



Claquage microonde par retournement temporel

Valentin Mazières

► To cite this version:

Valentin Mazières. Claquage microonde par retournement temporel. Electromagnétisme. Université Paul Sabatier - Toulouse III, 2020. Français. NNT : 2020TOU30230 . tel-03208800

HAL Id: tel-03208800

<https://tel.archives-ouvertes.fr/tel-03208800>

Submitted on 26 Apr 2021

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

En vue de l'obtention du

DOCTORAT DE L'UNIVERSITÉ DE TOULOUSE

Délivré par : *l'Université Toulouse 3 Paul Sabatier (UT3 Paul Sabatier)*

Présentée et soutenue le *Date de défense (27/11/2020)* par :

Valentin Mazières

Claquage Microonde par Retournement temporel

JURY

JULIEN DE ROSNY	Directeur de Recherche	Rapporteur
JEAN-PAUL BOOTH	Directeur de Recherche	Rapporteur
ÉLODIE RICHALOT	Professeur	Examinatrice
KREMENA MAKASHEVA	Directrice de Recherche	Examinatrice
LAURENT LIARD	Maître de conférence	Examineur
PHILIPPE POULIGUEN	Responsable du domaine des ondes acoustiques et radioélectriques DGA/AID	Examineur
OLIVIER PASCAL	Professeur	Directeur de thèse
ROMAIN PASCAUD	Professeur Associé	Co-directeur de thèse
CHRISTOPHE CALOZ	Professeur	Invité

École doctorale et spécialité :

GEET : Électromagnétisme et Systèmes Haute Fréquence

Unité de Recherche :

LAPLACE GRE - ISAE DEOS

Directeur(s) de Thèse :

Olivier Pascal et Romain Pascaud

Rapporteurs :

Julien De Rosny et Jean-Paul Booth

[...] science is always tentative, expecting that modifications in its present theories will sooner or later be found necessary, and aware that its method is one which is logically incapable of arriving at a complete and final demonstration. But in an advanced science the changes needed are generally only such as serve to give slightly greater accuracy ; the old theories remain serviceable where only rough approximations are concerned, but are found to fail when some new minuteness of observation becomes possible. Moreover, the technical inventions suggested by the old theories remain as evidence that they had a kind of practical truth up to a point. Science thus encourages abandonment of the search for absolute truth, and the substitution of what may be called “technical” truth, which belongs to any theory that can be successfully employed in inventions or in predicting the future. “Technical” truth is a matter of degree : a theory from which more successful inventions and predictions spring is truer than one which gives rise to fewer. “Knowledge” ceases to be a mental mirror of the universe, and becomes merely a practical tool in the manipulation of matter.— (Bertrand Russel - 1935)

Remerciements

Cette thèse a été effectuée conjointement au sein des groupes de recherche GRE (Groupe de Recherche en Électromagnétisme), GREPHE (Groupe de Recherche en Énergétique, Plasma Hors Équilibre) et MPP (Matériaux et Procédés Plasmas) du laboratoire LAPLACE (Laboratoire PLAsma et Conversion d'Énergie) de l'Université Paul Sabatier et du Département Électronique, Optro-nique et Signal (DEOS) de l'Institut Supérieur de l'Aéronautique et de l'Espace (ISAE-Supaéro). Ces travaux ont été co-financés par l'Université Paul Sabatier et la Direction Générale de l'Armement DGA/AID.

Je tiens tout d'abord à remercier Julien De Rosny, Directeur de Recherche à l'Institut Langevin, pour avoir accepté de rapporter cette thèse. Vos travaux autour du retournement temporel m'ont accompagné tout au long de ces trois années de thèse et ce fut un privilège d'avoir votre regard d'expert sur mes travaux.

Je tiens ensuite à remercier Jean-Paul Booth, Directeur de Recherche au LPP (Laboratoire de Physique des Plasmas), pour avoir accepté de rapporter cette thèse. Je vous remercie pour vos remarques et questions pertinentes lors de la soutenance.

Je remercie chaleureusement Kremena Makasheva, Directrice de recherche au laboratoire Laplace, pour avoir accepté de présider le jury de cette thèse.

Je remercie ensuite Elodie Richalot, Professeur à l'Université Paris-Est (ESYCOM), pour avoir accepté de faire partie du jury de cette thèse. Je vous remercie encore chaleureusement pour l'enthousiasme dont vous avez fait preuve ainsi que pour le regard intéressé que vous avez porté sur ces travaux.

Je remercie aussi Philippe Pouliguen, Responsable du domaine des ondes acoustiques et radioélectriques DGA/AID, pour l'intérêt manifeste qu'il a porté sur ces travaux durant ces trois années de thèse.

Le présent manuscrit rassemble les résultats obtenus durant trois années passées au laboratoire. Trois années durant lesquelles de nombreuses personnes ont contribué, plus ou moins directement, à l'aboutissement des travaux entrepris.

Je tiens tout d'abord à remercier les personnes non Toulousaines avec qui j'ai eu la chance de collaborer durant ces trois années de thèse.

Je remercie tout d'abord Luc Stafford, Professeur à l'Université de Montréal, pour l'intérêt qu'il a porté à mes travaux ainsi que la pédagogie dont il a fait preuve lors de nos différentes interactions.

Je remercie aussi Ali Al Ibrahim, alors doctorant à l'institut Pascal à Clermont-Ferrand lors de nos échanges et Pierre Bonnet, Professeur à l'Université Clermont Auvergne. Je suis vraiment fier de l'article que nous avons communément produit, fruit de nombreux échanges, notamment avec Ali. Je souhaite aussi te remercier Pierre pour l'enthousiasme dont tu as toujours fait preuve lors de nos échanges autour de ces travaux.

Je souhaite aussi remercier chaleureusement le CEA-Gramat pour le prêt d'un amplificateur 2 kW, indispensable à l'obtention de plasma sur notre dispositif expérimental. Les

résultats expérimentaux d'amorçage de plasmas présentés dans cette thèse n'auraient jamais vu le jour sans cet amplificateur. Un grand merci à Pierre Bruguière et à Jean Christophe Joly pour avoir suivi ces travaux, ainsi qu'à Jean-Luc Lavergne et Clovis Pouant, pour nous avoir apporté cet amplificateur au laboratoire et aidé à le mettre en place.

D'autres personnes, plus éloignées du contexte de ma thèse, ont aussi eu une influence significative sur le cheminement intellectuel que j'ai emprunté durant ces trois années.

Je pense tout d'abord à Jean-Pierre Boeuf, Directeur de Recherche au laboratoire Laplace, avec qui j'ai eu la chance d'interagir et qui a permis l'élaboration du code de simulation présenté dans ce ce manuscrit. Sans lui ce code n'aurait probablement pas vu le jour, alors je souhaite le remercier chaleureusement pour tout le temps qu'il a passé avec moi.

Une autre personne a aussi joué un rôle un clé à un moment où je me posais des questions autour du concept de chaotité : Florian Monsef, Maître de Conférences à l'Université Paris Sud. Les échanges que nous avons eus, tout d'abord lors du GDR Ondes à Gif-sur-Yvette et ensuite par mail, ont une une grande influence sur ma compréhension de ce concept de chaotité. Les résultats obtenus avec le code de simulation présentés dans ce manuscrit n'auraient probablement pas vu le jour sans ces explications précieuses. Un grand merci Florian pour l'intérêt que tu as porté à mes travaux, pour le temps que tu as pris pour discuter avec moi et pour avoir rendu ce concept de chaotité moins mystérieux.

Je souhaite ensuite remercier Christophe Caloz, Professeur à l'Université de Leuven, avec qui j'ai l'immense honneur de discuter depuis plus d'un an. Il est clair que les discussions que l'on a eues, notamment autour du fameux paradoxe thermodynamique, ont eu une énorme influence sur ma compréhension et ma vision de la physique. Je te remercie aussi énormément pour avoir accepté de faire partie de mon jury de thèse.

Je souhaite aussi remercier mes collègues du LAPLACE, que j'ai vraiment apprécié côtoyer durant ces années passées au laboratoire. Merci Abdel, Julien, Simo, Aurélien, Rafael, Jérôme, Thierry, Béatrice, Hugo... pour avoir rendu ces années si sympathiques.

Des immenses remerciements vont ensuite aux personnes qui m'ont encadrées durant ces trois années.

Je remercie d'abord profondément Laurent Liard et Simon Dap, qui ont d'abord été mes encadrants de stage de fin d'étude. Merci de m'avoir accueilli dans le laboratoire et merci pour votre soutien et votre aide durant ces années, sans lesquels les travaux de cette thèse n'auraient pas pu voir le jour. Je remercie aussi Richard Clergereaux pour son soutien et ses conseils toujours précieux. Un grand merci à tous les trois pour tout ce que vous m'avez appris et notamment pour m'avoir initié à la physique des plasmas.

Je tiens ensuite à présent à remercier profondément mon co-directeur de thèse, Romain Pascaud, pour son soutien infaillible ainsi que pour la bienveillance dont il fait preuve envers ses doctorants. Je te remercie sincèrement d'avoir accepté de faire partie de l'encadrement de cette thèse. Merci à la fois pour ton exigence scientifique ainsi que pour tes qualités humaines indéniables. J'ai beaucoup appris durant ces trois années, et ce fut un véritable plaisir de travailler à tes côtés, donc tout simplement merci pour tout !

Il est temps à présent pour moi de remercier Olivier Pascal, qui avant de devenir mon directeur de thèse, a d'abord été mon professeur d'électromagnétisme. Ce sont d'ailleurs tes

cours qui sont à l'origine de mon parcours scientifique, puisqu'ils ont éveillé en moi une envie de poursuivre dans ce domaine, tant ils m'ont marqué par leur justesse et leur rigueur et tant la passion qui en transparaissait était contagieuse. Il est clair que je ne serais pas là où je suis aujourd'hui sans avoir en premier lieu assisté à ces cours et sans avoir eu ensuite la chance de travailler à tes côtés. Tout ceci a joué un rôle fondamental dans la construction de la personne que je suis aujourd'hui, tant j'ai appris sur les plans scientifique et humain. Merci pour tout, pour la rigueur scientifique dont tu fais preuve, pour la vision que tu as de la recherche, pour ta curiosité contagieuse pour la science... Ce fût un honneur et un plaisir de travailler à tes côtés durant ces trois années.

Je pense enfin à mes amis et ma famille, notamment à mes parents, ma sœur, mon frère et mes super grands-parents (mamie Simone et papi Claude) que je remercie pour pour tout et sans qui je n'aurais pas pu arriver là où j'en suis aujourd'hui. Je remercie aussi bien sûr le superbe Thibault Douminge pour son soutien matériel, grâce auquel j'ai pu rédiger le présent manuscrit.

Je finirai par remercier ma chère et tendre Julie, qui m'a accompagné tout au long de ces trois années, les rendant si précieuses...

Captation vidéo de la soutenance

Le lecteur intéressé pourra visionner une captation vidéo de la soutenance de thèse, qui s'est déroulée le 27/11/2020 au laboratoire LAPLACE, en cliquant le lien suivant :

<https://youtu.be/XbaMGB55aDA>

Claquage microonde par Retournement Temporel



Présenté par	Jury	Date
Valentin Mazières	Julien de Rosny Rapporteur	27/11/2020
	Jean Paul Booth Rapporteur	
	Elodie Richalot Examinatrice	
	Kremena Makasheva Examinatrice	
	Philippe Pouliguen Examineur	
	Laurent Liard Examineur	
	Olivier Pascal Directeur de thèse	
	Romain Pascaud Co-directeur de thèse	
	Christophe Caloz Invité	



À toi, Théo

Table des matières

Introduction	1
Une brève histoire du plasma	1
Objectif de la thèse : La source plasma à retournement temporel	3
Organisation du manuscrit	4
1 Les plasmas microondes en cavité : Contexte et enjeux	7
1.1 Physique des plasmas froids hors équilibre d'argon	9
1.1.1 Échelle microscopique : Description des interactions	11
1.1.2 Échelle mésoscopique : Quasi-neutralité du plasma	17
1.1.3 Échelle macroscopique : Modélisation fluide	19
1.1.4 Échelle microscopique → Échelle macroscopique : Petit récapitulatif .	27
1.2 Les plasmas microondes en cavité : Limitations actuelles	29
1.2.1 Principe des décharges microondes en cavité	31
1.2.2 Les étapes de conception d'une cavité standard : Exemple	34
1.2.3 Limitations inhérentes aux méthodes standards de conception	36
1.3 La source plasma à retournement temporel : Solution ultime ?	39
1.3.1 Le retournement temporel : Chronologie des travaux fondateurs	40
1.3.2 La source plasma à retournement temporel : Les caractéristiques . . .	43
1.3.3 La source plasma à retournement temporel : Les avantages	46
2 La théorie autour du retournement temporel	49
2.1 Les origines du retournement temporel : La flèche du temps	51
2.1.1 Le temps, source d'inspiration et de fascination	51
2.1.2 La flèche du temps	53
2.1.3 Les prémisses du retournement temporel : Le paradoxe de réversibilité	55

2.1.4	Le principe des démons de Loschmidt appliqué aux ondes	58
2.2	Le retournement temporel en électromagnétisme	62
2.2.1	La causalité en électromagnétisme	62
2.2.2	La réversibilité en électromagnétisme	66
2.2.3	La réciprocité en électromagnétisme	71
2.2.4	Exemples	75
2.3	Le retournement temporel en pratique	79
2.3.1	Le retournement temporel en pratique : Réversible, réciproque, les deux ?	79
2.3.2	Le retournement temporel appliqué au contrôle plasma	82
3	Le retournement temporel en cavité réverbérante appliqué au contrôle plasma : Théorie et modélisation FDTD	87
3.1	Le retournement temporel mono-voie en cavité réverbérante	90
3.1.1	Principe du RT mono-voie en cavité réverbérante	90
3.1.2	Les cavités réverbérantes	95
3.1.3	La focalisation temporelle	100
3.1.4	La focalisation spatiale	111
3.1.5	Synthèse	114
3.2	Modélisation du retournement temporel en cavité réverbérante : Modèle FDTD 2D	116
3.2.1	Présentation du modèle FDTD 2D	116
3.2.2	Équivalence simulation - expérience	118
3.2.3	Étude des propriétés spatio-temporelles du retournement temporel simulé	120
3.3	Modélisation du retournement temporel en cavité pour l'amorçage de plasmas : Couplage Maxwell-fluide	125
3.3.1	Principe	125
3.3.2	Description du modèle fluide FDTD 2D	127
3.3.3	Équivalence simulation - expérience	128

3.3.4	Plasmas amorcés par retournement temporel en simulation	130
4	Les plasmas amorcés par retournement temporel : Résultats expérimentaux	141
4.1	Le retournement temporel expérimental : Dispositif et méthode	144
4.1.1	Dispositif expérimental	144
4.1.2	Le retournement temporel : Premier essai dans le domaine temporel .	148
4.1.3	Le retournement temporel : Développement dans le domaine fréquentiel	155
4.2	Résultats : Premiers plasmas amorcés par retournement temporel	161
4.2.1	Le dispositif expérimental destiné à l'amorçage des plasmas par RT .	161
4.2.2	Les propriétés du retournement temporel utilisé pour l'amorçage de plasmas	166
4.2.3	Le contrôle des plasmas par retournement temporel	168
4.3	Caractérisation des plasmas amorcés par retournement temporel	175
4.3.1	Dispositif expérimental	175
4.3.2	Caractérisation électrique	177
4.3.3	Caractérisation par imagerie rapide	182
4.3.4	La physique des plasmas amorcés par retournement temporel	187
	Conclusion et perspectives	193
	Conclusion	193
	Perspectives	195
	I - La physique des plasmas amorcés par retournement temporel : Études complémentaires	195
	II - Vers un pinceau plasma ondulatoire	198
	Publications	201
A	Détails sur la modélisation FDTD 2D	203
A.1	Discretisation des équations de Maxwell	203

A.2	Nombre de mode TE en cavité 2D	205
A.3	Équations du modèle plasma fluide	206
B	Comportement des antennes en émission et en réception	207
	Bibliographie	209

Introduction

- “Vous vous rendez compte, que le **temps** que la **lumière** met à atteindre notre œil, on voit que des choses passées.
- Hein ?
- Beh ouais si ça se trouve, je regarde l’image de ce **feu** alors qu’il est déjà éteint” [Ast05].

Ce qu’essaie maladroitement de dire le personnage de Perceval dans cet épisode de la série télévisée (humoristique) Kaamelott, c’est que la lumière émise par le feu qu’il est en train de regarder met un certain temps à atteindre son œil. La vision qu’il en a à un certain moment correspond donc en réalité à la vie passée de ce feu. Mais quel est le rapport avec les travaux de ce présent manuscrit me direz vous ? Dans cet exemple, le personnage de Perceval évoque en fait les deux domaines de la physique qui se rencontrent dans cette thèse : les ondes électromagnétiques et les plasmas. Il met en plus le doigt sur un concept qui est au cœur de l’originalité de ces travaux : le temps. En effet, l’objectif de cette thèse est de développer une source plasma microonde innovante et originale, basée sur le concept de “retournement temporel”.

Une brève histoire du plasma

Avant de discuter plus en détails des applications potentielles de cette source, il nous faut introduire le concept de plasma. Ce dernier est communément appelé le quatrième état de la matière. Il peut être grossièrement défini comme un gaz partiellement ionisé, *i.e.* un gaz auquel une énergie (électrique ou thermique) appliquée arrache les électrons des atomes ou des molécules. Ainsi, on retrouve dans un plasma des particules neutres (les atomes et molécules du gaz) et des particules chargées (des électrons et des ions). En fait, nous voyons tous du plasma tous les jours, puisque le soleil est un plasma. Et même par temps orageux, il nous est possible de voir et même d’entendre des plasmas avec la foudre. Et de tout temps, le plasma a fasciné l’espèce humaine. On le retrouve dans de nombreux contes et mythologies. Par exemple dans la mythologie Grecque le roi de l’Olympe Zeus est le dieu de la foudre, tout comme le dieu Thor dans la mythologie nordique. Dans des mythes plus contemporains, on retrouve dans l’univers Marvel la super-héroïne Storm capable aussi de contrôler la foudre, et on retrouve le plasma dans le Kamehameha de Dragon Ball ou dans le Rasengan de Naruto. Dans toutes ces représentations, le plasma est vu comme quelque chose de mystique et la personne qui arrive à le maîtriser possède alors des pouvoirs surnaturels. Cependant, le plasma n’est pas seulement associé à des pouvoirs surnaturels, il peut aussi être le fruit de technologies avancées. Par exemple, chez Marvel c’est Iron Man qui, grâce à ses connaissances scientifiques, construit une armure capable d’émettre des “rayons répulseurs”. Et ce mythe n’est pas très éloigné de la réalité. En effet, comme nous allons le voir le plasma n’est pas seulement un sujet de fascination pour les conteurs, il est aussi au cœur de nombreux dispositifs scientifiques.

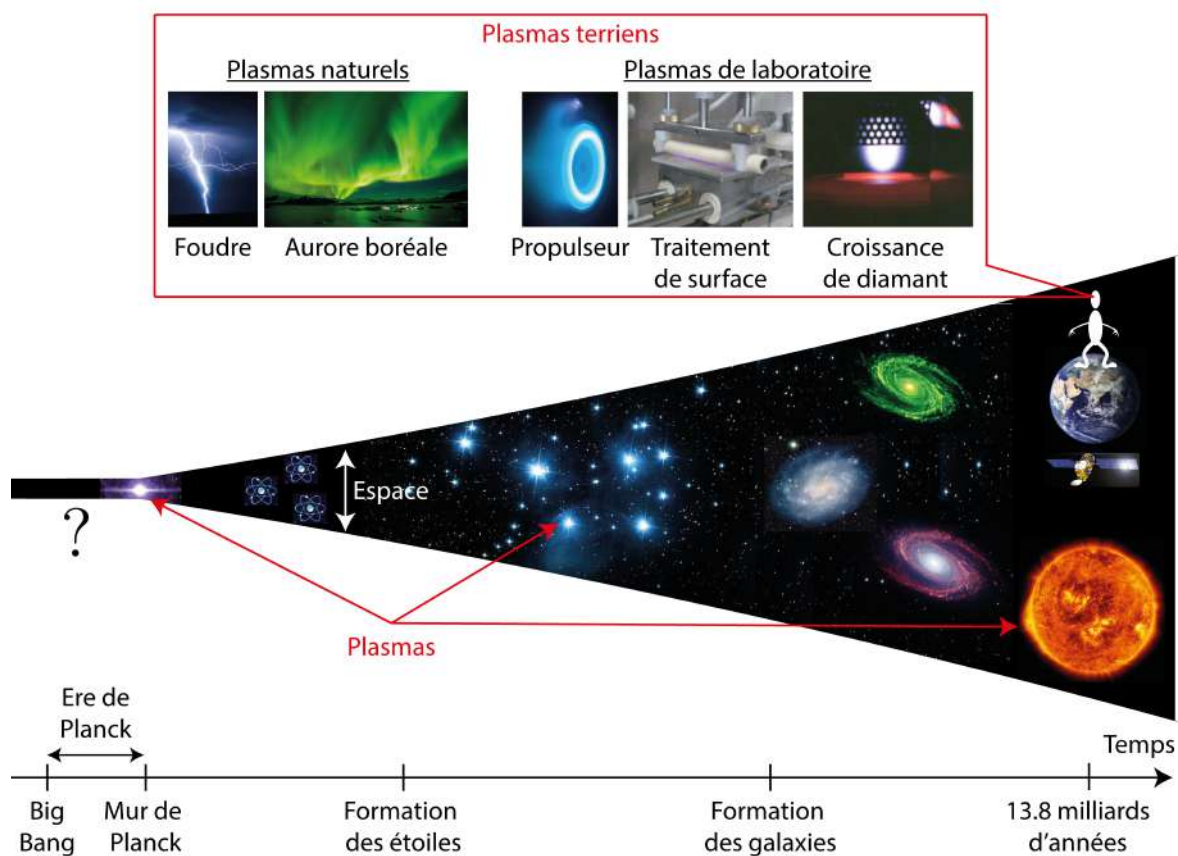


FIGURE I.1 – L’Histoire de l’univers selon les modèles cosmologiques standards actuels [Lid03] (les axes de temps et d’espace ne sont pas à l’échelle). Le plasma occupe une place importante dans les mécanismes de l’univers (les illustrations des plasmas de laboratoires sont reprises de [Boe16] et les autres, crédits images : NASA, NASA-JPL, ESO, CNES).

Les plasmas occupent une place importante dans la science en général et notamment en cosmologie. Les modèles cosmologiques standards actuels nous disent que l’univers est en expansion [Lid03]. De ce fait, en “remontant le temps”, on voit l’univers se contracter de plus en plus, jusqu’à devenir énormément dense et localisé. L’histoire de l’univers selon ces modèles est schématisée sur la Figure I.1. Ils sont capables de décrire ce qu’il s’est passé jusqu’à un instant nommé le “mur de Planck”. Pour des temps inférieurs, les théories de la relativité générale et de la mécanique quantique donnent des résultats contradictoires. Au moment du mur de Planck, l’univers est contracté dans un état très dense et très chaud : il est en fait un plasma. De ce plasma primordial, du fait de l’expansion de l’univers, sont ensuite nés les atomes et les molécules, puis les étoiles (qui sont des plasmas) et enfin les galaxies. Et c’est ainsi que 13.8 milliards d’années après le Big Bang¹, les conditions sur une

1. Le Big Bang correspond à la singularité gravitationnelle que l’on trouve en appliquant les équations de la relativité générale à l’univers et en remontant son histoire. Cependant, avant le mur de Planck (durant l’ère de Planck), l’univers est tellement dense et énergétique qu’il faut aussi prendre en compte les effets quantiques. Ces deux théories donnant des résultats incompatibles, il n’est ainsi pas possible, à l’heure actuelle, de conclure sur ce qui précède le mur de Planck.

certaine planète (nommée Terre) tournant autour d'une certaine étoile (nommée Soleil) dans un certain système planétaire (nommé système solaire) dans une certaine galaxie (nommée Voie lactée) sont réunies pour que la vie telle que nous la connaissons puisse apparaître. Et il est intéressant de noter que dans la matière de l'univers observable, jusqu'à 99 %, selon certaines estimations, se trouve être dans l'état plasma [GB05]. Sur Terre, à cause des conditions de température et de pression, ils sont moins présents naturellement mais on peut en apercevoir parfois lors d'orages comme nous l'avons déjà mentionné ou lors de phénomènes d'aurore polaire. On les retrouve cependant abondamment dans l'industrie et dans le domaine de la recherche scientifique. En effet, depuis les premières études en laboratoire sur les plasmas, de nombreuses applications ont émergé pour ces milieux aux propriétés physiques exotiques. Quelques illustrations d'applications sont proposées sur la Figure I.1. Les plasmas peuvent par exemple être utilisés dans l'espace pour propulser les satellites en accélérant les particules chargées du plasma sous l'effet d'un champ électrique [GK08]. Leurs propriétés lumineuses sont aussi exploitées dans certaines applications, telles que les éclairages néon ou les écrans plasmas. Enfin, de nombreuses applications exploitent les propriétés intéressantes de leur interaction avec des matériaux. Ils sont ainsi largement utilisés en traitement de surface [FM93]; [Gri94]; [LL05] pour changer les propriétés de la surface avec laquelle ils sont en contact. Ils sont aussi utilisés pour déposer des couches minces sur un matériau ou pour de la croissance de diamant [DW98]; [Su+14].

Objectif de la thèse : la source plasma à retournement temporel

Pour les applications où il est question d'interaction avec un matériau, il existe plusieurs technologies basées sur différents types de plasma. Les plus couramment utilisées sont celles basées sur des plasmas à couplage capacitif (CCP) ou inductif (ICP), les plasmas à résonance cyclotron électronique (ECR), ou les plasmas microondes [Sam+12]. Les décharges microondes sont étudiées de façon intensive car elles produisent des densités de plasma élevées à la fois dans des conditions de haute et de basse pression. Nous nous concentrerons dans cette thèse sur ces décharges. Elles sont généralement générées par des techniques dites modales. C'est à dire qu'elles exploitent le phénomène de résonance d'une cavité métallique soumise à un champ électromagnétique oscillant à une fréquence microonde. Nous verrons dans le chapitre 1 les limitations de ces techniques standards. La plus restrictive est la limite sur les dimensions des objets qu'il est possible de traiter. **La source plasma proposée dans cette thèse a pour objectif de surpasser ces techniques habituelles, en rendant possible le contrôle dynamique de la position des plasmas dans des cavités de grandes dimensions. Ainsi, les dimensions des objets à traiter ne sont plus limitées.**

Essayons à présent d'introduire rapidement le principe de la source plasma à retournement temporel (RT) développée dans cette thèse. Pour cela, commençons par le commencement. Pas celui de la thèse, celui de l'univers dont nous avons discuté précédemment, au sens du mur de Planck. Si on résume ce que nous disent les modèles cosmologiques actuels, en remontant le temps l'univers se contracte et devient localisé autour d'un point pour former un plasma. Ce que nous proposons dans cette thèse a de plus modestes ambitions : faire

“remonter le temps” aux ondes de façon à focaliser l’énergie électromagnétique dans le but d’amorcer un plasma local. Ici l’expression “remonter le temps” est entre guillemets parce qu’en pratique on ne remonte évidemment pas le temps, mais l’objectif est de faire revivre aux ondes leur vie passée de sorte à les focaliser à l’endroit initial de leur émission. Cette technique permettant de focaliser les ondes est appelée retournement temporel (RT). Elle a été conceptualisée dans les années 1990 par Mathias Fink et ses collègues. Une description des travaux fondateurs est proposée dans le chapitre 1 et le second chapitre se concentre sur la théorie autour du RT dans le domaine de l’électromagnétisme. Depuis sa découverte, le RT a trouvé de nombreuses applications dans divers domaines de la physique et notamment dans le domaine des microondes. Dans cette thèse, nous en proposons une nouvelle : le contrôle dynamique de la position de plasmas en cavité. **Tout au long de ce manuscrit, nous verrons en quoi cette source plasma à RT est en rupture avec les dispositifs actuels et pourquoi elle nous semble prometteuse pour les applications de traitement de surface notamment.**

L’objectif de cette thèse est de contribuer au développement de cette source plasma innovante. Nous reportons dans cette thèse la première étude sur ce genre de source. La compréhension des enjeux autour des sources plasmas microondes nécessite des développements relativement conséquents. Ces discussions font l’objet du chapitre 1. Les défis à relever que nous avons identifiés pour le développement de cette nouvelle source plasma sont listés à la fin de ce chapitre 1.

Organisation du manuscrit

Ce manuscrit est divisé en quatre chapitres, sommairement décrits ci-dessous :

Le premier chapitre présente le contexte et les enjeux autour des sources plasmas microondes en cavité. Tout d’abord, nous introduisons des notions théoriques essentielles à la compréhension de la physique des plasmas froids hors équilibre d’argon. Nous commençons par décrire les interactions au niveau microscopique avant de présenter les équations du modèle fluide permettant de décrire le comportement du plasma à l’échelle macroscopique. Ensuite, nous discutons des techniques standards de génération de plasmas microondes en cavité. Nous axons la discussion autour des limitations auxquelles elles sont soumises. Nous introduisons enfin le concept de la source plasma développée durant cette thèse. Nous expliquons en quoi cette nouvelle source permet un contrôle des plasmas en grandes cavités, là où c’était impossible par les techniques habituelles. Nous finissons ce chapitre en énonçant les défis à relever pour le développement de cette source innovante.

Le second chapitre se concentre sur la physique du RT dans le domaine de l’électromagnétisme. Afin d’appréhender les concepts théoriques présentés dans ce chapitre, une discussion autour du concept de temps est d’abord proposée. Nous présentons ensuite les discussions qui ont eu lieu entre Boltzmann et Loschmidt, à partir desquelles le RT a été pensé. Ensuite, nous discutons des concepts de causalité, de réversibilité et de réciprocité dans le domaine de l’électromagnétisme. Nous voyons que le RT ne viole pas la causalité, qu’il est basé sur

la réversibilité et, qu'en pratique, il est plus proche d'une expérience de réciprocité. Enfin nous discuterons du RT appliqué au contrôle plasma. Le point central est l'introduction d'un phénomène non-linéaire (claquage plasma) dans le processus de RT, qui lui est basé sur une physique linéaire.

Le troisième chapitre fait la transition entre le chapitre 2, très théorique, et le chapitre 4 présentant les résultats expérimentaux. Un outil se prête bien à cet objectif : la simulation. Dans un premier temps, les équations permettant de décrire le RT en cavité réverbérante selon son régime de recouvrement modal sont présentées. Nous voyons que le RT en cavité réverbérante présente des lobes secondaires temporels et spatiaux. Une propriété importante permet toutefois de réduire le niveau de ces lobes : la chaoticité de la cavité. Un code de simulation basé sur la méthode des différences finies dans le domaine temporel (ou FDTD pour Finite-difference time-domain [TH05]) est ensuite développé. Il est d'abord utilisé en codant uniquement les équations de Maxwell pour illustrer l'influence de la chaoticité de la cavité sur la qualité de la focalisation par RT. Le caractère essentiel de la chaoticité sur le contrôle des plasmas par RT est montré dans un second temps en couplant les équations de Maxwell à un modèle plasma fluide. Nous finissons ce chapitre en utilisant le code de simulation pour discuter des conditions d'amorçage des plasmas expérimentaux (régime pulsé + présence d'initiateurs).

Le quatrième chapitre présente les résultats expérimentaux obtenus durant cette thèse. Les premiers résultats de RT obtenus sur le dispositif expérimental avec les méthodes temporelles classiquement utilisées sont tout d'abord présentés. Ensuite, la méthode fréquentielle développée durant cette thèse est présentée et utilisée pour optimiser le RT sur le dispositif. De cette façon, le RT obtenu permet l'amorçage et le contrôle de la position des plasmas dans la cavité. Ces résultats constituent les premiers plasmas amorcés et contrôlés par RT. Nous finissons ensuite le chapitre par une caractérisation de ces plasmas amorcés par RT. Tout d'abord nous discutons de l'influence de la présence d'un plasma dans la cavité sur le comportement des ondes. Ensuite, à l'aide de mesures prises avec une caméra rapide, nous décrivons les mécanismes présents dans ces décharges plasmas. Une discussion décrivant la physique de ces plasmas est enfin proposée, à l'issue de laquelle l'importance de l'effet mémoire entre deux impulsions successives est identifiée.

Les plasmas microondes en cavité :

Contexte et enjeux

Sommaire

1.1	Physique des plasmas froids hors équilibre d'argon	9
1.1.1	Échelle microscopique : Description des interactions	11
1.1.2	Échelle mésoscopique : Quasi-neutralité du plasma	17
1.1.3	Échelle macroscopique : Modélisation fluide	19
1.1.4	Échelle microscopique → Échelle macroscopique : Petit récapitulatif . . .	27
1.2	Les plasmas microondes en cavité : Limitations actuelles	29
1.2.1	Principe des décharges microondes en cavité	31
1.2.2	Les étapes de conception d'une cavité standard : Exemple	34
1.2.3	Limitations inhérentes aux méthodes standards de conception	36
1.3	La source plasma à retournement temporel : Solution ultime?	39
1.3.1	Le retournement temporel : Chronologie des travaux fondateurs	40
1.3.2	La source plasma à retournement temporel : Les caractéristiques	43
1.3.3	La source plasma à retournement temporel : Les avantages	46

I would rather have questions that can't be answered than answers that can't be questioned.

Richard Feynman

L'objectif de ce chapitre est de présenter le principe de la source plasma microonde développée durant cette thèse. Elle repose sur le RT et contrairement aux sources plasmas existantes, elle rend possible un contrôle des plasmas dans les grandes cavités, contrôle quasi-indépendant de leur conception. Le développement expérimental de cette source innovante a fait l'objet de cette thèse. Les plasmas alors amorcés par RT en cavité réverbérante sont des plasmas froids hors équilibre d'argon. Nous nous concentrerons dans ce chapitre sur la description de la physique et des techniques usuellement utilisées pour ce genre de plasma. Cela nous permettra d'illustrer les avantages offerts par cette nouvelle source plasma à RT.

Dans une première partie, la physique générale des plasmas froids hors équilibre d'argon est présentée. Les plasmas sont des milieux dans lesquels il existe des interactions entre des

particules chargées et les neutres. Cela rend la description de ces milieux complexe, nécessitant une approche multi-physique et ce sur différentes échelles d'analyse. En effet, les champs électromagnétiques engendrent les mouvements des particules chargées, qui s'organisent alors, changeant ainsi la structure des champs électromagnétiques, ces derniers engendrent le mouvement de particules chargées et ainsi de suite... Afin de comprendre ces interactions complexes, nous commencerons par décrire les interactions entre les particules chargées au niveau microscopique. Nous ne considérerons pour notre plasma que les interactions prépondérantes, de type élastique et inélastique (excitation électronique ou ionisation). Ensuite, à une échelle mésoscopique, nous verrons pourquoi ce comportement des particules chargées se traduit par une quasi-neutralité du plasma. C'est à dire que le plasma est neutre au delà d'une certaine échelle et pour des phénomènes dynamiques d'une certaine durée. Nous verrons ensuite qu'une modélisation plasma fluide permet de décrire le plasma en manipulant des grandeurs macroscopiques. Ces grandeurs traduisent au niveau macroscopique les interactions qui ont lieu au niveau microscopique. Nous finirons par un petit récapitulatif retraçant la description du plasma de l'échelle microscopique à l'échelle macroscopique.

Dans une seconde partie, nous présenterons les limitations des techniques actuellement utilisées pour amorcer des plasmas microondes en cavité. Le principe des décharges plasmas microondes en cavité est d'abord rappelé. Nous verrons qu'il exploite l'intensification du champ électrique par résonance. Ensuite les différentes étapes nécessaires à la conception d'une source plasma microonde sont présentées sur un exemple. Nous verrons que cette conception commence par la conception de la cavité avec la prise en compte des éléments qu'elle doit contenir (fenêtres en quartz par exemple) de sorte à structurer le champ électromagnétique. Ensuite, la présence du plasma dans la cavité doit être prise en compte afin d'obtenir finalement la distribution spatiale de plasma souhaitée. Enfin, les différentes étapes de conception sont synthétisées afin d'illustrer les limitations des techniques basées sur ce principe. Nous verrons notamment que les principales limitations sont la forte dépendance de la distribution spatiale du plasma au design de la cavité ainsi que la restriction sur les dimensions du substrat à traiter.

Finalement dans une troisième partie, nous présenterons une alternative aux sources plasmas microondes standards : la source plasma à RT. Afin de faire sentir au lecteur le mécanisme du RT, nous présenterons tout d'abord rapidement la chronologie des travaux fondateurs du RT. Nous verrons qu'il permet la focalisation spatio-temporelle de l'énergie électromagnétique en cavité réverbérante. Ensuite nous nous concentrerons sur les caractéristiques de la source plasma à RT. Nous verrons qu'elle permet un contrôle temporel de la position des plasmas dans la cavité, grâce à un amorçage en régime transitoire. Pour que ces plasmas soient maintenus, il est nécessaire d'utiliser un régime pulsé. Nous discuterons rapidement de l'intérêt d'un tel régime, permettant notamment de restreindre le chauffage du gaz et des parois, qui sont des processus limitants. Finalement, nous présenterons la source plasma à RT. Nous verrons qu'elle est en rupture avec les sources plasmas microondes actuelles, en permettant le contrôle spatio-temporel des plasmas dans de grandes cavités, avec un contrôle quasi-indépendant du design de ces dernières.

1.1 Physique des plasmas froids hors équilibre d'argon

L'objectif de cette partie est de présenter la physique des plasmas étudiés dans cette thèse. Le domaine de la physique des plasmas est complexe et très vaste, et de nombreux plasmas existent avec des comportements complètement différents. La richesse de comportements des plasmas vient de la nature non-linéaire et auto-cohérente du couplage champs - plasmas [Rax05]. Cet effet est illustré sur la Figure 1.1. Les champs électromagnétiques ($\mathbf{E}(\mathbf{r}, t)$ et $\mathbf{B}(\mathbf{r}, t)$) engendrent des mouvements de particules chargées ($\mathbf{v}(\mathbf{r}, t)$) à travers la force de Lorentz, mouvements qui se traduisent par une réorganisation des charges présentes dans le plasma ($\rho(\mathbf{r}, t)$ et $\mathbf{J}(\mathbf{r}, t)$), qui sont les termes sources des équations de Maxwell. Cela engendre alors une modification des champs électromagnétiques, qui vont à leur tour engendrer des mouvements des particules chargées, et ainsi de suite... C'est donc l'interaction complexe entre ces mécanismes qui est à l'origine de la richesse des comportements des plasmas. Nous nous concentrerons donc dans cette partie sur la physique des plasmas permettant de décrire les plasmas amorcés par RT expérimentalement.

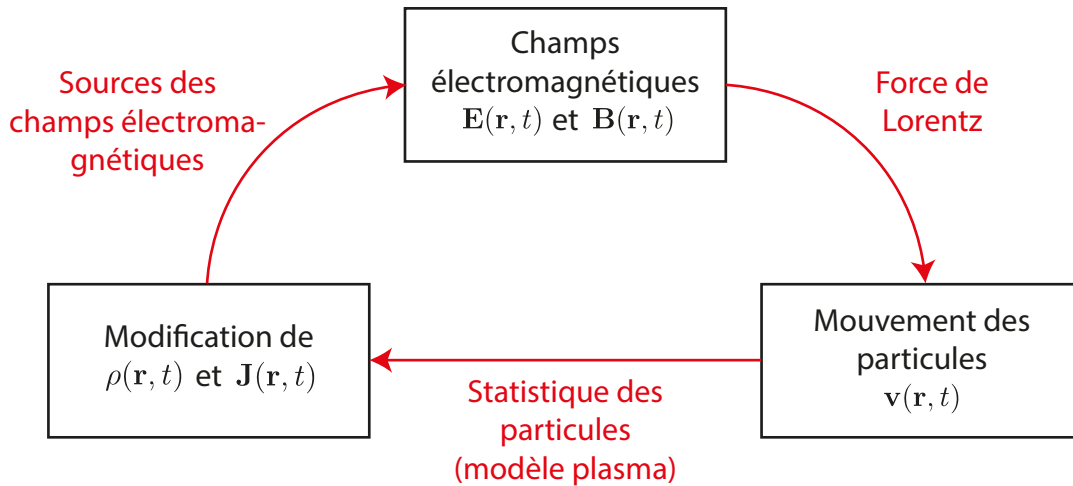


FIGURE 1.1 – Schéma représentant les liens entre les interactions électromagnétiques et le mouvement des particules dans un plasma (inspiré de [Rax05]).

Un milieu ionisé est un milieu constitué de neutres (de densité notée n_n) et d'espèces chargées, telles que les électrons (de densité notée n_e). Un milieu ionisé est un plasma quand la densité de particules chargées est suffisamment grande pour que le milieu reste macroscopiquement neutre (on parle en fait de quasi-neutralité du plasma, comme nous le verrons dans cette partie). On notera alors $n_e = n_i$ avec n_i la densité d'ions positifs. Comme nous l'avons dit, il existe différents types de plasma. Afin de classifier les plasmas, il est intéressant d'utiliser le degré d'ionisation, défini comme :

$$\delta_{ionisation} = \frac{n_e}{n_e + n_n} \quad (1.1)$$

Durant cette thèse, les plasmas amorcés par RT sont des plasmas froids, pour lesquels le

degré d'ionisation est typiquement $\delta_{ionisation} < 10^{-2}$. De plus, ces plasmas sont “hors équilibre thermodynamique”, *i.e.* les électrons possèdent une température beaucoup plus importante que celle des particules “lourdes”, ions et neutres. Par simplicité nous négligerons le mouvement de ces particules “lourdes”, et seul le mouvement des électrons est considéré. Dans les plasmas froids hors équilibre, les collisions entre particules chargées sont rares puisque le milieu est constitué essentiellement d'atomes ou de molécules neutres. Dans ces plasmas les collisions prépondérantes sont donc les collisions électron - neutre. Comme nous allons le voir, la physique de ces plasmas repose sur la description du type d'interaction électron - neutre. Pour finir, les plasmas amorcés par RT sont des plasmas d'argon. Nous nous concentrerons ainsi dans cette partie sur la description de la physique des **plasmas froids hors équilibre thermodynamique dans l'argon**.

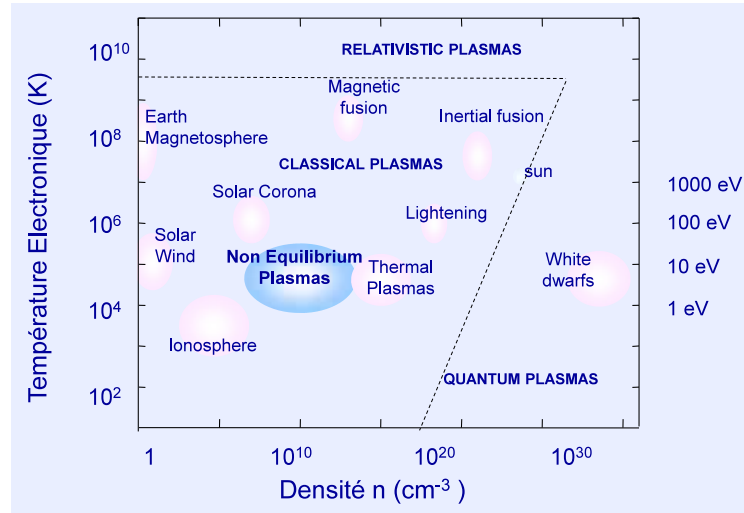


FIGURE 1.2 – Diagramme densité - température électronique (repris de [Boe16]).

Les différents types de plasmas sont rassemblés sur le diagramme de la Figure 1.2. Ils y sont classés en fonction de leur température électronique et de leur densité. Les plasmas qui nous intéressent dans cette thèse ont typiquement des températures électroniques de quelques eV ($1 \text{ eV} \leftrightarrow 11604 \text{ K}$) et des densités d'environ 10^{10} cm^{-3} .

Les plasmas sont des milieux constitués de particules chargées, qui interagissent entre elles au niveau microscopique. Cela donne à ces milieux certaines propriétés au niveau macroscopique. Afin de comprendre les mécanismes importants ayant lieu dans un plasma froid hors équilibre d'argon, nous commencerons par décrire les interactions entre les particules au niveau microscopique. Nous verrons que ces interactions peuvent être de nature très diverses, et dépendent fortement de la quantité d'énergie transmise entre les particules. Les niveaux d'énergie sont quantifiés. Nous verrons que cela se traduit par un comportement à seuil fortement non-linéaire. Nous passerons ensuite à un niveau d'analyse mésoscopique, nous permettant d'illustrer une propriété importante des plasmas : la quasi-neutralité. Nous verrons que cela se traduit par une différence de nature (propagation ou réflexion) de l'interaction onde - plasma en fonction de la fréquence de l'onde électromagnétique par rapport à la fréquence plasma. Pour finir, nous verrons par le biais d'un modèle fluide comment ces

interactions microscopiques se manifestent au niveau macroscopique. Nous verrons comment la modélisation permet de décrire le claquage du gaz ainsi que le comportement du plasma en présence d'ondes électromagnétiques.

1.1.1 Échelle microscopique : Description des interactions

La trajectoire d'un électron libre peut être modifiée sous l'action d'un champ électromagnétique. La force s'exerçant alors sur l'électron au point $\mathbf{r}$ au temps t peut s'écrire :

$$\mathbf{F}_L(\mathbf{r}, t) = q\mathbf{E}(\mathbf{r}, t) + q\mathbf{v} \wedge \mathbf{B}(\mathbf{r}, t) \quad (1.2)$$

avec $\mathbf{F}_L(\mathbf{r}, t)$ la force de Lorentz appliquée sur une particule, q la charge de la particule, $\mathbf{E}(\mathbf{r}, t)$ le champ électrique appliqué, $\mathbf{v}$ la vitesse de la particule et $\mathbf{B}(\mathbf{r}, t)$ le champ magnétique appliqué. Soumis à un champ électrique $\mathbf{E}$, les électrons, de charge $q = -e$, déjà présents initialement dans le gaz ou dans le plasma, seront alors mis en mouvement dans le sens opposé à la direction du champ électrique. En conséquence, ces électrons vont gagner de la quantité de mouvement et de l'énergie cinétique. Ils peuvent alors entrer en collision avec les molécules ou les atomes (neutres) du gaz. Un exemple de trajet d'un électron soumis à un champ électrique dans un gaz est représenté sur la Figure 1.3.

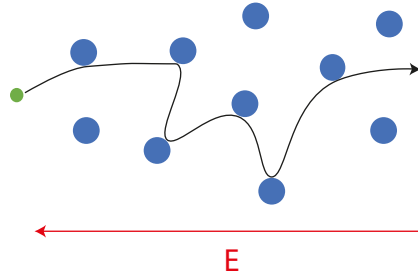


FIGURE 1.3 – Schéma représentant le trajet emprunté par un électron (en vert) soumis à un champ électrique dans un gaz (dont les atomes ou les molécules sont représentés en bleu).

On distingue alors la vitesse moyenne dirigée des électrons résultant des forces s'exerçant sur eux (ici $\mathbf{F}_L$), notée $\mathbf{u} = \overline{\mathbf{v}}$, de la vitesse d'agitation thermique résultant des collisions électron - particule, notée $\mathbf{v} = \mathbf{v} - \mathbf{u}$ et dont la moyenne du vecteur $\mathbf{v}$ est nulle : $\overline{\mathbf{v}} = \mathbf{0}$. On peut ensuite définir, à un instant donné t et un point donné $\mathbf{r}$, une énergie cinétique moyenne des électrons :

$$\mathcal{E}_e(\mathbf{r}, t) = \frac{1}{2} m_e \overline{v^2(\mathbf{r}, t)} = \frac{1}{2} m_e \overline{u^2(\mathbf{r}, t)} + \frac{1}{2} m_e \overline{v^2(\mathbf{r}, t)} \quad (1.3)$$

avec m_e la masse de l'électron en question et $v(\mathbf{r}, t) = |\mathbf{v}(\mathbf{r}, t)|$, $u(\mathbf{r}, t) = |\mathbf{u}(\mathbf{r}, t)|$ et $\mathcal{V}(\mathbf{r}, t) = |\mathbf{v}(\mathbf{r}, t)|$. L'énergie cinétique moyenne des électrons résulte donc du mouvement dirigé des électrons et de leur agitation thermique. L'énergie moyenne d'agitation des électrons est caractérisée par la température électronique $T_e(\mathbf{r}, t)$. En toute rigueur, elle caractérise un état

d'équilibre thermodynamique dans lequel la fonction de distribution des vitesses est Maxwellienne. Par abus de langage on parle également de température quand la distribution des vitesses n'est pas Maxwellienne et on la définit, pour chaque type de particule par :

$$\frac{3}{2}k_B T_e(\mathbf{r}, t) = \frac{1}{2}m_e \overline{v^2}(\mathbf{r}, t) \quad (1.4)$$

avec k_B la constante de Boltzmann.

Collision élastique	$e^- + X \rightarrow e^- + X$ 	$T_e \searrow$
Excitation électronique	$e^- + X \rightarrow e^- + X^* \rightarrow e^- + X + h\nu$ 	$T_e \searrow$ $I_l \nearrow$
Ionisation	$e^- + X \rightarrow X^+ + 2e^-$ 	$T_e \searrow$ $n_e \nearrow$

FIGURE 1.4 – Tableau synthétisant les différents types de collision électron - atome.

Selon la température T_e des électrons, il existe différents types de collisions entre un électron e^- et un atome, noté X . Les principaux types de collisions dans un gaz atomique (tel que l'argon) sont schématisés dans le tableau de la Figure 1.4.

- Lors d'une collision de type élastique, une partie de l'énergie de l'électron est cédée à l'atome sous la forme d'énergie cinétique.

- Lors d'une collision de type excitation électronique, l'électron participant à la collision cède de l'énergie à un électron de l'atome, qui passe alors d'un état d'énergie de niveau $\mathcal{E}_j$ à un niveau d'énergie plus élevé $\mathcal{E}_k$ que l'on appelle "état excité". L'atome ainsi excité est noté X^* . L'énergie transmise à l'électron de l'atome doit être égale à la différence d'énergie $\mathcal{E}_k - \mathcal{E}_j$. En se désexcitant au bout d'un certain temps, ce dernier émet un photon de longueur d'onde $\lambda = hc/(\mathcal{E}_k - \mathcal{E}_j)$, avec h la constante de Planck et c la vitesse de la lumière dans le vide. L'intensité lumineuse I_l augmente alors.

- Lors d'une collision de type ionisation (par impact électronique), l'énergie cédée à l'atome par l'électron incident est suffisante pour qu'un électron soit arraché de l'atome qui devient alors un ion positif noté X^+ . Cette opération permet donc la génération d'un électron (la densité n_e d'électrons augmente). Elle est essentielle dans le processus de génération d'un plasma.

Les états principaux de l'argon, de son état fondamental à ses états ionisés, sont représentés sur la Figure 1.5. Dans son état fondamental, la configuration électronique de l'argon est $1s^2 2s^2 2p^6 3s^2 3p^6$ (en notation de Paschen). Le niveau excité de plus faible énergie est alors

le $4s$, avec une énergie d'un peu plus de 11 eV. Il existe de nombreux niveaux d'énergies supérieures, correspondant aux différentes couches électroniques. Pour une différence d'énergie apportée à l'atome par l'électron de $\mathcal{E}_{ionisation} = \mathcal{E}_k - \mathcal{E}_j = 15.76$ eV l'électron est arraché à l'atome, il y a ionisation. La nature des interactions entre un électron et un atome est donc fortement dépendante de la valeur de l'énergie de l'électron incident. À chaque interaction est associée un niveau d'énergie bien défini. Les niveaux d'énergie possibles sont quantifiés. Nous verrons que ce comportement quantique des atomes au niveau microscopique est à l'origine du comportement à seuil du claquage d'un gaz par rapport au niveau du champ électrique appliqué, rendant ce phénomène fortement non-linéaire au niveau macroscopique.

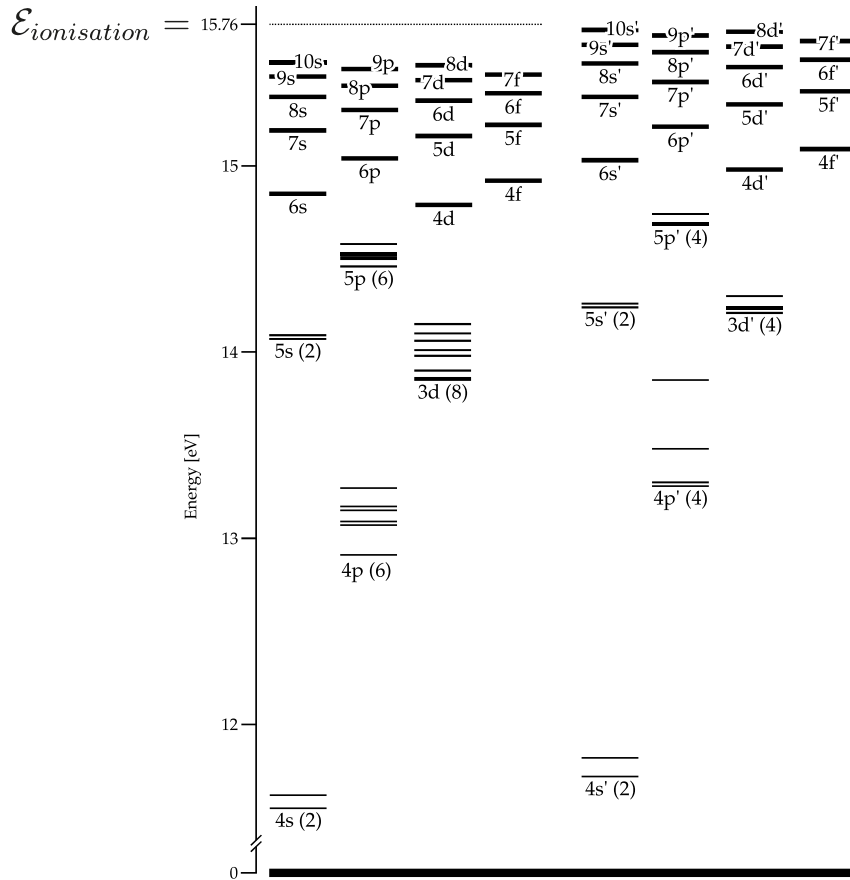


FIGURE 1.5 – Diagramme des états principaux de l'argon de son état fondamental à son état ionisé d'énergie $\mathcal{E}_{ionisation} = 15.76$ eV (repris de [Car13]).

Ainsi sous l'effet de champs électromagnétiques, différentes interactions électron - particule peuvent avoir lieu. Afin de décrire la probabilité qu'une collision ait lieu, on utilise le concept de section efficace, noté ici σ_{eff} . Pour comprendre ce concept, prenons par exemple un flux $\Gamma_F(x) = \Gamma_{F_0} = n_A v_A$ de particules de densité n_A et de vitesse v_A arrivant sur une cible au repos de surface S constituée de particules de densité n_B à la position x , comme représenté sur la Figure 1.6. La probabilité d'interaction $\mathbb{P}(interaction)$ entre les particules incidentes

et celles de la cible peut s'écrire [Alb15] :

$$\mathbb{P}(\text{interaction}) = \frac{S_{int}}{S} = \frac{(n_B S dx) \sigma_{eff}}{S} = n_B \sigma_{eff} dx \quad (1.5)$$

avec $n_B S dx$ le nombre de particules dans le volume $S dx$ et $S_{int} = (n_B S dx) \sigma_{eff}$ la surface d'interaction entre les particules cibles et les particules incidentes. La section efficace est donc homogène à une surface ($[m^2]$). Dans le cas simple où les particules sont représentées par des sphères élastiques dures et l'interaction entre les particules est une simple collision élastique, la section efficace devient alors égale à la surface de la sphère projetée sur la surface S et peut s'écrire simplement $\pi r_{sphère}^2$ avec $r_{sphère}$ le rayon de la sphère. Dans ce cas, la surface efficace correspond donc à une surface physique, la surface des sphères. Mais cette notion rend aussi compte des interactions plus complexes, de type inélastique (ionisation, excitation...). Cela permet de décrire avec un modèle classique ces interactions qui sont de nature quantique. Dans ce cas, la section efficace ne représente pas une surface physique réelle mais une surface fictive qu'auraient les particules dans une représentation classique. C'est pour cela que l'on parle de section **efficace**. La section efficace totale s'écrit alors comme la somme des sections efficaces des différentes interactions possibles, $\sigma_{eff}^{tot} = \sigma_{eff}^{el}(v_A) + \sigma_{eff}^i(v_A) + \sigma_{eff}^{ex}(v_A)$ avec $\sigma_{eff}^{el}(v_A)$ la section efficace des collisions de type élastique, $\sigma_{eff}^i(v_A)$ celle des collisions de type ionisation et $\sigma_{eff}^{ex}(v_A)$ celle des collisions de type excitation. On peut ajouter que les sections efficaces dépendent de la vitesse des particules incidentes.

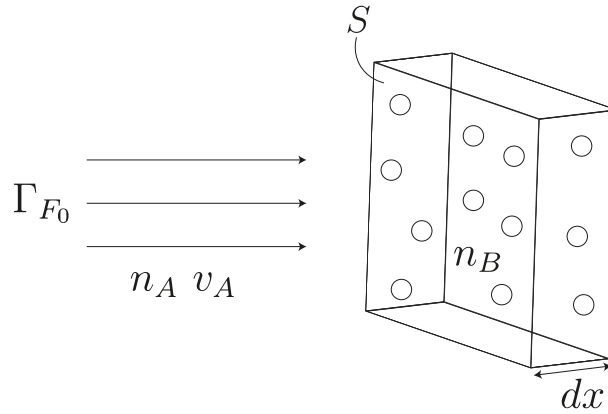


FIGURE 1.6 – Schéma représentant un flux de particules incidentes $\Gamma_F(x) = \Gamma_{F_0} = n_A v_A$ sur une cible au repos de surface S constituée de particules de densité n_B .

En utilisant l'équation 1.5, le flux sortant de la cible d'épaisseur dx peut s'écrire [Alb15] :

$$\Gamma_F(x + dx) = \Gamma_F(x) - \mathbb{P}(\text{interaction}) \cdot \Gamma_F(x) = \Gamma_F(x) (1 - n_B \sigma_{eff} dx) \quad (1.6)$$

et en notant que $\Gamma_F(x + dx) = \Gamma_F(x) + d\Gamma_F$, on trouve $d\Gamma_F(x) = -\Gamma_F(x) n_B \sigma_{eff} dx$ et donc [Alb15] :

$$\Gamma_F(x) = \Gamma_{F_0} e^{-x/\lambda_m} \quad (1.7)$$

avec $\lambda_m = 1/n_B \sigma_{eff}$ le libre parcours moyen, donnant la distance moyenne entre les collisions.

Il correspond, selon l'équation 1.7, à la distance sur laquelle le flux décroît d'une valeur $1/e^1$ de sa valeur initiale Γ_{F_0} . Plus la densité de particules cibles n_B est importante ou plus la section efficace est grande, plus les particules incidentes vont faire de collisions et donc plus ce libre parcours sera faible.

Le temps caractéristique τ_m entre collisions, pour une vitesse de particule incidente v_A peut alors s'écrire simplement $\tau_m = \lambda_m/v_A$. On trouve alors la fréquence de collision [LL05] :

$$\nu_m = \frac{1}{\tau_m} = n_B \sigma_{eff} v_A = n_B K \quad (1.8)$$

avec $K = \sigma_{eff} v_A$ (en $[m^3.s^{-1}]$) le taux de réaction. Et donc, selon l'équation 1.8, on retrouve bien que la fréquence de collision ν_m augmente avec la densité de particules cibles (dans le cas d'une collision électron - neutre, la fréquence de collision est donc proportionnelle à la densité de neutres et donc à la pression p) et la section efficace. De plus, cette fréquence de collision augmente aussi avec la vitesse des particules incidentes. Dans un cas pratique, les vitesses suivent une certaine distribution, la fréquence moyenne de collision s'exprime alors en fonction de la moyenne sur cette distribution.

Notons que : “La section efficace est une grandeur caractéristique primaire, c'est-à-dire issue de l'analyse microscopique de l'interaction ; le libre parcours moyen, la fréquence de collision et le taux de réaction sont des grandeurs secondaires permettant d'établir des bilans macroscopiques d'espèces” [Rax05]. Et il est important de comprendre que les sections efficaces dépendent de la vitesse des particules incidentes, donc de leur énergie $\mathcal{E}$. On notera alors $\sigma_{eff} = \sigma_{eff}(\mathcal{E})$. En pratique, l'énergie des électrons suit une distribution (la plus simple à étudier étant la Maxwellienne) et dans le cas d'une interaction électron - atome, les taux de réaction sont alors définis en intégrant sur la distribution. Ils sont alors définis en fonction de la température moyenne des électrons : $K = K(T_e)$.

Par exemple dans l'argon, les sections efficaces associées aux collisions électron - atome de types élastique, ionisation et excitation totale sont tracées en fonction de l'énergie de l'électron incident $\mathcal{E}$ sur la Figure 1.7(a). La section efficace de l'excitation totale est la somme des sections efficaces de tous les processus d'excitation de la forme $A_r + e^- \rightarrow A_r^* + e^-$. L'état excité A_r^* représente donc l'atome d'argon dans un des états représentés sur la Figure 1.5. On remarque tout d'abord la présence d'un minimum à faible énergie pour la section efficace élastique des électrons. Il est dû à une résonance mécanique quantique, que l'on nomme le minimum de Ramsauer [LL05]. Notons aussi que cette section efficace est toujours supérieure à celle d'ionisation ou d'excitation des électrons. Sur cette gamme d'énergie électronique, la nature des interactions électron - atome la plus probable est donc toujours la collision élastique. On retrouve sur la section efficace d'ionisation le seuil de $\mathcal{E}_{ionisation} = 15.76$ eV. En dessous de cette énergie, la section efficace, et donc la probabilité qu'une interaction électron - atome soit de nature ionisante, est nulle. Il faut que l'électron incident ait une énergie au moins supérieure à cette énergie seuil pour que la quantité d'énergie cédée à l'électron de l'atome corresponde à celle d'ionisation. De plus, on trouve une section efficace d'excitation

totale de même forme que celle d'ionisation, avec une énergie seuil un peu plus faible et une valeur maximale inférieure.

Les taux de réaction associés aux collisions électron - atome dans l'argon de type élastique, ionisation et excitation totale sont tracées en fonction de la température électronique moyenne sur la distribution T_e sur la Figure 1.7(b). Ces courbes traduisent un cas plus pratique. Elles sont lissées par l'intégration sur la distribution des températures électroniques. En dessous des niveaux d'énergie seuil d'ionisation et d'excitation, les taux de réaction décroissent exponentiellement avec T_e , traduisant la décroissance exponentielle du nombre d'électrons capables d'ioniser ou d'exciter l'atome [LL05]. On voit que même pour des températures électroniques moyennes de quelques eV, les taux de réaction d'ionisation et d'excitation sont non nuls.

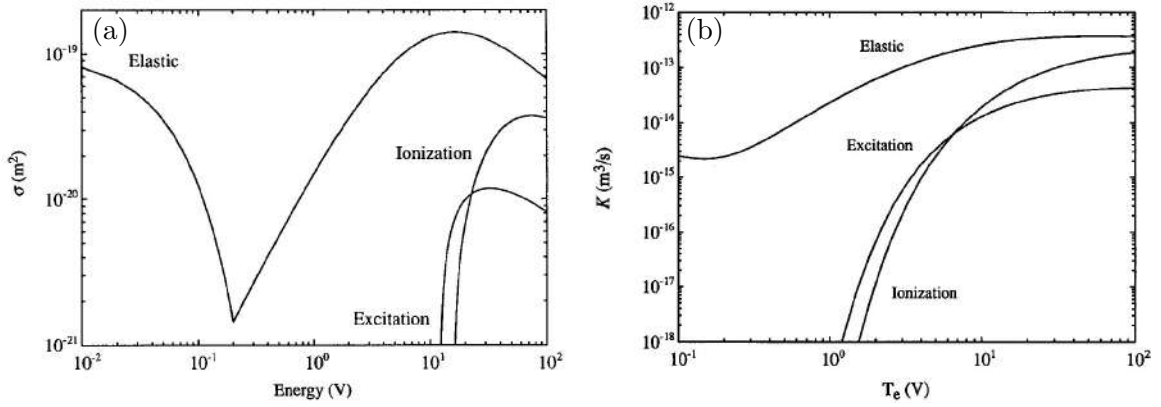


FIGURE 1.7 – (a) Sections efficaces et (b) taux de réaction élastique, d'ionisation et d'excitation pour les électrons dans l'argon (repris de [LL05]).

Pour que l'amorçage d'un plasma ait lieu, il faut que les électrons présents initialement dans le gaz soient suffisamment énergétiques pour donner naissance à de nouveaux électrons par ionisation. Lors de leurs collisions avec les particules du gaz, ces électrons vont perdre de l'énergie¹. Si cette perte d'énergie est trop importante, alors les électrons ne seront pas assez énergétiques pour que le processus d'ionisation ait lieu et il n'y aura pas de génération de plasma. Il faut donc que la perte d'énergie subie par les électrons soient compensée par le gain en énergie puisé dans le champ électrique. Il en résulte dans ce cas un milieu que l'on appelle plasma. Ce dernier émet de la lumière à travers le processus de désexcitation électronique que nous avons décrit précédemment et contient des espèces neutres et chargées (ions et électrons). En plus du champ électromagnétique appliqué au milieu de façon externe il peut aussi exister dans le plasma des différences de densité de charges spatiales. Cela entraîne la présence de forces s'exerçant sur les particules chargées. Ces forces sont à l'origine des mouvements des particules chargées au sein du plasma. La façon dont ces particules se comportent est conditionnée par les propriétés du plasma (densité n_e et température T_e

1. Il est aussi possible que l'électron gagne de l'énergie cinétique lors d'une collision avec un atome excité (qui se désexcite alors), collision que l'on appelle super-élastique [LL05].

électronique) et peut être caractérisée par longueur de Debye λ_D et la pulsation plasma ω_p , que nous allons présenter dans la sous-partie suivante.

1.1.2 Échelle mésoscopique : Quasi-neutralité du plasma

Dans un plasma, des interactions entre les particules ont donc lieu. Leur nature dépend de la température électronique et de la pression (équation 1.8). Des électrons et des ions positifs peuvent alors être créés par le processus d'ionisation. Ces deux particules, de charges électriques opposées, vont s'attirer l'une l'autre. Pourtant l'agitation thermique des électrons va limiter ce phénomène. En fait, "un plasma est soumis à deux tendances :

- une tendance de désordre due à l'agitation thermique, et
- une tendance à l'organisation due à l'aspect collectif que peut manifester l'interaction coulombienne.

Ces deux tendances permettent au plasma de rester sous forme ionisée tout en restant globalement neutre" [Rax05]. Pour être plus précis, elles créent des conditions s'opposant à la séparation de charges sur des distances supérieures à une certaine distance appelée la longueur de Debye (notée λ_D) et, pour des processus dynamiques, sur des temps plus longs qu'une certaine échelle de temps $\tau_p = 1/f_p$ (avec f_p la fréquence plasma). Le plasma peut être neutre sur un volume supérieur à λ_D^3 et sur une durée supérieure à τ_p . On parle ainsi de quasi-neutralité du plasma, puisque pour des volumes inférieurs à λ_D^3 et une dynamique plus rapide que τ_p une séparation de charges et donc un champ électrique peuvent apparaître [Rax05].

Intéressons nous d'abord à la distance à partir de laquelle un plasma peut être considéré comme neutre : la longueur de Debye. Cette longueur caractéristique est la conséquence du phénomène d'écrantage électrique. Il traduit le fait que si on place par exemple un ion chargé positivement dans un plasma, les électrons auront tendance à l'entourer comme représenté sur la Figure 1.8, de sorte à ce que par rapport aux autres particules du plasma l'ensemble ion + électron soit vu comme neutre. L'effet de la charge positive de l'ion est écranté par

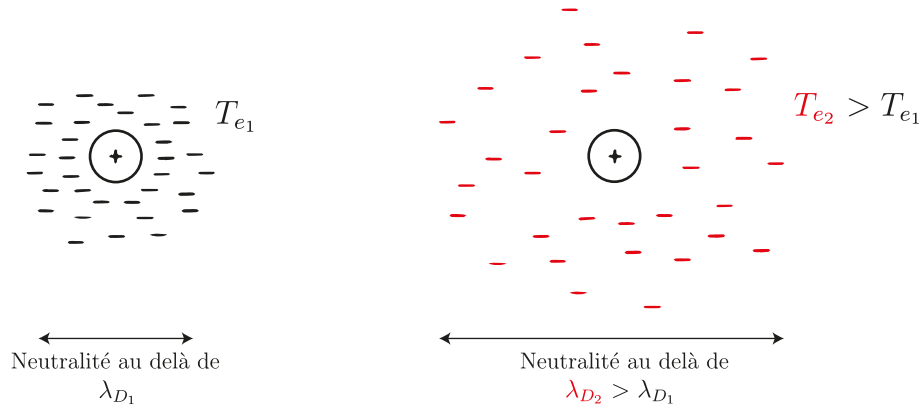


FIGURE 1.8 – Illustration de la longueur de Debye. Un ion chargé positivement est placée dans un plasma, à gauche avec une température électronique T_{e1} et à droite avec une température électronique $T_{e2} > T_{e1}$.

les électrons au delà d'une distance caractéristique nommé la longueur de Debye et noté λ_D . Cependant, l'agitation thermique des électrons va limiter ce mécanisme. En effet, plus la température électronique augmente (plus forte est l'agitation thermique), moins la trajectoire des électrons peut être modifiée par la présence d'ions. L'écrantage est alors moins efficace et la longueur de Debye augmente, comme représenté sur la Figure 1.8. Cet effet d'écrantage résulte alors de la compétition entre l'attraction des particules chargées et la température électronique. La longueur de Debye caractérise l'échelle typique sur laquelle la séparation de charge est autorisée. Elle s'écrit [LL05] ; [Rax05] ; [Che74] :

$$\lambda_D = \sqrt{\frac{\varepsilon_0 k_B T_e}{e^2 n_e}} \quad (1.9)$$

avec ε_0 la permittivité du vide et k_B la constante de Boltzmann.

On retrouve bien avec cette expression l'augmentation de la longueur de Debye avec la température électronique. De plus, elle dépend de l'inverse de la densité électronique n_e . En effet, plus cette densité est grande et plus la quantité d'électrons proche des ions sera grande. Finalement, cette longueur de Debye correspond à la distance au delà de laquelle un plasma peut être vu comme neutre. En deçà de cette distance, des séparations significatives de charge peuvent être observées. La longueur de Debye caractérise l'échelle typique sur laquelle la séparation de charge est autorisée.

Cependant, il peut aussi exister dans les plasmas des phénomènes rapides qui peuvent séparer les charges pendant une brève période que l'écrantage de Debye n'a pas le temps de résorber. Ces phénomènes sont appelés oscillations plasmas. Elles sont liées au mouvement collectif des électrons du plasma, qui peut alors briser la neutralité et induire une séparation de charges. Pour déterminer ce temps durant lequel une séparation de charges peut exister, prenons le cas simple d'un plasma de dimension finie constitué d'électrons froids ($T_e = 0$) et d'ions immobiles (avec $n_e = n_i$). Une perturbation unidimensionnelle (sous forme d'un excès de charges positives ou négatives) est alors introduite sur une distance $x(t)$ comme représenté sur la Figure 1.9. En utilisant le théorème de Gauss, le champ résultant de cette perturbation

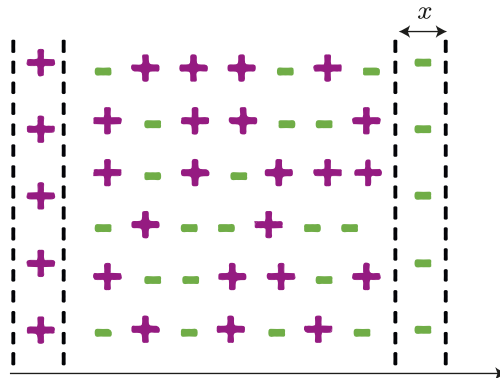


FIGURE 1.9 – Perturbation électronique dans un plasma [Boe16].

s'écrit $|\mathbf{E}| = en_e x / \varepsilon_0$ [LL05]. L'équation de transport pour les électrons, projetée sur l'axe x

s'écrit alors (en l'absence de champ magnétique) :

$$m_e \frac{d^2 x}{dt^2} = |\mathbf{F}_L| = -e|\mathbf{E}| = -\frac{e^2}{\varepsilon_0} n_e x \quad (1.10)$$

soit

$$\frac{d^2 x}{dt^2} + \omega_p^2 x = 0 \text{ avec } \omega_p = \sqrt{\frac{e^2 n_e}{\varepsilon_0 m_e}} \quad (1.11)$$

On trouve alors une équation typique d'oscillations non amorties de pulsation plasma ω_p . Selon cette équation 1.11, suite à la perturbation représentée sur la Figure 1.9, les électrons vont osciller indéfiniment autour d'une position d'équilibre avec une fréquence $f_p = \omega_p/2\pi$. La pulsation de ces oscillations dépend logiquement de l'inverse de la masse m_e de l'électron et augmente avec la densité électronique n_e (plus cette dernière sera grande et plus la force s'exerçant sur les électrons $|\mathbf{F}_L|$ sera grande). Cette équation 1.11 a été obtenue en supposant une température électronique nulle. Cela s'explique par le fait que ces oscillations sont extrêmement rapides, durant celles-ci l'équilibre thermodynamique local n'est pas assuré, la notion même de température n'est donc plus pertinente. Notons que dans un cas plus réaliste, ces oscillations seront amorties par collisions (et par le phénomène d'amortissement Landau) [LL05]. En fait la notion d'oscillation restera valable dans le cas $\nu_m \ll \omega_p$ [Rai91], *i.e.* dans le cas où les collisions n'amortissent pas trop rapidement ces oscillations.

Pour résumer, nous avons vu qu'un plasma peut être considéré comme neutre sur un volume supérieur à λ_D^3 et pour des phénomènes de dynamique inférieure à f_p , c'est ce que l'on appelle la quasi-neutralité du plasma. En deçà de la distance de Debye, des séparations de charges sont présentes à cause de la température électronique qui limite l'effet d'écrantage, et des processus plus lents que le temps caractéristique de l'oscillation plasma $\tau_p = 1/f_p$ laisseront au plasma le temps de se réorganiser afin de neutraliser la perturbation [Rax05]. Cette dernière propriété joue un rôle essentiel dans la nature de l'interaction entre une onde électromagnétique incidente et un plasma (tel que c'est le cas lors de l'amorçage de plasmas par RT). Pour une onde incidente oscillant à une fréquence bien inférieure à la fréquence plasma ($f_c \ll f_p$), le plasma aura le temps de se réorganiser de sorte à garder la neutralité et le champ électrique sera donc forcément nul dans le plasma. L'onde est alors réfléchiée par le plasma. Dans ce cas le plasma se comporte comme un conducteur. Pour une onde incidente oscillant à une fréquence bien supérieure à la fréquence plasma ($f_c \gg f_p$), le plasma n'a pas le temps de se réorganiser de sorte à garder la neutralité et un champ électrique peut donc exister dans le plasma. L'onde peut alors se propager, le plasma se comportant comme une diélectrique.

1.1.3 Échelle macroscopique : Modélisation fluide

Il existe plusieurs façons de modéliser les plasmas : particulaire, statistique et fluide. Afin de décrire les interactions onde - plasma, la plus intéressante est la modélisation fluide [LL05]. En effet, elle permet de manipuler des grandeurs macroscopiques telles que la densité ou la

vitesse moyenne des particules, qui, comme nous allons le voir, peuvent être ensuite facilement intégrées aux équations de Maxwell. Dans ce cas le plasma est traité comme un fluide constitué de particules chargées. Le problème se ramène donc à résoudre simultanément les équations de Maxwell et les équations de la mécanique des fluides adaptées au plasma.

Afin de modéliser le comportement des particules constituant un plasma suivant un modèle fluide, on introduit généralement le concept de fonction de distribution $f(\mathbf{r}, \mathbf{v}, t)$. C'est une modélisation cinétique du plasma, qui permet de décrire l'évolution temporelle des particules en fonction de leurs positions $\mathbf{r}$ et de leurs vitesses $\mathbf{v}$. Le produit $f(\mathbf{r}, \mathbf{v}, t)d^3rd^3v$ exprime alors le nombre de particules dont le centre est situé dans le domaine élémentaire compris entre les limites $x+dx$, $y+dy$ et $z+dz$ et dont les composantes des vitesses sont comprises entre les limites v_x+dv_x , v_y+dv_y et v_z+dv_z (avec $d^3r = dx dy dz$ et $d^3v = dv_x dv_y dv_z$). Sous l'action de forces macroscopiques, les particules peuvent entrer ou sortir du volume d^3rd^3v . La fonction de distribution $f(\mathbf{r}, \mathbf{v}, t)$ doit donc être continue. En faisant le bilan des flux d'électrons à travers les surfaces du volume d^3rd^3v , et en l'absence de collisions on trouve qu'elle vérifie [LL05] :

$$\frac{\partial f}{\partial t} + \mathbf{v}_e \cdot \nabla f + \mathbf{a}_e \cdot \nabla_{\mathbf{v}} f = 0 \quad (1.12)$$

avec $\nabla = (\mathbf{x}\partial/\partial x + \mathbf{y}\partial/\partial y + \mathbf{z}\partial/\partial z)$, $\nabla_{\mathbf{v}} = (\mathbf{x}\partial/\partial v_x + \mathbf{y}\partial/\partial v_y + \mathbf{z}\partial/\partial v_z)$ et $\mathbf{a}_e$ l'accélération des électrons. Cette équation est nommée l'équation de Vlasov ou équation de Boltzmann sans collisions [LL05]. En plus de pouvoir sortir ou entrer à travers les surfaces du volume, les particules peuvent soudainement apparaître ou disparaître suivant la nature des collisions. Cet effet est pris en compte en ajoutant un "terme de collisions" à la partie droite de l'équation 1.12, donnant alors l'équation de Boltzmann [LL05] ; [MP12] :

$$\frac{\partial f}{\partial t} + \mathbf{v}_e \cdot \nabla f + \frac{\mathbf{F}_L}{m_e} \cdot \nabla_{\mathbf{v}} f = \mathcal{S}(f) \quad (1.13)$$

avec ce terme de collision ² qui "traduit la variation, du fait des collisions élastiques et inélastiques, du nombre de particules dans l'élément de volume de l'espace des phases centré sur $\mathbf{r}$, $\mathbf{v}_e$ " [MP12].

À partir de cette équation de Boltzmann, il est possible de réduire considérablement la complexité des équations en intégrant sur les coordonnées de la vitesse, donnant alors des quantités macroscopiques. Pour obtenir ces quantités il faut déterminer les moments de l'équation de Boltzmann. La densité électronique est ainsi définie comme :

$$n_e(\mathbf{r}, t) = \iiint f d^3v \quad (1.14)$$

2. Habituellement on trouve ce terme de collision écrit comme $\frac{\partial f}{\partial t} \Big|_c$ [LL05], mais par soucis pédagogique nous avons choisi la notation utilisée dans [MP12] : $\mathcal{S}(f) = \frac{\partial f}{\partial t} \Big|_c$, où $\mathcal{S}(f)$ "désigne, de façon globale, l'opérateur de collisions".

1.1.3.1 Premier moment de Boltzmann : Équation de continuité

En intégrant tous les termes de l'équation 1.13 sur l'espace des vitesses, on obtient le premier moment de Boltzmann, correspondant à l'équation de continuité macroscopique :

$$\frac{\partial n_e}{\partial t} + \nabla \cdot (n_e \mathbf{u}_e) = s_e \quad (1.15)$$

avec $\mathbf{u}_e$ la vitesse moyenne des électrons et s_e le taux de production net d'électrons. Ce dernier inclut les différents types de collisions, tels que l'ionisation, l'attachement et les recombinaisons [CB10]. Dans l'argon, qui est un gaz atomique, il n'y a pas d'attachement. Le taux de production net d'électrons peut donc s'écrire : $s_e = \nu_i n_e - r_{ei} n_e^2$ (avec ν_i le taux d'ionisation et r_{ei} le coefficient de recombinaison électron - ion). L'équation 1.15 est une équation de diffusion. Le terme de gauche décrit l'évolution temporelle de la densité, il est associé à son évolution spatiale. La somme de ces deux termes doit être égale au nombre de particules créées ou perdues par collisions (partie droite). En fait la façon dont le plasma va diffuser dans l'espace dépend de la valeur de sa densité. Si cette densité est faible, alors les électrons pourront diffuser facilement dans l'espace, la diffusion est alors dite "libre" et caractérisée par le coefficient de diffusion libre D_e . Par contre si la densité est élevée, lorsque les électrons vont essayer de diffuser, la forte concentration de charges positives (dues aux ions positifs à l'endroit d'où diffusent les électrons) va se traduire par une force de rappel $|\mathbf{F}_L|$ importante qui empêche les électrons de diffuser. Le phénomène de diffusion est alors amoindri par rapport au cas libre, et on parle de diffusion ambipolaire, caractérisée par le coefficient de diffusion ambipolaire D_a . Pour rendre compte de cette diffusion du plasma dépendante de la densité, il est intéressant d'introduire un coefficient de diffusion effectif D_{eff} tel que [BCZ10] ; [CB10] :

$$\frac{\partial n_e}{\partial t} + \nabla \cdot (D_{eff} \nabla n_e) = s_e \quad (1.16)$$

avec :

$$D_{eff} = \frac{\xi D_e + D_a}{\xi + 1} \quad (1.17)$$

avec $\xi = \nu_i \tau_M$ et τ_M le temps de relaxation diélectrique défini comme $\tau_M = \varepsilon_0 / [en_e(\mu_e + \mu_i)]$, et μ_e et μ_i les mobilités respectivement des électrons et des ions.

Toutes ces grandeurs macroscopiques (n_e , D_{eff} , ν_i ...) dépendent de l'amplitude du champ électrique appliqué. Pour chaque configuration (gaz, pression, fréquence...) il existe un certain seuil, nommé champ de claquage $E_{claquage}$, à partir duquel il y a claquage du gaz et donc amorçage d'un plasma. Il faut pour cela que la production de particules chargées par ionisation soit supérieure aux pertes par diffusion et par recombinaison (et par attachement dans un gaz moléculaire). L'équation de continuité (équation 1.16) permet de quantifier la variation de la densité électronique sous l'effet de l'ionisation, des pertes par recombinaison (contenus

dans le terme de production s_e) et de la diffusion des particules chargées (traduit par le terme $\nabla \cdot (D_{eff} \nabla n_e)$). Le coefficient de diffusion effectif D_{eff} est une combinaison linéaire des coefficients de diffusion libre et ambipolaire, il rend compte du changement entre ces deux régimes en fonction de la densité électronique : il est égal au coefficient de diffusion libre pendant la phase de multiplication électronique et avant que le plasma quasi-neutre se forme ($D_{eff} \xrightarrow{n_e \rightarrow 0} D_e$), et une fois le plasma formé, ce coefficient doit être pris égal au coefficient de diffusion ambipolaire du plasma, qui est plus faible ($D_{eff} \xrightarrow{n_e \rightarrow \infty} D_a$). Ceci explique que le champ de claquage est souvent plus grand que le champ nécessaire à l'entretien du plasma.

Pour définir le claquage, on simplifie souvent le terme de diffusion de l'équation 1.16 en définissant une distance caractéristique de diffusion, notée Λ [MP12] (en négligeant la recombinaison) :

$$\frac{\partial n_e}{\partial t} = \left(\nu_i - \frac{D_{eff}}{\Lambda^2} \right) n_e \quad (1.18)$$

Le champ de claquage $E_{claquage}$ se déduit de cette équation en recherchant le champ électrique au-delà duquel la densité électronique croît, c'est à dire pour lequel :

$$\nu_i = \frac{D_{eff}}{\Lambda^2} \quad (1.19)$$

On doit alors, pour trouver la valeur du champ de claquage, estimer les pertes par diffusion. Si l'on est à basse pression et dans un volume limité par des parois, Λ peut être la longueur caractéristique de la cavité qui contient le plasma en formation. Si l'on est à haute pression et que le champ est inhomogène, par exemple beaucoup plus élevé au voisinage d'un objet pointu, Λ représente la dimension caractéristique de la région où le champ est supérieur au champ critique. Nous verrons dans le chapitre 4 que les plasmas sont amorcés expérimentalement avec des ondes électromagnétiques de fréquence centrale autour de $f_c = 2.4$ GHz. De plus, ces plasmas sont de forme sphérique, et de dimensions bien inférieures aux dimensions de la cavité. Le Λ s'exprime alors comme $\Lambda = L/\pi$ [Rai91] et nous prendrons comme première approximation $L = 5$ mm. Cela donne une distance caractéristique de diffusion d'environ $\Lambda \approx 1$ mm.

La Figure 1.10(a) montre le champ rms de claquage calculé à l'aide du logiciel Bolsig+ dans l'argon, pour différentes valeurs de longueur caractéristique Λ , avec une fréquence $f_c = 2.4$ GHz. Les différentes courbes possèdent les mêmes tendances. Elles passent par un minimum autour d'une pression comprise entre 1 et 2 torr. À cette pression, la quantité d'électrons générés par ionisation est plus importante. En deçà de cette pression, pour un même champ électrique appliqué les électrons ne font pas assez de collisions et n'ionisent pas assez le gaz. Le champ de claquage augmente alors. Au delà de cette pression, les électrons font des collisions trop "rapidement". Ils n'ont pas le temps de gagner assez d'énergie pour ioniser suffisamment

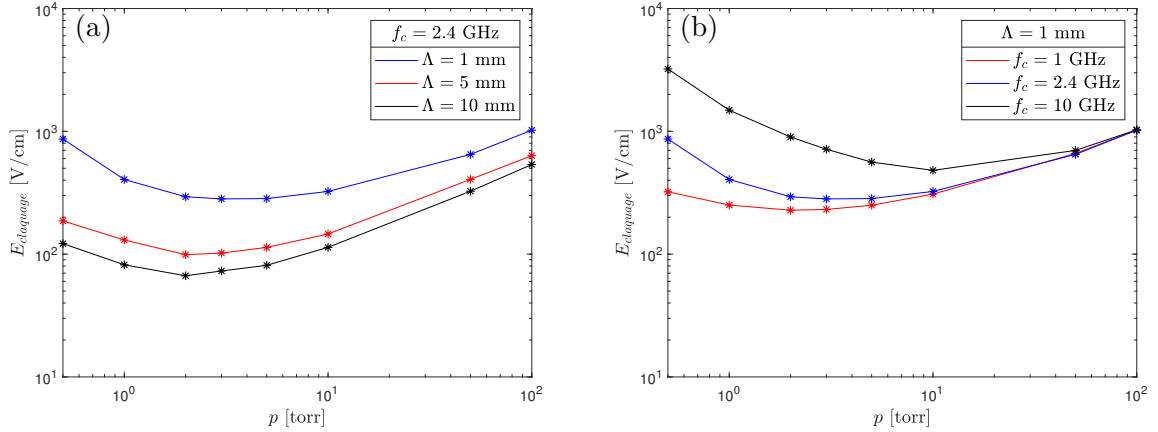


FIGURE 1.10 – (a) Champ de claquage rms dans l'argon en fonction de la pression calculé avec Bolsig+ pour plusieurs valeurs de la longueur de diffusion Λ avec une fréquence de 2.4 GHz. (b) Champ de claquage rms pour différentes fréquences f_c dans l'argon et pour une longueur de diffusion de 1 mm.

le gaz et le champ de claquage augmente. En outre, plus la longueur de diffusion est petite, plus la fréquence d'ionisation doit être élevée selon l'équation 1.19, ce qui se traduit par un champ de claquage plus important.

La Figure 1.10(b) présente le champ rms de claquage calculé à l'aide du logiciel Bolsig+ dans l'argon à longueur de diffusion Λ constante, pour différentes fréquences. Le minimum de champ de claquage est donc obtenu lorsque $\omega_c = 2\pi f_c \approx \nu_m$. À des pressions suffisamment élevées, $\omega_c \ll \nu_m$ et le champ de claquage ne dépend pas de la fréquence. En revanche, quand la pression diminue, le champ de claquage augmente avec la fréquence. Le champ de claquage minimum se situe à une pression telle que ν_m soit de l'ordre de ω_c , ce qui se vérifie également sur la Figure 1.10(b).

1.1.3.2 Deuxième moment de Boltzmann : Équation de transport de la quantité de mouvement

Pour le deuxième moment de Boltzmann, l'équation 1.13 est multipliée par $\mathbf{v}$ et intégrée sur l'espace des vitesses. On trouve alors [LL05] :

$$m_e n_e \left[\frac{\partial \mathbf{u}_e}{\partial t} + (\mathbf{u}_e \cdot \nabla) \mathbf{u}_e \right] = q n_e (\mathbf{E} + \mathbf{u}_e \wedge \mathbf{B}) - \nabla \cdot \overline{\overline{\mathbf{\Pi}}} + \mathbf{f}|_c \quad (1.20)$$

avec les deux termes $\partial \mathbf{u}_e / \partial t$ et $(\mathbf{u}_e \cdot \nabla) \mathbf{u}_e$ à gauche de l'équation qui représentent respectivement les accélérations dues à une variation temporelle et spatiale de la vitesse moyenne $\mathbf{u}_e$. Le terme spatial est un terme convectif, non-linéaire en $\mathbf{u}_e$, ce qui ne facilite pas la résolution de l'équation 1.20. Cependant, "sa contribution est souvent négligeable devant les autres

termes ; celle-ci demeure évidemment importante quand $\mathbf{u}_e$ est grand en valeur absolue ou lorsque son gradient est fort” [MP12]. Le terme $\overline{\Pi}$ représente le tenseur de pression. Dans la plupart des cas, ce dernier est isotrope on peut ainsi écrire : $\nabla \cdot \overline{\Pi} = \nabla \cdot p$. Le dernier terme de cette équation $\mathbf{f}|_c$ représente le taux de collisions de transfert d’énergie cinétique par unité de volume. Il peut s’écrire simplement, en négligeant les forces magnétiques et le mouvement des neutres, comme $\mathbf{f}|_c = -m_e n_e \nu_m \mathbf{u}_e$ [LL05]. On peut alors ré-écrire l’équation 1.20 sous la forme :

$$m_e n_e \frac{\partial \mathbf{u}_e}{\partial t} = q n_e \mathbf{E} - \nabla \cdot p - m_e n_e \nu_m \mathbf{u}_e \quad (1.21)$$

Ainsi, chaque nouveau moment de Boltzmann introduit une nouvelle inconnue. Le premier moment de Boltzmann, pouvant être interprété comme un bilan de particules, introduit la vitesse des particules. De même, le second moment de Boltzmann, pouvant être interprété comme un bilan de quantité de mouvement, introduit le tenseur de pression (on note au passage que l’équation d’évolution d’une grandeur physique représentée par un tenseur d’ordre k_{ordre} introduit une nouvelle grandeur physique représentée par un tenseur d’ordre $k_{ordre} + 1$). “Cette situation introduit donc nécessairement plus de variables que d’équations et ne peut conduire à un problème soluble que s’il existe d’autres relations entre les grandeurs physiques. Ces relations ne peuvent pas être données par l’approche fluide qui ne constitue donc pas une théorie fermée” [Rai13]. Généralement, la méthode pour fermer ce système d’équations consiste à utiliser des équations d’états thermodynamique afin de relier p et n_e . Pour une distribution Maxwellienne à l’équilibre, la relation isotherme est : $p = n_e k_B T$, ce qui donne $\nabla p = k_B T \nabla n_e$ [LL05]. Dans le cas simple dans lequel le plasma est uniforme (n_e constant $\rightarrow \nabla n_e = \mathbf{0}$) l’équation de transport de la quantité de mouvement des électrons devient :

$$\frac{\partial \mathbf{u}_e(\mathbf{r}, t)}{\partial t} = -\frac{e}{m_e} \mathbf{E}(\mathbf{r}, t) - \nu_m \mathbf{u}_e(\mathbf{r}, t) \quad (1.22)$$

Nous avons ainsi obtenu, à partir de l’équation de Boltzmann, les équations décrivant le comportement du plasma (sous certaines hypothèses) au niveau macroscopique. Nous verrons dans le chapitre 3 comment ces équations 1.16 et 1.22 (avec 1.17) sont discrétisées et intégrées dans un code de simulation FDTD dans lequel elles sont couplées aux équations de Maxwell afin de décrire l’amorçage des plasmas par RT.

1.1.3.3 Comportement électromagnétique du plasma : Permittivité équivalente

Pour finir, il est intéressant de construire la permittivité du plasma à partir des équations développées pour le modèle fluide. Cela permet de décrire le comportement du plasma de façon électromagnétique, et d’interpréter facilement la nature de l’interaction onde - plasma. Pour cela, il faut passer dans le domaine fréquentiel afin de simplifier les équations. L’équation

1.22 devient alors ³ (correspondant à la convention $e^{+j\omega t}$) :

$$\tilde{\mathbf{u}}_e(\mathbf{r}, \omega) = -\frac{e}{m_e(\nu_m + j\omega)} \tilde{\mathbf{E}}(\mathbf{r}, \omega) \quad (1.23)$$

Du point de vue électromagnétique, le plasma peut être modélisé par un courant de conduction des particules chargées dans le vide [LL05]. La masse des ions étant bien supérieure à celle des électrons, on négligera le courant des ions par rapport à celui des électrons, appelé $\tilde{\mathbf{J}}_e$. L'équation de Maxwell-Ampère en présence d'un plasma dans le domaine fréquentiel peut ainsi s'écrire :

$$\nabla \times \tilde{\mathbf{H}}(\mathbf{r}, \omega) = j\omega\epsilon_0 \tilde{\mathbf{E}}(\mathbf{r}, \omega) + \underbrace{\tilde{\mathbf{J}}_e(\mathbf{r}, \omega)}_{\text{Courant plasma}} = j\omega \underbrace{\epsilon_0 \tilde{\epsilon}_{pr}(\mathbf{r}, \omega)}_{\tilde{\epsilon}_p(\mathbf{r}, \omega)} \tilde{\mathbf{E}}(\mathbf{r}, \omega) = \underbrace{\tilde{\mathbf{J}}_T(\mathbf{r}, \omega)}_{\text{Courant total}} \quad (1.24)$$

En posant :

$$\tilde{\mathbf{J}}_e(\mathbf{r}, \omega) = \sigma_e(\mathbf{r}, \omega) \tilde{\mathbf{E}}(\mathbf{r}, \omega) = -en_e(\mathbf{r}, \omega) \tilde{\mathbf{u}}_e(\mathbf{r}, \omega) \quad (1.25)$$

avec $\sigma_e(\mathbf{r}, \omega)$ la conductivité du courant des électrons, on trouve la densité de courant total :

$$\tilde{\mathbf{J}}_T(\mathbf{r}, \omega) = j\omega\epsilon_0 \tilde{\mathbf{E}}(\mathbf{r}, \omega) + \tilde{\mathbf{J}}_e(\mathbf{r}, \omega) \quad (1.26)$$

$$\tilde{\mathbf{J}}_T(\mathbf{r}, \omega) = j\omega\epsilon_0 \tilde{\mathbf{E}}(\mathbf{r}, \omega) - en_e(\mathbf{r}, \omega) \tilde{\mathbf{u}}_e(\mathbf{r}, \omega) \quad (1.27)$$

$$\tilde{\mathbf{J}}_T(\mathbf{r}, \omega) = j\omega\epsilon_0 \underbrace{\left(1 + \frac{e^2 n_e(\mathbf{r}, \omega)}{j\omega\epsilon_0 m_e(\nu_m + j\omega)}\right)}_{\tilde{\epsilon}_{pr}(\mathbf{r}, \omega)} \tilde{\mathbf{E}}(\mathbf{r}, \omega) \quad (1.28)$$

avec $\tilde{\epsilon}_{pr}(\mathbf{r}, \omega)$ la permittivité équivalente du plasma, qui s'écrit (en utilisant l'équation 1.23) :

$$\tilde{\epsilon}_{pr}(\mathbf{r}, \omega) = \underbrace{1 - \frac{\omega_p^2}{\omega^2 + \nu_m^2}}_{\tilde{\epsilon}'_{pr}(\mathbf{r}, \omega)} - j \underbrace{\frac{\nu_m}{\omega} \frac{\omega_p^2}{\omega^2 + \nu_m^2}}_{\tilde{\epsilon}''_{pr}(\mathbf{r}, \omega)} \quad (1.29)$$

Afin d'illustrer la différence d'interaction onde - plasma en fonction de la fréquence plasma f_p dont nous avons discuté dans la partie précédente, intéressons nous en détail au comportement de cette permittivité plasma relative $\tilde{\epsilon}_{pr}$ dans les conditions expérimentales de cette thèse. Les expériences présentées par la suite ont été réalisées dans de l'argon à des pressions

3. Les quantités dans le domaine fréquentiel sont représentées avec un $\tilde{\cdot}$.

autour du torr avec des impulsions microondes autour d'une fréquence centrale $f_c = 2.4$ GHz. Dans ces conditions, la fréquence de collision électron - neutre ν_m est de l'ordre de quelques GHz. Typiquement elle varie entre 1 GHz pour des champ électriques faibles et 10 GHz pour de forts champs (données obtenues avec Bolsig + [HP05]). Prenons par simplicité une valeur intermédiaire $\nu_m = 5$ GHz (sa valeur ne change pas significativement le comportement de $\tilde{\varepsilon}_{pr}$). Dans ce cas, on peut tracer la partie réelle de la permittivité relative plasma $\tilde{\varepsilon}'_{pr}$ et sa partie imaginaire $\tilde{\varepsilon}''_{pr}$ en fonction de la densité électronique n_e ou de la fréquence plasma f_p , comme tracé sur la Figure 1.11 (la fréquence plasma et la densité électronique sont reliées par la relation 1.11).

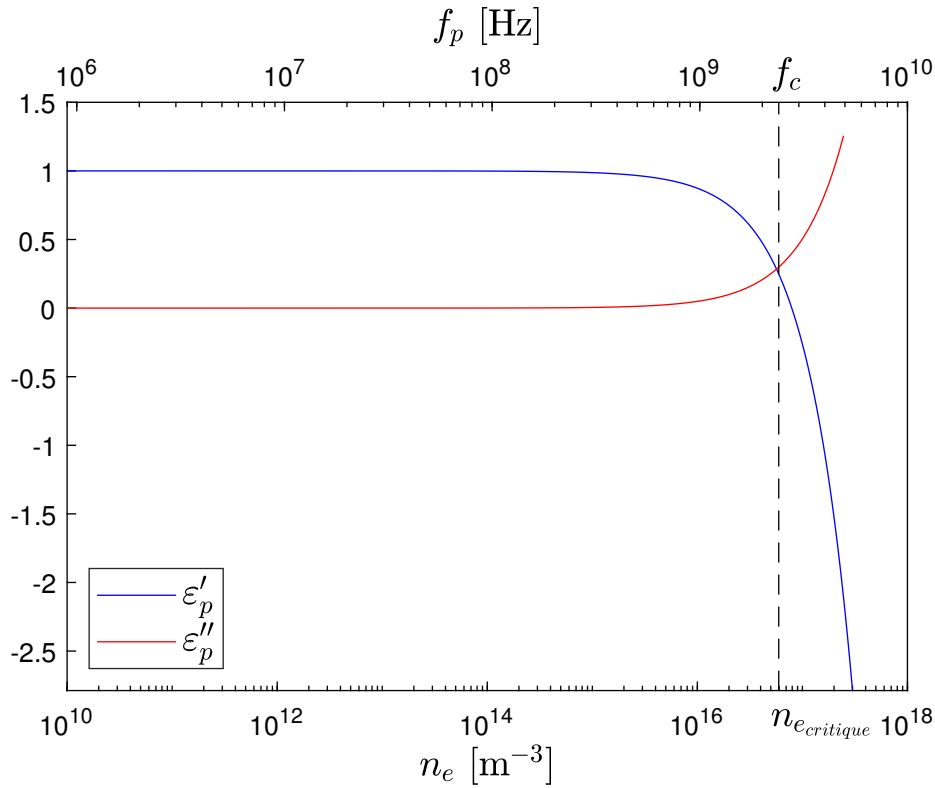


FIGURE 1.11 – Parties réelle et imaginaire de la permittivité relative du plasma en fonction de la densité électronique et de la fréquence plasma (pour une fréquence de collision électron - neutre $\nu_m = 5$ GHz et une pulsation $\omega_c = 2\pi f_c$).

Les courbes de la Figure 1.11 illustrent bien le comportement non-linéaire du plasma avec la densité électronique (et donc de la fréquence plasma). On remarque que pour des densités électroniques inférieures à environ 10^{16} m^{-3} la permittivité plasma relative est réelle et quasi-égale à 1. Cette dernière étant un facteur de proportionnalité par rapport à ε_0 , cela veut juste dire que pour des faibles densités électroniques, la présence des électrons ne modifie pas le comportement des ondes vis-à-vis d'une propagation dans le vide. En effet, la fréquence plasma reste inférieure à celle de l'onde et cette dernière se propage de la même façon que dans un gaz non-ionisé. Ce résultat se voit directement sur l'expression de la permittivité

plasma de l'équation 1.29.

En effet, pour la partie réelle $\tilde{\epsilon}'_{pr}$, la partie négative peut s'écrire $\omega_p^2/(\omega^2 + \nu_m^2) \simeq \omega_p^2/(\omega^2) = (e^2/\epsilon_0.m_e)/(n_{e_{critique}}^2).n_e \simeq 10^{-17}.n_e$. Cette partie négative indique la façon dont la partie réelle de la permittivité plasma s'éloigne de la permittivité du vide. On voit que pour que la densité électronique n_e ait un effet significatif sur la permittivité plasma à 2.4 GHz, il faut qu'elle soit de l'ordre de 10^{17} m^{-3} , comme observé sur la Figure 1.11. De plus on remarque un changement de signe de la partie réelle pour une densité critique $n_{e_{critique}}$ qui vérifie $f_p(n_{e_{critique}}) = f_c$. En dessous de cette densité critique, la partie réelle est positive et les ondes de fréquence f_c peuvent pénétrer dans le plasma. Plus la densité électronique augmente par rapport à cette valeur critique, plus les ondes sont réfléchies par le plasma.

Pour la partie imaginaire $\tilde{\epsilon}''_{pr}$ il y a le facteur $\nu_m/\omega \simeq 0.3$ en plus de $\omega_p^2/(\omega^2 + \nu_m^2)$. Ainsi, pour des densités électroniques supérieures à environ 10^{16} m^{-3} , la partie imaginaire augmente, traduisant des pertes (signe “-” sur la partie imaginaire).

Au final, la modélisation fluide d'un plasma permet d'exprimer sa permittivité par l'équation 1.29. Cette permittivité rend compte de l'interaction onde - plasma en fonction de la fréquence de l'onde (pulsation ω), de la fréquence plasma (pulsation ω_p) et de la fréquence de collision électron - neutre ν_m (qui dépend de la pression p). Elle correspond sous cette forme au modèle de Drüde [Dru00a]; [Dru00b]. Notons que ce modèle de Drüde a initialement été développé pour décrire le comportement des métaux, en considérant les électrons de ses couches supérieures comme libres. La fréquence plasma des métaux étant de l'ordre de la dizaine de milliers de THz [Rai57], cela explique pourquoi les ondes microondes (de l'ordre du GHz) sont réfléchies par les métaux. Nous verrons dans la partie suivante que cela permet de générer des ondes stationnaires dans les cavités métalliques, donnant lieu à un phénomène de résonance qui permet une intensification du champ électrique dans les cavités. De cette façon il est possible d'atteindre les seuils de claquage et d'amorcer des plasmas.

1.1.4 Échelle microscopique → Échelle macroscopique : Petit récapitulatif

Essayons de résumer la description des mécanismes à l'œuvre dans un plasma proposée dans cette partie. Lorsqu'un plasma est soumis à des champs électromagnétiques, les particules chargées du plasma vont alors se mettre en mouvement, changeant alors l'organisation des espèces chargées, qui vont en retour changer les champs électromagnétiques et ainsi de suite. Le lien entre ces mécanismes est représenté sur la Figure 1.1.

Commençons au niveau microscopique. Un plasma est un gaz contenant des particules neutres et chargées soumises à des forces et font ainsi des collisions, de différentes natures. Dans le cas où la force appliquée est de type électromagnétique, les particules chargées sont mises en mouvement (force de Lorentz). Les particules interagissent entre elles. La nature de ces interactions peut être élastique ou inélastique (ionisation, excitation...). Elle est fortement dépendante de l'énergie transférée entre les particules, les niveaux d'énergies étant quantifiés. L'interaction la plus importante est l'ionisation, qui permet d'arracher un électron à un atome. Les interactions se traduisent par une perte de l'énergie des électrons. Pour qu'il y ait un

plasma, il faut que cette perte d'énergie soit compensée par celle gagnée par les électrons issus de l'ionisation. Lors d'une collision de type excitation, un atome est excité sur un niveau supérieur d'énergie si la bonne quantité d'énergie lui est transmise. Cet atome excité va ensuite se désexciter en émettant un photon. Les plasmas émettent donc de la lumière. Les interactions de type ionisation et excitation sont des processus à seuil. Ils sont fortement non-linéaires. Notons que toutes ces interactions (élastique, ionisation, excitation...) se font avec un certain taux de réaction propre.

Ensuite, à un niveau mésoscopique, un plasma possède la propriété de quasi-neutralité. En fait, le plasma est considéré comme neutre sur des distances supérieures à la longueur de Debye et, pour des processus dynamiques, sur des temps plus lents que la période d'oscillation plasma $\tau_p = 1/f_p$. Sur des distances inférieures, l'effet d'écrantage de Debye n'est pas assez puissant pour vaincre l'agitation thermique des électrons, et le plasma n'est donc pas neutre. Pour des phénomènes plus lents que $\tau_p = 1/f_p$, le plasma a le temps de se réorganiser afin de neutraliser la perturbation. Il en résulte que la nature de l'interaction onde - plasma dépend de la fréquence de l'onde : si elle est inférieure à la fréquence plasma, le plasma se comporte comme un conducteur et l'onde est alors réfléchi par le plasma et, si elle est supérieure à la fréquence plasma, le plasma se comporte comme un diélectrique et l'onde peut alors se propager dans le plasma.

Pour finir, au niveau macroscopique il est intéressant d'utiliser une modélisation fluide du plasma pour rendre compte des interactions onde - plasma. Le plasma est alors décrit par des grandeurs macroscopiques, telle que la densité ou la vitesse moyenne des particules. Les moments de Boltzmann permettent alors de suivre le comportement du plasma au niveau macroscopique, en décrivant les phénomènes de diffusion (libre ou ambipolaire) et de collision (élastique ou inélastique). De plus, elles décrivent le changement d'état du milieu, passant de gaz à plasma. Ce changement d'état, appelé claquage, a lieu lorsque la production de particules chargées par ionisation est supérieure aux pertes par diffusion et par recombinaison (et par attachement dans un gaz moléculaire). Nous avons pu identifier que la valeur du champ de claquage dépend de la configuration (gaz, pression, fréquence...). Il est finalement possible de construire une permittivité du plasma, qui, couplée à ces moments de Boltzmann est en fait la manifestation au niveau macroscopique des mécanismes décrits précédemment à plus petite échelle.

1.2 Les plasmas microondes en cavité : Limitations actuelles

Les plasmas froids possèdent une physique riche et pluridisciplinaire. Ces plasmas peuvent avoir des propriétés différentes. De nombreuses applications, associées à chaque type de plasma, en découlent. “As a consequence, the trends in research, science and technology are difficult to follow and it is not easy to identify the major challenges of the field and their many sub-fields. Even for experts the road to the future is sometimes lost in the mist” [Sam+12]. Pour répondre à ce besoin de clarté, récemment de nombreux spécialistes se sont réunis et ont rédigé une feuille de route : “The 2012 Plasma Roadmap” [Sam+12].

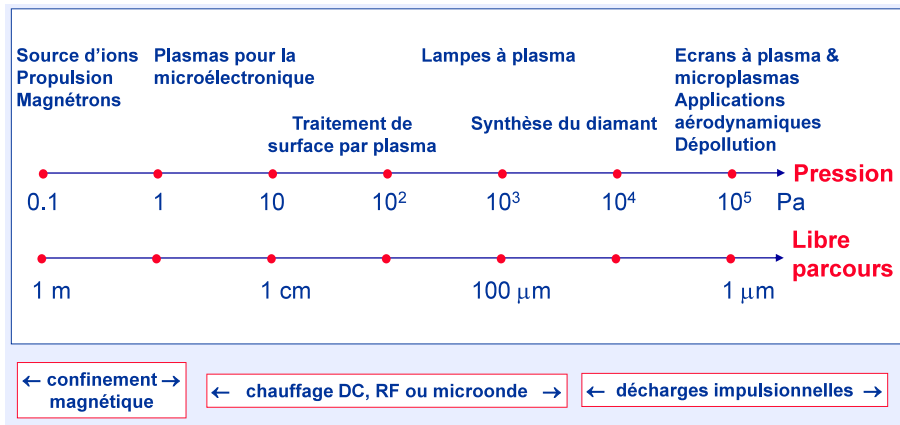


FIGURE 1.12 – Différentes applications des plasmas froids hors équilibre en fonction de la pression et du libre parcours (repris de [Boe16]).

Les applications principales des plasmas froids hors équilibre sont classées en fonction de la pression (et de l'ordre de grandeur du libre parcours correspondant) sur la Figure 1.12. Les décharges microondes font l'objet d'études intensives puisqu'elles permettent de produire des plasmas de haute densité à la fois à basse et haute pression. Elles occupent une place importante parmi les différentes applications des plasmas froids hors équilibre. Mais, indépendamment de l'application, les décharges microondes sont aussi des objets très intéressants pour des études plus fondamentales, puisqu'elles réunissent des phénomènes électrodynamiques, du domaine de la cinétique des plasmas et de la chimie des plasmas dans des conditions de non-équilibre et non-homogénéité (en général) [Leb15]. De plus, l'absence d'électrode dans ces décharges a l'avantage, comparé à des dispositifs à champ continu, d'éliminer la contamination du plasma et l'érosion des électrodes. Nous verrons cependant dans cette partie que l'absence totale d'érosion dans le réacteur nécessite de nombreuses étapes de conception. Notons aussi que la gamme des plasmas microondes et de leurs applications est très étendue. Le lecteur intéressé pourra se référer au papier de review de Yu A. Lebedev [Leb15], dans lequel il mentionne que tous les domaines clés identifiés dans la “Plasma Roadmap” [Sam+12] sont couverts par les décharges microondes.

Produits par ondes progressives comme dans le cas d'un surfatron [MZP79] ; [MZ91], par rayonnement de la puissance microonde avec un réseau d'antennes à fente [CS00] ; [Kor+96] ou en injectant les ondes dans une cavité métallique résonante [Asm+74] ; [HHG04], les plasmas

microondes sont contrôlés par l'équilibre entre [Asm89] :

- l'absorption d'énergie électromagnétique par le plasma
- les pertes de puissance dans le plasma (par collisions élastiques ou inélastiques et par chauffage du gaz ou des parois).

Une source plasma microonde est généralement constituée d'une source de puissance à ondes continues "continuous wave" (CW) à des fréquences de 0.915 ou 2.45 GHz, d'un applicateur plasma microonde, et d'un réacteur plasma [FM93] ; [Leb15]. Ce système peut être couplée à un champ magnétique statique tel que dans le cas des plasmas excités par résonance cyclotron électronique. Cela produit des plasmas plus uniformes et de densités supérieures [LPA97]. Comme nous allons le voir, l'applicateur plasma microonde ainsi que le réacteur plasma sont généralement conçus de sorte à optimiser la distribution spatiale du champ électrique qui est responsable du claquage et du maintien du plasma. En conséquence, de nombreuses sources plasmas microondes peuvent être trouvées dans la littérature [Leb15].

Notons que comme reporté dans [Leb15], il est difficile de classer les systèmes microondes plasma de façon catégorique. En effet, le problème est que les chercheurs conçoivent spécifiquement leurs dispositifs plasma de façon à répondre à un objectif. Il en résulte en fait de nombreux design possibles pour répondre à un objectif précis. De plus, comme nous allons le voir dans cette partie, de nombreuses étapes de conception des dispositifs sont nécessaires pour arriver à contrôler la distribution spatiale du plasma dans le réacteur.

Durant cette thèse, les plasmas ont été amorcés par retournement temporel expérimentalement dans une cavité métallique. L'intérêt est double, il est alors possible de contrôler la pression afin de se placer dans des conditions plus favorables de claquage, et la cavité est aussi nécessaire pour le processus du RT (nous reviendrons sur ce point au chapitre 3). Nous nous concentrerons donc dans cette partie sur les plasmas microondes en cavité. Dans ce cas le champ électrique est intensifié dans une cavité métallique, qui joue le rôle de réacteur plasma, en utilisant ses modes de résonances [FEB65] ; [Asm+74] ; [Sil+09] ; [Leb15]. Ce genre de plasma est largement utilisé pour le traitement de surface [FM93] ; [Gri94] ; [LL05], et pour la synthèse de diamant en utilisant le dépôt chimique en phase vapeur assisté par plasma [DW98]. Le design de la cavité occupe alors une partie importante dans la conception de cette source plasma, puisque l'optimisation de la géométrie de la cavité permet de façonner le champ électrique en son sein. Ainsi, la cavité peut être cylindrique [Li+14], ellipsoïdale [FWK98], ou même de forme plus complexe [Su+14]. Le point commun entre ces sources plasmas est le fait que lorsque le design de la cavité est réalisé, la position du plasma est fixée dans la cavité une fois pour toute. Nous verrons que nous proposons durant cette thèse une source plasma basée sur le RT, qui permet un contrôle dynamique de la position des plasmas dans les cavités quasi-indépendant de son design.

Notons aussi que la présence d'un plasma dans la cavité change la répartition spatiale du champ électromagnétique en son sein. Il est ainsi souvent nécessaire de prendre cet effet en compte au moment de la conception du réacteur. Il faut alors coupler au modèle électromagnétique, utilisé pour concevoir la cavité de sorte à structurer le champ électromagnétique, un modèle plasma rendant compte de la présence d'un plasma et de l'interaction onde - plasma. En plus de la prise en compte de cet effet, il est aussi nécessaire d'utiliser des éléments de

réglage (comme par exemple un piston, permettant d'ajuster le design de la cavité en présence du plasma). Sans eux, au moment de l'amorçage du plasma, la cavité devient désadaptée en impédance et la puissance réfléchie à son accès augmente alors, se traduisant par un couplage moindre, voire par l'insuccès de l'amorçage du plasma [Leb15].

Finalement, la conception d'un système microonde plasma requiert de nombreuses étapes de conception et la présence d'un dispositif de réglage (accord d'impédance). Nous rappellerons dans la partie suivante le principe des décharges microondes en cavité résonante. Nous verrons qu'il consiste à intensifier le champ électromagnétique dans la cavité par résonance, et nous verrons l'importance du design de la cavité (géométrie + dimensions) et du choix de la fréquence d'excitation sur la structure du champ dans la cavité. Ensuite, nous suivrons les différentes étapes de conception d'un système microonde plasma sur un exemple dont l'objectif est la croissance de diamant. L'importance des étapes de conception de la cavité ainsi que de la prise en compte du plasma sur la répartition spatiale du champ en son sein seront ainsi illustrées. Finalement, nous synthétiserons toutes les étapes de conception d'une source plasma microonde. Cela nous permettra d'identifier les limites des techniques classiquement utilisées. Dans la partie suivante nous verrons comment la source plasma à RT que nous proposons dans cette thèse permet de s'affranchir de toutes ces étapes de conception, en rendant possible un contrôle du plasma complètement indépendant du design de la cavité.

1.2.1 Principe des décharges microondes en cavité

Le principe des décharges microondes en cavité réverbérante consiste à utiliser le phénomène de résonance, donnant lieu à une intensification locale du champ électrique qui permet alors de claquer le gaz. Comme nous l'avons expliqué dans la partie précédente, pour qu'un claquage plasma ait lieu, il faut que les pertes d'électrons par diffusion et par recombinaison soient compensées par le gain d'électrons par ionisation. Dans le cas des décharges microondes, le champ électrique oscille à une fréquence fixe f_0 de l'ordre du GHz. En utilisant la notation complexe dans le but de simplifier les calculs, le champ électrique peut s'écrire dans ce cas $\mathfrak{E}(t) = \mathbf{E}_0 e^{j\omega_0 t}$ (avec $\omega_0 = 2\pi f_0$). Dans un plasma froid, le mouvement par agitation thermique des électrons peut être négligé devant celui engendré par l'application d'un champ électrique [MP06]. Les électrons vont alors osciller avec le champ (avec, potentiellement, un certain déphasage) et leur vitesse peut alors s'écrire : $\mathbf{v}_e(t) = \mathbf{v}_{e0} e^{j\omega_0 t}$ [MP06]. En introduisant ces complexes dans l'équation de quantité de mouvement 1.22, on trouve [MP06] :

$$\mathbf{v}_{e0} = -\frac{e}{m_e} \frac{1}{\nu_m + j\omega_0} \mathbf{E}_0 \quad (1.30)$$

Ensuite, en écrivant la densité de courant $\mathfrak{J}_e(t) = \mathbf{J}_{e0} e^{j\omega_0 t}$, on trouve :

$$\mathbf{J}_{e0} = -en_e \mathbf{v}_{e0} = \frac{e^2 n_e}{m_e} \frac{\nu_m - j\omega_0}{\nu_m^2 + \omega_0^2} \mathbf{E}_0 \quad (1.31)$$

La puissance absorbée par les électrons par unité de volume est alors donnée par [MP06] :

$$P_a = \frac{1}{2} \text{Re} [\mathfrak{E} \cdot \mathfrak{J}_e^*] = n_e \frac{e^2 E_0^2}{2m_e} \frac{\nu_m}{\nu_m^2 + \omega_0^2} \quad (1.32)$$

avec $E_0 = |\mathbf{E}_0|$. En l'absence de collision électron - neutre ($\nu_m = 0$), la puissance absorbée est nulle. Les électrons sont successivement accélérés et décélérés au cours d'un cycle (période d'oscillation) et ne gagnent pas d'énergie en moyenne. Les collisions transforment l'énergie dirigée due au champ électrique en énergie aléatoire et permettent le chauffage des électrons. Pour une fréquence de collision élevée $\nu_m \rightarrow \infty$, la puissance absorbée par les électrons tend vers 0. Dans ce cas les électrons n'ont pas le temps de gagner d'énergie, ils sont à peine accélérés qu'ils perdent déjà leur énergie par collision. Au final, la puissance absorbée par les électrons passe par un maximum pour $\nu_m = \omega_0$.

Selon cette formule 1.32, la puissance transmise aux électrons dépend de la fréquence de collision (donc de la pression p , selon l'équation 1.8), de l'amplitude du champ électrique E_0 et de la fréquence d'oscillation du champ. Pour les plasmas microondes, cette fréquence est de l'ordre du GHz. L'utilisation d'une cavité métallique permet alors le contrôle de la pression ainsi que l'intensification du champ électrique (par le phénomène de résonance) afin d'augmenter la puissance absorbée par les électrons de sorte à ce que l'énergie qu'il gagnent suffise à déclencher le processus d'ionisation. De cette façon, il est possible d'amorcer et de maintenir des plasmas microondes dans une cavité. La conception de la cavité est ainsi une étape incontournable et décisive. Elle conditionne à la fois l'intensification du champ électrique ainsi que sa répartition spatiale dans la cavité, et donc la position du plasma. Nous resterons dans cette partie à un niveau d'analyse très sommaire des cavités résonantes, l'objectif ici étant de comprendre les différentes étapes de conception d'une source plasma microonde standard à cavité. La théorie décrivant la réponse des cavités à une excitation sera développée en détail dans le chapitre 3.

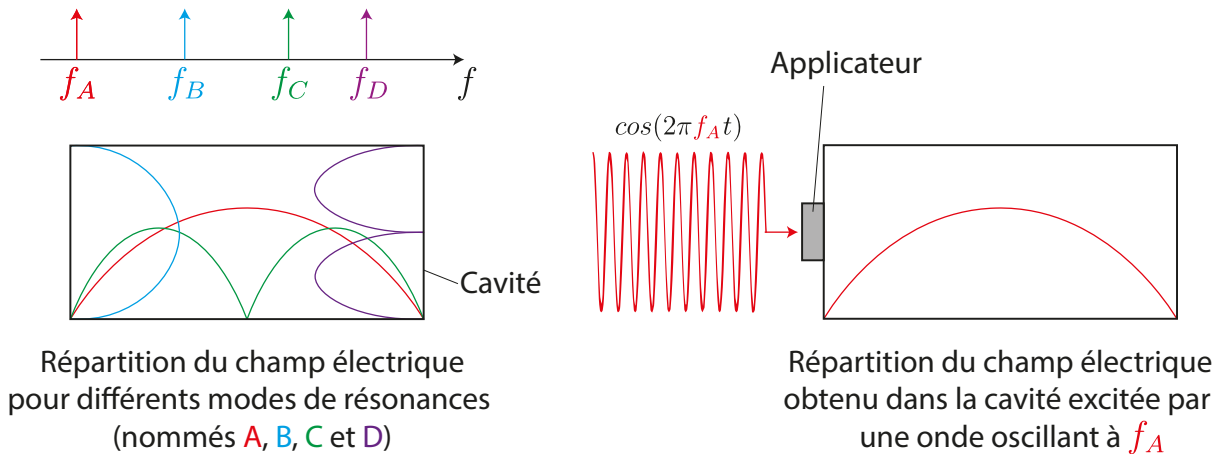


FIGURE 1.13 – (Gauche) “Répartition spatiale” du champ électrique pour différents modes (de fréquences différentes) dans une cavité. (Droite) Répartition spatiale du champ électrique obtenu dans la cavité excitée par une onde oscillant à la fréquence du mode A : f_A . (inspiré de [Sil19]).

Le principe de l'excitation d'une résonance de cavité est représenté sur la Figure 1.13 dans un cas sans perte (les modes sont alors localisés par des Dirac en fréquence). La façon dont le champ électrique s'organise dans la cavité dépend de la géométrie de cette dernière et de la fréquence d'oscillation. Plus la fréquence de résonance d'un mode est grande et plus le champ s'organise sur de petites dimensions (la longueur d'onde diminue). On voit par exemple que les variations spatiales du champ électrique du mode C sont plus grandes que celles du mode A ($\lambda_C < \lambda_A$ et $f_C > f_A$). Sur la droite de la Figure 1.13, un signal à la fréquence f_A est transmis à la cavité. Une onde stationnaire est ainsi créée (à cause des réflexions sur les parois de la cavité) et l'organisation spatiale du champ électrique correspond bien à celle du mode A.

Pour une cavité cylindrique la répartition spatiale du champ électrique des modes transverse électrique (TE) TE_{111} , TE_{211} et TE_{112} (champ électrique orthogonal à l'axe de la cavité) et transverse magnétique (TM) TM_{011} , et TM_{012} (champ magnétique orthogonal à l'axe de la cavité) est représentée sur la Figure 1.14. Les fréquences de résonance de ces modes sont tracées pour un diamètre constant de 150 mm en fonction de sa hauteur L . La répartition spatiale du champ reste la même pour chaque mode, mais plus la hauteur de la cavité augmente plus la fréquence associée à certains modes diminue. Par exemple à la fréquence typique de 2.45 GHz, pour que le champ s'organise dans la cavité selon la répartition du mode TE_{111} , la hauteur L de la cavité doit être d'exactement 70 mm. Pour une dimension deux fois plus grande ($L = 140$ mm) le champ sera organisé selon le mode TE_{112} . De plus si on s'intéresse maintenant à la fréquence des modes de résonance de cette cavité avec $L = 140$ mm, le nombre de modes que l'on peut exciter dans la bande de fréquence étudiée augmente fortement et ces derniers sont très rapprochés en fréquence.

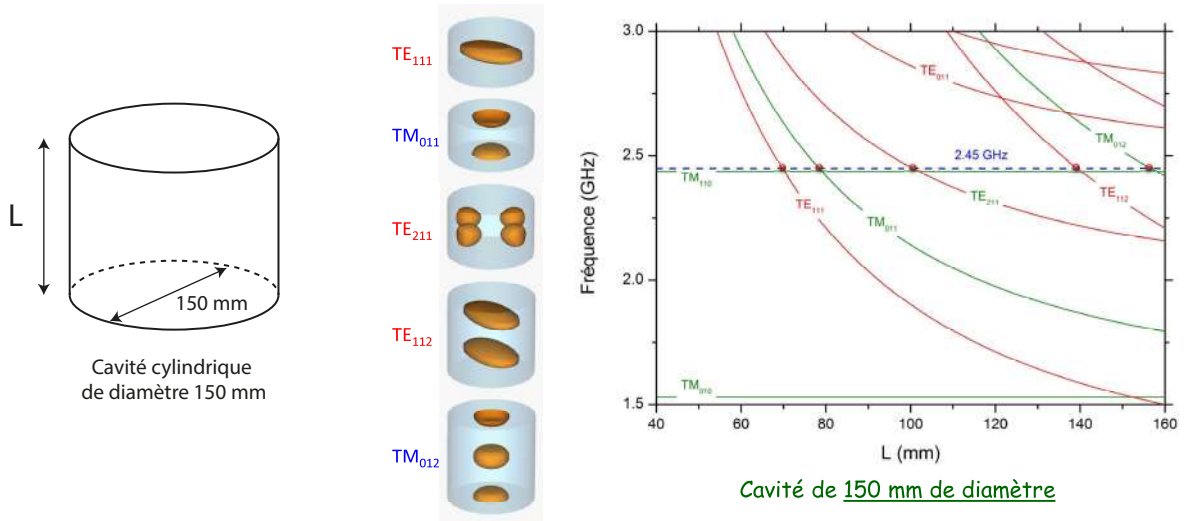


FIGURE 1.14 – Représentation de la répartition spatiale du champ électrique dans une cavité cylindrique des modes TE_{111} , TM_{011} , TE_{211} , TE_{112} et TM_{012} . Les courbes de la figure représentent les fréquences de ces modes en fonction de la hauteur L de la cavité cylindrique de diamètre 150 mm (reprise de [Sil19]).

Pour résumer, la structure du champ électrique dépend de la géométrie de la cavité ainsi que de la fréquence d'excitation. Pour une cavité excitée à une certaine fréquence (typiquement 2.45 GHz), plus les dimensions de la cavité augmentent, plus le nombre de modes potentiellement excités augmente. On dit alors que la cavité devient multi-modes et il est très difficile de contrôler la structure du champ dans ce cas (nous développerons dans le chapitre 3 les équations permettant de décrire le passage d'un régime mono-mode à un régime multi-mode). De plus, il ne suffit pas d'exciter la cavité à la fréquence du mode souhaité pour que le champ s'organise selon ce mode, la façon dont l'énergie microonde est transférée à la cavité est aussi essentielle. L'applicateur plasma microonde doit alors être conçu de sorte à ce qu'une des composantes du champ d'excitation (champ électrique ou magnétique) soit commune au champ du mode que l'on souhaite exciter.

Nous avons vu que la structuration du champ dans la cavité nécessite un travail de conception de l'ensemble applicateur + réacteur (cavité). Mais dans un cas plus réaliste, la fréquence de résonance d'un mode dépend aussi du contenu de la cavité. La présence d'une fenêtre (en quartz par exemple) et d'un plasma dans la cavité, que ce dernier se comporte comme un diélectrique ou comme un conducteur, va ainsi influencer sur la fréquence de résonance des modes ainsi que sur la structure spatiale du champ correspondante. Nous verrons dans la partie suivante, à travers un exemple de conception de cavité dans une application de croissance de diamant, que la présence du plasma (et des fenêtres) doit être prise en compte pour un contrôle efficace de la position du plasma dans la cavité.

1.2.2 Les étapes de conception d'une cavité standard : Exemple

Afin d'illustrer l'importance de la prise en compte de la présence d'un plasma sur le design de la cavité, nous allons suivre les étapes tirées de [Su+14] pour la conception d'une source plasma appliquée à la croissance de diamant. Elles sont rassemblées sur la Figure 1.15.

La première étape consiste à dimensionner l'ensemble applicateur + cavité de sorte à ce que la structure du champ soit intense au niveau du substrat à traiter pour une fréquence du champ $f_c = 2.45$ GHz. La source plasma ainsi conçue est représentée sur la première ligne de la Figure 1.15. Elle est constituée d'une cavité principale (réacteur plasma) dans laquelle le gaz et la pression sont contrôlés; d'un applicateur, séparé de la cavité principale par des fenêtres en quartz et restant à la pression atmosphérique dans l'air; d'un support sur lequel est disposé le substrat à traiter et d'un piston, permettant d'ajuster la hauteur de la cavité et donc la distribution du plasma durant le dépôt de diamant [Su+14]. Les résultats des simulations montrent que, excité à une fréquence $f_c = 2.45$ GHz, le champ est intense au niveau du substrat à traiter. On remarque cependant des zones de champ intense au niveau des coins supérieurs de la cavité, causant des plasmas parasites et donc des impuretés de carbone [Su+14].

Un nouveau design de la cavité est alors proposé en étape 2. La structure du champ est bien fidèle à celle souhaitée, *i.e.* un champ intense au niveau du substrat et des zones d'intensification non souhaitées réduites. Notons que le champ est aussi très intense à proximité

de l'applicateur. Cependant celui ci est laissé dans l'air à la pression atmosphérique pendant le processus de croissance de diamant. Cela permet d'éviter la génération de plasma à sa proximité. Notons aussi que le piston du design original a été enlevé, mais des mécanismes d'ajustement sont présents permettant de régler la hauteur de la cavité (et donc la distribution du plasma).

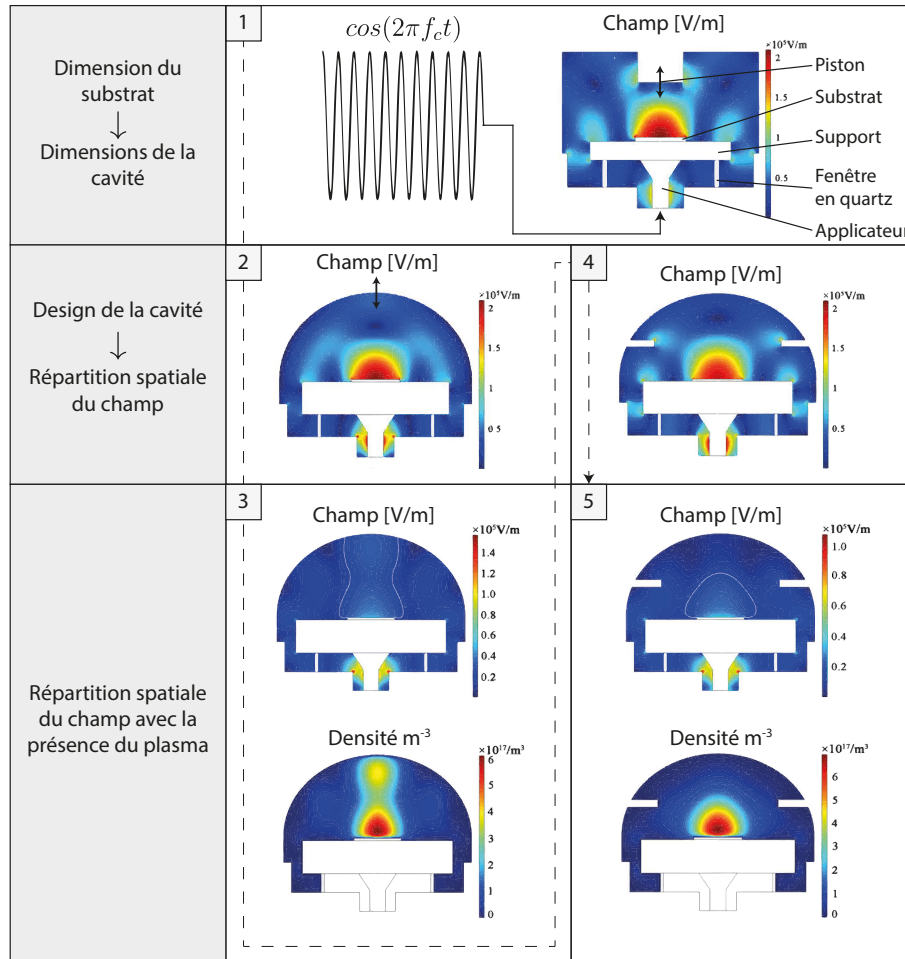


FIGURE 1.15 – Les différentes étapes de conception d'une cavité pour une application de croissance de diamant (repris de [Su+14]).

Les deux premières étapes de conception ont été réalisées uniquement en simulant un problème électromagnétique, *i.e.* en codant uniquement les équations de Maxwell. Lors de l'étape 3, les simulations sont réalisées en couplant ces équations au modèle de Drüde. Cela permet de décrire au niveau macroscopique les interactions onde - plasma. Dans ce cas, on remarque sur la structure du champ que la zone d'intensification (délimitée par des traits blancs) s'étend du substrat vers la paroi supérieure de la cavité. En conséquence, on remarque que la densité plasma obtenue dans cette cavité suit la même structure.

Afin de confiner le plasma au niveau du substrat, une nouvelle étape de conception est réalisée (étape 4) durant laquelle la cavité est "taillée" comme représenté sur la Figure 1.15.

Notons que ce changement de design ne modifie pas significativement la structure du champ dans la cavité (en comparant les étapes 2 et 4) tant que la présence de plasma n'est pas prise en compte dans les simulations du réacteur.

Finalement, l'étape 5 permet de valider le design de l'étape 4. En effet, en couplant les équations de Maxwell au modèle plasma fluide, on trouve une structure du champ intense uniquement au niveau du substrat, contrairement à la structure obtenue à l'étape 3. On observe alors que la position du plasma est bien localisée au niveau du substrat, comme souhaité.

1.2.3 Limitations inhérentes aux méthodes standards de conception

Le contrôle des plasmas microondes en cavité nécessite une grande sophistication de réalisation de la cavité et de disposer d'un système de couplage sélectif pour s'assurer que seul le mode d'intérêt sera excité. Il requiert également de prendre en compte la présence du plasma, qui va modifier l'organisation du champ dans la cavité. De nombreuses étapes de conception sont alors nécessaires, comme résumé sur la Figure 1.16.

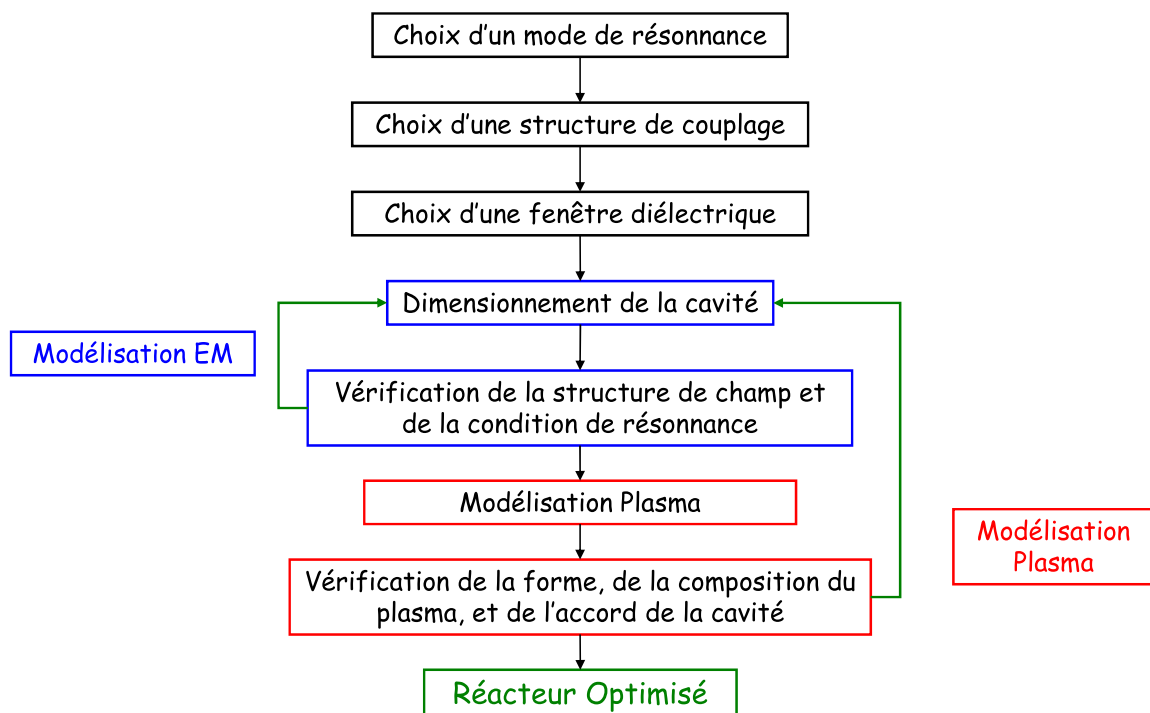


FIGURE 1.16 – Les étapes de la conception d'une source plasma microonde [Sil+09].

Tout d'abord, pour une certaine géométrie de cavité, le mode de résonance que l'on souhaite exciter (*i.e.* la structure globale du champ) est choisi. Il faut ensuite déterminer la structure de couplage (applicateur), permettant d'exciter le mode de résonance de la cavité et de transférer l'énergie microonde du générateur au plasma. Une des composantes du champ

d'excitation doit alors être commune au champ du mode que l'on souhaite exciter. Puis une fenêtre diélectrique peut être introduite dans la cavité, permettant d'isoler la partie de la cavité dans laquelle on souhaite générer un plasma. Cette étape permet d'éviter d'amorcer des plasmas parasites à proximité de l'applicateur ou dans des endroits de fort champ dans la cavité.

Ensuite, une modélisation électromagnétique de la cavité est réalisée. L'objectif est d'obtenir la structure du champ souhaitée, associée à un phénomène de résonance, en prenant en compte le système de couplage et la présence de diélectriques (fenêtres) dans la cavité.

L'étape suivante consiste à prendre en compte la présence du plasma dans la cavité sur la répartition du champ. On inspecte alors la forme et la composition du plasma. Si elles ne correspondent pas aux objectifs souhaités, comme c'est le cas par exemple dans l'exemple présenté dans la partie précédente, alors il faut revenir à l'étape modélisation électromagnétique. Enfin, lorsque la forme et la composition du plasma correspondent à l'objectif souhaité, on peut considérer que la source microonde est optimisée.

La précision de réalisation de la source plasma, ainsi que le couplage complexe entre le plasma et le champ électromagnétique rendent difficile la conception d'une source. Cela nécessite de nombreuses étapes de conception. De plus, une fois le design terminé la position du plasma est fixée une fois pour toute dans le réacteur. Même si des parois ajustables permettent parfois de modifier la cavité (comme dans l'exemple précédent), pour chaque nouvelle configuration, il est nécessaire de repasser par les étapes de conception que nous venons de décrire.

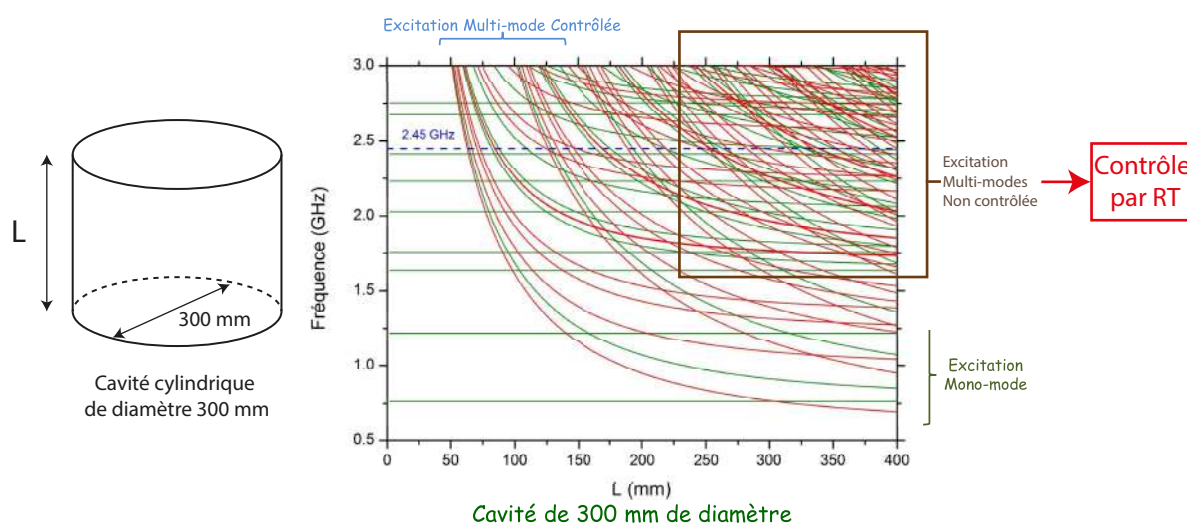


FIGURE 1.17 – Les courbes représentent les fréquences des modes en fonction de la hauteur L de la cavité cylindrique de diamètre 300 mm représentée à gauche (reprise de [Sil19]).

Enfin, notons que typiquement, la fréquence d'excitation est fixe : de 915 MHz ou 2.45 GHz [Leb15]. La dimension du substrat à traiter est alors limitée. En effet, pour ces fréquences,

augmenter la taille de la cavité rend la cavité multi-modale. Dans ce cas, la structure du champ dans la cavité n'est pas contrôlable par ces techniques modales. Cet effet est illustré sur la Figure 1.17. La cavité est la même cavité cylindrique que précédemment (Figure 1.14) mais avec un diamètre plus grand (de 300 mm). Dans ce cas, on remarque que pour des fréquences supérieures à environ 1.5 GHz et une dimension L supérieure à environ 200 mm, il n'est pas possible d'exciter seulement un unique mode. Même avec un système de couplage sélectif et un design optimisé, les fréquences des modes sont trop proches les unes des autres. Notons en plus que la source microonde d'excitation possède une certaine largeur spectrale (typiquement 50 MHz pour un magnétron [Vya+16]). Il n'est donc pas possible de contrôler la structure du champ dans la cavité avec les techniques modales classiques. C'est un problème qui a été identifié dans la feuille de route que nous avons déjà mentionnée, dans laquelle il est indiqué : “electromagnetic effects make design of conventional ICP and microwave sources challenging at larger dimensions. [...] there is an increased need to explore and develop plasma technologies that are more readily scalable to larger dimensions” [Sam+12].

Dans un article de review de 2014 portant sur les sources plasmas microondes, il est clairement spécifié, en parlant des sources plasmas en cavité multi-mode, que “the electromagnetic aspects of these microwave enhanced processing systems need to be addressed in a radically different manner to reach successful applications” [Stu+14]. Et c'est ce que nous proposons dans cette thèse. Comme nous allons le voir en détail dans le chapitre 3, le contrôle de la structure transitoire du champ et donc des plasmas est possible par RT dans les cavités à grande dimension. Nous verrons que le caractère multi-modal de la cavité est même nécessaire pour un bon contrôle des plasmas par RT.

L'objectif de la thèse consiste à éliminer toutes les étapes de conception représentées sur la Figure 1.16, en rendant possible le contrôle dynamique de la position des plasmas dans des cavités de géométrie quelconque et de grandes dimensions. Nous verrons dans la partie suivante que pour arriver à cet objectif, il est intéressant de travailler sur la forme d'onde que l'on transmet à la cavité.

1.3 La source plasma à retournement temporel : Solution ultime ?

L'objectif de la thèse est de proposer une nouvelle façon de contrôler les plasmas dans des cavités, afin de s'affranchir de toutes les étapes présentées dans la partie précédente et de rendre le contrôle des plasmas en cavité multi-mode possible. En regardant les techniques classiquement utilisées pour les sources plasmas microondes (présentées dans la partie précédente), on s'aperçoit qu'elles reposent sur l'excitation en régime CW (ou pulsé) d'un mode de résonance de la cavité. Le signal transmis à la cavité est alors un simple signal sinusoïdal et le plasma est maintenu en structurant la cavité de sorte à ce que la répartition du champ électrique maintienne la plasma à l'endroit souhaité.

Cependant, il se trouve que la forme d'onde du signal transmis à la cavité détermine le comportement des ondes en son sein. On peut alors se demander s'il ne serait pas judicieux de jouer sur cette forme d'onde afin de contrôler l'organisation spatio-temporelle du champ dans la cavité, et donc du plasma. L'idée de la nouvelle source plasma développée durant cette thèse repose sur ce principe, et nécessite donc la manipulation de signaux plus complexes. Son principe est représenté sur la Figure 1.18. La forme du signal $s_{transmis}(t)$ transmis à la cavité détermine la façon dont l'énergie électromagnétique se répartit dynamiquement dans la cavité. En contrôlant astucieusement la forme de ce signal, il est alors possible de contrôler la position du plasma dans le temps, comme un "pinceau plasma ondulatoire". C'est comme si les *ondes* jouaient le rôle d'un *pinceau* qui déplace la position du *plasma*.

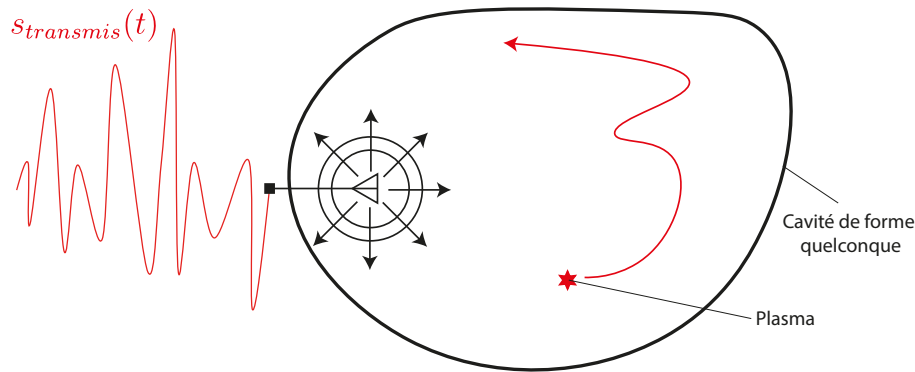


FIGURE 1.18 – Principe de la source plasma développée dans cette thèse : le “pinceau plasma ondulatoire”. La position du plasma dans la cavité est contrôlée par la forme du signal transmis à la cavité $s_{transmis}(t)$, qui structure dynamiquement le champ dans la cavité.

Pour atteindre cet objectif, il est nécessaire d'utiliser une technique élaborée de contrôle des ondes, permettant le contrôle spatio-temporel de l'organisation du champ dans la cavité. Cette technique doit permettre non seulement le contrôle du champ mais elle doit aussi comprendre un phénomène d'intensification du champ à l'endroit et au moment souhaités pour amorcer le plasma. Nous allons voir durant cette thèse que le RT répond tout à fait à ces critères et permet bien d'amorcer et de contrôler les plasmas en cavité. C'est une technique développée dans les années 1990 qui permet de focaliser spatio-temporellement

les ondes dans un environnement réverbérant. Nous ferons dans la partie suivante un bref historique des travaux fondateurs de cette technique. Cela permettra de faire sentir au lecteur les mécanismes sur lesquels elle est basée (une étude plus détaillée de la théorie autour du RT est proposée au chapitre 2) et de comprendre ce qu'elle peut faire en terme de focalisation de l'énergie. Nous verrons alors que cette technique requiert la manipulation de signaux de courte durée et de formes complexes. Nous discuterons ensuite des caractéristiques de la source plasma à RT, qui présente la possibilité de contrôler dans le temps la position des plasmas (amorçage transitoire). Finalement, le principe et les avantages a priori de la source plasma à RT développée durant cette thèse sont présentés.

1.3.1 Le retournement temporel : Chronologie des travaux fondateurs

L'objectif de cette partie est de faire un bref rappel sur les travaux fondateurs qui ont été menés à propos du RT, dont la paternité revient à Mathias Fink, qui a dirigé toutes les thèses présentées dans cette partie. Nous reviendrons dans la suite de cette thèse plus en détails sur la physique du RT. En effet le chapitre 2 de cette thèse réalise une étude des concepts théoriques sur lesquels est basé le RT dans le domaine microonde et dans le chapitre 3 la théorie autour de sa mise en application en cavité réverbérante est ensuite développée.

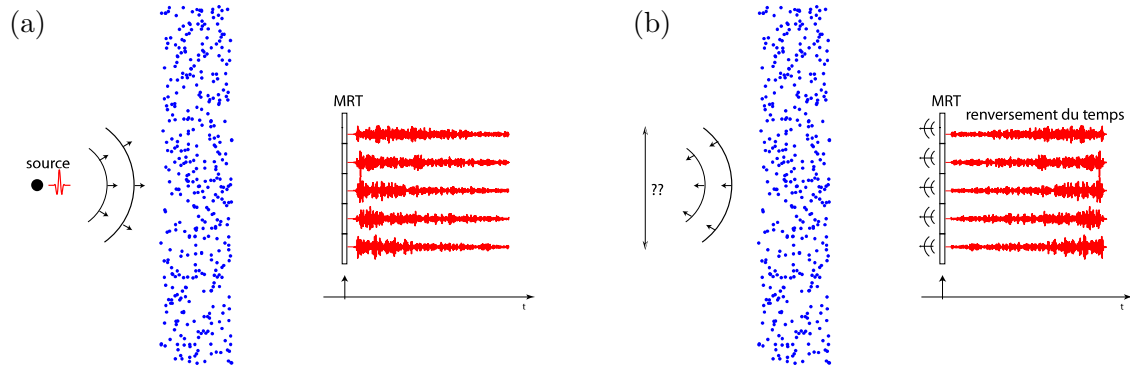


FIGURE 1.19 – (a) Phase d'enregistrement du champ. (b) Phase de réémission du champ retourné temporellement (repris de [Lem11]).

Les premières études sur l'exploitation des propriétés spatio-temporelles du RT ont été menées dans les années 1990 par Claire Prada dans le domaine acoustique dans le cadre de ses travaux de thèse, soutenue en 1991 [Pra91] (on y trouve, en 1989, la première occurrence du terme “time reversal”, en tant que dispositif permettant de focaliser les ondes [Fin+89]). Ces études ont vu le jour avec la nécessité croissante de contrôler les ondes dans des milieux complexes. À l'époque, la technique de conjugaison de phase (qui avait d'abord été développée en optique puis transposée en acoustique) permettait de contrôler et de focaliser spatialement des ondes dans des milieux complexes. Elle nécessite l'utilisation de nombreux transducteurs et fonctionne uniquement avec des ondes quasi monochromatiques. Les travaux de thèse de Claire Prada proposent d'étendre cette faculté de focalisation à des signaux de courte durée (*i.e.* des signaux larges bandes) en utilisant le RT. Le RT est constitué de deux étapes. La première consiste à “sonder” un milieu de propagation complexe en mesurant l'information de propagation des ondes dans ce milieu. Cette information est ensuite retournée temporellement

et transmise au milieu, donnant lieu à une focalisation spatio-temporelle. Depuis ces travaux pionniers, de nombreuses études ont été réalisées sur le RT pour comprendre les mécanismes et exploiter au mieux ses propriétés de focalisation spatio-temporelle.

Durant la thèse d'Arnaud Derode (soutenue en 1994) le RT a été appliqué dans le domaine acoustique à des milieux complexes dits désordonnés [Der94] ; [DRF95]. Son principe est illustré sur la Figure 1.19. Le milieu de propagation est constitué de nombreuses tiges (représentées en bleu). Une impulsion est tout d'abord émise à l'extérieur de la forêt de tiges (Figure 1.19(a)). Les ondes qui se propagent, sont diffusées par ce milieu complexe et sont mesurées par un réseau de transducteurs (au nombre de 96 dans cette expérience), appelé le miroir à retournement temporel (MRT). Lors de la deuxième étape, l'information de la propagation des ondes ainsi mesurée est retournée temporellement et transmise par les transducteurs. Il en résulte une focalisation spatio-temporelle au niveau de l'émission initiale de l'impulsion. Ainsi, il est possible de contrôler le comportement des ondes en milieu désordonné, sans connaissance a priori du milieu, qui peut être de forme et de géométrie complexes.

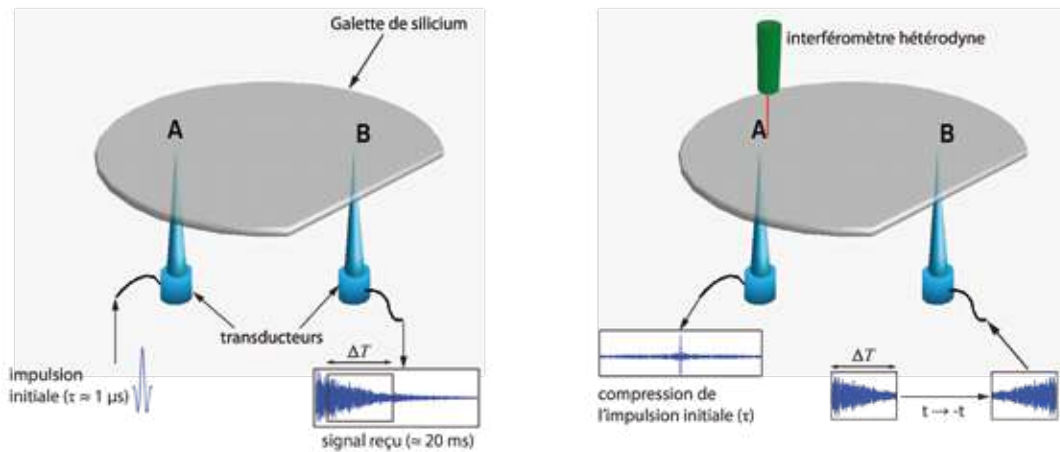


FIGURE 1.20 – Les deux étapes d'une expérience de retournement temporel mono-voie en cavité (repris de [Lem11]).

Ensuite, Carsten Draeger a montré durant sa thèse (soutenue en 1997) la possibilité de réduire considérablement le nombre de transducteurs dans le MRT à un unique transducteur, en tirant profit des multi-réflexions des ondes en environnement réverbérant [DRA97]. Ces premières expériences de RT en cavité réverbérante ont été réalisées dans le domaine acoustique avec pour cavité le disque de silicium tronqué représenté sur la Figure 1.20. Lors de la première phase, une impulsion de $\tau = 1 \mu s$ est transmise au niveau du transducteur A. L'information contenue dans le signal reçu au niveau du transducteur B s'étale sur une durée d'environ 20 ms, considérablement plus longue à cause des réflexions des ondes sur les parois de la cavité. Cette information (ou une partie de durée ΔT) est retournée temporellement et transmise au niveau du transducteur B. Il en résulte une focalisation spatio-temporelle de durée égale à celle de l'impulsion initiale τ et de dimension égale à la demi-longueur d'onde. On retrouve bien cette caractéristique sur la cartographie du champ au moment de la focalisation (pour $t = 0 \mu s$) de la Figure 1.21. Le disque tronqué dans lequel ces expériences ont eu

lieu constitue en fait une cavité chaotique, dont l'importance sur la qualité de la focalisation a d'abord été évoquée par Carsten Draeger [DF97]. Nous reviendrons sur l'importance de la chaoticité sur le contrôle des plasmas par RT dans le chapitre 3.

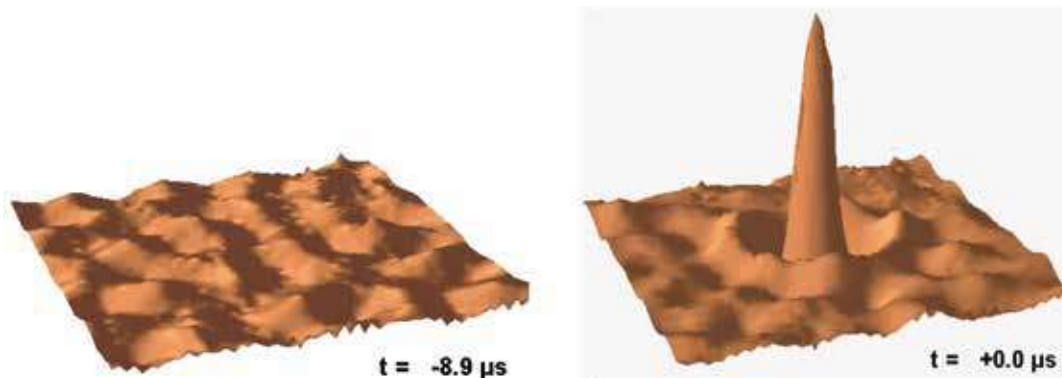


FIGURE 1.21 – Cartographie du déplacement à la surface de la galette de silicium à un instant quelconque (à gauche) et à l'instant de la focalisation temporelle (à droite). (repris de [Lem11]).

En regardant la première phase des expériences en cavité réalisées par Carsten Draeger de la Figure 1.20(a), on pressent que l'information de la propagation des ondes partout dans la cavité est contenue dans le signal mesuré au niveau du transducteur B. C'est entre autre sur ce point que s'est concentré Julien de Rosny durant sa thèse (soutenue en 2000) [Ros00]. Il a alors analysé finement dans les domaines temporel et fréquentiel l'information contenue dans ce signal, en utilisant notamment le terme de “grains d'information” introduit par Arnaud Derode *et al.* [DTF99]. Ce terme désigne la quantité d'information décorrélée contenue dans ce signal, et conditionne la qualité de la focalisation par RT. On trouve aussi une étude du comportement de ces “grains” en fonction de la chaoticité de la cavité. Dans ce présent manuscrit et notamment dans les développements effectués au chapitre 3, nous nous appuyerons amplement sur les développements réalisés par Julien de Rosny, qui permettent une compréhension fine des mécanismes avec cette vision duale temps - fréquence.

Les propriétés de focalisation spatiale et temporelle du RT ont ensuite été démontrées dans le domaine microonde par Geoffroy Lerozey durant sa thèse (soutenue en 2006) [Ler06] ; [Ler+04]. Il a en particulier montré la possibilité de focaliser spatio-temporellement l'énergie électromagnétique en cavité métallique sur des durées égales à celle de l'impulsion et sur des dimensions égales à la demi-longueur d'onde. De plus, le dispositif expérimental de contrôle des ondes qu'il a développé permet la génération de signaux hautes fréquences en manipulant seulement l'information en bande de base, en s'appuyant sur une technique de modulation. Suite à ces expériences, Matthieu Davy a étudié durant sa thèse (soutenue en 2010) la possibilité d'utiliser le RT dans l'espoir de développer un “bazooka électromagnétique”. Pour ce faire, il a utilisé une cavité ouverte pour focaliser les ondes à l'extérieur de celle-ci [Dav10]. Depuis ces premières expériences, le RT a été appliqué dans le domaine microonde pour des applications de communication sans fils [Kha10] ; [LY19], de localisation de source [FA11], et de transfert de puissance sans fils [Ibr+16] ; [Li+19]. Notons aussi que le RT a été transposé au domaine de l'optique [McC+11]. Nous verrons dans le chapitre 4 de ce manuscrit que

pour la nouvelle application que nous en proposons, à savoir le contrôle plasma par RT, notre dispositif expérimental s'inspire largement du dispositif développé par Geoffroy Lerosey.

Pour finir, notons que dans les travaux plus récents tels les thèses de Fabrice Lemoult (soutenue en 2011) [Lem11] et celle de Matthieu Dupré (soutenue en 2015) [Dup15] portant sur l'amélioration du contrôle des ondes en milieux complexes, le concept de "grains d'information" est englobé dans un concept plus général : celui de "degrés de liberté". Nous verrons dans le développement du chapitre 3 de ce manuscrit que nous avons aussi préféré parler de degré de liberté.

1.3.2 La source plasma à retournement temporel : Les caractéristiques

Nous avons vu dans la partie précédente que le retournement temporel permet de focaliser spatio-temporellement l'énergie électromagnétique sur une durée égale à celle de l'impulsion initiale et sur une dimension égale à la demi-longueur d'onde. L'objectif de cette thèse est de développer une source plasma basée sur ce principe, dont les avantages par rapport aux techniques actuelles sont présentés dans la partie suivante. Il faut cependant noter que l'utilisation d'impulsions de courtes durées pour le RT conduit à une durée d'application du champ électromagnétique très courte (de l'ordre de la nanoseconde dans nos expériences). On parlera ainsi d'amorçage transitoire (par opposition à l'amorçage en CW présenté dans la première section de ce chapitre). Cela permet, en plus du contrôle de la position du plasma à l'endroit de la focalisation, de contrôler la position du plasma dans le temps. Le maintien des plasmas ainsi obtenus est rendu possible par un fonctionnement en régime pulsé. Nous explorerons tout d'abord dans cette partie les conséquences de cet amorçage transitoire puis nous discuterons des caractéristiques du régime pulsé.

1.3.2.1 Amorçage en régime transitoire

Nous avons discuté dans la première section de ce chapitre du claquage microonde en régime permanent. Or, le claquage du gaz est par définition un phénomène qui se produit en un temps donné. Le champ de claquage estimé dans la première section de ce chapitre constitue donc un champ minimal. L'augmentation du champ électrique appliqué va engendrer une diminution du temps de claquage. De ce fait, la valeur du champ de claquage dépend de la durée d'application du champ.

Le temps de claquage est en général défini comme la somme d'un retard statistique (présence d'électrons libres au moment où le champ est appliqué) et d'un temps de formation du plasma. Nous ne nous intéressons ici qu'au temps de formation du plasma.

Les expériences d'amorçage de plasmas par RT présentées dans cette thèse ont été réalisées dans de l'argon à des pressions de l'ordre du torr avec des impulsions microondes autour d'une fréquence centrale $f_c = 2.4$ GHz et de durée initiale $\tau = 8$ ns. Le champ électrique est donc intensifié à l'endroit et au moment de la focalisation sur une dimension de $\lambda_c/2 = 6$ cm

pendant une durée de $\tau = 8$ ns. Le principe de l'amorçage des plasmas par RT consiste à ce que ce champ électrique de focalisation soit suffisant pour amorcer un plasma par les processus de chauffage des électrons et d'ionisation que nous avons décrits dans la première section de ce chapitre. On peut alors parler d'amorçage transitoire. Nous présenterons dans cette partie les propriétés du plasma dans ces conditions. Les données plasma sont obtenues avec Bolsig + [HP05].

Afin de décrire les premiers instants de l'amorçage d'un plasma par RT, le terme de diffusion de l'équation de bilan de particules développée dans la première section de ce chapitre (équation 1.16) peut être négligé en première approximation. On trouve alors, pour une densité initiale n_0 d'électrons une évolution exponentielle de la densité :

$$n_e(t) = n_0 \cdot e^{\nu_i t} \quad (1.33)$$

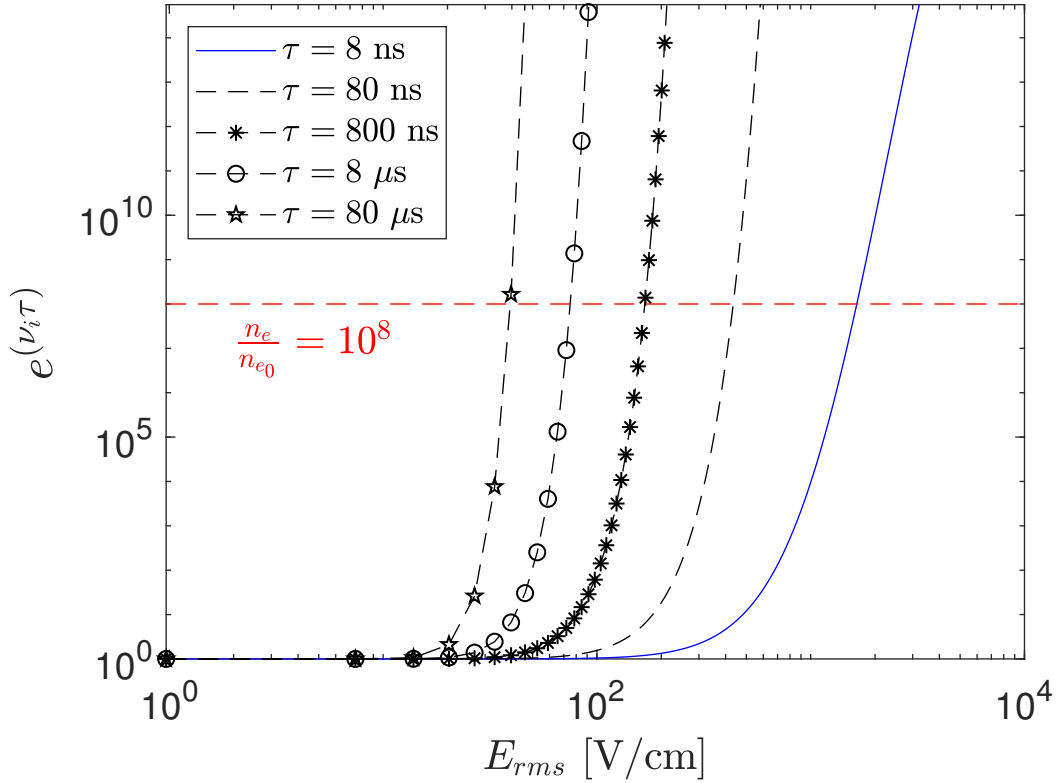


FIGURE 1.22 – Valeur de l'exponentielle $e^{\nu_i \tau}$.

Pour qu'un plasma soit généré, il faut que $n_e(t + \Delta t) > n_e(t)$ et donc que $\tau > 1/\nu_i$. Autrement dit, il faut que la durée d'application du champ soit plus longue que le temps caractéristique entre deux ionisations. En effet, si ce n'est pas le cas, les électrons n'ont même pas le temps d'ioniser les atomes que le champ est cessé d'être appliqué. Plus précisément, la

multiplication critique au-delà de laquelle on considère qu'il y a claquage est en général (et de façon assez empirique) prise égale à $n_e/n_{e0} = 10^8$ [GR56].

Comme nous l'avons vu dans la première section de ce chapitre, la fréquence d'ionisation dépend du gaz, de sa pression p et de la vitesse des électrons (et donc du niveau du champ électrique). Dans l'argon à $p = 1$ torr, la valeur de l'exponentielle $e^{\nu_i \tau}$ est tracée sur la Figure 1.22 pour différents τ en fonction de l'amplitude du champ microonde. On retrouve ici clairement le caractère non-linéaire de l'amorçage du plasma avec le champ électrique. En dessous d'un certain champ électrique, la valeur de l'exponentielle de l'équation 1.33 n'est pas assez élevée pour augmenter significativement n_e de sorte à obtenir l'avalanche électronique. Pour $\tau = 8$ ns, on observe un changement de régime entre 100 et 1000 V/cm et une valeur du champ pour obtenir une multiplication électronique de 10^8 d'environ 1500 V/cm. De plus, on voit que le claquage du plasma dépend du temps d'application du champ τ . Plus ce temps est long et plus le niveau du champ requis pour claquer le gaz est faible. Ces courbes sont en accord avec les résultats expérimentaux disponibles dans la littérature [Fel66]. Nous verrons dans le chapitre suivant comment cette propriété de non linéarité du claquage du gaz se traduit sur le processus du RT, qui est lui basé sur une physique linéaire.

1.3.2.2 Intérêt des plasmas pulsés nanosecondes

Nous venons donc de voir que les plasmas amorcés par RT sont des plasmas transitoires, *i.e.* avec une durée d'application du champ très courte d'environ environ 8 ns dans notre cas. Nous verrons au chapitre 4 que les plasmas amorcés par RT expérimentalement sont en fait maintenus en régime pulsé, c'est à dire que la focalisation spatio-temporelle de l'énergie microonde par RT est répétée dans le temps à une certaine fréquence de répétition f_{rep} . L'utilisation d'une source plasma microonde pulsée possède plusieurs avantages par rapport aux sources plasmas microondes CW. En effet, alors que les collisions inélastiques sont nécessaires pour l'amorçage d'un plasma, le chauffage du gaz et des parois sont des processus limitants. L'utilisation de puissance pulsée permet de limiter favorablement ces pertes. La modulation temporelle de la puissance d'injection offre ainsi différentes capacités de contrôle des caractéristiques des plasmas [Rou+94]. L'interruption de la puissance entre des impulsions répétitives permet, par exemple, de restreindre le chauffage du gaz [Car+14] ou de contrôler des processus chimiques présents dans le volume plasma et sur les surfaces [Mar09] ; [Has+01] ; [HMK02] ; [Cun+12] ; [Bou+06]. Dans ce cas, l'intensité du champ électrique, la durée de l'impulsion τ , la fréquence de répétition f_{rep} , les temps de montée t_r et de descente, ainsi que le rapport cyclique D , défini comme $D = \tau f_{rep}$ sont les principaux paramètres qui contrôlent la production et le maintien du plasma. Des investigations minutieuses sur les plasmas microondes pulsés ont été menées sur des sources exploitant des ondes de surface. Par exemple, Hubner *et al.* et Carbone *et al.* ont étudié des décharges microondes pulsées avec τ descendant jusqu'à 1 μ s, t_r descendant jusqu'à quelques dizaines de nanosecondes et $D > 10$ % [Hü+12] ; [Car+14]. Comparé à ces travaux, nous verrons dans le chapitre 4 que dans le cas des plasmas amorcés par RT, le rapport cyclique, le temps de montée et la durée d'application du champ sont bien plus faibles. Cela rend cette nouvelle source plasma microonde pulsée très originale du point de vue de la physique des mécanismes ayant lieu durant une impulsion et entre deux

impulsions consécutives, comme nous le verrons à la fin du chapitre 4.

1.3.3 La source plasma à retournement temporel : Les avantages

Finalement, nous avons vu que le RT permet de focaliser spatio-temporellement l'énergie électromagnétique dans une cavité réverbérante. Nous proposons dans cette thèse une nouvelle source plasma tirant profit de cette propriété : la source plasma à RT. Son principe est représenté sur la Figure 1.23. Il nécessite deux phases. Lors de la première phase, l'information de la propagation entre plusieurs endroits dans la cavité (notés $\mathbf{r}_1$, $\mathbf{r}_2$ et $\mathbf{r}_3$), et une antenne est mesurée. Les signaux ainsi mesurés sont notés $r_1(t)$, $r_2(t)$ et $r_3(t)$. Ces signaux sont ensuite retournés temporellement avant d'être transmis à la cavité par l'antenne lors de la seconde phase. De l'émission des signaux $r_1(-t)$, $r_2(-t)$ ou $r_3(-t)$ résulte une focalisation spatio-temporelle respectivement aux points $\mathbf{r}_1$, $\mathbf{r}_2$ et $\mathbf{r}_3$. Le principe de la source plasma à RT consiste à obtenir un plasma à l'endroit et à l'instant de la focalisation, que l'on peut déplacer dans la cavité en changeant le signal transmis. Par exemple sur la Figure 1.23(b) ces trois signaux sont transmis successivement à l'antenne. Le plasma "passera" alors du point $\mathbf{r}_1$ au temps t_{rt1} , puis au point $\mathbf{r}_2$ au temps t_{rt2} et enfin au point $\mathbf{r}_3$ au temps t_{rt3} . Dans ce cas la position du plasma est entièrement contrôlée par la forme d'onde transmise à la cavité. Le RT est donc un bon candidat pour répondre aux objectifs de la thèse et développer le "pinceau plasma ondulatoire".

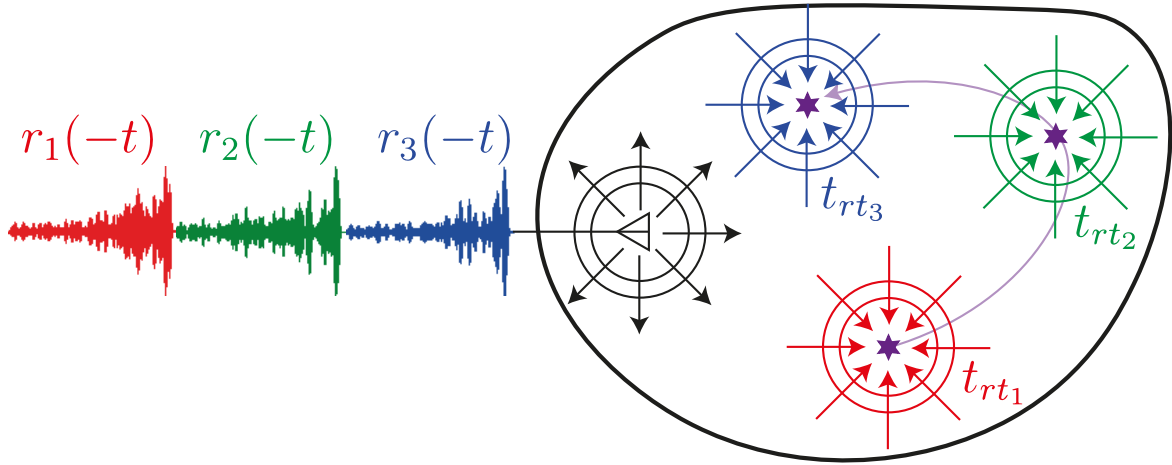


FIGURE 1.23 – Principe de la source plasma à RT. L'émission successive des trois signaux retournés temporellement $r_1(-t)$ - $r_2(-t)$ - $r_3(-t)$ permet d'obtenir un plasma (représenté par une étoile violette) à un temps t_{rt1} au point $\mathbf{r}_1$, puis au temps t_{rt2} au point $\mathbf{r}_2$ et enfin à un temps t_{rt3} au point $\mathbf{r}_3$ (avec t_{rt1} , t_{rt2} et t_{rt3} les instants de la focalisation par RT).

Finalement cette source plasma possède théoriquement les avantages suivants :

- ✓ Contrôle de la position du plasma quasi-indépendant du design de la cavité.
- ✓ Contrôle d'autant meilleur que les dimensions de la cavité augmentent (de plus le contrôle est d'autant meilleur que le désordre dans la cavité est important).
- ✓ Dimensions du substrat non limitées.

- ✓ Contrôle **dans l'espace et dans le temps** du plasma.
- ✓ Possibilité d'utiliser n'importe quelle cavité pour plusieurs applications.

L'objectif de la thèse est de développer expérimentalement cette source plasma à RT. Pour cela, il est nécessaire de présenter la théorie autour du RT, afin de comprendre tout d'abord sur quels principes physiques le RT repose et ensuite comment l'introduction d'un milieu non-linéaire (le plasma) va affecter ce processus, basé sur une physique linéaire. Cette étude fait l'objet du chapitre 2. Nous verrons ensuite dans le chapitre 3, à l'aide d'un code de simulation FDTD, comment le RT se met en place en cavité réverbérante, et nous illustrerons l'importance de la chaoticité de la cavité sur le contrôle plasma par RT. Finalement, dans le chapitre 4 les plasmas obtenus avec la source plasma à RT développée durant cette thèse sont présentés.

CHAPITRE 1 : Ce que qu'il faut retenir

Nous avons vu dans ce chapitre que les techniques actuelles de conception des sources de plasma à microondes possèdent des limitations. Tout d'abord, leur conception demande une grande précision sur la géométrie du dispositif (cavité + applicateur). De plus, de nombreuses étapes sont nécessaires, durant lesquelles il faut développer des outils décrivant l'influence du plasma sur le comportement du champ électromagnétique dans la cavité. Enfin, pour une certaine fréquence d'excitation, le volume de la cavité est limité, et en conséquence la dimension du substrat à traiter aussi.

La source plasma à RT que nous proposons dans cette thèse est en rupture avec les dispositifs actuels. Elle permet de s'affranchir de toutes les étapes de conception habituellement requises en offrant un contrôle dynamique et adaptable de la position du plasma. En outre, elle permet ce contrôle en régime multimodal, et surtout il est d'autant meilleur que le volume de la cavité est grand. Cette nouvelle source rend alors possible le traitement de substrats de grandes dimensions.

Pour toutes ces raisons, cette nouvelle source constitue une potentielle révolution dans le domaine des plasmas microondes et des procédés associés. Cette thèse reporte la première étude sur ce genre de source. Comme pour tout travail pionnier, il y a donc de nombreux défis à relever. Nous avons tenté de les identifier en les rassemblant selon trois axes principaux :

- 1) Études théoriques.
 - (a) Comprendre le RT en présence de milieux linéaires ou non-linéaires (plasmas).
 - (b) Identifier les propriétés du système qui fixent les performances en amorçage et en contrôle des plasmas.
- 2) Premières études expérimentales de la source avec du matériel générique.
 - (a) Proposer une architecture de banc expérimental permettant, à la fois la manipulation de formes d'ondes complexes pour le RT et le contrôle de l'environnement (gaz, pression) pour l'amorçage des plasmas.
 - (b) Choisir des initiateurs pour abaisser les niveaux de puissance requis et mesurer les signaux nécessaires.
 - (c) Contrôler l'amorçage et l'entretien (pulsé) des plasmas par RT sur des initiateurs.
 - (d) Étudier la physique (ns) de ce nouveau type de décharge.
- 3) Développement de la source (pinceau plasma ondulatoire) : vers du matériel dédié.
 - (a) Développer un prototype de source autorisant une étude expérimentale dans des conditions optimales.
 - (b) Décrocher le plasma de l'initiateur puis contrôler le plasma partout dans la cavité.
 - (c) Étudier l'influence de la présence d'un plasma dans la cavité sur le RT (non-linéaire).
 - (d) Étudier les procédés qui pourront exploiter cette nouvelle source.

Cette thèse contribue à répondre à ces défis. Nous verrons à la fin de ce manuscrit, à l'issue du bilan des travaux réalisés, les défis auxquels nous avons répondu et ceux qui restent à explorer.

La théorie autour du retournement temporel

Sommaire

2.1	Les origines du retournement temporel : La flèche du temps	51
2.1.1	Le temps, source d'inspiration et de fascination	51
2.1.2	La flèche du temps	53
2.1.3	Les prémisses du retournement temporel : Le paradoxe de réversibilité . .	55
2.1.4	Le principe des démons de Loschmidt appliqué aux ondes	58
2.2	Le retournement temporel en électromagnétisme	62
2.2.1	La causalité en électromagnétisme	62
2.2.2	La réversibilité en électromagnétisme	66
2.2.3	La réciprocité en électromagnétisme	71
2.2.4	Exemples	75
2.3	Le retournement temporel en pratique	79
2.3.1	Le retournement temporel en pratique : Réversible, réciproque, les deux ?	79
2.3.2	Le retournement temporel appliqué au contrôle plasma	82

Time is the worst place, so to speak, to get lost in, as Arthur Dent could testify, having been lost in both time and space a good deal. At least being lost in space kept you busy.

Douglas Adams

Nous avons vu dans le premier chapitre en quoi l'utilisation de signaux d'excitation sinusoïdaux possède des limites en terme de contrôle spatio-temporel d'un plasma microonde. L'objectif principal de la thèse consiste donc à utiliser des techniques basées sur des formes d'ondes plus complexes, permettant un meilleur contrôle de ces plasmas. L'idée est d'utiliser une technique appelée RT. Mais qu'est ce que le Retournement Temporel ? Sur quels principes physiques repose-t-il ? Ce chapitre a pour objectif de répondre, au moins partiellement, à ces questions.

Tout d'abord dans la section 2.1, nous allons plonger dans le temps. Il est en effet évident qu'avant même de pouvoir prétendre comprendre ce qu'est le RT, il faut d'abord être au clair sur ce qu'est le *temps*, pour pouvoir espérer le retourner. Nous verrons dans une première partie, à travers quelques exemples, que ce concept est très compliqué à définir et qu'il est bien mystérieux. Puis nous verrons que les physiciens distinguent ce qu'ils ont appelé le cours du temps de la flèche du temps. Il est important de comprendre la distinction entre ces deux concepts, puisque le RT prétend pouvoir retourner la flèche du temps et non pas son cours. Le RT n'est ainsi pas un moyen de remonter le temps tel que son nom pourrait peut être le laisser penser et ne remet pas en question le principe de causalité (du moins pas directement). Il permet "juste" de faire se dérouler des phénomènes en inverse. Cette capacité de certains phénomènes à se dérouler identiquement dans un sens ou dans l'autre est appelée *réversibilité*. Nous verrons que pour comprendre ce principe, il est intéressant de "retourner" plus d'un siècle en arrière, aux discussions de l'époque entre Boltzmann et Loschmidt. C'est en effet à partir de ces discussions que des physiciens français, à la fin du XX^{ème} siècle, ont développé ce principe du RT. À la fin de cette partie, nous verrons enfin rapidement, comment le RT se met en place en pratique.

Ensuite dans la section 2.2, nous étudierons comment le concept de réversibilité se traduit dans le domaine de l'électromagnétisme. En pratique, il est souvent plus intéressant de parler de "réversibilité restreinte". Nous montrerons ensuite que le RT s'appuie non seulement sur la réversibilité des équations de Maxwell, mais aussi sur un autre concept, la *réciprocité*. Dans la plupart des travaux autour du RT, ce dernier concept est souvent invoqué mais rarement explicité de façon claire. Une explication intuitive de ce concept ainsi que son lien avec la réversibilité sont ensuite proposés à travers des exemples simples. Finalement, nous verrons comment ces principes s'appliquent sur le cas pratique de deux antennes dipôles. Nous en concluons que dans ce cas, le RT est quelque part plus lié de la notion de réciprocité qu'à celle de réversibilité.

Finalement, dans la section 2.3.2 nous présenterons le principe de l'amorçage de plasmas par RT. Nous verrons alors comment les concepts de réversibilité, réversibilité restreinte et réciprocité présentés dans la section précédente (section 2.2) s'appliquent dans ce cas. Nous parlerons de réversibilité et de réciprocité transitoires, conséquences de la non-linéarité du claquage du gaz.

2.1 Les origines du retournement temporel : La flèche du temps

2.1.1 Le temps, source d'inspiration et de fascination

“Qu’est-ce donc que le temps ? Si personne ne m’interroge, je le sais ; si je veux répondre à cette demande, je l’ignore” écrivait déjà Saint Augustin entre 397 et 401 [Aug97]. Et même encore de nos jours, ce terme qui nous paraît tous familier est pourtant impossible à définir sans inéluctablement utiliser la notion de temps. Citons par exemple Jean Giono : “Le temps, c’est ce qui passe quand rien ne se passe”. Dans cette définition, la notion de “passer” contient la notion de temporalité. Et ce problème se retrouve dès lors que l’on s’essaie à définir le temps. Ainsi pour Blaise Pascal, le temps est un mot primitif, c’est à dire qu’on ne peut pas le dériver d’un concept plus profond que lui [Pas04]. Pour lui, il est impossible et même inutile de définir le temps [Pas04]. Cela n’a pas empêché les physiciens mais aussi les philosophes, écrivains et artistes de se frotter depuis des siècles à cette question aussi épineuse qu’est la question de la nature du temps. Afin de saisir la notion de temps, nous proposons de voir dans cette première sous-partie comment le temps est vu dans différents domaines, tels que la poésie, la philosophie ou la peinture.

Souvent le temps est vu comme une chose sur laquelle on n’a pas d’emprise, qu’on ne peut pas maîtriser. Cela se traduit par exemple chez Georges Brassens par “Il est mort, c’était le bon temps” dans sa chanson “Le temps passé”. On ne pourra jamais retrouver ce bon temps puisqu’il est mort. Cette incapacité à agir face au temps amène souvent les auteurs à en avoir une vision sombre, fataliste même. Par exemple pour Delmore Schwartz le temps nous consume littéralement, avec “Time is the fire in which we burn” [Sch20]. On peut aussi citer la fameuse chanson “Mistral Gagnant”, dans laquelle Renaud décrit le temps comme un “assassin qui emporte avec lui les rires des enfants”. Mais c’est sûrement Jacques Brel, dans sa chanson “les vieux”, qui a le mieux dépeint cette vision du temps. Chez lui le temps se matérialise par cette terrible “pendule d’argent”, “qui dit oui, qui dit non” traduisant notre impuissance face à elle, avant de finir en disant qu’elle “nous attend” pour nous rappeler que personne ne peut y échapper. Notre impuissance face à cette horloge ardente et assassine pousse certains à préférer ne pas en parler, comme Robert Penn Warren dans son poème “Tell Me a Story” : “The name of the story will be Time, But you must not pronounce its name”. De cette façon peut être, l’auteur espère échapper à son emprise.

D’autres conceptions du temps sont parfois proposées dans la littérature ou dans les œuvres cinématographiques. On peut citer par exemple la série “True detective”, dans laquelle pour le personnage de Rust le temps est circulaire et nous sommes condamnés à revivre toujours exactement la même chose pour l’éternité (si tant est que cette notion ait du sens dans cette représentation circulaire du temps). Idée que l’on retrouve dans le film “Un jour sans fin” ou même déjà bien plus tôt avec le mythe de Sisyphe de la mythologie grecque. Nous verrons dans la partie suivante que les physiciens rejettent en général cette vision du temps, en se basant sur le principe de causalité. Il y a aussi des récits dans lesquels les voyages dans le temps sont rendus possibles. Cependant, ce principe de causalité est tellement ancré dans notre façon de concevoir le monde que cela donne presque systématiquement lieu à des paradoxes logiques.

Il ne nous est donc pas seulement impossible à définir, penser le monde avec un temps différent ou voire même sans temps nous demande des efforts d'imagination extrêmes. La question de l'origine du temps, c'est à dire de savoir si le temps a eu une origine et de comprendre ce que cela pourrait signifier en est un parfait exemple, comme le soulève Etienne Klein notamment dans sa conférence "D'où vient le temps?" [Kle10]. En effet, si le temps a une origine, c'est que quelque chose, qui n'est pas le temps, l'a précédé. Dit autrement il y a eu quelque chose *avant* le temps. Et donc "cette chose est par définition antérieure au temps. Mais dès lors que la notion d'antériorité n'a de sens qu'au sein d'une chronologie, qui nécessite la notion de temps pour lui donner sens, comment une chose distincte du temps et lui donnant naissance, pourrait-elle être antérieure au temps?" [Kle10]. On voit là toute la difficulté que constitue cette question.

Dans cette conférence, Étienne Klein poursuit en décrivant le temps comme le processus qui remplace systématiquement l'instant présent par un autre instant présent, différent du précédent. Et paradoxalement ce processus, lui, est invariant par rapport au temps. C'est à dire que le moteur du temps, le moteur de ce renouvellement de l'instant présent ne change pas dans le temps, n'est pas affecté par le temps. Et ce processus est indispensable à la vie telle que nous la connaissons. C'est en effet grâce à ce processus, immuable et infatigable, que l'on peut expérimenter la vie tel que nous le faisons. Sans lui, comment pourrait-on grandir, évoluer, vieillir? Pourrait-on encore parler de vie? C'est peut être là que réside la réponse à la nature du temps, peut être que tout simplement "Time is Nature's way to keep everything from happening all at once" [Whe89].



FIGURE 2.1 – “The Persistence of Memory” - Salvador Dali.

Enfin, comme pour rajouter une difficulté supplémentaire à notre quête pour donner du sens au temps, notre perception de ce dernier dépend de notre état psychologique. Par exemple, le “comportement” du temps nous paraît beaucoup plus erratique lorsque nous sommes en état de rêve. C’est souvent l’interprétation qui est donnée du célèbre tableau de Salvador Dali présenté sur la Figure 2.1. La fonte des horloges représenterait l’impuissance du temps durant l’état de rêve. Le temps paraît ainsi moins rigide, comme malléable. Une autre interprétation est souvent aussi proposée pour expliquer ce tableau. Il symboliserait

la théorie de la relativité restreinte élaborée par Albert Einstein en 1905. Les horloges, en fondant, perdraient de leur pouvoir et de leur caractère absolu, illustrant le concept de temps propre de la relativité restreinte¹, théorie qui a révolutionné notre conception du temps. Cependant, comme nous allons le voir dans les parties suivantes, les physiciens n’ont pas attendu ce célèbre savant pour essayer de percer les mystères du temps.

2.1.2 La flèche du temps

Nous allons nous intéresser dans la suite de cette partie au temps des physiciens et plus particulièrement à la “flèche du temps”, nécessaire comme nous allons le voir pour décrire certains phénomènes physiques.

L’objectif premier des physiciens consiste à trouver les lois qui régissent l’univers dans lequel nous vivons. Et dans cette réalité que nous expérimentons tous les jours, se déroulent des phénomènes irréversibles. Par exemple, lorsque l’on verse une goutte de lait dans son café, le milieu obtenu est un mélange de café et de lait, du café au lait. Mais à partir de ce dernier, on ne voit jamais le café et le lait se séparer naturellement en reformant la goutte de lait à la surface du café. On pourrait donc s’attendre à ce que cette irréversibilité se retrouve naturellement inscrite dans les équations et les modèles développés par les physiciens. Cependant, ce n’est pas le cas. En effet, les lois fondamentales qui régissent la dynamique des particules nous disent qu’il n’y a pas de raison de privilégier l’un ou l’autre de ces scénarios. Ce paradoxe a suscité un vif intérêt des physiciens, qui sont venus à développer des théories pour en expliquer la cause.

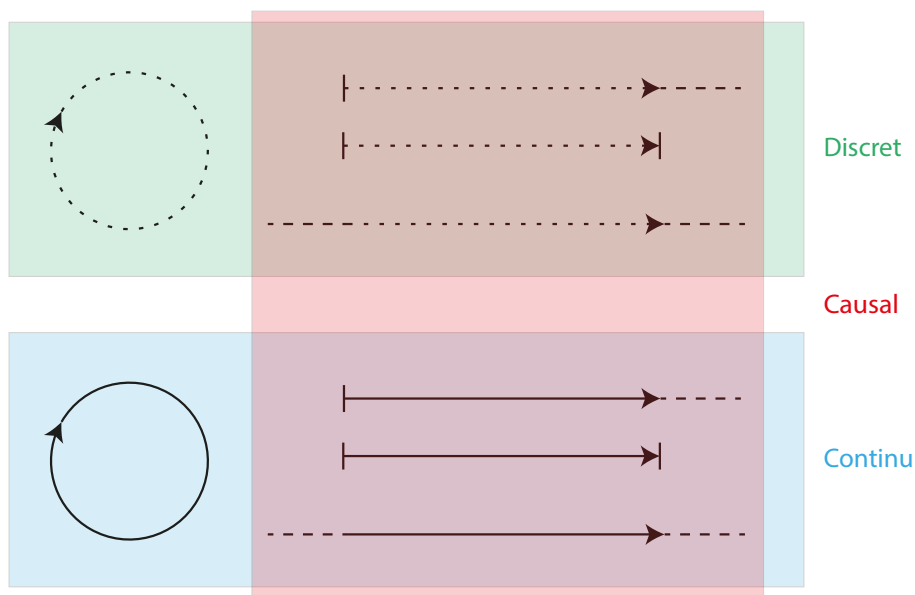


FIGURE 2.2 – Quelques représentations du cours du temps.

1. Cette interprétation a été réfutée par Salvador Dali, qui a expliqué que les horloges molles avait pour inspiration la perception surréaliste d’un camembert fondant au soleil [Wik20].

Avant de rentrer plus en détail sur l'origine de cette irréversibilité, il faut faire attention à la distinction qui est faite en physique entre la flèche et le cours du temps. Ce dernier traduit le passage du temps, le remplacement systématique de l'instant présent par un autre. En un sens, il est le temps lui-même [Kle07]. Il existe plusieurs possibilités pour le représenter. Il peut être de nature discrète ou continue, de forme circulaire ou linéaire. Il peut avoir ou non une origine et/ou une fin, comme illustré sur la Figure 2.2. Les problèmes conceptuels que posent la question de savoir si le temps a eu une origine ont été rapidement évoqués dans la partie précédente (on peut facilement comprendre que des problèmes similaires apparaissent pour concevoir une fin du temps, fin prévue par exemple par les modèles cosmologiques du "Big Rip" [CKW03]). Qu'il soit discret ou continu, soumis au principe de causalité, qui stipule qu'un effet ne peut pas précéder sa cause, ce cours du temps devient irréversible et ne peut pas être circulaire. Ce principe de causalité est l'un des plus fondamentaux de la physique. Et pourtant, récemment, Huw Price a avancé des arguments en faveur d'une possible rétrocausalité² de certains phénomènes quantiques [Pri12]. Il faut cependant rester prudent, ces développements n'ont de sens que pour certains phénomènes quantiques et ne sont pas encore matures, le titre de l'article de Huw Price est d'ailleurs explicite à ce sujet : "Does Time-Symmetry Imply Retrocausality? How the Quantum World Says "Maybe" " [Pri12]. Comme l'a dit Gilles Cohen-Tannoudji, "le principe de causalité sera sans doute un des derniers auxquels les sciences renonceront un jour". Et nous verrons dans les parties suivantes que ce principe est fondamental dans la description des phénomènes en électromagnétisme. La question de la représentation discrète ou continue du temps (ou de l'espace) a aussi beaucoup animé les discussions entre les physiciens et les philosophes. Certains, comme le philosophe Zénon, ont eu du mal à accepter cette représentation infinie. Il argumentait par exemple qu'on "ne peut traverser un nombre infini de points en temps fini", mais comme l'explique Bertrand Russell "l'idée qu'un nombre infini d'instantanés constitue un temps infiniment long n'est pas vraie" [RD02]. Ainsi le cours du temps est généralement représenté par une ligne continue (donc respectant la causalité), terminée par une flèche pour indiquer qu'il a un sens. Il est irréversible dans le sens où chaque point de la ligne ne peut être traversé qu'une fois (avec les mots d'Étienne Klein cela donne "Le facteur temps ne sonne jamais deux fois", titre de l'un de ses livres).

Paradoxalement, cette flèche ne représente pas ce qu'on appelle en physique la "flèche du temps" [Kle07]. Cette dernière est une propriété de certains phénomènes physiques se déroulant au cours du temps. Dans son livre, Étienne Klein écrit : "la plupart de ceux [les phénomènes] qui existent à notre échelle subissent au cours du temps des transformations irréversibles qui les empêchent à tout jamais de retrouver leur état initial. [...] Leur dynamique est temporellement "fléchée" " [Kle07]. Pour revenir à l'exemple de la goutte de lait dans le café, la flèche du temps indique le sens unique selon lequel le phénomène se déroule naturellement, c'est à dire la goutte qui se mélange au café. Cette flèche traduit donc l'irréversibilité du phénomène, comme illustré sur la Figure 2.3. Elle rend ainsi compte du fait qu'on ne verra jamais le lait se dissocier du café pour reformer la goutte initiale, scénario pourtant en accord avec les lois fondamentales de la dynamique des particules.

2. Le lecteur intéressé pourra aussi aller voir les théories développées à ce sujet en philosophie, avec par exemple les travaux Michel Anthony Eardley Dummett : "Can an Effect Precede its Cause?" [Dum54].

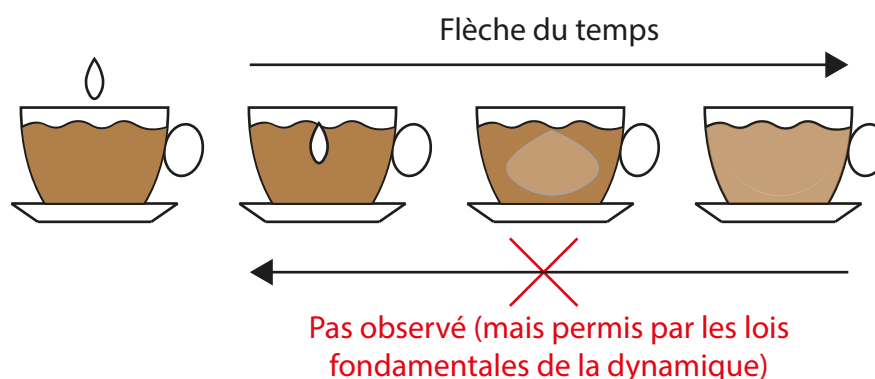


FIGURE 2.3 – Illustration de la flèche du temps.

Ainsi, au sein du cours du temps, se produisent des phénomènes irréversibles, donc soumis à la flèche du temps, alors que les équations qui les décrivent sont réversibles. Cette apparente incohérence a suscité un vif intérêt chez de nombreux physiciens, qui ont dû développer de nouvelles théories pour expliquer l'origine de cette flèche du temps. Nous verrons dans la partie suivante la solution proposée par Ludwig Boltzmann, qui a eu une influence considérable sur ces questions. Il est intéressant pour cela de présenter les discussions qu'il a eu avec son ami Johann Josef Loschmidt et en particulier le fameux *paradoxe de réversibilité*. Ce dernier constitue la base théorique à partir de laquelle le RT a par la suite été élaboré.

2.1.3 Les prémisses du retournement temporel : Le paradoxe de réversibilité

Au XIX^{ème} siècle, la physique fondamentale se trouve dans une position particulière puisque deux courants de pensée s'affrontent : les énergétistes et les atomistes [Kle07] ; [Dug59]. Ce genre de situation est symptomatique de l'aube des révolutions scientifiques [Kuh62]. C'est une situation dans laquelle le paradigme en place ne permet pas de répondre à certaines questions sur les observations physiques. Il en résulte une remise en cause de ce paradigme et l'élaboration d'un nouveau, permettant de rendre compte des phénomènes que l'ancien ne pouvait pas expliquer [Kuh62]. Ainsi durant le XIX^{ème} siècle, un des problèmes auquel les physiciens faisaient face était l'irréversibilité observée de certains phénomènes macroscopiques alors que la dynamique au niveau microscopique de ces mêmes phénomènes était régie par des lois de la mécanique de Newton, qui sont réversibles.

C'est dans ce contexte que Boltzmann proposa pour la première fois, avec le théorème H, une interprétation du second principe de la thermodynamique dans le cadre de la mécanique statistique [Dug59]. Cette fonction H rend compte de la façon dont les systèmes isolés tendent naturellement vers leur état d'équilibre. C'est une vision statistique de l'évolution d'un système de particules. Ce théorème H explique pourquoi, pour un système isolé comportant un nombre assez grand de particules, à partir d'équations au niveau microscopiques réversibles, son état macroscopique tend irréversiblement vers un état moins organisé. Ainsi, Boltzmann

arrive à montrer comment, en partant des équations réversibles de la mécanique régissant la dynamique des particules, on arrive à des phénomènes irréversible à l'échelle macroscopique. Il en conclut que : “le théorème H peut être considéré comme l'explication microscopique des phénomènes macroscopiques. Il rend en effet compte de l'émergence d'une flèche du temps à partir d'équations réversibles par rapport à la variable temps” [Kle07].

Dès lors, il s'est heurté à une opposition de nombreux scientifiques, énergétistes et atomistes confondus [Dug59]. En effet, les scientifiques en général sont les premiers à résister aux nouvelles théories [Bar61]. Cela peut paraître surprenant au premier abord, mais c'est en fait nécessaire au fonctionnement même de la science [Kuh62]. Les scientifiques ne peuvent pas rejeter le paradigme dans lequel ils ont évolué d'un revers de main pour en accepter aussitôt un nouveau, “ils ne pourraient pas le faire et rester pourtant des scientifiques” [Kuh62].

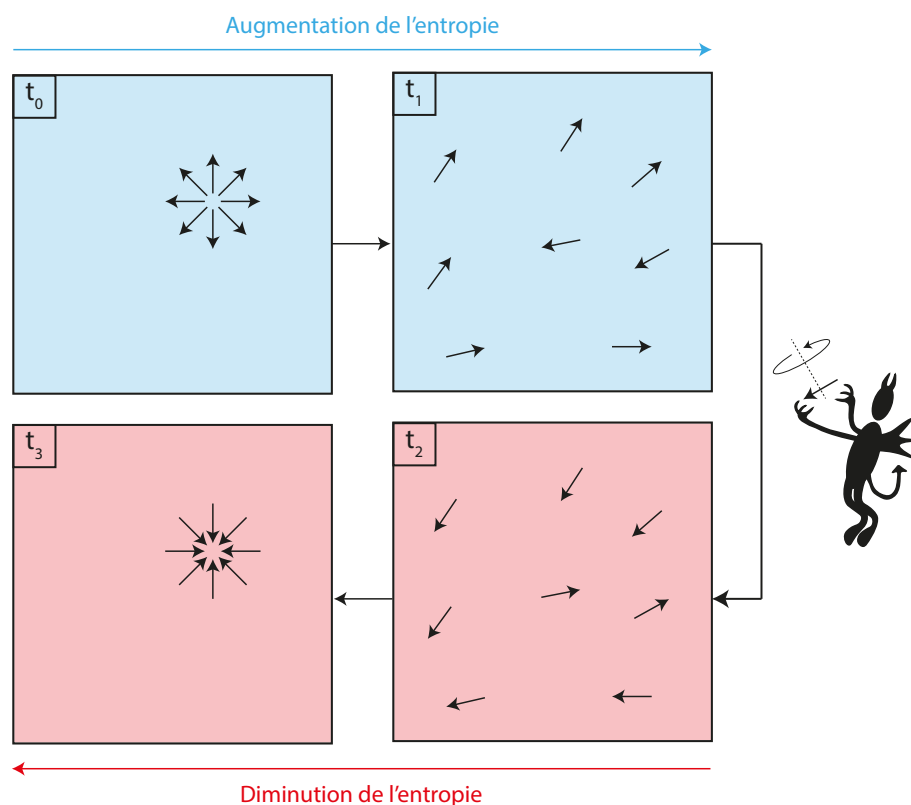


FIGURE 2.4 – Illustration du paradoxe de réversibilité, dans sa forme revisitée par Mathias Fink [Fin15].

Parmi les scientifiques qui lui ont fait face, Loschmidt, qui était un ami de Boltzmann, lui a consacré pas moins de quatre mémoires à la critique de la théorie cinétique [Dug59]. La plus célèbre de ses objections est l'*objection de réversibilité*, appelée aussi *paradoxe de réversibilité* [Dug59] ou *paradoxe de Loschmidt* [Fin09]. Il est illustré sur la Figure 2.4 dans sa forme revisitée par Mathias Fink [Fin15]. Au temps t_0 , un certain nombre de molécules sont placées dans un enceinte fermée dont les parois sont absolument lisses et élastiques. Elles se trouvent alors dans un état “organisé”, possédant toutes la même vitesse initiale.

Au bout d'un certain temps, au temps t_1 , ces molécules auront subies des collisions entre elles et avec la paroi et auront atteint un état plus désorganisé qu'au temps t_0 (la répartition des vitesses sera proche de l'état d'équilibre, *i.e.* proche d'une répartition Maxwellienne). Ainsi l'entropie³, qui caractérise le désordre d'un système, augmente entre les temps t_0 et t_1 , en accord avec le second principe de la thermodynamique. Au temps t_1 le système est figé et des démons, les démons de Loschmidt, viennent inverser le sens du vecteur vitesse de chaque molécule, donnant au temps t_2 les nouvelles conditions initiales du système. Ainsi le processus va se dérouler de façon inverse entre les temps t_2 et t_3 , c'est à dire que les molécules vont revivre leur vie passée, refaisant les collisions dans le sens inverse, pour revenir à l'état organisé initial du temps t_0 . Entre les temps t_2 et t_3 , le système devient donc plus "organisé" et donc l'entropie décroît (de façon exactement opposée à son augmentation entre les temps t_1 et t_2). En se basant sur cette expérience de pensée, Loschmidt prétend réfuter le théorème H de Boltzmann.

Dans sa réponse, Boltzmann concéda que l'entropie diminuait bien entre les temps t_2 et t_3 . Cependant il rétorqua que "cela n'est nullement en contradiction avec ce qui a été démontré auparavant ; car la supposition, faite alors, que la répartition ne présentait aucune organisation moléculaire n'est pas réalisé ici puisque, après l'inversion exacte de toutes les vitesses, chaque molécule ne heurtera pas les autres conformément aux lois de probabilités, mais suivant une loi très particulière que l'on peut calculer à l'avance" [Bol+02]. Cela a été rapporté de façon peut être plus claire par Émile Borel "Il n'y a là qu'une contradiction apparente, car si le raisonnement que nous avons fait est bien applicable dans le cas d'un mouvement naturel quelconque, il ne l'est pas dans le cas du mouvement inverse, qui n'est d'ailleurs pas *physiquement* possible" [Bor25]. La subtilité réside dans le mot "quelconque". C'est à dire que la théorie de Boltzmann rend compte de l'évolution de systèmes isolés dont les molécules évoluent de façon aléatoire et non pas de ceux dont les conditions initiales sont contrôlées de sorte à orienter les molécules vers une certaine organisation spatiale. Sa théorie n'exclut pas par ailleurs la possibilité que le système tende vers un état plus organisé, elle prétend juste que c'est hautement improbable. On peut s'en convaincre facilement en remarquant qu'il est hautement improbable que le système de la Figure 2.4 puisse se retrouver dans l'état au temps t_3 de lui même, sans intervention extérieure. Boltzmann ajouta aussi qu'après le temps t_3 , l'entropie du système augmentera bien, les particules étant alors de nouveau soumise à un mouvement quelconque [Bol+02]. Finalement cette objection émise par Loschmidt, qui a obligé Boltzmann à pousser encore plus loin sa réflexion, a exercé une influence déterminante sur sa pensée. Sa réponse l'a en effet finalement amené à reformuler le second principe de la thermodynamique avec une vision probabiliste, connu sous sa fameuse forme⁴ $S = k_B \ln(\Omega)$ [Dug59] ; [Kle07].

Ainsi Boltzmann propose une explication sur l'origine de cette fameuse flèche du temps. Elle serait donc une manifestation au niveau macroscopique du fait que certaines configurations des phénomènes microscopiques - eux réversibles - soient privilégiées par rapport à

3. Par soucis de simplicité, ici entropie = - fonction H.

4. Cela permet parfaitement d'illustrer en quoi la réaction de défiance des scientifiques à l'égard de nouvelles théories demande à ses défenseurs de redoubler d'arguments et ainsi de pousser encore plus loin les théories qu'ils élaborent, comme expliqué par Thomas Samuel Khun [Kuh62].

d'autres, puisque beaucoup plus probables. Cependant cette explication proposée par Boltzmann n'a pas fait l'unanimité chez les physiciens. Max Planck par exemple, bien qu'anti-énergétiste, n'a jamais été convaincu par cette interprétation car "il demeurerait persuadé que la loi de l'entropie était absolue, et donc que son fondement était non probabiliste" [Kle07]. Les travaux qu'il a alors entrepris sur le rayonnement corps noir dans le but d'aboutir à une nouvelle compréhension de l'irréversibilité thermodynamique, l'ont finalement amené, "dans un acte de désespoir", à imaginer que "les échanges d'énergie entre le rayonnement et la matière ne peuvent se faire que par des paquets discontinus d'énergie, les quantas" [Kle07]. Ces travaux ont déclenché une véritable révolution en physique, qui a abouti à l'élaboration de la physique quantique [Kle07].

Une chose est sûre, les concepts de temps et de (d'ir)réversibilité demeurent encore, au moins en partie, mystérieux et ils n'ont pas fini de fasciner les physiciens. Et c'est presque un siècle plus tard vers la fin du XX^{ème} siècle, qu'un groupe de physicien français, en exhumant ces vieilles discussions, en est arrivé à développer le concept de retournement temporel.

2.1.4 Le principe des démons de Loschmidt appliqué aux ondes

Vers la fin du XX^{ème} siècle, Mathias Fink s'intéresse à ces discussions autour du paradoxe de réversibilité et il se demande comment faire en pratique pour réaliser l'opération de renversement des vitesses afin d'obtenir le déroulé inverse d'un phénomène. On comprend aisément que l'expérience de pensée proposée par Loschmidt n'est pas réalisable en pratique. En effet, on voit difficilement comment on pourrait stopper toutes les molécules d'un gaz, mesurer leurs vitesses et renvoyer chaque molécule dans la direction opposée. Mais l'idée de Mathias Fink consiste à transposer le principe des démons de Loschmidt dans le domaine des ondes, en suivant le raisonnement suivant.

Prenons l'équation bien connue de propagation d'ondes dans un milieu sans pertes et sans sources :

$$\left\{ \frac{\partial^2}{\partial x^2} + \frac{\partial^2}{\partial y^2} + \frac{\partial^2}{\partial z^2} - \frac{1}{c^2(\mathbf{r})} \frac{\partial^2}{\partial t^2} \right\} \Phi(\mathbf{r}, t) = 0 \quad (2.1)$$

avec x , y et z qui représentent les trois variables d'espace, t celle du temps, $c(\mathbf{r})$ la vitesse de propagation de l'onde et Φ son champ (qui peut être de nature acoustique ou électromagnétique). D'un point de vue mathématique, cette équation est réversible, comme le traduit la seule présence d'une dérivée seconde par rapport au temps. En effet la double dérivation d'une fonction $f(-t)$ donne une fonction de même signe $f''(-t)$ et donc si la fonction $f(t)$ est solution de l'équation 2.1 alors $f(-t)$ l'est elle aussi.

Pour comprendre comment cette réversibilité se traduit en pratique, il est intéressant d'utiliser les fonctions de Green, que l'on peut par exemple écrire comme $G(\mathbf{r}, \mathbf{r}_0, t)$ pour la fonction de Green causale. Ces fonctions correspondent par définition à la réponse du milieu à

une impulsion de Dirac spatiale et temporelle, émise en un point $\mathbf{r}_0$ au temps $t = 0$. La seule connaissance de cette fonction permet ainsi, par le principe de superposition, de connaître la réponse du milieu pour n'importe quelle source spatio-temporelle. Pour l'équation 2.1, qui est réversible, il existe deux solutions, causale et anti-causale, comme illustré sur la Figure 2.5. La fonction causale correspond à un scénario dans lequel les ondes sont divergentes et se propagent du point source $\mathbf{r}_0$ à partir du temps $t = 0$. La fonction anti-causale correspond elle au scénario inverse, dans lequel les ondes convergent vers le point $\mathbf{r}_0$ depuis le passé et arrivent à ce point au temps $t = 0$. Soumis au principe de causalité, il ne reste qu'une solution, la fonction de Green causale, qui représente la solution que l'on observe physiquement. La fonction de Green anti-causale est bien mathématiquement solution de l'équation 2.1, mais on ne l'observe pas en pratique, puisqu'elle viole le principe de causalité.

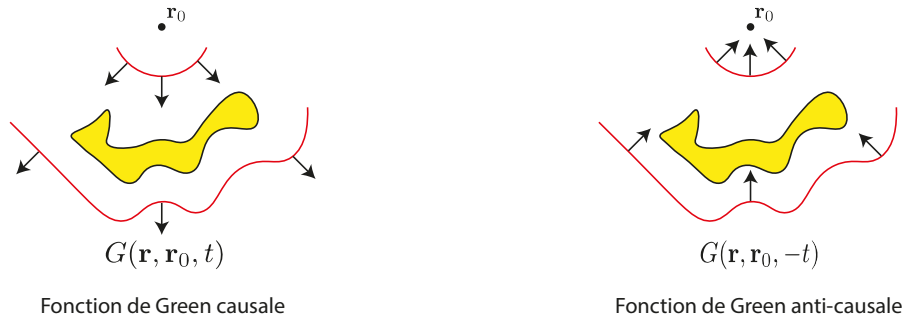


FIGURE 2.5 – Les solutions causale et anti-causale de l'équation d'onde 2.1 [Fin15].

Cependant puisque ce scénario dans lequel les ondes convergent vers la source est possible selon les équations de la physique, on peut se demander comment le réaliser. Il faut pour cela remarquer que dans la fonction $G(\mathbf{r}, \mathbf{r}_0, -t)$, la variable t évolue des temps les plus négatifs (du passé) vers l'instant présent $t = 0$. Pour pouvoir réaliser ce scénario dans un univers causal, il va falloir faire de ce passé l'instant présent, de sorte à ce que les ondes convergent vers la source à un instant futur. Ainsi, ce n'est pas la fonction anti-causale $G(\mathbf{r}, \mathbf{r}_0, -t)$ que l'on va chercher à créer, mais la fonction $G(\mathbf{r}, \mathbf{r}_0, T - t)$ qui est, elle, causale (avec T la durée de l'expérience).

Mathias Fink propose pour ce faire deux options [Fin15]. La première, la plus évidente, consiste à inverser le sens de propagation d'une onde divergente (comme par exemple celle décrite par la fonction de Green causale de la Figure 2.5) en tout point de l'espace à un instant fixé. Elle ressemble fortement à l'expérience de renversement des vitesses proposée par Loschmidt. Cependant, même si c'est en se basant sur cette idée que le principe du RT a été d'abord pensé, ce n'est pas celle qui a été développée au début des travaux sur le RT. La deuxième option consiste à inverser le sens de propagation de l'onde sur une surface pendant un certain intervalle de temps (contrairement à la première qui consiste à inverser le sens de propagation de l'onde en tout point spatial à un instant). Elle tire parti du fait qu'il est possible de connaître l'évolution temporelle d'une onde dans l'espace uniquement à partir de la connaissance de son évolution temporelle sur une surface. Ce principe est généralement appelé principe de Huygens. Il stipule que l'évolution d'un champ dans un certain volume est complètement déterminée par les composantes tangentielles de ce champ au niveau de la

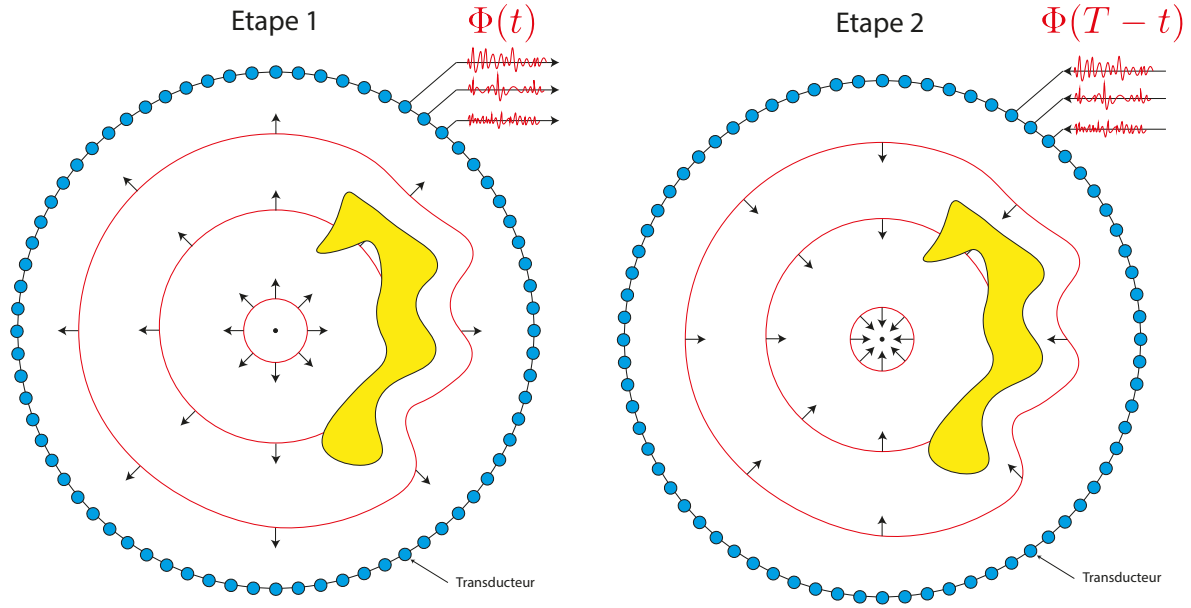


FIGURE 2.6 – Illustration du principe de la cavité à RT.

surface contenant ce volume [Kon00]. Et cette deuxième option semble peut être plus facile à mettre en place. Elle a initialement été pensée de la façon suivante. Un certain nombre de transducteurs réversibles sont placés sur une surface sphérique entourant une source, comme illustré sur la Figure 2.6. Ce dispositif, développé par Didier Cassereau et Mathias Fink est appelé cavité à RT [CF92]. Durant la première étape, ces transducteurs mesurent l'évolution temporelle du champ $\Phi(t)$ issue de la propagation des ondes émises par une source et stockent cette information. La deuxième étape consiste à renvoyer par ces capteurs, alors utilisés comme sources, cette information retournée temporellement $\Phi(T-t)$. De cette façon, on assiste durant la deuxième étape au scénario inverse de celui de la phase 1. Enfin presque. Parce qu'il manque une certaine information, présente lors de la première étape, mais qui ne l'est pas lors de la deuxième. Cette information correspond à la partie de l'énergie électromagnétique qui ne s'est pas propagée de la source et qui n'est donc pas mesurée par les transducteurs, celle des ondes évanescentes. De ce fait, les ondes focalisent lors de la deuxième étape sur une dimension égale à la limite de diffraction (*i.e.* $\lambda/2$). De plus, au moment de cette focalisation, la source de l'étape 1 n'émettant pas un signal inversé, les ondes continuent ainsi de se propager au delà de la source.

La mise en place du dispositif présenté sur la Figure 2.6 est compliquée en pratique. Elle nécessiterait le déploiement d'un très grand nombre de transducteurs. En effet, d'après le critère de Shannon [Jer77], un champ oscillant à une fréquence spatiale $1/\lambda$ est échantillonné sans perte d'information si la distance séparant les transducteurs est inférieure ou égale à $\lambda/2$. De plus, il faudrait synchroniser tous ces transducteurs. On se confronterait aussi rapidement à des problèmes de couplage entre les transducteurs, pouvant rendre impossible l'opération. C'est ainsi que, comme nous l'avons évoqué au chapitre précédent, Carsten Draeger a développé le concept de miroir à RT (dans le domaine acoustique) en environnement réverbérant [DRA97] ; [DF97]. Son principe est expliqué en détail dans le chapitre suivant.

De plus, nous avons dit que c'est parce que l'équation 2.1 est réversible que la fonction de Green anti-causale existait, permettant d'imaginer le concept de la cavité à RT de la Figure 2.6. Cependant, cette équation est valable pour un cas sans pertes et en pratique, tous les milieux présentent des pertes. Par ailleurs, si on regarde le chemin emprunté par les ondes de la source vers chaque transducteur lors de l'étape 1, on remarque que ce chemin doit être le même lors de la phase inverse, afin que toutes les ondes retournent bien à la source. N'est-ce pas là la définition de la réciprocité ? Et donc le milieu doit-il être réversible ou réciproque (ou les deux) ? Y a-t-il un lien entre ces deux concepts ? La section suivante se concentre ainsi sur la compréhension de ces concepts dans le domaine de l'électromagnétisme.

2.2 Le retournement temporel en électromagnétisme

Dans cette section, les principes sur lesquels repose le RT en électromagnétisme sont présentés. Tout d'abord, la question de la réversibilité des équations de Maxwell est discutée. Ensuite, le lien entre réversibilité et réciprocité est abordé à travers des exemples simples, dans le but de donner au lecteur une vision intuitive de ces concepts. Notons que ces concepts théoriques décrivent le comportement de milieux mais ils peuvent être étendus à la notion de système, comprenant des milieux et/ou des dispositifs et des composants [Cal+18]. Nous verrons dans cette section que nous parlerons selon le cas considéré de milieu ou de système.

2.2.1 La causalité en électromagnétisme

Nous allons tout d'abord présenter les équations de Maxwell dans un contexte général. Nous verrons que pour décrire les milieux d'un point de vue électromagnétique, la causalité joue un rôle important. Prenons les équations de James Clerk Maxwell dans un milieu LTI (linéaire et invariant dans le temps) hétérogène et sans source [Kem11] :

$$\nabla \cdot \mathbf{D}(\mathbf{r}, t) = 0 \quad (2.2)$$

$$\nabla \cdot \mathbf{B}(\mathbf{r}, t) = 0 \quad (2.3)$$

$$\nabla \times \mathbf{E}(\mathbf{r}, t) = -\frac{\partial \mathbf{B}(\mathbf{r}, t)}{\partial t} \quad (2.4)$$

$$\nabla \times \mathbf{H}(\mathbf{r}, t) = \frac{\partial \mathbf{D}(\mathbf{r}, t)}{\partial t} \quad (2.5)$$

dans lesquelles les vecteurs déplacement électrique $\mathbf{D}$ et excitation magnétique $\mathbf{H}$ sont décrits dans le domaine fréquentiel⁵ par⁶ :

$$\tilde{\mathbf{D}}(\mathbf{r}, \omega) = \tilde{\tilde{\epsilon}}(\mathbf{r}, \omega) \cdot \tilde{\mathbf{E}}(\mathbf{r}, \omega) \quad (2.6)$$

$$\tilde{\mathbf{B}}(\mathbf{r}, \omega) = \tilde{\tilde{\mu}}(\mathbf{r}, \omega) \cdot \tilde{\mathbf{H}}(\mathbf{r}, \omega) \quad (2.7)$$

avec $\tilde{\tilde{\epsilon}}(\mathbf{r}, \omega)$ et $\tilde{\tilde{\mu}}(\mathbf{r}, \omega)$ les tenseurs complexes de permittivité diélectrique et de perméabilité magnétique du milieu. Ces relations sont obtenues dans le domaine temporel en prenant leur transformée de Fourier inverse :

$$\mathbf{D}(\mathbf{r}, t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} [\tilde{\tilde{\epsilon}}(\mathbf{r}, \omega) \cdot \tilde{\mathbf{E}}(\mathbf{r}, \omega)] \cdot e^{j\omega t} d\omega \quad (2.8)$$

$$\mathbf{B}(\mathbf{r}, t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} [\tilde{\tilde{\mu}}(\mathbf{r}, \omega) \cdot \tilde{\mathbf{H}}(\mathbf{r}, \omega)] \cdot e^{j\omega t} d\omega \quad (2.9)$$

5. Rappel : Les quantités dans le domaine fréquentiel sont représentées avec un $\tilde{\cdot}$.

6. Par soucis de simplicité, on se place dans le cas homoisotropique, c'est à dire que les tenseurs de couplage magnetic-to-electric et electric-to-magnetic sont pris égaux à zéros.

Ce qui donne des produits de convolution [Gar67] :

$$\mathbf{D}(\mathbf{r}, t) = \int_{-\infty}^t \overline{\overline{g}}_{\varepsilon}(\mathbf{r}, \gamma) \cdot \mathbf{E}(\mathbf{r}, t - \gamma) d\gamma \quad (2.10)$$

$$\mathbf{B}(\mathbf{r}, t) = \int_{-\infty}^t \overline{\overline{g}}_{\mu}(\mathbf{r}, \gamma) \cdot \mathbf{H}(\mathbf{r}, t - \gamma) d\gamma \quad (2.11)$$

où :

$$\overline{\overline{g}}_{\varepsilon}(\mathbf{r}, t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \tilde{\tilde{\varepsilon}}(\mathbf{r}, \omega) \cdot e^{j\omega t} d\omega \quad (2.12)$$

$$\overline{\overline{g}}_{\mu}(\mathbf{r}, t) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \tilde{\tilde{\mu}}(\mathbf{r}, \omega) \cdot e^{j\omega t} d\omega \quad (2.13)$$

Les équations 2.10 et 2.11 traduisent la façon avec laquelle un système réagit à une certaine excitation. Par exemple, l'équation 2.10 traduit la réponse $\mathbf{D}(\mathbf{r}, t)$ d'un système à une excitation $\mathbf{E}(\mathbf{r}, t)$, à travers la réponse impulsionnelle du milieu $\overline{\overline{g}}_{\varepsilon}(\mathbf{r}, t)$. La permittivité d'un milieu peut ainsi être vue comme une fonction de transfert entre une entrée $\mathbf{E}(\mathbf{r}, t)$ et une sortie $\mathbf{D}(\mathbf{r}, t)$ (même raisonnement pour la perméabilité).

Pour que les tenseurs $\overline{\overline{g}}_{\varepsilon}(\mathbf{r}, t)$ et $\overline{\overline{g}}_{\mu}(\mathbf{r}, t)$ représentent un milieu physique, il faut qu'ils aient les propriétés suivantes [Gar67] :

1) Ils doivent respecter le principe de causalité, qui dans ce cas se traduit par le fait que la réponse d'un milieu ne peut pas précéder l'excitation : $\overline{\overline{g}}_{\varepsilon}(\mathbf{r}, t) = \overline{\overline{g}}_{\mu}(\mathbf{r}, t) = 0$ pour $t < 0$. On voit ici une illustration du caractère essentiel du principe de causalité en électromagnétisme : s'il n'est pas respecté les milieux que l'on décrit ne correspondent pas à une réalité physique.

2) Ils doivent être bornés : $\int_{-\infty}^{\infty} |\overline{\overline{g}}_{\varepsilon}(\mathbf{r}, t)| dt < \infty$ et $\int_{-\infty}^{\infty} |\overline{\overline{g}}_{\mu}(\mathbf{r}, t)| dt < \infty$, traduisant le fait que la réponse d'un milieu à une impulsion (Dirac) ne peut pas durer infiniment.

Pour comprendre plus en détails l'origine physique de ce lien entre les réponses ($\mathbf{D}$, $\mathbf{B}$) et les excitations ($\mathbf{E}$, $\mathbf{H}$) à travers un produit de convolution, il faut s'intéresser à la façon dont se comporte la matière au niveau microscopique. Les atomes ou les molécules constituant un matériau diélectrique se déforment sous l'effet d'une excitation électromagnétique. Cette déformation résulte de la force électromagnétique à laquelle sont soumises les particules chargées. Il en résulte une orientation des moments dipolaires ainsi créés avec l'excitation électromagnétique. Cette capacité d'un milieu diélectrique à se polariser est prise en compte par les susceptibilités électriques et magnétiques, notées respectivement $\tilde{\tilde{\chi}}_e$ et $\tilde{\tilde{\chi}}_m$. Dans un milieu diélectrique, la permittivité et la perméabilité peuvent alors s'écrire :

$$\tilde{\tilde{\varepsilon}}(\mathbf{r}, \omega) = \varepsilon_0(1 + \tilde{\tilde{\chi}}_e(\mathbf{r}, \omega)) \quad (2.14)$$

$$\tilde{\tilde{\mu}}(\mathbf{r}, \omega) = \mu_0(1 + \tilde{\tilde{\chi}}_m(\mathbf{r}, \omega)) \quad (2.15)$$

Les équations 2.10 et 2.11 décrivent la façon dont les ondes électromagnétiques se comportent au niveau macroscopique dans un milieu. Ces ondes résultent en fait de la superposition d'ondes issues de l'interaction entre l'onde et les particules chargées des atomes [Kra27]. Ces dernières, en oscillant avec le champ électromagnétique, génèrent des ondes sphériques secondaires, qui sont elles mêmes soumises aux équations de Maxwell dans le vide [Kra27]. Ce phénomène est décrit par les équations 2.14 et 2.15, puisque les propriétés du milieu résultent de la somme de l'onde se propageant dans le vide (décrite par ε_0 et μ_0) et de la façon dont le matériau réagit à cette onde (décrite par les susceptibilités). En intégrant ces équations dans les équations 2.6 et 2.7 on trouve les relations constitutives d'un milieu diélectrique dans le domaine temporel [Jac98] :

$$\mathbf{D}(\mathbf{r}, t) = \varepsilon_0 \cdot \mathbf{E}(\mathbf{r}, t) + \varepsilon_0 \cdot \int_{-\infty}^t \overline{\overline{G_{\chi_e}}}(\mathbf{r}, \gamma) \cdot \mathbf{E}(\mathbf{r}, t - \gamma) d\gamma \quad (2.16)$$

$$\mathbf{B}(\mathbf{r}, t) = \mu_0 \cdot \mathbf{H}(\mathbf{r}, t) + \mu_0 \cdot \int_{-\infty}^t \overline{\overline{G_{\chi_m}}}(\mathbf{r}, \gamma) \cdot \mathbf{H}(\mathbf{r}, t - \gamma) d\gamma \quad (2.17)$$

avec $\overline{\overline{G_{\chi_e}}}(\mathbf{r}, t)$ et $\overline{\overline{G_{\chi_m}}}(\mathbf{r}, t)$, les transformées de Fourier inverse des susceptibilités, soumises aux propriétés 1) et 2) présentées ci dessus.

Ainsi la réponse d'un matériau à une excitation électromagnétique résulte de deux contributions. L'une correspond à la réponse du vide qui est instantanée (terme de gauche de la somme dans les équations 2.16 et 2.17). L'autre correspond à la façon dont le matériau (et plus précisément ses atomes ou molécules) réagit à cette excitation. Si la susceptibilité (électrique ou magnétique) est constante en fonction de la fréquence, alors $\overline{\overline{G_{\chi_e}}}(\mathbf{r}, \gamma) \propto \delta(\gamma)$ et $\overline{\overline{G_{\chi_m}}}(\mathbf{r}, \gamma) \propto \delta(\gamma)$. Cela se traduit par une relation de proportionnalité entre les susceptibilités et les champs dans le domaine temporel, comme dans le domaine fréquentiel. Dans ce cas la réponse est instantanée. Si la susceptibilité n'est pas constante, c'est à dire si le milieu est dispersif, ce qui est le cas des matériaux réels (voir suite de cette partie), alors la partie de droite de la somme dans ces équations traduit le fait que la réponse d'un milieu au temps t dépend des champs $\mathbf{E}$ et $\mathbf{H}$ aux temps précédents et peut être retardée par rapport à l'instant d'application de l'excitation.

On ne peut pas s'empêcher de remarquer, en regardant les équations 2.16 et 2.17, que la réponse du vide à une excitation n'est pas zéro. Si ce que nous venons de dire pour expliquer l'origine de la réponse d'un matériau est vrai, comment le vide peut il avoir une réponse non nulle ? Cette question représente sûrement le plus gros problème conceptuel auquel se heurte la théorie de l'électromagnétisme Maxwellienne. Maxwell lui même, lorsqu'il a rassemblé les travaux de Gauss, Faraday, Ampère et d'autres dans son article "A dynamical theory of the electromagnetic field" [Max65] n'envisageait pas la possibilité pour les ondes de propager dans le vide⁷. Pour lui les ondes électromagnétiques étaient dues aux déplacements des particules constituant la matière [Kuh62]. Comme beaucoup de physiciens de son époque, il était donc convaincu de l'existence d'un éther, appelé "Luminiferous aether" dans ce cas, qui permet-

7. Il est intéressant de noter que Maxwell, tout comme Boltzmann avec sa théorie de la mécanique statistique, s'est heurté à l'opposition des scientifiques de l'époque [Kuh62].

taut aux ondes électromagnétiques de se propager [Kuh62]. Vers la fin du XIX^{ème} siècle, les physiciens ont petit à petit abandonné l'idée de l'existence d'un éther. Cela a donc posé des problèmes pour comprendre comment les ondes peuvent se propager dans le vide. C'est peut être dans les travaux d'Einstein que résident des pistes de solution à ce problème. Dans une conférence de 1909, basée sur ses fameux papiers de 1905, Einstein commence sa présentation en rappelant que d'après la théorie Maxwellienne, "light waves appear to be essentially an aggregate of states of a hypothetical medium, the ether, which is present everywhere even in the absence of radiation" [Ein09]. Puis il continue en expliquant que "the theory of [special] relativity has thus changed our views on the nature of light insofar as it does not conceive of light as a sequence of state of a hypothetical medium, but rather as something having an independence existence just like matter". Il termine sa conférence en suggérant, pour la première fois, que la lumière puisse être *à la fois* de nature ondulatoire et particulaire (les discussions autour de cet étonnant comportement de la lumière donneront ensuite naissance à la physique quantique). Et donc si la lumière est une particule, il n'y a pas de problème pour comprendre qu'elle puisse se propager dans le vide. Néanmoins, ces arguments, si importants qu'ils soient du point de vue du développement de la science, ne constituent pas une explication sur la nature profonde des phénomènes électromagnétiques, la dualité onde-particule de la lumière ne fait même que complexifier leur interprétation⁸. Ainsi, l'électromagnétisme réserve encore sûrement bien des surprises aux physiciens, comme l'écrivait déjà Thomas Samuel Kuhn en 1962 : "Il se peut qu'un renversement [...] soit en cours dans la théorie électromagnétique. [...] il se peut qu'un jour nous arrivions à savoir ce qu'est un déplacement électrique" [Kuh62].

Les relations constitutives des milieux dans le domaine temporel (équations 2.10 et 2.11 ou 2.16 et 2.17) sont en général difficiles à manipuler puisqu'elles contiennent des produits de convolution. Il en résulte que les développements sont souvent faits dans le domaine fréquentiel puisque le produit de convolution devient un simple produit scalaire (équations 2.6 et 2.7). Comme nous l'avons déjà signalé, pour un milieu non dispersif c'est encore plus simple puisque le produit de convolution laisse directement sa place au produit scalaire dans le domaine temporel (et fréquentiel). Il faut cependant faire attention, en réalité tous les matériaux sont dispersifs, puisque la polarisabilité des atomes dépend de la fréquence du champ exciteur⁹. De plus, pour qu'un matériau soit "physique", il faut que sa permittivité et sa perméabilité souscrivent aux relations de Kramers et Kronig, qui assurent le respect des conditions 1) et 2) présentées ci-dessus [Gar02]. Ces relations impliquent que tout matériau réel est dissipatif [Ker87]. Lorsque les propriétés d'un matériau ne satisfont pas à ces équations de Kramers et Kronig, cela donne lieu à des paradoxes [Gar02] ; [Gar67]. Ce raisonnement peut être résumé simplement par

$$Causal \Rightarrow Dispersif \text{ (et donc réel)} \Rightarrow Dissipatif \quad (2.18)$$

Cela illustre clairement l'importance de la causalité dans la description des phénomènes électromagnétiques. Nous allons voir dans la partie suivante que les équations de Maxwell, qui

8. À ce sujet, plus récemment, les résultats des travaux d'Alain Aspect [ADR82] ; [Jac+07] sont troublants et illustrent parfaitement les "problèmes conceptuels de compréhension de la mécanique quantique" [Asp19].

9. En pratique les propriétés d'un milieu (permittivité ou perméabilité) peuvent être considérées constantes sur la bande de fréquence de l'excitation [Ish17] (si celle-ci est assez éloignée d'une résonance).

ont intégré le principe de causalité, sont réversibles. Cela revient quelque part aux discussions du début de ce chapitre sur la différence entre le cours du temps et la flèche du temps.

2.2.2 La réversibilité en électromagnétisme

Nous allons discuter dans cette partie de ce que veut dire la réversibilité en électromagnétisme. Nous allons voir que pour ce faire, nous allons discuter de la façon dont cette réversibilité se traduit en théorie et nous finirons par présenter une version plus pratique : la “réversibilité restreinte” [Cal+18] ; [Asa+20].

Mathématiquement, la réversibilité (ou symétrie par retournement temporel¹⁰) peut être représentée par un opérateur $\mathcal{R}$, qui appliqué à l’état $\Psi\{\mathbf{r}, t\}$ d’un processus (amplitude, phase, fréquence (temporelle et spatiale), polarisation...) s’écrit : $\mathcal{R}\{\Psi\{\mathbf{r}, t\}\} = \Psi_r\{\mathbf{r}, -t\}$. L’état $\Psi_r\{\mathbf{r}, -t\}$ correspond à l’état du système au temps $-t$ à la position $\mathbf{r}$ après avoir été soumis à l’opération de symétrie par retournement temporel. Le système est réversible si pour tout temps t et toute position $\mathbf{r}$ la relation suivante reste valide :

$$\mathcal{R}\{\Psi\{\mathbf{r}, t\}\} = \Psi_r\{\mathbf{r}, -t\} = \Psi\{\mathbf{r}, t\} \quad (2.19)$$

Une représentation du “trajet” suivi par $\Psi\{\mathbf{r}, t\}$ pour un processus se déroulant entre $t = 0$ et $t = T$ est présenté sur la Figure 2.7. Pour que ce processus soit réversible, il faut que pour toute position $\mathbf{r}$, $\Psi_r\{\mathbf{r}, -t\}$ soit égal à $\Psi\{\mathbf{r}, t\}$ pour tout temps t , c’est à dire que le trajet soit symétrique par rapport au temps $t = 0$, comme représenté par le trajet bleu. Pour que ce processus soit non-réversible (trajet vert), il suffit que l’état $\Psi_r\{\mathbf{r}, -t\}$ diffère de $\Psi\{\mathbf{r}, t\}$ à une position $\mathbf{r}$ et à instant t .

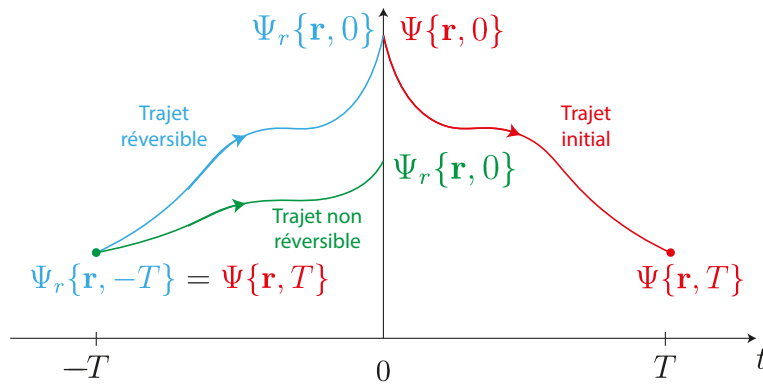


FIGURE 2.7 – Illustration de la réversibilité (repris de [Cal+18]).

Notons que l’introduction de temps négatifs sur la représentation de la Figure 2.7 ne viole pas la causalité. Le choix de positionner l’instant $t = 0$ a été choisi ainsi sur cette

10. Ici le terme “retournement temporel” n’est pas utilisé tout à fait dans le même sens que “RT”.

représentation pour faciliter la distinction entre les processus réversibles et irréversibles. L'état d'un processus réversible sera symétrique par rapport à l'axe vertical en $t = 0$ et sera donc facilement identifiable sur cette représentation.

Comme nous l'avons déjà dit, la plupart des lois physiques fondamentales sont symétriques par retournement temporel (*i.e.* réversibles). Cependant, les quantités physiques impliquées dans ces lois physiques, notées $f(\mathbf{r}, t)$, peuvent être anti-symétriques ou symétriques par retournement temporel [Cal+18] ; [Asa+20], c'est à dire :

$$\mathcal{R}\{f(\mathbf{r}, t)\} = f_r(\mathbf{r}, -t) = \pm f(\mathbf{r}, -t) \quad (2.20)$$

Par exemple en électromagnétisme les vecteurs $\mathbf{E}$ et $\mathbf{D}$ sont des vecteurs vrais et les vecteurs $\mathbf{H}$ et $\mathbf{B}$ sont des pseudo-vecteurs, ce qui se traduit par [Asa+20] :

$$\mathcal{R}\{\mathbf{E}(\mathbf{r}, t)\} = \mathbf{E}(\mathbf{r}, -t) \quad (2.21)$$

$$\mathcal{R}\{\mathbf{D}(\mathbf{r}, t)\} = \mathbf{D}(\mathbf{r}, -t) \quad (2.22)$$

$$\mathcal{R}\{\mathbf{H}(\mathbf{r}, t)\} = -\mathbf{H}(\mathbf{r}, -t) \quad (2.23)$$

$$\mathcal{R}\{\mathbf{B}(\mathbf{r}, t)\} = -\mathbf{B}(\mathbf{r}, -t) \quad (2.24)$$

En appliquant l'opérateur $\mathcal{R}$ aux équations de Maxwell (2.2,2.3,2.4,refMaxwell4) et en suivant les équations précédentes, on retrouve les champs inversés temporellement soumis à ces équations :

$$\nabla \cdot \mathbf{D}(\mathbf{r}, -t) = 0 \quad (2.25)$$

$$\nabla \cdot \mathbf{B}(\mathbf{r}, -t) = 0 \quad (2.26)$$

$$\nabla \times \mathbf{E}(\mathbf{r}, -t) = -\frac{\partial \mathbf{B}(\mathbf{r}, -t)}{\partial t} \quad (2.27)$$

$$\nabla \times \mathbf{H}(\mathbf{r}, -t) = \frac{\partial \mathbf{D}(\mathbf{r}, -t)}{\partial t} \quad (2.28)$$

Cette réversibilité des équations de Maxwell indique que “si *toutes* les quantités électromagnétiques d'un système sont retournées temporellement (suivant les règles des équations 2.21, 2.22, 2.23, 2.24), alors le système retourné temporellement aura la même solution électromagnétique, quel que soit la complexité du système” [Cal+18]. Il faut noter que par souci pédagogique, le raisonnement est effectué sur les équations de Maxwell sans source. Soumise à l'opérateur $\mathcal{R}$, une source, par exemple une source de courant $\mathbf{J}(\mathbf{r}, t)$ devient alors $-\mathbf{J}(\mathbf{r}, -t)$ [Cal+18] ; [Asa+20]. Cela traduit le fait que la source doit être inversée pour que les équations de Maxwell soient réversibles. Nous retrouverons ce point lors de la mise en pratique du RT

entre deux antennes dans la section suivante.

Ainsi tous les systèmes électromagnétiques sont en théorie réversibles. Cependant il faut se questionner sur la signification des équations 2.22 et 2.24. Quelle est la signification de l'inversion temporel de la réponse d'un milieu ? Encore une fois, les développements des équations dans le domaine fréquentiel permet une vision plus simple de ce que cela implique. Nous allons donc poursuivre le développement de cette partie dans le domaine fréquentiel. La forme des équations que nous allons dériver sera donc valable ; soit pour une excitation de forme monochromatique, c'est à dire qu'elles nous renseignent sur le comportement d'un système excité à une seule fréquence ; soit pour un système non dispersif (avec les réserves évoquées précédemment sur ces milieux). Il est important de garder à l'esprit que pour des excitations large bande, tel que c'est le cas pour un RT à un seul transducteur dans le MRT par exemple, les équations peuvent se traduire dans le domaine temporel de façon beaucoup plus complexes. Dans le domaine fréquentiel, l'opérateur $\mathcal{R}$ se traduit par une opération de conjugaison [Cal+18] ; [Asa+20] :

$$\tilde{\mathcal{R}}\{\tilde{f}(\mathbf{r}, \omega)\} = \pm \tilde{f}^*(\mathbf{r}, \omega) \quad (2.29)$$

Par exemple en électromagnétisme, $\tilde{\mathcal{R}}\{\tilde{\mathbf{D}}(\mathbf{r}, \omega)\} = \tilde{\mathbf{D}}^*(\mathbf{r}, \omega)$ et $\tilde{\mathcal{R}}\{\tilde{\mathbf{B}}(\mathbf{r}, \omega)\} = -\tilde{\mathbf{B}}^*(\mathbf{r}, \omega)$. Les équations constitutives du milieu soumises à l'opérateur $\tilde{\mathcal{R}}$ s'écrivent ainsi dans le domaine fréquentiel :

$$\tilde{\mathbf{D}}^*(\mathbf{r}, \omega) = \tilde{\tilde{\epsilon}}^*(\mathbf{r}, \omega) \cdot \tilde{\mathbf{E}}^*(\mathbf{r}, \omega) \quad (2.30)$$

$$\tilde{\mathbf{B}}^*(\mathbf{r}, \omega) = \tilde{\tilde{\mu}}^*(\mathbf{r}, \omega) \cdot \tilde{\mathbf{H}}^*(\mathbf{r}, \omega) \quad (2.31)$$

Ainsi, l'opération de retournement temporel que l'on a appliquée aux équations de Maxwell, donnant les équations 2.25, 2.26, 2.27 et 2.28, implique que le milieu dans lequel se propagent les ondes après ce retournement possède les propriétés conjuguées du milieu initial. De plus, ces propriétés du milieu peuvent aussi dépendre d'un biais extérieur $\mathbf{F}_0$, qui peut être par exemple un champ magnétique stationnaire appliqué à un plasma ou un ferrite. Par souci de clarté, nous écrirons dans ce qui suit la permittivité et la perméabilité comme dépendant uniquement de ce biais (nous omettrons leur dépendance spatiale et fréquentielle). Cela veut dire que pour que les quantités électromagnétiques inversées temporellement soient solutions des équations de Maxwell, il faut que le milieu ait subi la transformation [Cal+18] :

$$\tilde{\mathcal{R}}\{\tilde{\tilde{\epsilon}}(\mathbf{F}_0)\} = \tilde{\tilde{\epsilon}}^*(-\mathbf{F}_0) \quad (2.32)$$

$$\tilde{\mathcal{R}}\{\tilde{\tilde{\mu}}(\mathbf{F}_0)\} = \tilde{\tilde{\mu}}^*(-\mathbf{F}_0) \quad (2.33)$$

Pour qu'un milieu soit réversible il faut donc que lors du retournement temporel des quantités électromagnétiques, la direction du biais soit inversée et ses propriétés conjuguées.

On remarque tout d'abord que pour un milieu isotrope (les tenseurs de permittivité et perméabilité deviennent des scalaires) sans biais ($\mathbf{F}_0 = \mathbf{0}$) et sans perte (qui se traduit dans le cas scalaire par : $\text{Im}[\tilde{\varepsilon}(\mathbf{r}, \omega)] = 0$ et $\text{Im}[\tilde{\mu}(\mathbf{r}, \omega)] = 0$), le milieu de propagation est le même avant et après retournement temporel. Il en résulte qu'un **milieu isotrope sans biais et sans perte est réversible**.

Pour comprendre la signification de ces relations 2.32 et 2.33, prenons tout d'abord un milieu sans biais ($\mathbf{F}_0 = \mathbf{0}$). Dans ce cas, pour que les équations de Maxwell soient réversibles, ces relations impliquent la conjugaison des propriétés du milieu. Pour comprendre ce que cela implique, par simplicité, prenons un milieu diélectrique isotrope, dont la permittivité relative est décrite par un scalaire complexe $\varepsilon_r = \varepsilon'_r - j\varepsilon''_r$ et la perméabilité par un scalaire égal à μ_0 . Le champ $\mathbf{E}$ d'une onde plane se propageant dans ce milieu le long de l'axe z vers les z positifs peut se mettre sous la forme :

$$\mathbf{E}(\mathbf{r}, t) = \text{Re} \left[\mathbf{E}_0 \cdot e^{-jk_z z} \cdot e^{-i\omega t} \right] = \text{Re} \left[\mathbf{E}_0 \cdot e^{-j(k'_z - jk''_z)z} \cdot e^{-i\omega t} \right] = \text{Re} \left[\mathbf{E}_0 \cdot e^{-jk'_z z} \cdot e^{-k''_z z} \cdot e^{-j\omega t} \right] \quad (2.34)$$

avec $\mathbf{E}_0$ un vecteur uniforme et k_z la composante selon z du vecteur d'onde complexe qui s'exprime comme : $k_z = k'_z - jk''_z = k_0 \cdot \sqrt{\varepsilon' - j\varepsilon''}$. La composante $e^{-jk'_z z}$ traduit l'oscillation du champ à la fréquence spatiale k_z et la composante $e^{-k''_z z}$ traduit les pertes dans le milieu par une décroissance exponentielle de l'enveloppe du champ dans le sens des z positifs.

Appliquons ensuite l'opérateur $\tilde{\mathcal{R}}$ à cette équation 2.34. La propagation va alors s'effectuer dans le sens inverse, ce qui est traduit par $\tilde{\mathcal{R}}\{\mathbf{k}(\mathbf{r}, \omega)\} = -\mathbf{k}^*(\mathbf{r}, \omega)$. Le vecteur d'onde s'écrit ainsi $k_z = -(k'_z + jk''_z)$ et on trouve :

$$\mathcal{R}\{\mathbf{E}(\mathbf{r}, t)\} = \text{Re} \left[\mathbf{E}_0 \cdot e^{jk'_z z} \cdot e^{k''_z z} \cdot e^{j\omega t} \right] \quad (2.35)$$

Dans ce cas le sens de propagation est inversé (signe + dans $e^{jk'_z z}$) et les ondes revivent le scénario en sens inverse. Pour cela les pertes se sont transformées en gain. Ce gain est traduit par une croissance exponentielle $e^{k''_z z}$, qui produit exactement l'effet inverse que les pertes dans le sens initial de propagation. De cette façon, si on regarde le phénomène lors du sens initial, on va voir l'onde se propager dans le milieu vers les z positifs en subissant des pertes selon une décroissance exponentielle. Dans le sens inversé, l'onde se propage dans le sens des z négatifs et gagne de l'énergie selon une croissance exponentielle. On retrouve le scénario inverse, le système est donc bien réversible.

Prenons pour finir le cas d'un système sans perte dont les propriétés dépendent d'un biais $\mathbf{F}_0$. Ces équations nous disent que pour que le système soit réversible, le sens du biais doit

être inversé entre les deux sens de propagation.

Donc pour qu'un système soit réversible, au sens théorique du terme, il faudrait changer la nature du milieu entre la propagation dans un sens et dans l'autre. Pour le cas avec biais ce dernier doit être inversé et pour le cas à perte, il faut transformer les pertes en gain (ou vice-versa). On comprend aisément que ces opérations ne s'opèrent pas naturellement. De plus, on voit difficilement comment on pourrait réaliser l'opération de transformation de perte en gain en pratique. En effet, lors de la propagation des ondes dans le premier sens de propagation, les ondes vont subir des pertes dues à leur propagation dans ce milieu. Le milieu chauffe ainsi suite à l'énergie cédée à ses particules. Ensuite, lorsque que l'on renverse le sens de propagation, les ondes vont subir une nouvelle fois ces pertes avant de revenir à leur état initial. On ne verra jamais, naturellement, les particules chauffées lors du premier passage des ondes se remettre à osciller de sorte à générer des ondes électromagnétiques, qui se traduirait par un gain lors du second passage des ondes dans le sens inverse. En fait cela revient quelque part au paradoxe de Loschmidt : le mouvement des particules au niveau macroscopique est irréversible et donc une fois l'énergie transmise à ces dernières on ne verra jamais le processus inverse, pourtant en accord avec les équations régissant leur dynamique, se traduisant ici par la conjugaison des propriétés du milieu. Autrement dit, les pertes, telles qu'elles se présentent à nous au niveau macroscopique, sont irréversibles.

Ainsi, la réversibilité prise selon cette définition n'a pas vraiment de sens en pratique puisque le système doit être altéré entre les deux sens de propagation. “In such a case, the time-reversal experiment is irrelevant, since it compares apples and pears” [Cal+18]. C'est pourquoi on parle plutôt de “réversibilité restreinte” [Cal+18] ; [Asa+20]. Elle consiste à tester la réversibilité d'un système sans opérer la conjugaison des propriétés du milieu des équations 2.32 et 2.33. Elle intègre quand même la possibilité de changer ou pas l'orientation du biais $\mathbf{F}_0$, puisque cette opération est facilement réalisable (contrairement au changement des pertes en gain). Les relations de réversibilités restreintes peuvent s'écrire :

$$\mathcal{R}\{\tilde{\tilde{\epsilon}}(\mathbf{F}_0)\} = \tilde{\tilde{\epsilon}}(-\mathbf{F}_0) \quad (2.36)$$

$$\mathcal{R}\{\tilde{\tilde{\mu}}(\mathbf{F}_0)\} = \tilde{\tilde{\mu}}(-\mathbf{F}_0) \quad (2.37)$$

Les équations 2.32 et 2.33 font ainsi apparaître les deux origines possibles de la non-réversibilité d'un système : les pertes ou un biais d'orientation constante. Cependant, ces deux mécanismes semblent différer. En effet, dans le cas d'un système à pertes, lors des deux phases aller et retour, les ondes peuvent se propager “de la même façon” contrairement au cas avec biais pour lequel les ondes peuvent se propager dans un sens et complètement différemment dans l'autre (principe utilisé par exemple dans les isolateurs à base de ferrite magnétisé). La différenciation entre ces deux phénomènes requiert un autre concept, appelé *réciprocité*, comme développé dans la partie suivante.

2.2.3 La réciprocité en électromagnétisme

Le mot “réciprocité” veut étymologiquement dire “se déplacer de la même façon en avant et en arrière” [Cal+18]. Il se traduit en pratique par le fait qu’un système possède la même réponse lorsque l’émetteur et le récepteur sont inter-changés. C’est une propriété très importante pour la mise en place pratique des dispositifs à RT, puisqu’elle permet de ne pas avoir à échanger émetteur(s) et récepteur(s) entre les deux étapes d’un RT. Cependant elle est souvent mal comprise et confondue parfois avec la réversibilité. C’est pour cela que récemment, des auteurs ont proposé des articles dont l’objectif est de donner une perspective globale de ce que sont ces concepts, afin d’éviter leur confusion [Cal+18] ; [Asa+20].

Pour évaluer la réciprocité, il faut comparer la façon dont les ondes “voient le milieu” dans un sens de propagation et dans l’autre. Il est donc évident que dans les développements qui suivent, le milieu doit rester le même dans un sens et dans l’autre. En fait, ce concept est étroitement lié à celui de réversibilité restreinte. Il en constitue une version plus “pratique”, puisque tandis que la réversibilité restreinte nécessite l’inversion temporelle des ondes dans tous l’espace, le principe de réciprocité peut être facilement mis en place : il suffit d’inter-changer l’émetteur et le récepteur. Ces deux concepts font la même chose : ils “sondent” le milieu et donnent des informations sur son comportement entre deux sens de propagation inverses.

Par souci de clarté, dans cette partie, nous omettrons les dépendances spatiale et fréquentielle des champs électriques et magnétiques, de la permittivité et de la perméabilité et des sources. Prenons les équations de Maxwell dans le domaine fréquentiel en présence d’une source $\tilde{\mathbf{J}}_1$:

$$\nabla \times \tilde{\mathbf{E}}_1 = -j\omega\tilde{\mathbf{B}}_1 \quad (2.38)$$

$$\nabla \times \tilde{\mathbf{H}}_1 = j\omega\tilde{\mathbf{D}}_1 + \tilde{\mathbf{J}}_1 \quad (2.39)$$

et en présence d’une autre source $\tilde{\mathbf{J}}_2$:

$$\nabla \times \tilde{\mathbf{E}}_2 = -j\omega\tilde{\mathbf{B}}_2 \quad (2.40)$$

$$\nabla \times \tilde{\mathbf{H}}_2 = j\omega\tilde{\mathbf{D}}_2 + \tilde{\mathbf{J}}_2 \quad (2.41)$$

En soustrayant l’équation 2.38 multipliée par $\tilde{\mathbf{H}}_2$ à l’équation 2.41 multipliée par $\tilde{\mathbf{E}}_1$ et de même avec les équations 2.39 et 2.40 :

$$\tilde{\mathbf{E}}_1 \cdot \nabla \times \tilde{\mathbf{H}}_2 - \tilde{\mathbf{H}}_2 \cdot \nabla \times \tilde{\mathbf{E}}_1 = j\omega\tilde{\mathbf{H}}_2 \cdot \tilde{\mathbf{B}}_1 + j\omega\tilde{\mathbf{E}}_1 \cdot \tilde{\mathbf{D}}_2 + \tilde{\mathbf{E}}_1 \cdot \tilde{\mathbf{J}}_2 \quad (2.42)$$

$$\tilde{\mathbf{E}}_2 \cdot \nabla \times \tilde{\mathbf{H}}_1 - \tilde{\mathbf{H}}_1 \cdot \nabla \times \tilde{\mathbf{E}}_2 = j\omega\tilde{\mathbf{H}}_1 \cdot \tilde{\mathbf{B}}_2 + j\omega\tilde{\mathbf{E}}_2 \cdot \tilde{\mathbf{D}}_1 + \tilde{\mathbf{E}}_2 \cdot \tilde{\mathbf{J}}_1 \quad (2.43)$$

Puis, en utilisant l'identité $\mathbf{A} \cdot \nabla \times \mathbf{B} - \mathbf{B} \cdot \nabla \times \mathbf{A} = -\nabla \cdot (\mathbf{A} \times \mathbf{B})$ on trouve :

$$-\nabla \cdot (\tilde{\mathbf{E}}_1 \times \tilde{\mathbf{H}}_2) = j\omega \tilde{\mathbf{H}}_2 \cdot \tilde{\mathbf{B}}_1 + j\omega \tilde{\mathbf{E}}_1 \cdot \tilde{\mathbf{D}}_2 + \tilde{\mathbf{E}}_1 \cdot \tilde{\mathbf{J}}_2 \quad (2.44)$$

$$-\nabla \cdot (\tilde{\mathbf{E}}_2 \times \tilde{\mathbf{H}}_1) = j\omega \tilde{\mathbf{H}}_1 \cdot \tilde{\mathbf{B}}_2 + j\omega \tilde{\mathbf{E}}_2 \cdot \tilde{\mathbf{D}}_1 + \tilde{\mathbf{E}}_2 \cdot \tilde{\mathbf{J}}_1 \quad (2.45)$$

En soustrayant ces deux équations on trouve :

$$\tilde{\mathbf{E}}_2 \cdot \tilde{\mathbf{J}}_1 - \tilde{\mathbf{E}}_1 \cdot \tilde{\mathbf{J}}_2 = \nabla \cdot (\tilde{\mathbf{E}}_1 \times \tilde{\mathbf{H}}_2 - \tilde{\mathbf{E}}_2 \times \tilde{\mathbf{H}}_1) - j\omega (\tilde{\mathbf{E}}_2 \cdot \tilde{\mathbf{D}}_1 - \tilde{\mathbf{E}}_1 \cdot \tilde{\mathbf{D}}_2 + \tilde{\mathbf{H}}_1 \cdot \tilde{\mathbf{B}}_2 - \tilde{\mathbf{H}}_2 \cdot \tilde{\mathbf{B}}_1) \quad (2.46)$$

Et en intégrant sur le volume V formé par la surface S et en appliquant le théorème de la divergence :

$$\begin{aligned} & \iiint_V \tilde{\mathbf{E}}_2 \cdot \tilde{\mathbf{J}}_1 dv - \iiint_V \tilde{\mathbf{E}}_1 \cdot \tilde{\mathbf{J}}_2 dv \\ &= \oint_S (\tilde{\mathbf{E}}_1 \times \tilde{\mathbf{H}}_2 - \tilde{\mathbf{E}}_2 \times \tilde{\mathbf{H}}_1) \cdot \hat{\mathbf{n}} ds \\ & - j\omega \iiint_V (\tilde{\mathbf{E}}_2 \cdot \tilde{\mathbf{D}}_1 - \tilde{\mathbf{E}}_1 \cdot \tilde{\mathbf{D}}_2 + \tilde{\mathbf{H}}_1 \cdot \tilde{\mathbf{B}}_2 - \tilde{\mathbf{H}}_2 \cdot \tilde{\mathbf{B}}_1) dv \end{aligned} \quad (2.47)$$

Cette équation, appelée le *Lorentz lemma* [Asa+20], est valable pour n'importe quel milieu, réciproque ou non. Ensuite, il faut remarquer que l'intégrale de surface disparaît puisque la surface peut toujours être étendue à l'infini (par rapport aux émetteurs) et dans ce cas les champs électriques et magnétiques sont liés par la relation $\hat{\mathbf{n}} \times \tilde{\mathbf{E}}_{1,2} = \eta \tilde{\mathbf{H}}_{1,2}$ (avec η inchangé pour les deux émetteurs) [Cal+18] ; [Asa+20]. On trouve ainsi :

$$\iiint_V \tilde{\mathbf{E}}_2 \cdot \tilde{\mathbf{J}}_1 dv - \iiint_V \tilde{\mathbf{E}}_1 \cdot \tilde{\mathbf{J}}_2 dv = -j\omega \iiint_V (\tilde{\mathbf{E}}_2 \cdot \tilde{\mathbf{D}}_1 - \tilde{\mathbf{E}}_1 \cdot \tilde{\mathbf{D}}_2 + \tilde{\mathbf{H}}_1 \cdot \tilde{\mathbf{B}}_2 - \tilde{\mathbf{H}}_2 \cdot \tilde{\mathbf{B}}_1) dv \quad (2.48)$$

Si le membre de droite de l'équation 2.48 est égal à zéro, on obtient donc que le membre de gauche est égal à zéro et on retrouve ainsi la relation de réciprocité de Lorentz :

$$\iiint_V \tilde{\mathbf{E}}_2 \cdot \tilde{\mathbf{J}}_1 dv = \iiint_V \tilde{\mathbf{E}}_1 \cdot \tilde{\mathbf{J}}_2 dv \quad (2.49)$$

Cette relation stipule que le champ électrique, généré par un émetteur, et mesuré au

niveau d'un récepteur est le même si le rôle émetteur - récepteur est échangé. Elle peut aussi s'expliquer avec le concept de réaction introduit par Victor Henry Rumsey [Rum54]. Et on voit que l'équation 2.48 nous donne la condition selon laquelle cette relation est valable :

$$\iiint_V \left(\tilde{\mathbf{E}}_2 \cdot \tilde{\mathbf{D}}_1 - \tilde{\mathbf{E}}_1 \cdot \tilde{\mathbf{D}}_2 + \tilde{\mathbf{H}}_1 \cdot \tilde{\mathbf{B}}_2 - \tilde{\mathbf{H}}_2 \cdot \tilde{\mathbf{B}}_1 \right) dv = 0 \quad (2.50)$$

Ensuite, il faut remarquer qu'entre le cas 1 et 2, les ondes se propagent dans des directions opposées entre l'émetteur et le récepteur. Le milieu de propagation est le même dans ces deux cas, on laissera juste la possibilité de changer le sens du biais lors du cas 2, ce qui permettra d'englober tous les cas possibles :

$$\tilde{\tilde{\epsilon}}_1 = \tilde{\tilde{\epsilon}}(\mathbf{F}_0) \quad (2.51)$$

$$\tilde{\tilde{\mu}}_1 = \tilde{\tilde{\mu}}(\mathbf{F}_0) \quad (2.52)$$

$$\tilde{\tilde{\epsilon}}_2 = \tilde{\tilde{\epsilon}}(\pm \mathbf{F}_0) \quad (2.53)$$

$$\tilde{\tilde{\mu}}_2 = \tilde{\tilde{\mu}}(\pm \mathbf{F}_0) \quad (2.54)$$

En utilisant les relations constitutives du milieu, on trouve :

$$\iiint_V \left(\tilde{\mathbf{E}}_2 \cdot \tilde{\tilde{\epsilon}}(\mathbf{F}_0) \cdot \tilde{\mathbf{E}}_1 - \tilde{\mathbf{E}}_1 \cdot \tilde{\tilde{\epsilon}}(\pm \mathbf{F}_0) \cdot \tilde{\mathbf{E}}_2 + \tilde{\mathbf{H}}_1 \cdot \tilde{\tilde{\mu}}(\mathbf{F}_0) \cdot \tilde{\mathbf{H}}_2 - \tilde{\mathbf{H}}_2 \cdot \tilde{\tilde{\mu}}(\pm \mathbf{F}_0) \cdot \tilde{\mathbf{H}}_1 \right) dv = 0 \quad (2.55)$$

Et en utilisant l'identité $\mathbf{a} \cdot \overline{\overline{\mathbf{X}}} \cdot \mathbf{b} = (\mathbf{a} \cdot \overline{\overline{\mathbf{X}}} \cdot \mathbf{b})^T = \mathbf{b} \cdot \overline{\overline{\mathbf{X}}}^T \cdot \mathbf{a}$ on trouve finalement :

$$\iiint_V \left(\tilde{\mathbf{E}}_2 \cdot \left[\tilde{\tilde{\epsilon}}(\mathbf{F}_0) - \tilde{\tilde{\epsilon}}^T(\pm \mathbf{F}_0) \right] \cdot \tilde{\mathbf{E}}_1 + \tilde{\mathbf{H}}_1 \cdot \left[\tilde{\tilde{\mu}}(\mathbf{F}_0) - \tilde{\tilde{\mu}}^T(\pm \mathbf{F}_0) \right] \cdot \tilde{\mathbf{H}}_2 \right) dv = 0 \quad (2.56)$$

Cette équation représente la forme généralisée du théorème de réciprocité de Lorentz pour un milieu LTI anisotrope [Cal+18]. Pour qu'elle soit vraie pour tout champ, on doit avoir :

$$\tilde{\tilde{\epsilon}}(\mathbf{F}_0) = \tilde{\tilde{\epsilon}}^T(\pm \mathbf{F}_0) \quad (2.57)$$

$$\tilde{\tilde{\mu}}(\mathbf{F}_0) = \tilde{\tilde{\mu}}^T(\pm \mathbf{F}_0) \quad (2.58)$$

Pour qu'un milieu soit réciproque, il faut que ces relations soient vérifiées. Dans le cas où

le biais n'est pas modifié ($\tilde{\tilde{\epsilon}}_2 = \tilde{\tilde{\epsilon}}(\mathbf{F}_0)$ et $\tilde{\tilde{\mu}}_2 = \tilde{\tilde{\mu}}(\mathbf{F}_0)$), il y aura toujours inégalité de l'une de ces deux relations, traduisant la non réciprocity du milieu [Cal+18]. Prenons par exemple un plasma sans perte magnétisé par un biais $\mathbf{F}_0 = \mathbf{H}_0$. Sa perméabilité correspond à celle du vide et sa permittivité peut s'écrire [Ish17] :

$$\tilde{\tilde{\epsilon}}(\pm\mathbf{H}_0) = \begin{bmatrix} \epsilon & jq(\pm\mathbf{H}_0) & 0 \\ -jq(\pm\mathbf{H}_0) & \epsilon & 0 \\ 0 & 0 & \epsilon_z \end{bmatrix} \quad (2.59)$$

Avec :

$$q(\mathbf{H}_0) = +\frac{\omega_c\omega_p^2}{\omega(\omega_c^2 - \omega^2)} \text{ lorsque } \mathbf{H}_0 \text{ est orienté selon } +z \quad (2.60)$$

$$q(-\mathbf{H}_0) = -\frac{\omega_c\omega_p^2}{\omega(\omega_c^2 - \omega^2)} \text{ lorsque } \mathbf{H}_0 \text{ est orienté selon } -z \quad (2.61)$$

$$\epsilon = 1 + \frac{\omega_p^2}{\omega_c^2 - \omega^2} \quad (2.62)$$

$$\epsilon_z = 1 + \frac{\omega_p^2}{\omega^2} \quad (2.63)$$

et $\omega_c \propto \mathbf{H}_0$ la pulsation cyclotronique et ω_p la pulsation plasma.

On peut aisément voir que l'on trouve une égalité en inversant le sens du biais ($\tilde{\tilde{\epsilon}}(\mathbf{H}_0) = \tilde{\tilde{\epsilon}}^T(-\mathbf{H}_0)$) et une inégalité si le biais reste inchangé ($\tilde{\tilde{\epsilon}}(\mathbf{H}_0) \neq \tilde{\tilde{\epsilon}}^T(\mathbf{H}_0)$). Le même raisonnement vaut pour la perméabilité d'un ferrite magnétisé. Plus généralement un milieu possédant un biais (ferrite ou plasma magnétisés par exemple) est non-réciproque si ce biais reste inchangé, c'est à dire si on ne modifie pas le système. On a donc toujours :

$$\tilde{\tilde{\epsilon}}(\mathbf{F}_0) \neq \tilde{\tilde{\epsilon}}^T(\mathbf{F}_0) \quad (2.64)$$

$$\tilde{\tilde{\mu}}(\mathbf{F}_0) \neq \tilde{\tilde{\mu}}^T(\mathbf{F}_0) \quad (2.65)$$

Il suffit que l'une ou l'autre de ses relations soit vraies pour que le milieu soit non-réciproque. Et si ce biais est inversé entre les deux sens de propagation, on trouve

$$\tilde{\tilde{\epsilon}}(\mathbf{F}_0) = \tilde{\tilde{\epsilon}}^T(-\mathbf{F}_0) \quad (2.66)$$

$$\tilde{\tilde{\mu}}(\mathbf{F}_0) = \tilde{\tilde{\mu}}^T(-\mathbf{F}_0) \quad (2.67)$$

Si l'une de ces deux équations n'est pas respectée, le milieu est non-réciproque [Cal+18].

Pour un milieu sans biais ($\mathbf{F}_0 = 0$) on retrouve à partir des équations 2.57 et 2.58 les relations classiques de réciprocité [Kon00] :

$$\tilde{\tilde{\epsilon}} = \tilde{\tilde{\epsilon}}^T \quad (2.68)$$

$$\tilde{\tilde{\mu}} = \tilde{\tilde{\mu}}^T \quad (2.69)$$

Une conséquence directe de ces relations est qu'un système LTI isotrope (sans biais) est réciproque. De plus les systèmes anisotropes dont les tenseurs de permittivité et de perméabilité sont symétriques sont réciproques. Dans la plupart des dispositifs, c'est la présence d'un biais (inchangé) qui rend un milieu non-réciproque. En revenant par exemple au plasma sans perte magnétisé (décrit par l'équation 2.59) on remarque que si on annule le biais ($\mathbf{F}_0 = \mathbf{H}_0 = \mathbf{0}$), la permittivité se réduit à un scalaire et ce milieu est donc réciproque (dans ce cas le biais est à l'origine à la fois de l'anisotropie et de la non-réciprocité).

Pour finir, il faut remarquer que dans les équations 2.66 et 2.67, qui sont les conditions générales de réciprocité (en présence d'un biais), les tenseurs transposés correspondent en fait aux propriétés du milieu 2 (équations 2.53 et 2.54). Ils traduisent donc la propagation des ondes dans le sens opposé au cas 1. On retrouve ainsi une similitude forte avec les équations 2.36 et 2.37 traduisant une opération de réversibilité restreinte.

La différence entre réversibilité (au sens strict) et réciprocité apparaît clairement en regardant leurs conditions de validité sur les propriétés du milieu. En effet, pour la réversibilité, l'opération de conjugaison (équations 2.32-2.33) signifie une "inversion temporelle du milieu" (la conjugaison correspond dans le domaine temporel à changer t en $-t$) et pour la réciprocité, l'opération de transposition (équations 2.66-2.67) signifie une "inversion spatiale du milieu". Cependant pour la réversibilité restreinte, l'opération de conjugaison est abandonnée, résultant en un procédé très proche de la réciprocité. Afin de comprendre les liens entre ces concepts, ils sont appliqués dans la partie suivante à des cas simples.

2.2.4 Exemples

L'objectif de cette partie est d'illustrer sur des exemples simples comment se traduisent les concepts de réversibilité, réversibilité restreinte et réciprocité. Nous proposerons à la fin de chapitre d'utiliser ces concepts pour décrire l'amorçage d'un plasma par RT. Nous allons pour cela nous intéresser à trois systèmes LTI nommé A, B et C, aux propriétés différentes :

	Sans pertes	Symétrie des tenseurs de permittivité et de perméabilité
Système A	✓	✓
Système B	✗	✓
Système C	✓	✗

Les propriétés sans pertes des systèmes A et C englobent les milieux théoriques sans perte (nous avons vu précédemment qu'un milieu causal est à pertes) et les milieux réels à pertes négligeables (par exemple, pour certaines bande de fréquences microondes, l'air peut être considéré comme un milieu sans perte). L'asymétrie d'un des tenseurs (de permittivité ou de perméabilité) du système C est une conséquence de la présence d'un biais $\mathbf{F}_0$.

D'après les discussions précédentes, un système est réversible au sens strict du terme si en "inversant" le milieu (*i.e.* conjugué et inversion du biais) et les quantités électromagnétiques, le processus se déroule dans le sens inverse. Cette opération est représentée sur la première ligne de la Figure 2.8. Lors de la phase retournée, les pertes deviennent un gain et le sens du biais est inversé. On retrouve que cette opération n'a pas vraiment d'intérêt, les trois systèmes étant réversibles. De plus, leur propriétés sont changées lors de la phase retournée, donc ce n'est finalement pas vraiment le même système que l'on étudie.

	Système A : Sans perte Tenseur symétrique	Système B : Avec pertes Tenseur symétrique	Système C : Sans perte Tenseur non-symétrique
Phase retournée avec inversion du milieu			
Réversibilité	✓	✓	✓
Phase retournée sans inversion du milieu			
Réversibilité restreinte	✓	✗	✓
Inversion émetteur - récepteur			
Réciprocité	✓	✓	✗

FIGURE 2.8 – Illustration de la réversibilité et de la réciprocité (inspiré de [Cal+18] et [Asa+20]).

La deuxième ligne de la Figure 2.8, correspondant à une opération de réversibilité res-

treinte, est plus intéressante. Dans ce cas là les trois systèmes restent inchangés entre les deux phases (on choisit ici quand même d'inverser le sens du biais par souci pédagogique, comme expliqué plus bas). Le système A est bien réversible aussi au sens de la réversibilité restreinte puisqu'il est sans pertes (et sans biais). Dans ce cas, réversibilité (équations 2.32 et 2.33) et réversibilité restreinte (équations 2.36 et 2.37) sont équivalentes. Pour le système B, les champs électromagnétiques subissent les pertes deux fois (une fois entre les états $\Psi\{\mathbf{r}, 0\}$ et $\Psi\{\mathbf{r}, T\}$ et une deuxième fois entre les états $\Psi_r\{\mathbf{r}, -T\}$ et $\Psi_r\{\mathbf{r}, 0\}$). Dans ce cas $\Psi\{\mathbf{r}, t\} \neq \Psi_r\{\mathbf{r}, -t\}$ (sauf au temps $t = T$) et donc ce système n'est pas réversible au sens de la réversibilité restreinte. On comprend ici pourquoi un système à pertes ne peut pas être réversible au sens restreint. Les ondes subissent les pertes durant le trajet direct et on leur demande ensuite de repasser une nouvelle fois dans le système, subissant de nouveau les pertes qui sont irréversibles. Pour le système C, si le biais est inversé entre les deux phases, alors les relations 2.36-2.37 sont vérifiées et ce milieu est réversible au sens restreint.

La troisième ligne de la Figure 2.8 correspond à une opération d'inversion émetteur - récepteur entre deux points 1 et 2, respectivement aux positions $\mathbf{r} = \mathbf{r}_1$ et $\mathbf{r} = \mathbf{r}_2$. Le système dans l'état $\Psi_1\{\mathbf{r}_1, 0\}$ au point 1 donne au point 2 l'état $\Psi_{1 \rightarrow 2}\{\mathbf{r}_2, t\}$ (à un temps t). Si le système est réciproque, placer ce système dans le même état initial au point 2 ($\Psi_2\{\mathbf{r}_2, 0\} = \Psi_1\{\mathbf{r}_1, 0\}$) donne le même état au point 1 ($\Psi_{1 \leftarrow 2}\{\mathbf{r}_1, t\} = \Psi_{1 \rightarrow 2}\{\mathbf{r}_2, t\}$) pour tout temps t . Autrement dit, exciter le système au point 1 et regarder au point 2 ou l'exciter de la même façon au point 2 et regarder au point 1 donne exactement la même chose *à l'endroit auquel on regarde*. Par contre les choses peuvent être complètement différentes partout ailleurs dans le système. Cette différence entre les deux phases est représentée par le trait plein de 1 vers 2 différent du trait pointillé de 2 vers 1. Nous reviendrons sur ce point dans la partie suivante (voir Figure 2.9). Pour le système A, la symétrie des tenseurs de permittivité et de perméabilité se traduit par le fait que les chemins empruntés par ces dernières sont les mêmes dans un sens (de 1 vers 2) que dans l'autre (de 2 vers 1) et donc le système est réciproque. L'ajout de pertes dans le système B ne change pas ce processus, les tenseurs étant toujours symétriques. On voit là une différence entre la réversibilité restreinte et la réciprocité. La première fait subir aux ondes les pertes deux fois et donc un système à pertes sera forcément non réversible au sens restreint. Dans la deuxième l'état initial est le même dans un sens ou dans l'autre, et ainsi, si les tenseurs sont symétriques, les ondes subissent les pertes de la même façon dans un sens et dans l'autre. Pour le système C, le biais reste inchangé (on garde le même milieu), les équations 2.66-2.67 ne sont donc pas respectées et le milieu est donc non-réciproque. Ce qui se traduit par $\Psi_{1 \leftarrow 2}\{\mathbf{r}_1, t\} \neq \Psi_{1 \rightarrow 2}\{\mathbf{r}_2, t\}$.

Ces discussions peuvent être résumées simplement par :

$$\text{Pertes non négligeables} \Rightarrow \text{Non - Réversibilité (Restreinte)} \quad (2.70)$$

$$\text{Tenseur asymétrique} \Rightarrow \text{Non - Réciprocité} \quad (2.71)$$

Pour faire une analogie, on pourrait penser à la propagation d'eau (représentant les ondes) dans un tuyau (représentant un guide d'ondes). Un tuyau percé est non-réversible (au sens

restreint) puisque l'eau perdue par les fuites durant le trajet initial ne re-rentre pas dans le tuyau lors de l'inversion du sens de propagation de l'eau (et la quantité d'eau perdue par les fuites sera doublée à l'issue des deux phases). Par contre il est bien réciproque puisque si on envoie une certaine quantité d'eau dans un sens et qu'on regarde quelle quantité a réussi à traverser, on trouvera la même quantité indépendamment du sens dans lequel on fait couler l'eau. En introduisant une pompe créant un courant (représentant un biais) dans un des deux sens de propagation, le système devient non-réciproque. En effet, dans le sens de ce courant l'eau pourrait se propager sans difficulté mais dans le sens inverse l'eau serait ralentie par le courant induit par la pompe. L'eau resterait alors plus longtemps dans le tuyau et donc beaucoup plus d'eau serait perdue par les fuites. C'est exactement ce principe qui est utilisé dans les isolateurs à ferrite (le biais est dans ce cas un champ magnétique stationnaire). Les ondes peuvent passer dans un sens et sont "perdues" par les pertes dans l'autre.

Il est important de noter que le choix d'inverser le biais lors d'une opération de réversibilité restreinte (suivant les équations 2.36-2.37) a été fait pour des raisons pédagogiques, afin d'illustrer les deux propriétés ci-dessus. Cependant en pratique, on ne change pas le biais, le système reste le même et donc généralement les systèmes possédant un tenseur asymétrique (comme le système C) sont non-réversibles au sens restreint (en plus d'être non-réciproques).

Finalement, pour le milieu B, les opérations de réversibilité restreinte et de réciprocité ne semblent pas si différentes. En effet, le système est le même dans les deux cas et de plus l'opération de réversion temporelle se traduit en fait par la propagation des ondes dans le sens inverse, faisant penser à l'expérience de réciprocité. En fait, sa seule différence avec cette dernière semble être le fait que l'état du système n'est pas réinitialisé lors de la propagation dans le sens inverse. En un sens c'est vrai, mais seulement si on parle de ce qu'il se passe entre les points 1 et 2. Pour comprendre la différence entre ces deux concepts, nous proposons dans la section suivante de s'intéresser au RT en pratique. Nous verrons pour finir comment il peut être appliqué à l'amorçage de plasmas.

2.3 Le retournement temporel en pratique

2.3.1 Le retournement temporel en pratique : Réversible, réciproque, les deux ?

Pour comprendre la nuance entre la réciprocité et la réversibilité restreinte intéressons nous à ce que donnent ces concepts sur un cas pratique : la transmission entre deux antennes dipôles. Il correspond en fait à un RT avec une seule antenne dans le MRT, comme c'est souvent le cas pour les dispositifs basés sur le RT en électromagnétisme [Ler+04]. Nous verrons que le RT ressemble en pratique plus à une expérience de réciprocité que de réversibilité (restreinte ou pas).

Prenons un milieu réciproque d'impédance η , c'est à dire décrit par des tenseurs symétriques, comme c'est généralement fait en pratique dans les dispositifs à RT. Dans ce milieu nous disposons de deux antennes dipôles numérotées ① et ②. Lors de la phase initiale, le dipôle ① est excité par une tension $V_1(t)$ donnant un courant $I_1(t)$ qui lui même produit un champ rayonné. Le dipôle ②, utilisé en réception, capte une partie de l'énergie ainsi rayonnée par le dipôle ①, correspondant à celle qui emprunte les chemins possibles entre les deux dipôles. Il en résulte un courant $I_{1\rightarrow 2}(t)$ sur le dipôle ②. Cette étape initiale est représentée sur la première ligne de la Figure 2.9.

Lors d'une opération de réversibilité, tout est inversé : les ondes reviennent dans le sens inverse et convergent vers la source et les pertes du milieu et les résistances se transforment en gain ; la source devient $V_1(-t)$ (voir la première ligne de la Figure 2.9). Comme nous l'avons déjà dit cette opération n'est pas très intéressante, puisque difficilement réalisable en pratique.

Lors d'une opération de réversibilité restreinte, seulement la source et le sens de propagation des ondes sont inversés. On voit sur la Figure 2.9 toute la difficulté que représente la mise en place de cette opération en pratique. Pour pouvoir retourner toutes les ondes il faudrait entourer le dipôle ① de capteurs, ressemblant en fait à une cavité à RT, dont nous avons déjà parlé dans la section précédente (cela deviendrait similaire à ce qui est présenté en Figure 2.6). De plus pour faire une vraie expérience de réversibilité restreinte, il faut aussi retourner la source initiale ($V_1(t)$ devient $V_1(-t)$). En pratique, cette opération a été appelée puits à RT. Elle permet de focaliser l'énergie sur des dimensions inférieures à $\lambda/2$, du fait de l'annulation de l'onde divergente succédant la focalisation dans une expérience classique de RT [RF01]. Nous verrons dans les perspectives de cette thèse en quoi les plasmas amorcés par RT peuvent peut être jouer le rôle de puits à RT.

Pour la réciprocité, c'est le dipôle ② qui est excité par une tension $V_2(t)$, donnant un courant $I_2(t)$ qui lui même produit un champ rayonné (dernière ligne de la Figure 2.9). Le dipôle ①, en réception ici, capte une partie de l'énergie ainsi rayonnée par le dipôle ②, correspondant à celle qui emprunte les chemins possibles entre les deux dipôles (les mêmes que lors de la phase initiale). Il en résulte un courant $I_{1\leftarrow 2}(t)$ sur le dipôle ①. Puisque le système est réciproque, la relation de réciprocité 2.50 s'applique et elle devient simplement dans ce cas $V_1(t)I_{1\leftarrow 2}(t) = V_2(t)I_{1\rightarrow 2}(t)$ [Kon00]. Dans cette relation, les tensions correspondent aux

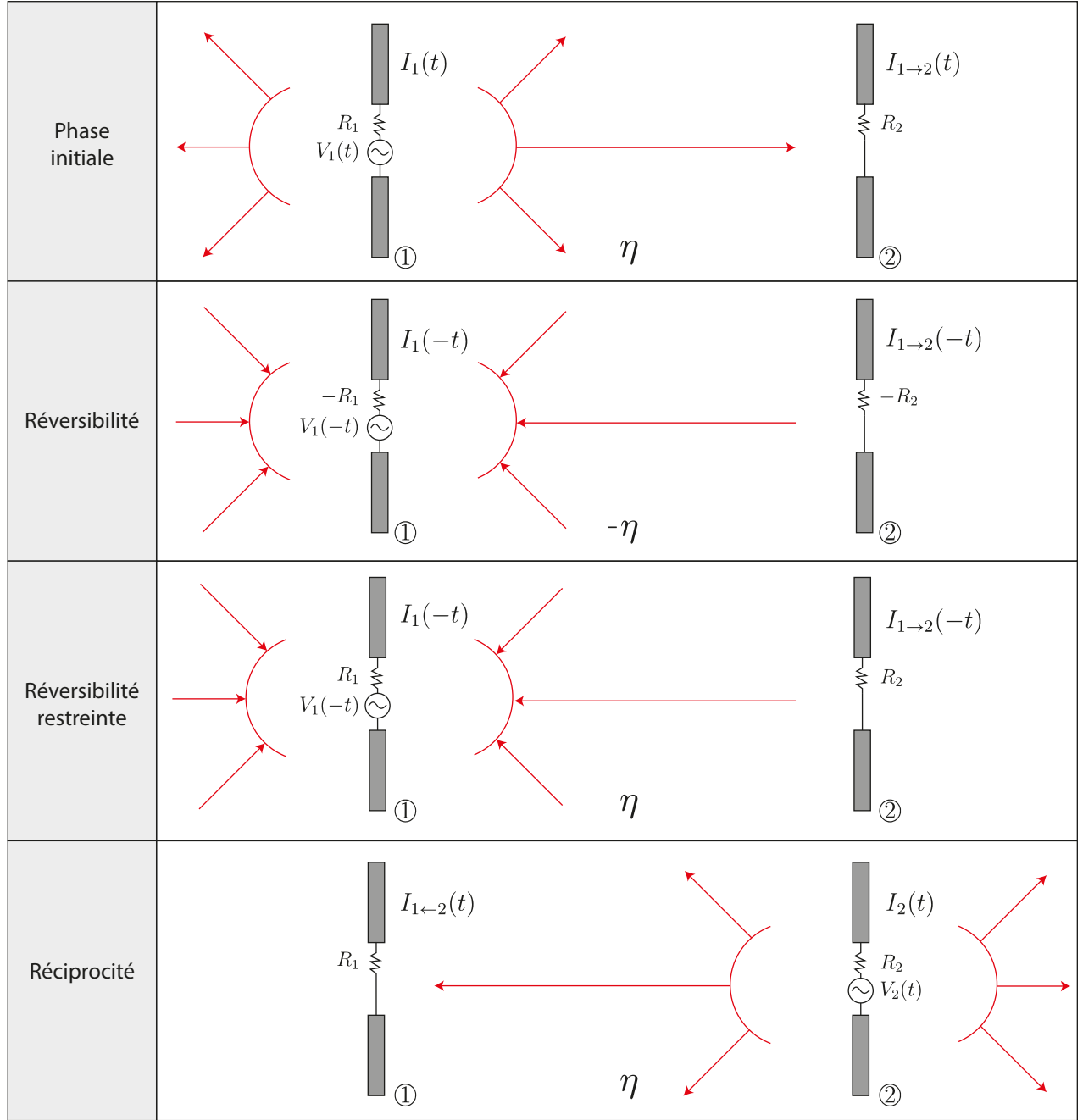


FIGURE 2.9 – Illustration de la réversibilité et de la réciprocité entre deux antennes dipôles (inspiré de [Asa+20]).

excitations et les courants aux effets de ces excitations. On voit par exemple qu'une même excitation $V_1(t) = V_2(t)$ donne bien les mêmes effets $I_{1\leftarrow 2}(t) = I_{1\rightarrow 2}(t)$. Dit autrement, ce système possède la même réponse lorsque les rôles d'émission et de réception sont inter-changés. Et pourtant, les deux scénarios (phase initiale et phase réciproque) semblent substantiellement différents, la répartition du champ entre les deux n'a pas grand chose à voir. La raison

est double : tout d'abord lors de la première phase (initiale), seulement une partie de l'énergie rayonnée par le dipôle ① est captée par le dipôle ②, ensuite lors de la deuxième phase (réciprocité) le dipôle ② rayonne l'énergie dans toutes les directions, pas seulement dans la direction dans laquelle il l'avait captée lors de la phase initiale. En fait ce qui est réciproque, c'est l'effet d'une source *au niveau* d'un récepteur, qui est le même si les rôles sont inversés.

Et il est important de comprendre que la réciprocité est valable pour deux antennes différentes, aux diagrammes de rayonnement complètement différents (contrairement aux deux dipôles de la Figure 2.9). Supposons par exemple deux antennes telles que le chemin entre les deux antennes correspond à un maximum de gain pour l'une et un minimum pour l'autre. Lorsque l'antenne à fort gain est excitée, une grande quantité d'énergie électromagnétique est donc rayonnée (par rapport à celle de l'excitation) le long de cette direction. Cette énergie arrive à la seconde antenne mais très peu est prélevée par celle-ci à cause de son faible gain (puisque à une certaine fréquence, une antenne rayonne de la même façon que ce qu'elle capte). Lors de l'expérience inverse, très peu d'énergie est rayonnée par l'antenne mais une grande quantité de celle-ci est prélevée par l'autre. Au final, ces deux scénarios donnent les mêmes bilans de transmission.

En fait, le concept de cavité à RT élaboré par Mathias Fink fait plus penser à une expérience de réciprocité que de réversibilité restreinte (ou de réversibilité). En effet, en pratique on ne contrôle pas la direction dans laquelle les ondes sont émises. En fait, la cavité à RT de la Figure 2.6 se rapprocherait plus de l'expérience de réciprocité de la Figure 2.9 avec plusieurs dipôles, entourant le dipôle ①, utilisés en réception lors de la phase initiale.

Finalement, le concept de réversibilité est peut être plus fondamental que celui de réciprocité, qui lui a plus de sens en pratique. On pourrait ainsi résumer le raisonnement qui a amené au développement du RT en pratique à partir de considérations théoriques comme suit :

1. Les équations de Maxwell sont réversibles.
2. Cependant cette réversibilité suppose l'inversion du milieu, ce qui n'est pas très intéressant en pratique donc on va lui préférer un concept plus pratique : la réversibilité restreinte.
3. Mettre en place cette dernière suppose de pouvoir inverser le champ dans toutes les directions de la phase initiale et d'inverser la source.
4. L'inversion de la source, même si elle est utile dans certains cas, n'est pas facile à mettre en place et limite les applications puisqu'il faut être capable de la réémettre à l'endroit de la focalisation. La réversibilité restreinte suppose de plus la capacité de contrôler la direction dans laquelle les ondes sont transmises lors de la deuxième étape.
5. Ainsi en pratique, on va finalement être amené à réaliser des dispositifs qui sont plus proches d'une expérience de réciprocité. Et ces dispositifs nécessitent deux phases :
 - Phase initiale : émission d'une impulsion au niveau d'une antenne. La partie de l'énergie électromagnétique ainsi rayonnée qui est capable d'être captée par l'antenne (ou les antennes) du MRT est enregistrée.

- Phase inverse : l'information recueillie lors de la première phase est inversée temporellement et émise par le MRT dans toutes les directions (selon son diagramme de rayonnement). Si le système est réciproque (tenseurs symétriques), les ondes vont pouvoir prendre les chemins inverses entre les deux antennes. Ainsi, le long de ces chemins et donc finalement au niveau de l'antenne initiale d'émission. On verra se passer un processus proche de l'impulsion inverse, se traduisant par une focalisation spatio-temporelle. Les paramètres influant sur la qualité de cette focalisation lors d'un processus de RT sont présentés dans le chapitre suivant.

Nous allons voir dans la partie suivante comment ces concepts se traduisent dans le cas d'amorçage de plasmas par RT.

2.3.2 Le retournement temporel appliqué au contrôle plasma

Nous allons nous intéresser dans cette partie aux plasmas amorcés par RT. L'objectif n'est pas de proposer des développements théoriques complets sur ces derniers, mais d'avantage une discussion et des pistes d'interprétation. Nous allons tout d'abord rappeler le principe, qui consiste à utiliser la focalisation de l'énergie électromagnétique par RT pour amorcer des plasmas locaux au moment et à l'endroit de cette focalisation. Ensuite, puisque le plasma est un milieu non-linéaire, nous discuterons de l'influence de cette dernière propriété sur le bon fonctionnement du processus de RT (puisque'il suppose la linéarité du système pour fonctionner). De plus, ce plasma n'est présent qu'à la fin de la seconde phase du RT, en conséquence le milieu de propagation diffère entre la première et la deuxième phase du RT. Nous verrons ainsi comment ce changement se traduit sur les concepts de réversibilité et de réciprocité. Notons que nous parlerons ici de réversibilité et de réciprocité du système vu dans son ensemble, en ce sens que les propriétés de ce système évoluent avec le temps en conséquence de l'apparition d'un plasma au moment de la focalisation. Les concepts de réversibilité et de réciprocité introduits précédemment sont classiquement utilisés sur des systèmes LTI, bien que récemment ces derniers sont largement étendus à des systèmes plus exotiques, tel que des milieux non-linéaire ou modulé dans le temps [Cal+18] ; [Asa+20]. Nous proposons dans cette partie d'utiliser ces concepts pour décrire l'amorçage (non-linéaire) d'un plasma par RT.

Le principe de l'amorçage d'un plasma d'argon par RT est présenté sur la Figure 2.10. Lors de la phase initiale (courbe rouge), une impulsion est émise au temps $t = 0$ (état $\Psi\{\mathbf{r}, 0\}$) dans un gaz (décrit par ε_0, μ_0) à une pression p_{init} choisie de sorte à ce que le champ électrique des ondes soit inférieur au champ de claquage (en conséquence le milieu restera décrit par ε_0, μ_0). Les ondes ainsi créées divergent par rapport au point source jusqu'au temps T (état $\Psi\{\mathbf{r}, T\}$). Lors de la phase inverse (courbe bleue), le sens de propagation des ondes est inversé (état $\Psi_r\{\mathbf{r}, -T\}$) et les ondes convergent vers la source, du temps $t = -T$ au temps $t = 0$. Si la pression du gaz est la même que lors de la phase initiale p_{init} , le champ électrique restera inférieur au champ de claquage et les propriétés du milieu seront les mêmes (ε_0, μ_0). Dans ce cas les ondes convergent bien vers la source jusqu'au moment $t = 0$, la relation $\Psi_r\{\mathbf{r}, -t\} = \Psi\{\mathbf{r}, t\}$ est toujours vérifiée (courbe bleue) et donc le système est réversible (au sens général ou non).

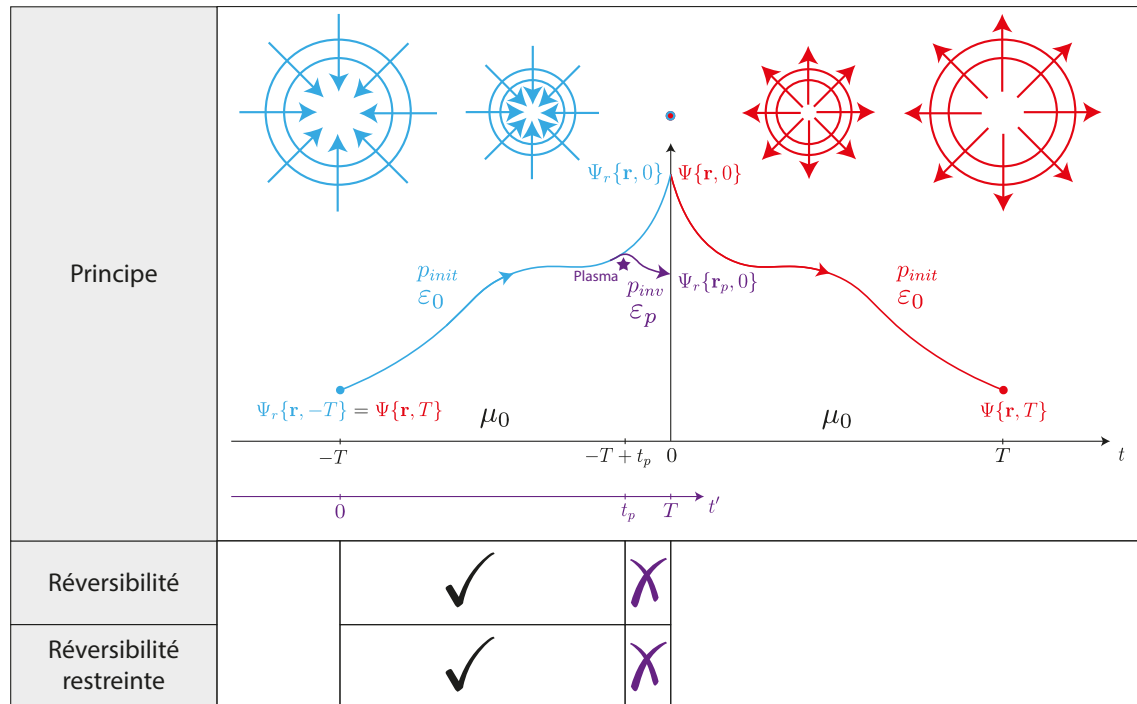


FIGURE 2.10 – Illustration du principe de l’amorçage d’un plasma par RT.

Le principe de l’amorçage d’un plasma par RT consiste à utiliser l’énergie électromagnétique de la focalisation pour claquer le gaz. Comme nous l’avons vu dans le premier chapitre, il faut pour cela que la norme du champ électrique dépasse un certain seuil : le champ de claquage $E_{claquage}$. Ce dernier dépend notamment de la pression du gaz. Pour un champ oscillant à une fréquence donnée f_0 , il existe une pression p_{opt} pour laquelle le transfert d’énergie aux électrons est optimal et donc le champ claquage minimal. Pour amorcer un plasma par RT, lors de la phase inverse, il existe donc deux démarches possibles : soit le niveau du champ électrique des ondes est augmenté soit la pression est changée pour être plus proche de la pression optimale que lors de la phase initiale. Par souci pédagogique nous allons présenter le principe de l’amorçage d’un plasma par RT avec la seconde démarche, mais un raisonnement similaire peut être fait avec la première. Cette seconde démarche consiste donc à changer la pression de la phase inverse par rapport à celle de la phase initiale de sorte à ce qu’à cette pression le champ électrique des ondes convergentes au moment de la focalisation soit suffisant pour claquer le gaz. Cette pression est appelée pression de la phase inverse, notée p_{inv} . L’état $\Psi_r\{\mathbf{r}, -t\}$ correspondant à cette opération est représenté en violet. Par souci de clarté, une nouvelle origine des temps est choisie pour décrire l’opération inverse, noté t' (cours du temps violet). À cette pression, le champ électrique de la focalisation à l’endroit initial de la source est suffisant pour claquer le gaz à un temps noté $t' = t_p$. À ce moment là, un plasma est généré à l’endroit initial de la source et l’état du système n’est donc plus le même. À partir de ce temps ($t' > t_p$), à cet endroit $\mathbf{r} = \mathbf{r}_p$ le milieu est décrit par la permittivité plasma $\varepsilon_p \neq \varepsilon_0$ et μ_0 . Le système est donc réversible (au sens restreint ou non) jusqu’à ce temps t_p , le milieu étant décrit par ε_0, μ_0 comme lors de la phase initiale. Par contre après ce temps t_p , l’état du système n’est plus symétrique par rapport au temps $t = 0$. En effet au point $\mathbf{r} = \mathbf{r}_p$ le milieu

change localement par rapport à la phase initiale, *i.e.* $\Psi_r\{\mathbf{r}_p, -t\} \neq \Psi\{\mathbf{r}_p, t\}$ pour $t' > t_p$. En conséquence de la non-linéarité présente lors de la phase inverse, le système n'est alors plus réversible (au sens restreint ou non) à partir de ce temps $t' = t_p$. Nous proposons alors de parler de **réversibilité transitoire**.

Nous avons vu dans ce chapitre qu'en pratique une expérience de RT ressemble plus à une expérience de réciprocité. Nous avons vu qu'un système composé de deux antennes est réciproque si le signal mesuré sur une antenne suite à l'excitation de l'autre antenne est le même lorsque les rôles émetteur - récepteur sont inversés. Pour une expérience de RT, cette propriété est essentielle puisqu'elle assure que les ondes empruntent les mêmes chemins entre les deux antennes lors de la phase initiale et de la phase inverse. Lors de l'amorçage de plasmas par RT, la question de savoir si le système est réciproque est complexe. En fait tout dépend de ce que nous entendons par "système" et surtout comment il est considéré dans le temps. En effet, le système au sens de milieu de propagation évolue dans le temps puisqu'un plasma est généré au moment de la focalisation. Cependant, si on considère le milieu dans lequel se propagent les ondes lors de la phase initiale, ce dernier est bien réciproque puisque c'est tout simplement de l'air. Lors de la phase inverse, un plasma est amorcé au moment et à l'endroit de la focalisation par RT. Si on considère le système cavité/antennes + plasma ainsi obtenu, il est aussi réciproque, en ce sens qu'une fois le plasma présent, le signal mesuré sur une antenne suite à l'excitation de l'autre antenne est le même lorsque les rôles émetteur - récepteur sont inversés. Par contre du point de vue du processus de RT, le chemin emprunté par les ondes lors de la phase inverse change au moment et à l'endroit de la focalisation par rapport à la phase initiale. En ce sens, nous proposons de parler dans notre cas de "**réciprocité transitoire**". Le système que l'on considère alors est dynamique, *i.e.* ses propriétés évoluent avec le temps, modifiant dans le temps le comportement des ondes dans le système. Et de la même façon que pour la réversibilité transitoire, on peut alors considérer ce système comme réciproque uniquement jusqu'au temps $t' = t_p$.

Ainsi dans le cas de plasmas amorcés par RT, on peut parler de réversibilité (au sens restreint ou non) et de réciprocité *transitoires*, conséquence de la non-linéarité du système lors de la phase inverse. En effet cette non-linéarité a pour conséquence un changement local de propriété du milieu : la permittivité change de celle du vide ε_0 à la permittivité plasma ε_p (dont nous avons discuté au chapitre précédent) à la position $\mathbf{r}_p$ à l'instant t_p . En fait c'est la non-linéarité du claquage du gaz qui permet de contrôler l'instant et la position des plasmas amorcés par RT. Il n'y a qu'au moment (vers t_p) et à l'endroit ($\mathbf{r}_p$) de la focalisation que le champ est assez élevé pour faire changer significativement la permittivité du milieu. Cette non-linéarité n'est pas présente lors de la phase initiale (pour le même niveau de champ électrique). Le fait que la réversibilité (au sens restreint ou non) et la réciprocité soient transitoires est donc une conséquence du changement de pression lors de la phase inverse qui rend le milieu non-linéaire.

Nous pouvons aussi noter qu'une autre façon de présenter l'amorçage de plasmas par RT aurait pu être la suivante (première démarche présentée plus haut) : la pression serait la même lors des deux phases (initiale et inverse) mais ce serait le niveau du champ électrique porté par les ondes qui serait sensiblement augmenté entre la phase initiale et la phase inverse. De

cette façon lors de la phase initiale les ondes se propagent dans un milieu linéaire décrit par ε_0, μ_0 et lors de la phase inverse le milieu devient non-linéaire à cause du niveau élevé du champ électrique. Cette façon de présenter est équivalente à celle présentée sur la Figure 2.10, l'élément clé étant que le milieu devient non-linéaire lors de la phase inverse.

Afin de générer en pratique des plasmas par RT, nous avons vu que le contrôle de la pression est un critère important, puisqu'il permet de se placer dans des conditions optimales de claquage. C'est pourquoi les plasmas par RT sont générés dans une cavité métallique dans laquelle le gaz et la pression sont contrôlés. Cette cavité est aussi indispensable pour pouvoir faire du RT avec un fort niveau de champ électrique en pratique (fort niveau requis pour claquer le gaz). En effet, le dispositif de cavité à RT (présenté sur la Figure 2.6) est très compliqué, voire impossible à mettre en place à cause du nombre trop important de transducteurs requis. C'est pourquoi un dispositif plus pratique, appelé le miroir à RT a été développé en environnement réverbérant par Carsten Draeger [DRA97]; [DF97]. Il permet de diminuer significativement le nombre de transducteurs nécessaires à une opération de RT, jusqu'à même n'en nécessiter qu'un seul dans une cavité chaotique. Son principe est expliqué dans le chapitre suivant.

CHAPITRE 2 : Ce qu'il faut retenir

L'objectif de ce chapitre était de comprendre les concepts théoriques sur lesquels reposent le RT. Ainsi nous avons pu discuter en fin de chapitre de l'influence de l'amorçage d'un plasma (non-linéaire) sur le processus de RT. Nous pouvons résumer cette discussion suivant la liste suivante :

- Le raisonnement qui a permis de construire théoriquement le RT s'appuie sur les discussions entre Boltzmann et Loschmidt autour du paradoxe apparent entre la réversibilité des phénomènes à l'échelle microscopique et l'irréversibilité de ces phénomènes à l'échelle macroscopique.
- Le RT en électromagnétisme trouve son origine dans la réversibilité des équations de Maxwell.
- La causalité n'est pas remise en cause par le RT.
- La réversibilité au sens théorique nécessite une "inversion temporelle" du milieu de propagation.
- On lui préfère alors un concept qui a plus de sens pratique : la réversibilité restreinte.
- Le RT repose aussi sur la réciprocité (en général conséquence de la réversibilité).
- En pratique, les expériences de RT sont plus proches d'une expérience de réciprocité que de réversibilité. En conséquence lors d'un processus de RT en cavité, des lobes secondaires temporels et spatiaux sont présents. La chapitre suivant se concentre sur l'étude de l'influence des propriétés de la cavité sur ces lobes.
- L'amorçage d'un plasma par RT nécessite un changement de comportement du milieu entre la phase initiale et la phase inverse du RT : soit par une augmentation du niveau du champ électrique, soit par un changement de pression. L'objectif de l'une ou l'autre de ces opérations est de rendre le milieu non-linéaire, avec la génération d'un plasma au moment et à l'endroit de la focalisation par RT. On peut alors parler de réversibilité et de réciprocité transitoires.

Le retournement temporel en cavité réverbérante appliqué au contrôle plasma : Théorie et modélisation FDTD

Sommaire

3.1	Le retournement temporel mono-voie en cavité réverbérante	90
3.1.1	Principe du RT mono-voie en cavité réverbérante	90
3.1.2	Les cavités réverbérantes	95
3.1.3	La focalisation temporelle	100
3.1.4	La focalisation spatiale	111
3.1.5	Synthèse	114
3.2	Modélisation du retournement temporel en cavité réverbérante : Modèle FDTD 2D	116
3.2.1	Présentation du modèle FDTD 2D	116
3.2.2	Équivalence simulation - expérience	118
3.2.3	Étude des propriétés spatio-temporelles du retournement temporel simulé	120
3.3	Modélisation du retournement temporel en cavité pour l'amorçage de plasmas : Couplage Maxwell-fluide	125
3.3.1	Principe	125
3.3.2	Description du modèle fluide FDTD 2D	127
3.3.3	Équivalence simulation - expérience	128
3.3.4	Plasmas amorcés par retournement temporel en simulation	130

Every path is the right path. Everything
could have been anything else and it
would have just as much meaning.

Mr Nobody

Ce chapitre a pour objectif de faire la transition entre le chapitre précédent, très théorique et le suivant présentant les résultats expérimentaux obtenus durant cette thèse. Il s'agit de

comprendre comment les concepts théoriques développés précédemment se traduisent en pratique. Nous utiliserons pour cela un outil répondant parfaitement à ce besoin : la simulation numérique. Nous verrons en effet que la simulation permet de visualiser et de comprendre les mécanismes se déroulant en pratique lors d'un RT ainsi que lors de l'amorçage de plasmas par RT.

Nous avons vu dans le chapitre précédent que le RT était plus proche d'une expérience de réciprocité que de réversibilité. En effet, lors de la phase inverse du RT, le milieu n'est pas inversé temporellement (pas de conjugaison des permittivité et perméabilité), la source initiale n'est pas inversée et les ondes sont émises dans toutes les directions, pas seulement selon celles le long desquelles elles ont été captées durant la phase initiale. De plus, la mise en place de la cavité à RT est compliquée, puisque cela nécessite le déploiement et la synchronisation d'un très grand nombre de transducteurs pour entourer la source. Il existe cependant une solution pour contourner ce problème. Elle consiste à tirer profit de la diffusion multiple en milieux de propagation complexes. Dans ce cas, lors de la phase initiale du RT, les ondes vont "rester" un certain temps dans le milieu en conséquence de leur diffusion. L'information mesurée par les transducteurs s'étale ainsi dans le temps sur une durée plus longue que celle de l'impulsion initiale. Lors de la phase inverse, chaque transducteur ré-émet cette information retournée temporellement, de sorte à ce que les ondes "subissent" le milieu de la même façon mais dans le sens de propagation inverse. Les ondes convergent ainsi vers la source initiale. Dans ce cas un nombre limité de transducteurs est suffisant pour obtenir une bonne focalisation. Ils constituent ce qui a été appelé un Miroir à Retournement Temporel (MRT). Les premières expériences de RT en milieu complexe ont été réalisées en acoustique par Arnaud Derode *et al.* avec 96 transducteurs dans le MRT [DRF95]. Puis Carsten Draeger *et al.* ont montré, toujours en acoustique, que le nombre de transducteurs nécessaires pouvait être drastiquement réduit en environnement réverbérant, jusqu'à ne nécessiter qu'un seul transducteur (on parle alors de RT mono-voie) [DF97]. Dans ce cas, les ondes issues de l'émission d'une impulsion durant la phase initiale sont stockées dans la cavité pendant un temps considérablement long par rapport à la durée de l'impulsion, à cause de leur réflexion sur les parois. Dans le domaine de l'électromagnétisme, Geoffroy Lerosey *et al.* ont montré le bon fonctionnement du RT mono-voie en cavité métallique [Ler+04].

Afin d'amorcer des plasmas par RT, l'utilisation d'une cavité métallique possède alors deux avantages : elle permet, comme nous venons de le dire, de n'utiliser qu'un seul transducteur dans le MRT et de contrôler le gaz et sa pression afin de se placer dans des conditions favorables pour obtenir un claquage dans l'enceinte (ce qui correspond au changement de pression lors de la phase inverse du RT présenté à la fin du chapitre précédent). En revanche, le RT mono-voie en cavité réverbérante a pour conséquence la présence de nombreux lobes secondaires temporels et spatiaux entourant le pic de RT [DF97], comme nous allons le voir dans ce chapitre. L'amorçage de plasmas par RT consiste alors à ce que le champ électrique obtenu au moment et à l'endroit de la focalisation par RT dans la cavité soit suffisant pour claquer localement le gaz. Pour un contrôle efficace de la position du claquage, il faut cependant s'assurer que le champ électrique soit inférieur à celui de la focalisation partout ailleurs dans la cavité et pendant toute la durée de l'expérience (sinon un plasma serait amorcé à un autre endroit et/ou à un autre instant que ceux désirés, *i.e.* l'endroit de la focalisation). Comme

nous allons le voir dans ce chapitre, il existe une propriété des cavités qui permet d'assurer ce dernier point : la chaotité.

Dans un premier temps, les équations permettant de décrire le RT mono-voie en environnement réverbérant sont présentées. Nous l'utiliserons dans le chapitre suivant pour l'amorçage expérimental de plasmas par RT. Pour cela, nous verrons qu'il est nécessaire de distinguer différents régimes de comportement des cavités en fonction de la fréquence à laquelle elles sont excitées. Nous montrerons, en utilisant les équations développées par Julien De Rosny [Ros00], l'influence de la chaotité de la cavité sur la qualité de la focalisation temporelle par RT. Nous finirons par discuter de son influence sur la qualité de la focalisation spatiale.

Ensuite, le code de simulation basé sur la méthode des différences finies dans le domaine temporel (ou FDTD pour Finite-difference time-domain [TH05]) en deux dimensions (2D) développé durant cette thèse est présenté. Cette méthode est intrinsèquement dispersive et elle n'est donc généralement pas recommandée pour la simulation de phénomènes en cavité réverbérante sur des temps longs tels que ceux d'un RT. Cependant, le processus de RT permet de compenser la dispersion subie par les ondes lors de la phase initiale [DF99]. De plus nous discuterons de ses avantages quand à l'observation des champs électromagnétiques, et plus précisément l'uniformité et l'isotropie de l'énergie électromagnétique induits par une cavité à géométrie chaotique. En conséquence, nous verrons que le RT en cavité chaotique possède des propriétés de focalisation spatio-temporelle supérieures à celles d'un RT en cavité régulière, comme prévu par les équations développées dans la première partie de ce chapitre.

Enfin, le modèle FDTD 2D est couplé avec un modèle plasma fluide afin de décrire l'amorçage de plasmas par RT en cavité réverbérante. Nous montrerons le bon contrôle de la position du plasma par RT dans la cavité chaotique, contrairement à la cavité régulière (dans les conditions dans lesquelles sont réalisées les simulations). En effet dans cette dernière, l'organisation spatiale régulière de l'énergie électromagnétique a pour conséquence la présence de lieux d'intensification du champ électrique dans la cavité, qui suffisent à amorcer des plasmas non désirés à d'autres endroits que celui de la focalisation par RT. Cela permet d'illustrer clairement le caractère essentiel de la chaotité sur le contrôle plasma par RT. Nous montrerons qu'une cavité chaotique permet de contrôler efficacement la position des plasmas par RT, dont la dimension est de l'ordre de la demi-longueur d'onde.

L'objectif de ce chapitre étant de préparer le lecteur aux résultats expérimentaux présentés dans le chapitre suivant, les paramètres de simulation du modèle électromagnétique (dimensions de la cavité, fréquence des ondes et durée de l'impulsion) et ceux du modèle fluide (pression, gaz...) sont choisis de sorte à ce que les simulations se déroulent dans des conditions similaires à celles des expériences présentées au chapitre suivant. Ainsi, en comparant les performances du RT obtenues expérimentalement et celles obtenues en simulation dans les cavités régulière et chaotique, nous pourrions conclure sur la chaotité de la cavité expérimentale utilisée et donc sur le contrôle des plasmas en son sein.

3.1 Le retournement temporel mono-voie en cavité réverbérante

Comme mentionné précédemment, la mise en place d’une expérience de RT nécessite en pratique l’utilisation d’un MRT. Dans notre cas, un seul transducteur est utilisé dans le MRT (RT mono-voie) et les expériences sont réalisées en cavité réverbérante (nécessaire à la fois pour le RT mono-voie et pour le contrôle de la pression). Cela a pour conséquence la présence de nombreux lobes secondaires temporels et spatiaux entourant le pic de RT [DF97], limitant la qualité de la focalisation. Pour caractériser la qualité d’un RT, on utilise généralement les concepts de contrastes temporel [Alb+19] ; [Ros00] ; [Ler06] et spatial [Alb+19]. Le contrôle des plasmas par RT dépend de la qualité du RT, *i.e.* de ces contrastes. Nous allons dans cette section développer les équations permettant de décrire le RT mono-voie en cavité réverbérante afin de comprendre quels sont les paramètres qui permettent d’augmenter ces contrastes, et donc d’améliorer le contrôle des plasmas par RT.

3.1.1 Principe du RT mono-voie en cavité réverbérante

Le problème considéré ici est celui d’un miroir à RT plongé dans un environnement diffusant ou réverbérant. Dans ce cas, on peut tirer profit des multiples réflexions auxquelles sont soumises les ondes pour réduire considérablement le nombre de transducteurs nécessaires à l’opération de RT, comparé à celui nécessaire pour la cavité à RT présentée au chapitre précédent (voir Figure 2.6). Nous allons illustrer ce principe dans le cas mono-voie (un seul transducteur dans le MRT) à l’aide de schémas dans une cavité quelconque, représentée sur la Figure 3.1. Deux transducteurs nommés ① et ② y sont positionnés. Le milieu est supposé réciproque dans la cavité. Afin de comprendre comment le RT mono-voie en cavité réverbérante fonctionne, nous allons décrire la phase initiale et les phases inverses sur la Figure 3.2. L’évolution temporelle du signal (émis ou mesuré selon les cas) au niveau du transducteur ① y est représentée en dessous du schéma de la cavité et celui au niveau du transducteur ② en dessus.

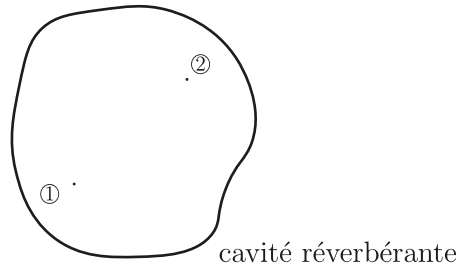


FIGURE 3.1 – Schéma de la cavité utilisée dans cette partie pour illustrer le RT mono-voie en cavité réverbérante.

Lors de la phase initiale (ligne du haut du tableau la Figure 3.2), une impulsion $e(t)$ de courte durée est émise à un temps t_0 au niveau du transducteur ①, utilisé dans ce cas comme émetteur. Les ondes vont alors se propager dans la cavité et y être stockées pendant

un temps supérieur à la durée de l'impulsion à cause des réflexions des ondes sur les parois. Le transducteur ②, utilisé comme récepteur, mesure alors la réponse à cette impulsion. Le signal ainsi mesuré est appelé réponse impulsionnelle¹, noté $r(t)$. Le niveau et le signe du signal mesuré dépendent du "chemin" emprunté par les ondes entre les deux transducteurs, à savoir le nombre de réflexions subies par les ondes et la longueur du chemin. Afin de comprendre les mécanismes d'un RT en pratique, trois chemins possibles entre les deux transducteurs sont représentés sur la Figure 3.2. Le trajet direct est représenté en pointillés triangles, celui le long duquel les ondes ont subi une réflexion en pointillés ronds et deux réflexions en pointillés carrés. Au temps t_1 le transducteur ② mesure le signal issu du trajet direct (trajet balistique) entre les deux transducteurs (représenté en pointillés triangles). Les ondes ont donc mis un temps $t_1 - t_0$ pour se propager de ① vers ② (sans réflexion). L'amplitude du lobe (représenté en rouge) alors mesuré est élevée puisque les ondes n'ont subies aucune réflexion. Au temps t_2 les ondes ayant empruntées le chemin en pointillés ronds sont mesurées par le transducteur ② (représenté en bleu). Puisqu'elles ont subi une réflexion sur les parois, leur signe est inversé et leur amplitude amoindrie (la durée de ce chemin est de $t_2 - t_0$). Finalement, au temps t_3 , les ondes ayant suivi le chemin en pointillés carrés arrivent au niveau du transducteur ② (représenté en vert), avec le même signe que celui du trajet balistique et une amplitude encore plus faible (puisqu'elles ont subi deux réflexions). L'amplitude mesurée sur la réponse impulsionnelle $r(t)$ suit ainsi une décroissance exponentielle avec le temps, à cause du nombre croissant des réflexions subies par les ondes dans la cavité. En pratique, les ondes n'arrêtent pas de se propager au niveau du transducteur ② mais par souci de clarté, nous omettons sur le schéma les chemins empruntés par ces ondes après avoir été mesurées par ce dernier.

Lors de la phase inverse, la réponse impulsionnelle retournée temporellement, notée $r(-t)$, est émise au niveau du transducteur ②, utilisé alors en émission (milieu du tableau de la Figure 3.2). L'objectif de cette opération est d'obtenir un pic de focalisation au niveau du transducteur ①, utilisé maintenant comme récepteur. Les ondes issues de chaque lobe du signal $r(-t)$ peuvent être suivies indépendamment. En fait, au niveau du transducteur ② chaque lobe est émis dans toutes les directions et non pas uniquement le long des directions selon lesquelles il a été capté lors de la phase initiale. Chaque lobe émis peut alors être vu comme une impulsion excitant la cavité de la même façon que lors de la phase initiale. Au temps t_0 est émis au niveau du transducteur ② le lobe mesuré en dernier lors de la phase initiale, avec la même amplitude et le même signe. Les chemins empruntés par les ondes ainsi que les lobes ainsi mesurés sur le transducteur ① sont représentées en vert. Ces chemins sont les mêmes que lors de la phase initiale, mais ils sont empruntés par les ondes dans le sens inverse. On retrouve alors sur le signal mesuré sur le transducteur ① la même forme que la réponse impulsionnelle $r(t)$ mesurée lors de la phase initiale (avec une amplitude plus faible), *i.e.* on retrouve les lobes aux temps t_1 , t_2 et t_3 (représentés en vert). Le même mécanisme a lieu lors de l'émission des deux autres lobes de $r(-t)$. Les chemins et les lobes mesurés des second et troisième lobes de $r(-t)$ sont représentés respectivement en bleu et rouge. Le lobe représenté en bleu est émis à un temps noté t'_0 . En notant $t'_1 = t'_0 + t_1$, $t'_2 = t'_0 + t_2$ et $t'_3 = t'_0 + t_3$, on retrouve bien les lobes aux temps t'_1 , t'_2 et t'_3 , avec un signe opposé. Enfin,

1. Par souci de clarté, nous appellerons réponse impulsionnelle - qui est par définition la réponse d'un système à une impulsion de Dirac - la réponse du système à une impulsion d'une certaine durée, décrivant un cas plus pratique.

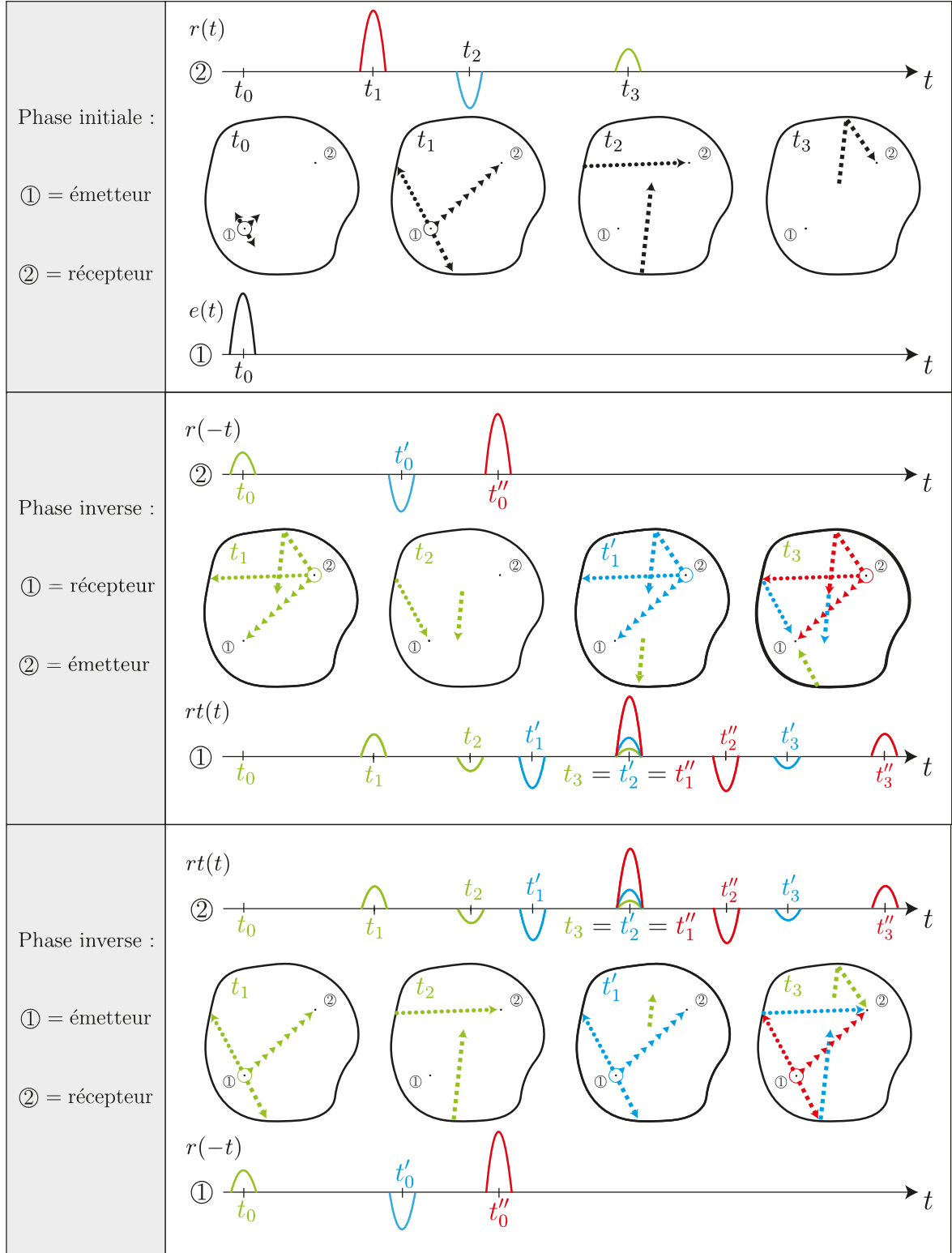


FIGURE 3.2 – Illustration du principe du miroir à RT mono-voie.

le lobe représenté en rouge est émis à un temps noté t''_0 . En notant $t''_1 = t''_0 + t_1$, $t''_2 = t''_0 + t_2$ et $t''_3 = t''_0 + t_3$, on retrouve bien les lobes aux temps t''_1 , t''_2 et t''_3 . Sur le signal mesuré sur le transducteur ①, noté $rt(t)$, nous obtenons au niveau de ce transducteur une sommation en phase des lobes au temps $t_3 = t'_2 = t''_1$, traduisant une focalisation par RT.

Nous venons d'illustrer l'origine de la présence de lobes secondaires temporels autour du pic de RT à l'endroit de la focalisation. En regardant la répartition spatiale des ondes au moment de la focalisation (au temps $t_3 = t'_2 = t''_1$), nous pouvons aussi remarquer la présence d'ondes dans la cavité ne participant pas à la focalisation, traduisant la présence de lobes secondaires spatiaux. Et ceci est aussi le cas pour les temps précédant la focalisation. Selon ce que nous venons de dire, les lobes secondaires temporels et spatiaux ont la même origine : l'émission des ondes lors de la phase inverse dans toutes les directions. Nous pouvons aussi noter que même si les ondes étaient émises le long de la direction selon laquelle elles ont été captées lors de la phase initiale, nous assisterions à la présence de lobes secondaires temporels, conséquence du fait que lors de la phase initiale les ondes peuvent, après un certain nombre de réflexions, "passer" au niveau du transducteur d'émission avant d'être mesurées par le transducteur en réception. Par exemple, lors de la phase initiale représentée sur la Figure 3.2, il existe plusieurs chemins (non représentés sur ce schéma) selon lesquels les ondes issues de l'émission de l'impulsion vont se réfléchir sur les parois de la cavité, repasser au niveau du transducteur ① avant d'être mesurées par le transducteur ②. Cela aura pour conséquence la présence de lobes secondaires temporels lors de la phase inverse, même si les ondes sont émises le long des directions selon lesquelles elles ont été captées lors de la phase initiale. Finalement, la présence de lobes secondaires possède deux origines : l'utilisation d'un environnement réverbérant et l'émission des ondes lors de la phase inverse dans toutes les directions.

La comparaison entre les configurations spatiales des ondes dans les phases initiale et inverse montre clairement que les choses sont différentes entre ces deux phases. Lors de la phase inverse, on peut alors difficilement parler de "film passé à l'envers" (image souvent utilisée pour expliquer le RT). En fait lors de cette phase, on assiste plutôt à la superposition du film de la phase initiale passé à l'envers avec d'autres films non présents initialement. Du point de vue des ondes, le principe de superposition nous dit que ce que l'on voit ce sont les ondes de la phase initiale dont le sens de propagation est inversé, interférant avec les ondes issues de l'émission dans toutes les directions lors de la phase inverse (qui n'étaient pas présentes lors de la phase initiale). Cela illustre bien les conclusions du chapitre précédent : en pratique le RT ressemble plus à une expérience de réciprocité que de réversibilité.

On voit que c'est la réciprocité du système qui est importante lors d'une expérience de RT. Elle assure en effet que les ondes se propagent de la même façon le long des chemins possibles entre les deux transducteurs. En conséquence, comme nous l'avons vu dans le chapitre précédent, un échange des rôles des transducteurs lors de la phase inverse donne le même signal $rt(t)$ au niveau du transducteur utilisé en réception. Cette opération est illustrée sur la dernière ligne du tableau de la Figure 3.2. Dans ce cas, le signal $r(-t)$ est émis au niveau du transducteur ①. Les ondes vont alors prendre les mêmes chemins possibles entre les transducteurs que lors de la phase inverse (ligne du milieu), mais dans le sens inverse. Au niveau du

transducteur ②, le même signal $rt(t)$ est alors mesuré. En effet, il suffit que les lobes de $r(-t)$ soient envoyées aux bons moments et avec le bon signe pour qu'au temps $t_3 = t'_2 = t''_1$ les ondes se somment en phase au niveau du transducteur en réception. On peut aussi remarquer que l'émission de l'impulsion $e(t)$ au niveau du transducteur ② aurait donné la même réponse impulsienne $r(t)$ au niveau du transducteur ①. D'un point de vue pratique cette propriété est très intéressante puisqu'elle permet de ne pas avoir à échanger le rôle récepteur - émetteur des deux transducteurs utilisés lors d'un RT (l'émission de la réponse impulsienne retournée se fait sur la même antenne que celle qui avait émise l'impulsion initiale). Par contre nous pouvons remarquer que la répartition spatiale des ondes est complètement différente entre les deux phases.

Comme nous l'avons mentionné dans le chapitre précédent, nous pouvons aussi noter que les ondes convergentes vers le point de focalisation vont continuer de se propager puisque la source n'est pas inversée au niveau du transducteur sur lequel la focalisation a lieu. À l'endroit (et au moment) de la focalisation, interfèrent alors une onde convergente et une divergente, ayant pour conséquence une dimension minimale de la focalisation égale à la demi-longueur d'onde. Notons au passage que lors de l'amorçage de plasmas au moment et à l'endroit de la focalisation, une certaine quantité de puissance contenue dans les ondes est transmise au gaz. Le plasma joue alors potentiellement le rôle d'un puits électromagnétique à RT, diminuant la dimension spatiale de la focalisation par RT. Nous rediscuterons de ce point dans les perspectives de cette thèse.

Nous avons ainsi expliqué pourquoi en pratique le pic de RT est entouré de lobes secondaires temporels et spatiaux. Généralement, ces lobes sont appelés bruits temporel [Alb+19] ; [Ros00] ; [Ler06] et spatial [Alb+19]. La qualité d'un RT va donc dépendre de sa capacité à discriminer le pic de RT par rapport au bruit. Comme mentionné précédemment, pour caractériser la qualité d'un RT on utilise généralement les concepts de contrastes temporel [Alb+19] ; [Ros00] ; [Ler06] et spatial [Alb+19]. En regardant le schéma de la Figure 3.2 on peut facilement soupçonner que ces contrastes dépendent des caractéristiques de la cavité (pertes, dimensions par rapport à la longueur d'onde, géométrie...). Avant de développer les équations de ces contrastes, il faut d'abord comprendre comment les propriétés de la cavité influent sur le comportement des ondes. Pour cela, notons que l'ensemble transducteur/cavité d'un dispositif est un système de transmission linéaire, continu et stationnaire (en l'absence de plasma). On peut donc assimiler cet ensemble à un filtre, et donc par définition à un système de convolution [Cot]. Nous noterons ainsi $h(t)$ la réponse d'un système entre deux transducteurs à une impulsion de Dirac $\delta(t)$. Nous allons voir dans la suite qu'il est utile d'étudier le comportement de l'ensemble transducteur/cavité dans le domaine fréquentiel. La réponse spectrale du système $H(f)$ est alors généralement appelée fonction de transfert de l'ensemble transducteur/cavité. Nous verrons ensuite comment les propriétés de la cavité influent sur la forme de cette fonction de transfert $H(f)$ et enfin comment cela se traduit sur les contrastes temporel et spatial du RT.

3.1.2 Les cavités réverbérantes

Avant de présenter les équations des contrastes temporel et spatial, il faut tout d'abord comprendre comment une cavité réagit à une excitation. La réponse de la cavité dépend bien entendu de ses propriétés mais aussi de l'excitation. Ces propriétés déterminent la façon dont l'énergie électromagnétique se répartit dans la cavité et conditionne donc la qualité de la focalisation par RT. Le raisonnement développé dans cette partie est réalisé en régime permanent et permet de comprendre le comportement des ondes dans la cavité en fonction de ces propriétés, à travers la fonction de transfert $H(f)$. Nous verrons ensuite dans les parties suivantes comment la fonction de transfert d'une cavité est sollicitée en régime transitoire lors d'un RT.

Pour commencer, notons que le champ électrique $\mathbf{E}(\mathbf{r}, f)$ obtenu au point $\mathbf{r}$ (contenu dans la cavité) en excitant une cavité par une source arbitraire $\mathbf{J}(\mathbf{r}_0, f)$ localisée en $\mathbf{r}_0$ peut s'écrire comme la somme des champs propres $\mathbf{E}_n$ associés aux résonances propres f_n [Gro14] :

$$\mathbf{E}(\mathbf{r}, f) = \sum_{n=1}^{\infty} a_n(f) \mathbf{E}_n(\mathbf{r}) \quad (3.1)$$

Dans le cas sans perte, les coefficients $a_n(f)$ deviennent des Dirac aux fréquences $f = f_n$ (comme nous l'avons vu au chapitre 1). Dans un cas plus réaliste, la présence de pertes se traduit par une certaine largeur spectrale des coefficients $a_n(f)$, notée Γ_n . La qualité d'une résonance est généralement caractérisée par le facteur de qualité de la cavité à la fréquence de résonance, défini comme $Q_n = f_n/\Gamma_n$. Il correspond au rapport entre l'énergie stockée et l'énergie dissipée (notamment par les parois). En première approximation, l'énergie stockée est proportionnelle à l'inverse de la longueur d'onde au cube (plus la longueur d'onde diminue et plus la cavité devient grande pour les ondes, l'énergie stockée est donc plus grande) et l'énergie dissipée est proportionnelle à l'inverse de la longueur d'onde au carré (puisque c'est au niveau des parois de la cavité que l'énergie est dissipée). Dans ces conditions le facteur de qualité est proportionnel à la fréquence f_n : $Q_n \propto (1/\lambda_n^3)/(1/\lambda_n^2) \propto 1/\lambda_n = f_n/c$. La largeur spectrale doit alors être constante, comme obtenu expérimentalement dans [MC14]. On prendra ainsi $\Gamma_n = \Gamma$. Mathématiquement, l'équation 3.1 stipule que le champ électrique $\mathbf{E}(\mathbf{r}, f)$ résulte de la contribution de tous les modes de résonance (somme infinie). En pratique les contributions $a_n(f)$ des modes dont la fréquence f_n est trop éloignée de f sont négligeables. Ainsi seuls certains modes, appelés modes significatifs [MC14], assez proches de la fréquence d'excitation f , influent significativement sur le champ électrique $\mathbf{E}(\mathbf{r}, f)$. Ces modes significatifs dépendent du régime modal de la cavité. De façon simple, ils peuvent être quantifiés par un paramètre sans dimension d , appelé recouvrement modal moyen, défini par [Gro14] :

$$d = \frac{\Gamma}{\delta f} \quad (3.2)$$

où δf représente l'espacement moyen entre les fréquences de résonances des modes. Et il existe trois régimes de recouvrement modal, qui dépendent des pertes (*i.e.* de Γ) et de l'espacement moyen entre les fréquences des modes δf (*i.e.* de la fréquence d'excitation). Ils sont représentés sur la Figure 3.3 :

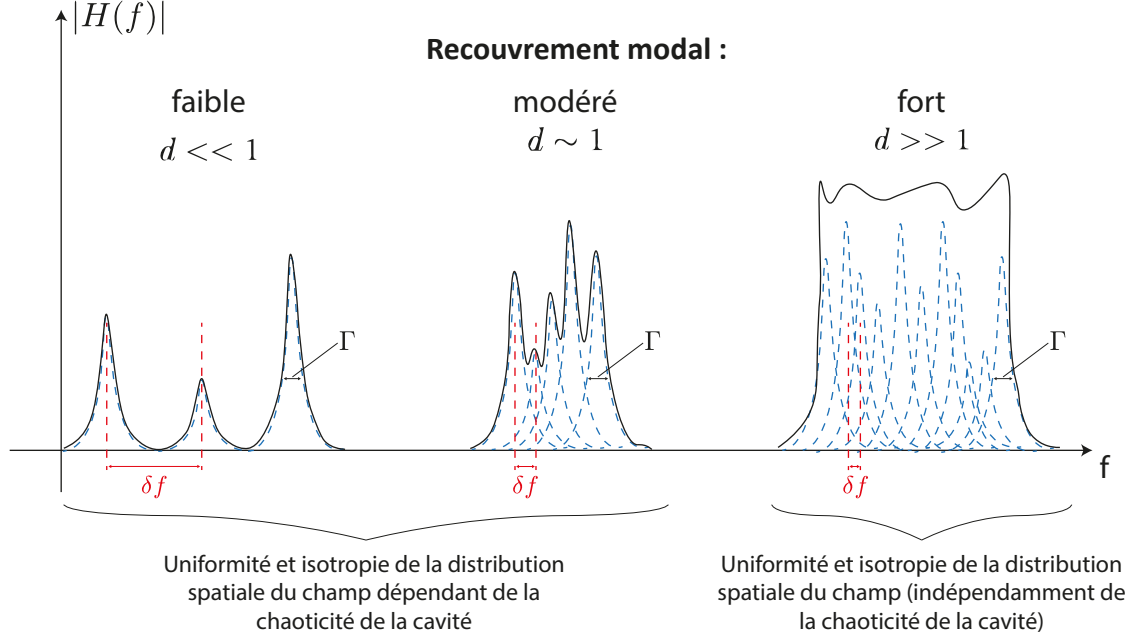


FIGURE 3.3 – Illustration des différents régimes de recouvrement modal (inspiré de [Gro14]).

- À basse fréquence (*i.e.* proche de la fréquence fondamentale de la cavité), les modes sont très éloignés les uns des autres ($d \ll 1$) et peuvent être vus comme indépendants les uns des autres. Le recouvrement modal est alors faible (cela équivaut au cas mono-modal dont nous avons discuté au chapitre 1). La réponse à une excitation ressemble à la réponse sans perte : $\mathbf{E}(\mathbf{r}, f) \propto \mathbf{E}_n(\mathbf{r})$ pour $f = f_n$ et est proche de 0 sinon. La distribution spatiale du champ est donc très similaire à celle du champ propre $\mathbf{E}_n(\mathbf{r})$.
- À haute fréquence, l'espacement entre les modes est inférieur à la largeur des modes ($d \gg 1$), le recouvrement modal est donc fort (cela équivaut au cas multi-modal dont nous avons discuté au chapitre 1). Les modes se recouvrent les uns les autres et la réponse $\mathbf{E}(\mathbf{r}, f)$ à une excitation dépend d'un grand nombre de modes. Le nombre de modes significatifs est important. Dans ce cas la distribution spatiale de chaque champ propre n'a pas d'importance et le champ résultant peut être assimilé à une superposition aléatoire d'ondes planes. “La réponse est alors statistiquement uniforme et isotrope quelle que soit [...] la configuration de la cavité réverbérante” [Gro14].
- Entre ces deux régimes ($d \sim 1$) le recouvrement modal est modéré. C'est à dire que l'espacement entre les modes est de l'ordre de la largeur spectrale des modes. Dans ce cas seulement quelques modes contribuent à la réponse pour une excitation. Elle correspond à la somme des distributions spatiales des champs propres de ces modes. “Dans ce régime, la réponse n'est pas forcément statistiquement uniforme et isotrope pour toutes

les configurations [polarisation, position des sources] de la cavité réverbérante et toutes les fréquences d'excitation" [Gro14].

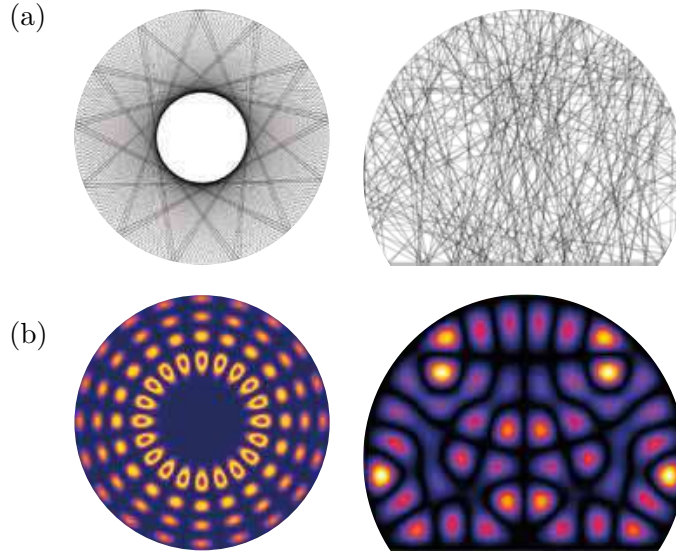


FIGURE 3.4 – Comparaison entre une cavité régulière (gauche) et une chaotique (droite). (a) Illustration de la trajectoire suivie par une particule dans les deux cavités (appelés dans ce cas billards régulier et chaotique) [Mic09]. (b) Illustration de la répartition spatiale de l'intensité du champ d'un mode dans les deux cavités [Mic09].

Dans les régimes de recouvrement faible et modéré, l'uniformité et l'isotropie de la distribution spatiale du champ dans la cavité dépendent de la chaoticité de la cavité, déterminée par sa géométrie. Pour comprendre intuitivement la différence entre une cavité régulière et une chaotique, il est intéressant de considérer le comportement des ondes dans le cadre de l'optique géométrique (longueur d'onde très petite devant la dimension du milieu). Dans ce cas le problème peut se ramener à un problème de dynamique classique de particule se mouvant dans la cavité. Par exemple, les trajectoires suivies par une particule dans une cavité régulière et une chaotique sont représentées sur la Figure 3.4(a). Dans la cavité régulière, la particule ne passera jamais par le centre du disque alors que dans la cavité chaotique la particule parcourt tout l'espace de la cavité (au bout d'un temps assez long). Dans la cavité régulière, il y a donc des endroits qui sont beaucoup plus empruntés par la particule, contrairement à la cavité chaotique dans laquelle il n'y pas d'endroit privilégié par la particule (la chaoticité de la cavité se traduit aussi par une dynamique soumise à une extrême sensibilité aux conditions initiales). Dans le domaine ondulatoire, on ne peut plus parler de propagation de l'énergie suivant des rayons lumineux. La chaoticité se traduit alors par la façon dont l'énergie est répartie spatialement. Les répartitions spatiales du champ d'un mode pour les deux mêmes cavités sont représentées sur la Figure 3.4(b). On remarque une corrélation entre les endroits de la cavité empruntés par la particule et la répartition spatiale du champ. Pour la cavité régulière, on remarque une structuration spatiale du champ très régulière, avec des zones de forte concentration. Pour la cavité chaotique, l'énergie est répartie de façon beaucoup plus

uniforme. De plus le champ possède la propriété d'isotropie, c'est à dire que le champ mesuré en un point résulte d'ondes venant de toutes les directions (ce n'est pas le cas pour la cavité régulière).

Les cavités réverbérantes en électromagnétisme ont été largement utilisées et étudiées dans le domaine de la compatibilité électromagnétique (CEM) pour tester la réaction de dispositifs à des agressions électromagnétiques [61003]. Dans ce cas des propriétés d'uniformité et d'isotropie de la distribution spatiale du champ électromagnétique dans la cavité sont ainsi souvent requises. Nous montrerons par la suite que dans le cas d'un RT mono-voie en cavité réverbérante ces propriétés sont aussi souhaitées pour assurer une bonne qualité de la focalisation spatio-temporelle. Suivant ce que nous venons de dire, il suffit donc de se placer dans un régime de recouvrement fort. Ce régime est atteint à partir d'une certaine fréquence, dépendant du volume de la cavité. Il est donc intéressant d'étendre ces propriétés d'uniformité et d'isotropie au régime modéré pour pouvoir étendre la bande de fréquence sur laquelle la cavité est utilisable. Dans le domaine CEM cette extension est généralement réalisée par une opération de brassage, le plus souvent de façon mécanique à l'aide d'un brasseur de modes. Ce brassage des modes est censé permettre d'obtenir statistiquement un champ uniforme et isotrope dans la cavité, en changeant la position du brasseur par exemple. Récemment, des études ont été menées pour décrire plus finement le comportement du champ électromagnétique dans les cavités réverbérantes et notamment autour du recouvrement modal modéré [Gro14] ; [Sel14]. Elles montrent que les propriétés d'uniformité et d'isotropie de la distribution spatiale du champ électromagnétique peuvent être obtenues en rendant la cavité chaotique. Pour ce faire, il faut supprimer les symétries de la géométrie de la cavité de sorte à "mélanger" les ondes. Dans une cavité parallélépipédique ce peut être fait par exemple en plaçant des demi-sphères sur les surfaces en regard [Gro14] ; [Sel14]. Nous verrons dans le chapitre suivant que nous avons opté pour cette solution pour notre dispositif expérimental. Dans la section suivante de ce chapitre, nous étudierons en simulation le RT en cavité 2D régulière et en cavité chaotique (le disque tronqué présenté sur la Figure 3.4) dans des conditions similaires aux conditions expérimentales et nous verrons dans la section d'après comment la chaoticité est déterminante pour le contrôle des plasmas par RT (en régime de recouvrement modéré).

On peut aussi remarquer que pour étendre la propriété de recouvrement fort aux plus basses fréquences il suffit d'augmenter la largeur spectrale des modes Γ , c'est à dire augmenter les pertes. Cependant ce n'est pas souhaitable puisque cela se traduit par une diminution du temps de stockage de l'énergie dans la cavité et donc une diminution de l'effet de concentration du champ électromagnétique dans la cavité, effet que l'on cherche généralement à exploiter en utilisant une cavité réverbérante.

Pour caractériser le degré de recouvrement modal d'une cavité, il est utile de connaître l'évolution en fréquence du nombre de ses modes propres. Il est facile de trouver une expression pour les fréquences des modes propres dans le cas canonique d'une cavité parallélépipédique remplie d'air et dont les parois sont sans pertes (conductivité infinie). En effet dans ce cas, la divergence du champ électrique $\mathbf{E}$ et celle du champ magnétique $\mathbf{H}$ sont nulles en tout point de la cavité ($\nabla \cdot \mathbf{E}(\mathbf{r}) = 0$ et $\nabla \cdot \mathbf{H}(\mathbf{r}) = 0$) et les conditions aux limites imposent que les composantes tangentielles du champ électrique $\mathbf{E}$ soient nulles au niveau des parois. La

solution des équations de Maxwell dans la cavité est une solution de type onde stationnaire satisfaisant l'équation de propagation $\Delta \mathbf{E} - \varepsilon_0 \mu_0 \partial^2 \mathbf{E} / \partial t^2 = 0$, on trouve ainsi les fréquences f_{lmn} auxquelles le champ peut exister dans la cavité [Hil98] :

$$f_{lmn} = \frac{c_0}{2} \sqrt{\left(\frac{l}{L_x}\right)^2 + \left(\frac{m}{L_y}\right)^2 + \left(\frac{n}{L_z}\right)^2} \quad (3.3)$$

avec L_x , L_y et L_z les dimensions de la cavité et l , m et n des entiers positifs² (et non tous nuls). Pour compter le nombre de modes, la dégénérescence des modes doit être prise en compte, c'est à dire que lorsque aucun des indices l, m, n n'est nul, il y a deux types de modes, transverse électrique TE_{lmn} et transverse magnétique TM_{lmn} , pour chaque fréquence propre [Hil98].

Pour les cavités aux géométries plus complexes, la détermination analytique des fréquences propres peut devenir très compliquée. Dans ce cas il est utile d'utiliser la formule de Weyl qui donne une approximation du nombre de modes jusqu'à une fréquence f en fonction du volume V de la cavité [Hil98] ; [LCM83] :

$$N_w(f) = \frac{8\pi V}{3c_0^3} f^3 \quad (3.4)$$

Dans le cas simple parallélépipédique sans pertes cette équation est une version lissée de l'équation 3.3 (cette dernière suit une forme d'escalier). Sur une bande passante $\Delta f \ll f$ au voisinage d'une fréquence centrale f , il y a statistiquement un nombre $N_{\Delta f}(f)$ de modes [Dup15] :

$$N_{\Delta f}(f) = \frac{dN_w(f)}{df} \Delta f = \frac{8\pi V}{c_0^3} f^2 \Delta f \quad (3.5)$$

Nous pouvons ainsi exprimer un paramètre important dans la caractérisation de la qualité d'une focalisation par RT : le **temps de Heisenberg** t_H . Il est défini comme l'inverse de l'écart fréquentiel entre les modes de la cavité [Ler06] et peut donc s'écrire :

$$t_H(f) = \frac{1}{\delta f(f)} = \frac{N_{\Delta f}(f)}{\Delta f} = \frac{8\pi V}{c_0^3} f^2 \quad (3.6)$$

Notons que ce temps de Heisenberg est aussi appelé densité modale [Ler+06] ; [MC14].

2. Attention ici m ne représente pas la masse des particules et n ne représente pas la densité.

Nous verrons dans la prochaine partie le rôle de ce paramètre sur la qualité d'une focalisation par RT. En appliquant l'équation 3.5 sur une bande passante égale à la largeur spectrale des modes $\Delta f = \Gamma$, on trouve une expression pour le nombre de modes significatifs d se chevauchant sur cette bande spectrale :

$$d(f) = \frac{8\pi V}{c_0^3} f^2 \Gamma \quad (3.7)$$

Dans le domaine acoustique, Manfred Robert Schroeder a proposé un nombre de modes se chevauchant dans la largeur spectrale égal à trois comme un critère à partir duquel le recouvrement modal peut être considéré comme fort : $d \geq 3$ [Sch87]. Même si ce critère est arbitraire et demande une analyse statistique plus fine pour trouver un critère plus pertinent comme développé dans [Coz11] ; [MC14], nous en resterons dans cette thèse à cette première approximation. Nous verrons dans la suite de ce chapitre comment ce critère s'applique aux cavités étudiées en simulation, et dans le chapitre suivant comment il s'applique sur la cavité expérimentale et ses conséquences sur les propriétés d'uniformité et d'isotropie du champ électromagnétique dans la cavité.

Nous allons voir dans les prochaines parties de cette section comment le régime modal et les propriétés d'uniformité et d'isotropie du champ dans la cavité influent sur la qualité de la focalisation spatio-temporelle par RT.

3.1.3 La focalisation temporelle

L'objectif de cette partie est de développer et comprendre l'équation du contraste temporel d'un RT. Pour cela les développements sont présentés dans les domaines temporel et fréquentiel. Nous avons dit précédemment que l'ensemble transducteur/cavité peut être assimilé à un filtre, et donc par définition à un système de convolution [Cot]. Ainsi en notant $h(t)$ la réponse d'un système LTI (entre deux transducteurs) à une impulsion de Dirac $\delta(t)$ et $e(t)$ le signal en entrée, le signal de sortie sera obtenu par l'équation [Cot] :

$$r(t) = e(t) * h(t) \quad (3.8)$$

Comme nous l'avons vu, ce signal est appelé réponse impulsionnelle. Et donc, par le théorème de Plancherel, dans le domaine spectral on aura :

$$R(f) = E(f).H(f) = |E(f)|.|H(f)|.e^{j(\phi_E(f)+\phi_H(f))} = |R(f)|.e^{j\phi_R(f)} \quad (3.9)$$

avec $r(t) \xrightarrow{TF} R(f)$ et $h(t) \xrightarrow{TF} H(f)$ où $H(f)$ la fonction de transfert du système an-

tennes/cavité ($\xrightarrow{TF}$ représentant la transformée de Fourier) et $\phi_E(f)$ et $\phi_H(f)$ respectivement la phase du signal d'entrée et la phase de la fonction de transfert $H(f)$, $\phi_R(f) = \phi_E(f) + \phi_H(f)$ la phase de la réponse de la cavité et $|R(f)| = |E(f)| \cdot |H(f)|$ la norme de la réponse de la cavité.

Lors de la deuxième phase du RT, cette réponse est retournée temporellement et transmise à la cavité au niveau d'un des deux transducteurs utilisés lors de la phase initiale (comme illustré sur la Figure 3.2). Si le système est réciproque, la cavité est alors aussi décrite par $h(t)$. Le signal $rt(t)$ alors obtenu s'écrit :

$$rt(t) = r(-t) * h(t) = e(-t) * h(-t) * h(t) \quad (3.10)$$

Ce qui s'écrit dans le domaine fréquentiel :

$$RT(f) = R^*(f) \cdot H(f) = |R(f)| \cdot |H(f)| \cdot e^{j(-\phi_R(f) + \phi_H(f))} = |E(f)| \cdot |H(f)|^2 \cdot e^{j(-\phi_E(f))} \quad (3.11)$$

et donc :

$$RT(f) = E^*(f) \cdot |H(f)|^2 \quad (3.12)$$

Selon l'équation 3.12, l'inversion temporelle du signal $R(f)$ avant de le transmettre à la cavité permet de compenser le déphasage induit par la cavité. Elle traduit le fait que lors d'une expérience de RT, le signal obtenu à l'endroit de l'émission de l'impulsion initiale $E(f)$ correspond à l'impulsion retournée temporellement $E^*(f)$ multipliée par le carré du module de la fonction de transfert entre le point de mesure et ce point là. Le carré traduit le fait que les ondes se propagent deux fois dans la cavité, une fois lors de la phase initiale et une autre lors de la phase inverse, le signal est donc filtré deux fois par la cavité.

En pratique, l'impulsion initiale $e(t)$ possède une bande passante finie $\Delta f = [f_{min}; f_{max}]$. Selon l'équation 3.12, pour obtenir au point initial de l'émission de la source exactement l'inverse temporel de l'impulsion initiale, il faut que la norme de la fonction de transfert de la cavité $|H(f)|$ soit constante sur Δf . Or nous avons vu dans la partie précédente que la fonction de transfert d'une cavité réverbérante est structurée par ses modes de résonance, ayant pour conséquence la présence de maximums et de minimums en fréquence. Comme nous allons le voir, dans le domaine temporel cela se traduit par la présence de lobes secondaires. Généralement le contraste temporel C_t , correspondant au rapport entre l'intensité du pic au moment de la compression par RT et l'intensité du bruit temporel, est utilisé pour caractériser la qualité de la compression temporelle obtenue par RT. Ce contraste est proportionnel au produit des degrés de liberté spatiaux N_s et temporels N_t [Dup15] :

$$C_t \propto N_{liberté} = N_s \cdot N_t \quad (3.13)$$

avec les degrés de liberté spatiaux qui correspondent au nombre de transducteurs dans le MRT $N_s = N_{ant}$.

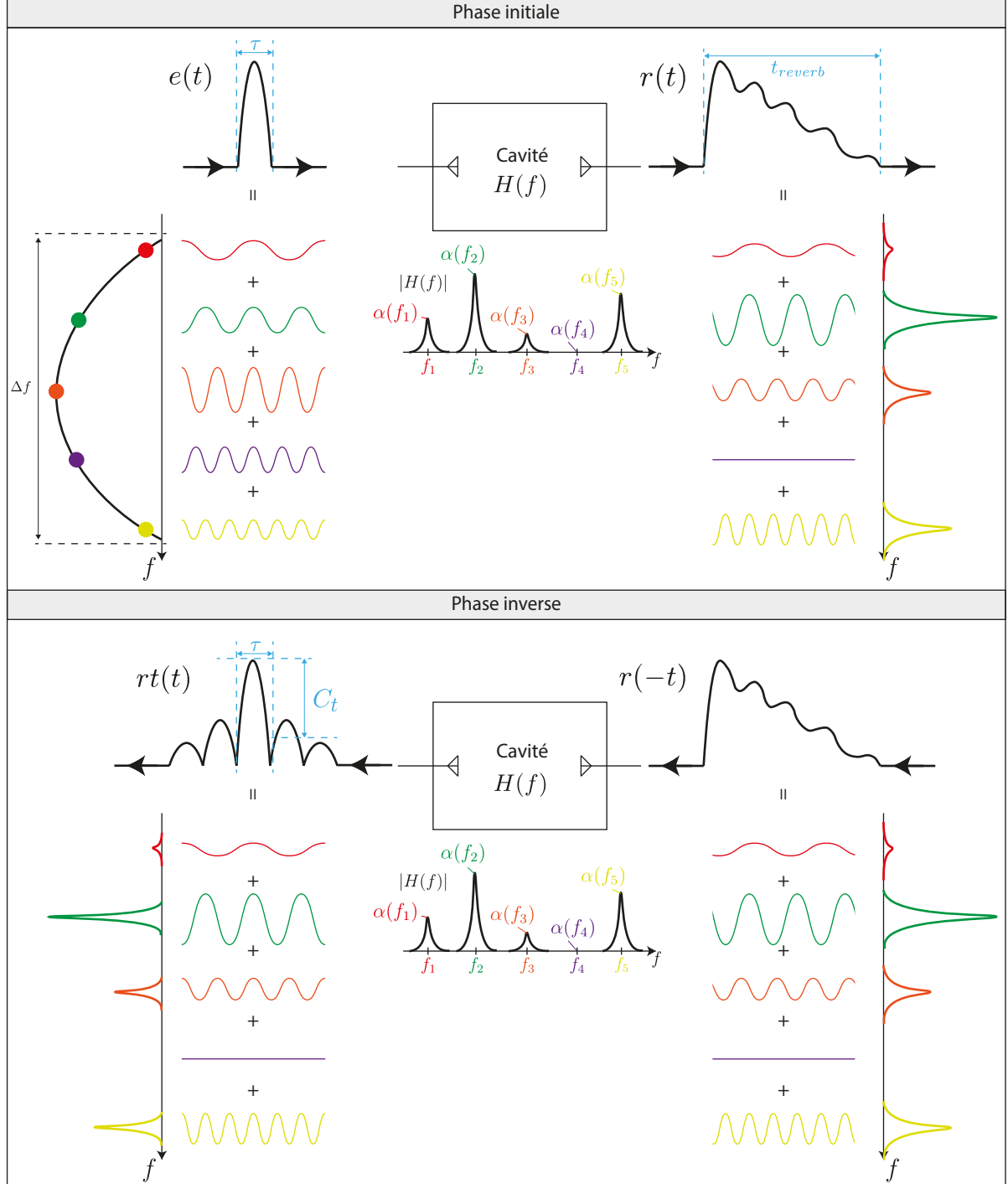


FIGURE 3.5 – Schéma d'un RT mono-voie en cavité.

Dans le cas d'un RT mono-voie qui nous intéresse ici, le nombre d'antennes dans le MRT est égal à $N_{ant} = 1$. D'après l'équation 3.13, on voit que la seule façon d'avoir un bon contraste

consiste à avoir un nombre de degrés de liberté temporels important. Ce nombre de degrés de liberté temporel N_t est lié aux modes de résonance de la cavité. En effet comme nous l'avons vu précédemment, dans le cas d'une cavité réverbérante, le module de la fonction de transfert $|H(f)|$ ne sera pas constant, les conditions aux limites imposant des fréquences auxquelles les ondes entrent en résonance. Prenons une impulsion $e(t)$ de courte durée donnant un Δf grand devant la largeur des résonances de la cavité. Dans le domaine fréquentiel, l'impulsion $e(t)$ correspond à la somme de sinusoïdes en phase au moment de l'impulsion, ces sinusoïdes possédant des amplitudes telles que leur somme s'annule exactement en dehors de la durée τ de l'impulsion, comme représenté sur la Figure 3.5. Le module $|H(f)|$ de la fonction de transfert de la cavité possède alors des amplitudes $\alpha(f_n)$ différentes à chaque fréquence de résonance f_n . La réponse de la cavité à cette impulsion $r(t)$ sera beaucoup plus longue que l'impulsion $e(t)$ puisque l'énergie est stockée pendant un certain temps dans la cavité. Ce temps est généralement appelé temps de réverbération de la cavité, noté ici t_{reverb} . Dans le domaine fréquentiel, cette réponse $r(t)$ correspond au signal $E(f)$ filtré par la cavité (*i.e.* $R(f) = E(f).H(f)$) et donc l'amplitude et la phase des sinusoïdes sont modifiées par rapport à l'impulsion initiale, modification qui dépend de la forme du complexe $H(f)$ (son module influe sur l'amplitude et son argument sur la phase comme traduit par l'équation 3.9)³. Dans ce cas, les sinusoïdes ne sont pas sommés en phase et le signal $r(t)$ décroît exponentiellement car les ondes perdent de l'énergie à chaque réflexion sur les parois de la cavité, comme nous l'avons illustré précédemment. Lors de la phase inverse, le signal obtenu $rt(t)$ résulte du filtrage de $r(-t)$ par la cavité. La phase du signal $rt(t)$ correspond à celle de l'impulsion initiale $e(t)$ avec un signe moins (équation 3.11), les sinusoïdes sont donc sommés en phases et on obtient un pic de durée τ . Cependant ces sinusoïdes ne sont pas sommés avec la bonne amplitude et certains ont même disparu par rapport à ceux dont la somme donne l'impulsion $e(t)$. Cela traduit le fait que le signal $rt(t)$ ne s'annule pas avant ni après le pic de durée τ . Il en résulte la présence de lobes secondaires autour du pic. Ainsi, on voit sur le schéma de la Figure 3.5 que le contraste temporel dépend de la façon dont les modes contenus dans $|H(f)|$ modifient les amplitudes des signaux dans le domaine fréquentiel. Et la forme de $|H(f)|$ dépend des propriétés de la cavité, telles que le régime de recouvrement modal ou la chaoticité.

L'expression du contraste temporel C_t en fonction des propriétés de la cavité a été développée par Julien de Rosny dans [Ros00] et mise sous une forme plus générale par Nicolas Quieffin dans [Qui04]. Afin de ne pas surcharger ce manuscrit, nous nous contenterons d'exploiter la formule que ce dernier a construite. L'objectif ici est de comprendre les paramètres qui influent sur le contraste temporel. La formule de ce contraste sera donc donnée dans sa forme développée dans [Qui04] et nous nous concentrerons dans la suite de ce chapitre sur l'explication des termes présents dans cette formule. Nous verrons qu'il faut introduire deux nouveaux paramètres, le kurtosis k_u et le facteur $g_{repulsion}$, et nous verrons dans la suite de cette section à travers des exemples simples comment ces paramètres influent sur ce contraste. L'expression du contraste temporel développée dans [Qui04] est :

3. Dans une cavité LTI passive, il est évident que $|H(f)|$ doit toujours être inférieur à 1. Par souci pédagogique, sur le schéma de la Figure 3.5, les amplitudes des modes pour $R(f)$ ne sont pas représentés avec la même échelle de l'axe des ordonnées que pour $E(f)$, donnant l'impression que $|H(f)|$ aurait pu être supérieur à 1, ce qui n'est bien sûr pas le cas.

$$C_t = \frac{I_{pic}}{I_{bruit}} = \frac{rt^2(t_{rt})}{rt^2(t \neq t_{rt})} = \frac{4\sqrt{\pi}t_H\Delta fshc^2\left(\frac{\Delta T}{t_{reverb}}\right)}{\frac{t_H}{\Delta T}shc\left(\frac{2\Delta T}{t_{reverb}}\right) + k_ushc^2\left(\frac{\Delta T}{t_{reverb}}\right) - g_{repulsion}} \quad (3.14)$$

avec :

- I_{pic} l'intensité du pic de l'impulsion au moment de la focalisation, définie comme $I_{pic} = rt^2(t_{rt})$, et I_{bruit} l'intensité du bruit :

$$I_{bruit} = rt^2(t \neq t_{rt}) = \frac{1}{t_a - t_b} \int_{t_a}^{t_b} rt^2(t)dt \quad (3.15)$$

où “les bornes t_a et t_b sont choisies dans une zone qui évite le pic de focalisation” [Ros00].

- ΔT la durée de la fenêtre temporelle appliquée à la réponse impulsionnelle avant de la retourner temporellement et la ré-émettre, comme illustré sur la Figure 3.6.
- t_{reverb} le temps de réverbération.
- t_H le temps de Heisenberg de la cavité, il est par définition égal à l'inverse de l'écart moyen entre les fréquences de modes δf contenus dans Δf : $t_H = 1/\delta f$
- $k_u = \langle \alpha^4(f_n) \rangle / \langle \alpha^2(f_n) \rangle^2$ le kurtosis des amplitudes $\alpha(f_n)$ des pics de résonance du spectre associées à la fréquence des modes f_n , c'est à dire “le moment d'ordre 4 sur le moment d'ordre 2 au carré des amplitudes des pics de résonances du spectre” [Ros00] aussi appelé “facteur d'aplatissement (« flatness factor ») de la distribution d'amplitude des modes de vibration de la cavité” [Qui04]. Par simplicité il sera simplement nommé kurtosis dans ce manuscrit.
- $g_{repulsion}$ un facteur qui rend compte de l'effet de la répulsion des fréquences propres du système (qui est nul pour une cavité régulière).

Comme mentionné précédemment, dans le cas d'un RT mono-voie le contraste temporel C_t est proportionnel au nombre de degrés de liberté temporels N_t . Nous allons voir qu'il est intéressant d'exprimer ce dernier dans les domaines temporel et fréquentiel. En effet, cette vision duale permet une compréhension fine des mécanismes ayant lieu à chaque étape du RT. Il faut alors tout d'abord distinguer deux cas, $\Delta T < t_{reverb}$ et $\Delta T > t_{reverb}$ représentés sur la Figure 3.6. La réponse impulsionnelle $r(t)$ est représentée par une simple décroissance exponentielle (à cause des pertes lors des réflexions des ondes dans la cavité). Dans le cas $\Delta T < t_{reverb}$, le signal contenu dans la fenêtre ΔT ne contient pas la décroissance exponentielle. En supposant le signal contenu dans ΔT comme environ constant, sa transformée de Fourier est un sinus cardinal de largeur spectrale à mi-hauteur $\Gamma = 1/\Delta T$. Par contre dans le cas $\Delta T > t_{reverb}$, le signal contenu dans la fenêtre ΔT contient la décroissance exponentielle. Sa transformée de Fourier est alors une Lorentzienne de largeur spectrale $\Gamma = 1/(\pi\tau_{exp})$. Afin de garder une symétrie temps - fréquence, il faut que les durées dans le domaine temporel soient inversement proportionnelles aux largeurs spectrales dans le domaine fréquentiel. Nous

définissons alors le temps de réverbération comme :

$$t_{reverb} = \pi \tau_{exp} = \pi \sqrt{\frac{\int t^2 \cdot r^2(t) dt}{\int r^2(t) dt}} \quad (3.16)$$

avec τ_{exp} le temps pendant lequel la plus grande partie de l'énergie contenue dans la réponse impulsionnelle est concentrée. Il correspond à la constante de temps de l'exponentielle décroissante suivie par l'amplitude de la réponse impulsionnelle, c'est à dire au temps au bout duquel cette amplitude a diminué d'environ 63 % de sa valeur maximale. Le facteur π est rajouté ici par rapport à l'expression utilisée par G. Lerosey [Ler06]. Il permet d'assurer que ce temps de réverbération corresponde bien à l'inverse de la largeur spectrale des modes : $t_{reverb} = 1/\Gamma$. Nous verrons dans la suite de cette partie que cela permet d'exprimer le nombre de degrés de liberté temporels N_t dans les domaines temporel et fréquentiel de façon symétrique.

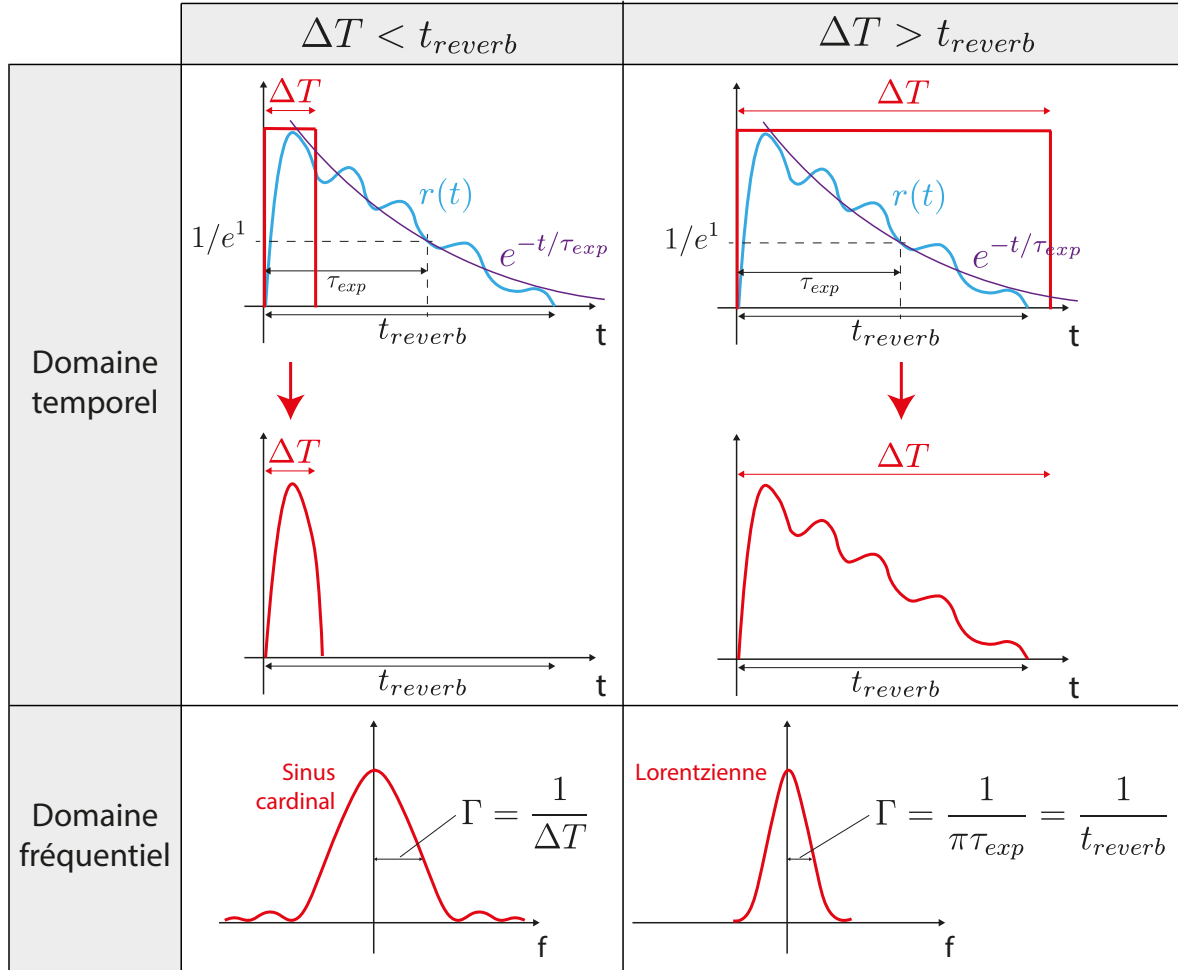


FIGURE 3.6 – Illustration de la largeur spectrale Γ obtenue dans deux cas : $\Delta T < t_{reverb}$ et $\Delta T > t_{reverb}$.

Pour résumer, le contraste temporel C_t représente le rapport entre la puissance du pic de focalisation et la puissance moyenne du bruit autour de ce pic. Il rend compte de la qualité

de la focalisation temporelle d'un processus de RT en fonction :

- de la durée de l'impulsion initiale (à travers Δf)
- de la durée de la fenêtre temporelle ΔT par rapport au temps de réverbération t_{reverb}
- de la longueur d'onde par rapport au volume de la cavité (à travers t_H)
- des pertes du système (parois de la cavité + antennes) à travers t_{reverb}
- de la chaotité de la cavité (à travers k_u et $g_{repulsion}$)

On remarque tout d'abord que le contraste est directement proportionnel à la bande passante Δf . Nous avons dit que le contrôle des plasmas par RT augmente avec les contrastes temporel et spatial. **Le contrôle des plasmas par RT est ainsi d'autant meilleur que la bande passante est grande.**

Nous allons ensuite étudier l'influence sur le contraste tout d'abord du rapport entre le temps de Heisenberg et le temps de réverbération (à travers le comportement limite du contraste temporel) et ensuite de la chaotité la cavité.

3.1.3.1 Comportement limite du contraste temporel

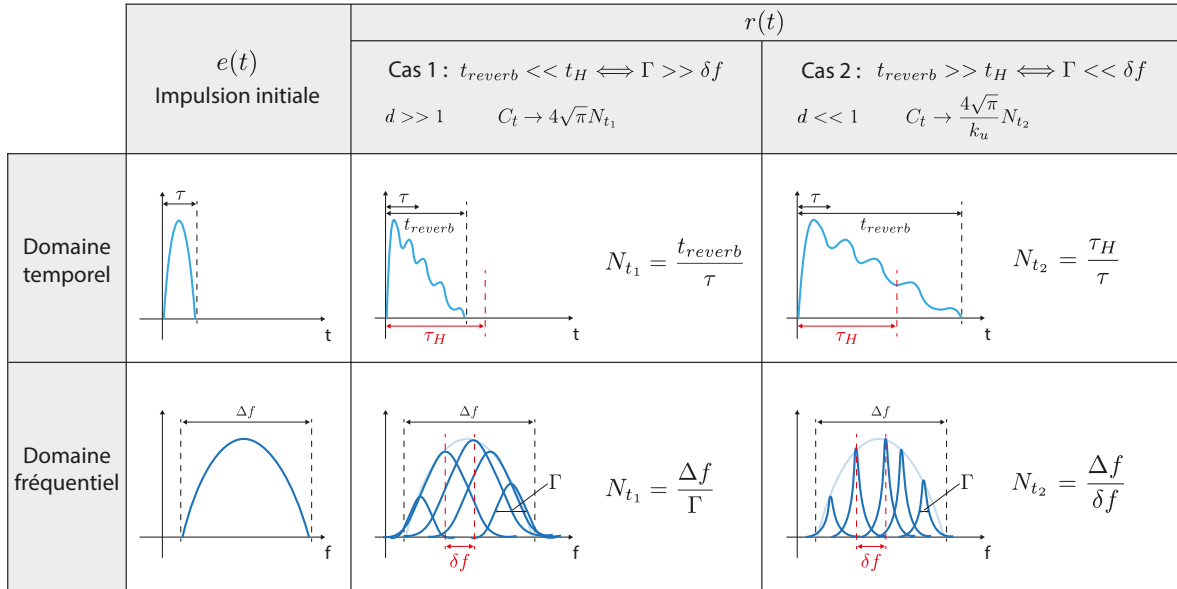


FIGURE 3.7 – Illustration des degrés de libertés temporels dans deux régimes (cas 1 et 2).

Nous nous intéressons ici au comportement limite du contraste temporel : pour $t_{reverb} \ll t_H$ et pour $t_{reverb} \gg t_H$, comme illustré respectivement sur les cas 1 et 2 de la Figure 3.7. Dans les deux cas on suppose ici que la durée de la fenêtre temporelle ΔT est supérieure ou égale au temps de réverbération, c'est à dire que l'on se contente juste de retourner temporellement et de ré-émettre la réponse impulsionnelle mesurée. De plus sur ce schéma les signaux dans le domaine temporel représentent l'enveloppe de signaux modulés. Lorsque le temps de réverbération est inférieur au temps de Heisenberg (cas 1) la largeur spectrale des

modes Γ est supérieure à l'espacement entre les modes δf et les modes se recouvrent les uns les autres : les modes ne sont pas résolus et le régime modal est donc fort ($d \gg 1$). Ainsi la quantité d'information contenue dans la réponse impulsionnelle $r(t)$ est inférieure au nombre de modes dans la bande passante ($\Delta f / \delta f$). Le nombre de degrés de liberté, traduisant la quantité d'information décorrélée contenue dans la réponse impulsionnelle $r(t)$ s'écrit donc $N_{t_1} = t_{\text{reverb}} / \tau = \Delta f / \Gamma$. Nous avons dit précédemment que le contraste était proportionnel au nombre de degrés de liberté, proportionnalité que l'on retrouve bien en prenant le cas limite $t_{\text{reverb}} \ll t_H$ de l'équation du contraste 3.14 : $C_t \rightarrow 4\sqrt{\pi} t_{\text{reverb}} \Delta f$ [Qui04]. Dans ce cas 1 la chaotité de la cavité n'a pas d'influence sur le contraste temporel. Lorsque le temps de réverbération est supérieur au temps de Heisenberg (cas 2) la largeur spectrale des modes Γ est inférieure à l'espacement entre les modes δf et donc les modes sont résolus (recouvrement modal faible $d \ll 1$). Par contre il n'est pas possible d'avoir plus d'information décorrélée que le nombre des modes contenus dans la bande passante Δf . Dans ce cas le nombre de degrés de liberté correspond au nombre des modes dans la bande passante et s'écrit $N_{t_2} = t_H / \tau = \Delta f / \delta f = N_{\Delta f}$. En prenant la limite $t_{\text{reverb}} \gg t_H$ on retrouve bien que le contraste lui est proportionnel : $C_t \rightarrow 4\sqrt{\pi} t_H \Delta f / k_u$ [Qui04]. Dans ce cas 2, l'influence de la chaotité de la cavité est prise en compte par le kurtosis k_u (qui sera proche de 3 pour une cavité chaotique et supérieur dans une cavité régulière).

Ces discussions peuvent être résumées en traçant l'évolution du contraste en fonction du temps de réverbération pour une certaine cavité excitée sur une bande passante Δf constante autour d'une fréquence centrale f_c , comme illustré sur la Figure 3.8(a). Dans ce cas le nombre des modes propres et leurs fréquences sont les mêmes dans les cas 1 et 2 et donc le temps de Heisenberg est aussi le même dans les deux cas. Ce qui change entre le cas 1 et 2, c'est une diminution des pertes (qui pourrait par exemple être réalisé par un dépôt d'une surface très conductrice sur les parois de la cavité) se traduisant par une augmentation du nombre de degrés de liberté N_t . Plus simplement, une étude similaire peut être faite dans une cavité à faible perte ($t_{\text{reverb}} \gg t_H$), en diminuant progressivement la durée ΔT de la fenêtre temporelle appliquée sur la réponse impulsionnelle avant de la retourner et de la ré-émettre [DAF99] ; [Ros00] ; [Qui04], comme cela sera fait en simulation dans la section suivante.

Il est aussi intéressant de tracer l'évolution du contraste en fonction du temps de Heisenberg pour une certaine cavité excitée sur une bande passante Δf , constante, avec des pertes constantes, comme illustré sur la Figure 3.8(b). Dans ce cas, en augmentant la fréquence centrale f_c , le contraste temporel augmente en passant du régime de recouvrement modal faible à celui modéré comme illustré sur la Figure 3.8(b) (sous l'hypothèse de pertes constantes en fréquence *i.e.* $t_{\text{reverb}} = 2/\Gamma$ constant). À partir du moment où le régime de recouvrement modal fort est atteint, les modes se recouvrent les uns les autres et, sous l'hypothèse d'une largeur spectrale des modes constante, le contraste va saturer puisqu'il sera proportionnel à $N_{t_1} = t_{\text{reverb}} / \tau = \Delta f / \Gamma$. Une augmentation en fréquence, même si elle se traduit par une augmentation du nombre de modes dans la bande passante ($\searrow \delta f$), ne permettra pas d'augmenter le contraste puisque ces modes seront recouverts par les autres modes et ne constitueront pas une information décorrélée supplémentaire. Autrement dit, pour une certaine cavité excitée sur une bande passante et sous l'hypothèse $\Gamma_n = \Gamma$, le régime de recouvrement modal modéré ($d \sim 1 \leftrightarrow \delta f \sim \Gamma$) correspond au seuil à partir duquel le contraste sature en fonction de la

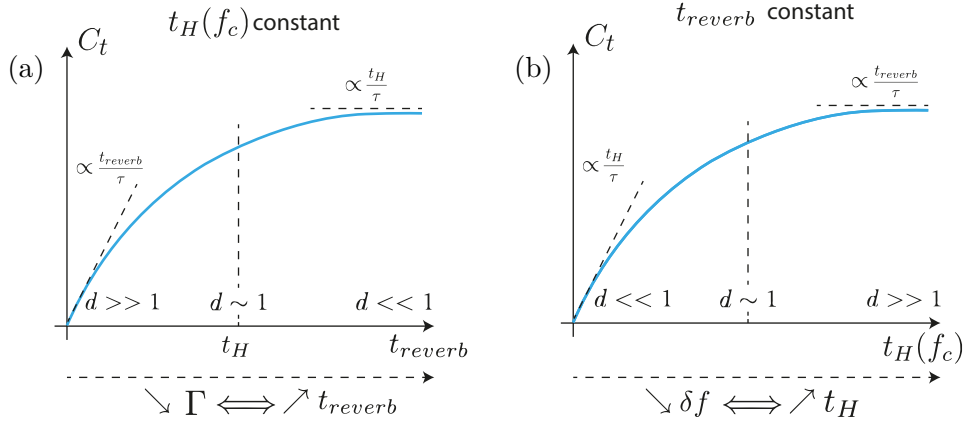


FIGURE 3.8 – Illustration de l'évolution du contraste pour une cavité de bande passante Δf constante et (a) de fréquence centrale f_c , en fonction des pertes (t_{reverb}) (b) de pertes constantes ($t_{reverb} = 1/\Gamma$ constant), en fonction de la fréquence centrale f_c .

fréquence.

Pour une cavité donnée excitée sur une certaine bande passante, il est alors possible d'améliorer le contrôle des plasmas par RT :

- soit à fréquence d'excitation constante, en diminuant les pertes aux parois, qui se traduit par une augmentation du temps de réverbération, jusqu'à saturation du contraste temporel lorsque les modes sont résolus
- soit à pertes constantes, en augmentant la fréquence centrale d'excitation, qui se traduit par une augmentation du temps de Heisenberg, jusqu'à ce que les modes se recouvrent les uns les autres
- soit en couplant ces deux solutions.

3.1.3.2 Influence de la chaoticité sur le contraste temporel

Afin de comprendre plus en détails l'équation du contraste temporel 3.14, intéressons nous maintenant à l'influence de la chaoticité sur ce dernier. Comme nous l'avons dit, celle-ci a une influence seulement en régime de recouvrement modal faible ou modéré (voir Figure 3.3). Pour un régime de recouvrement fort, il y a tellement de modes se chevauchant les uns les autres dans la bande passante que le champ résultant possède des propriétés d'uniformité et d'isotropie indépendamment de la chaoticité de la cavité. Le contraste temporel tend alors vers $C_t \rightarrow 4\sqrt{\pi}t_{reverb}\Delta f$ (cas 1 de la Figure 3.7). Par contre à plus faible recouvrement modal, la chaoticité de la cavité influe sur le contraste temporel, à travers le facteur $g_{repulsion}$ et le kurtosis k_u .

Le facteur $g_{repulsion}$ traduit le degré de répulsion entre les modes, phénomène appelé "level repulsion". Pour le caractériser il est intéressant de tracer la densité de probabilité de l'espacement spectral entre deux modes successifs $P(\delta f)$ (ou distribution des écarts entre plus

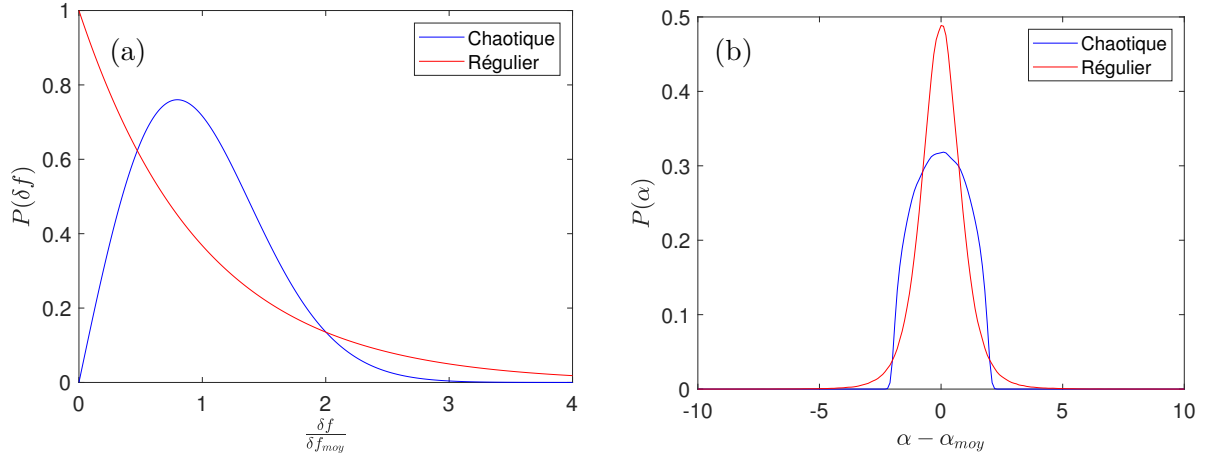


FIGURE 3.9 – Comparaison des propriétés statistiques d’une cavité régulière et chaotique. (a) Densité de probabilité de l’espacement spectral entre deux modes successifs $P(\delta f)$. (b) Densité de probabilité de l’amplitude des modes $P(\alpha)$ (pour les deux cas la variance est égale à 1).

proches voisins). Cette quantité statistique permet de caractériser la chaoticité d’un système [Gro14] ; [SS90]. Elle est tracée pour une cavité chaotique et une autre régulière sur la Figure 3.9 [Gro14] ; [Ros00]. Cette densité de probabilité $P(\delta f)$ suit une loi de Poisson pour une cavité régulière. Plus l’espacement spectral δf diminue et plus la densité de probabilité $P(\delta f)$ augmente, **c’est comme si les modes s’attirent les uns les autres** (la présence de modes dégénérées est prise en compte par la forte probabilité autour d’un espacement spectral nul). Dans ce cas le facteur $g_{repulsion}$ est nul. Pour une cavité chaotique, la densité de probabilité $P(\delta f)$ suit une loi de Wigner. Dans ce cas la probabilité de trouver deux modes très proches spectralement l’un de l’autre tend vers zéro. C’est ce qu’on appelle le “level repulsion”, **c’est comme si deux modes successifs n’aimaient pas rester trop proches l’un de l’autre** (on peut aussi remarquer que deux modes successifs n’aiment pas, non plus, rester trop loin l’un de l’autre). Ainsi l’espacement spectral δf entre les modes a tendance à être concentré autour d’une valeur moyenne δf_{moy} . Plus la densité de probabilité $P(\delta f)$ se rapprochera de celle de Wigner, plus la cavité sera chaotique [Gro14] et plus le facteur $g_{repulsion}$ sera important.

Le kurtosis k_u lui rend compte de la façon dont les amplitudes $\alpha(f_n)$ des modes sont distribuées. C’est un paramètre utilisé en statistique qui traduit l’importance de la “queue” d’une densité de probabilité, c’est à dire du “poids” de cette queue comparé au centre de la distribution. La densité de probabilité de l’amplitude des modes $P(\alpha)$ est tracée pour des cavités chaotique et régulière sur la Figure 3.9(b) (la moyenne et la variance sont les mêmes dans les deux cas). Pour une cavité chaotique, les amplitudes des modes restent concentrées proche d’une valeur moyenne, le kurtosis est alors proche de 3. Dans le cas $k_u = 3$ la distribution des amplitudes des modes suit une loi gaussienne [Gro14], et il peut même chuter jusqu’à $k_u = 2$ en prenant en compte les pertes [LW00]. Pour une cavité régulière, les amplitudes des modes sont réparties de façon moins uniforme que le cas chaotique (le kurtosis augmente). **En fait la probabilité $P(\alpha)$ de trouver des modes avec une amplitude très éloignée de la**

valeur moyenne est supérieure dans la cavité régulière. La “queue” de la distribution est plus élevée.

Pour comprendre simplement comment ces propriétés (répulsion et kurtosis) influent sur le contraste, les signaux obtenus par RT $rt(t)$ avec une impulsion initiale $e(t)$ (de bande passante Δf) pour trois cavités aux chaoticités différentes sont tracés sur la Figure 3.10. Ces trois cavités sont excitées en régime de recouvrement modal modéré, puisque nous avons expliqué dans la partie précédente que la chaoticité n’a pas beaucoup d’influence sur la qualité de la focalisation spatio-temporelle en régime de recouvrement modal fort. Cette différence de chaoticité se traduit par des formes différentes du module de leur fonction de transfert $|H(f)|$, paramètre déterminant la forme du signal de RT comme le montre l’équation 3.12. Ces trois cavités sont excitées avec une impulsion $e(t)$ d’une durée τ modulée à une fréquence f_c . Elles sont donc excitées avec la même bande passante Δf . Pour ces trois cavités, le temps de Heisenberg t_H (nombre de modes sur la bande passante) et le temps de réverbération t_{reverb} (*i.e.* largeur spectrale des modes) sont les mêmes.

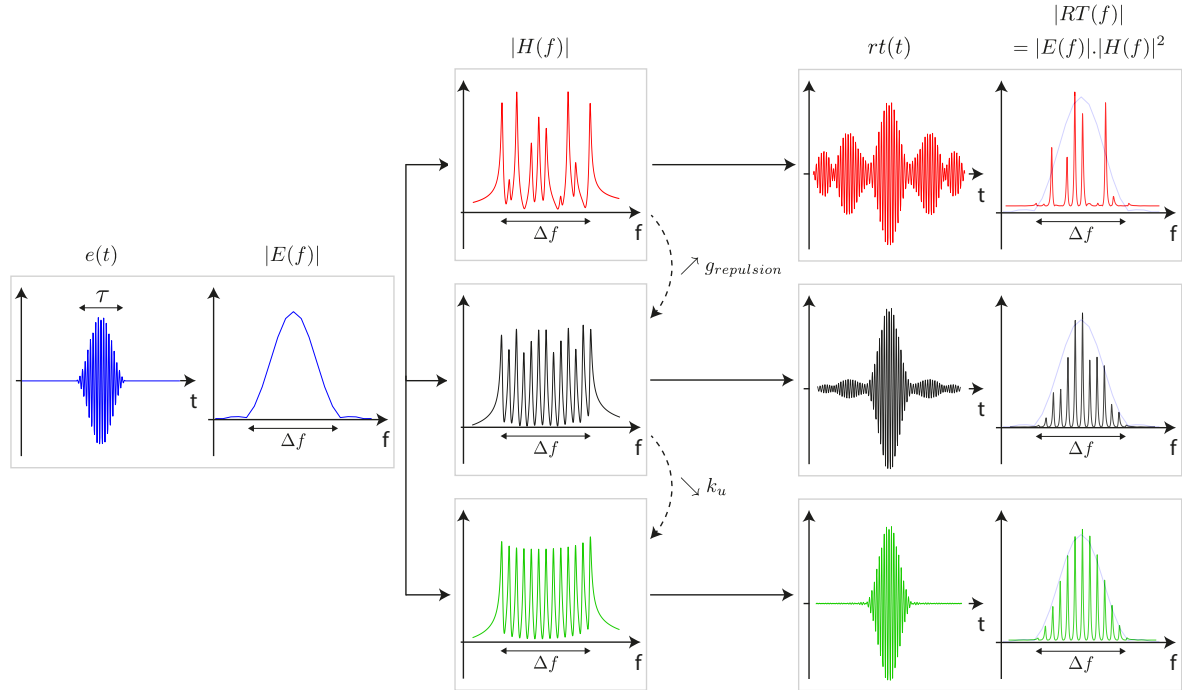


FIGURE 3.10 – Illustration de l’influence de la chaoticité de la cavité (répulsion et kurtosis), dont dépend la forme de $|H(f)|$, sur le contraste du RT.

La première cavité (en rouge) représenterait plutôt une cavité régulière, l’espacement spectral entre les modes δf n’étant pas constant (suivant par exemple l’équation 3.3 pour un parallépipède) et la différence entre les niveaux des modes étant grande. Dans ce cas $grepulsion$, qui traduit la façon dont les modes se repoussent les uns des autres, est nul puisqu’il y a des modes dégénérés (plusieurs modes à la même fréquence) et le kurtosis ku , qui traduit la façon dont les amplitudes des modes sont distribuées est grand (par rapport au cas où ces amplitudes sont réparties suivant une loi gaussienne pour lequel $ku = 3$). Dans ce cas le signal obtenu par RT $rt(t)$ présente des lobes secondaires temporels importants. On peut voir

directement l'origine de ces lobes dans le domaine fréquentiel sur le signal $|RT(f)|$. L'objectif du RT est de reconstruire le plus fidèlement possible l'impulsion initiale $|E(f)|$ et on voit bien sur $|RT(f)|$ qu'il manque beaucoup de fréquences et que les amplitudes des modes diffèrent de celles de $E(f)$. Cela se traduit dans le domaine temporel par la présence de lobes secondaires importants.

La seconde cavité (en noir) représente une cavité dont les fréquences des modes sont réparties de manière plus homogène que celles de la première cavité (en rouge), comme si les modes dégénérés avaient été séparés les uns des autres. Dans ce cas le facteur $g_{repulsion}$ est supérieur à zéro (et le kurtosis a diminué par rapport à la première cavité). Le contraste augmente selon la relation 3.14, comme on le voit sur la Figure 3.10. L'information fréquentielle contenue dans le signal $|RT(f)|$ obtenu par RT est mieux répartie en fréquence et les lobes secondaires sont ainsi diminués dans le domaine temporel.

La troisième cavité (en vert) représente une cavité dont les amplitudes des modes sont réparties de façon plus uniforme que la deuxième cavité (et que la première). Cela se traduit par un kurtosis moins élevé, pouvant diminuer jusqu'à $k_u = 2$ pour une cavité chaotique (le facteur $g_{repulsion}$ est lui le même que pour la deuxième cavité). Selon la relation 3.14, le contraste est supérieur pour cette troisième cavité par rapport à la deuxième (et à la première). Cela se voit clairement sur le signal $|RT(f)|$, ce signal ressemble en fait à une version échantillonnée de l'impulsion initiale $|E(f)|$.

Pour conclure nous avons vu que le contraste temporel est proportionnel au nombre de degrés de liberté temporel dans le cas d'un RT mono-voie en cavité réverbérante. Ce nombre de degrés de liberté correspond à la quantité d'information décorrélée contenue dans la bande passante. De plus nous avons vu que le contraste est d'autant meilleur que la cavité est chaotique (en régime de recouvrement faible ou modéré). **Le contrôle des plasmas par RT sera donc d'autant meilleur que la cavité est chaotique.** Nous allons voir dans la partie suivante en quoi cette propriété de chaotité est essentielle pour avoir une bonne focalisation spatiale.

3.1.4 La focalisation spatiale

Nous avons étudié dans la partie précédente la qualité de la focalisation temporelle caractérisée par le contraste temporel C_t . Ce contraste s'applique au signal de RT $rt(t)$ mesuré à l'endroit de la focalisation, que nous noterons $\mathbf{r}_{rt}$. Par souci de clarté nous avons omis dans la partie précédente la dépendance spatiale des signaux puisque nous nous intéressons seulement à ce qu'il se passe à l'endroit de la focalisation. De façon plus générale, on peut écrire que $rt(t) = s_{rt}(\mathbf{r}_{rt}, t)$ avec :

$$s_{rt}(\mathbf{r}, t) = r(-t) * h_s(\mathbf{r}, t) = e(-t) * h(-t) * h_s(\mathbf{r}, t) \quad (3.17)$$

avec $h_s(\mathbf{r}, t)$ la fonction de transfert entre le point d'émission du signal $r(-t)$ et n'importe quel point $\mathbf{r}$ contenu dans la cavité. En se plaçant à l'endroit de la focalisation, c'est à dire pour $\mathbf{r} = \mathbf{r}_{rt}$, on a donc $h(t) = h_s(\mathbf{r}_{rt}, t)$ et l'équation 3.17 revient bien à l'équation 3.10 que nous avons utilisée pour l'étude du contraste temporel. Cette relation s'écrit dans le domaine fréquentiel :

$$S_{rt}(\mathbf{r}, f) = R^*(f) \cdot H_s(\mathbf{r}, f) = E^*(f) \cdot |H(f)| \cdot |H_s(\mathbf{r}, f)| \cdot e^{j(\phi_{H_s}(\mathbf{r}, f) - \phi_H(f))} \quad (3.18)$$

En se plaçant à l'endroit de la focalisation ($\mathbf{r} = \mathbf{r}_{rt}$), le module devient $|H(f)| \cdot |H_s(\mathbf{r}, f)| = |H(f)|^2$, la phase s'annule et on retrouve bien l'équation 3.12 correspondant au signal de RT $RT(f)$. La première différence à remarquer entre le signal obtenu à l'endroit de la focalisation $RT(f)$ et celui ailleurs dans la cavité $S_{rt}(\mathbf{r} \neq \mathbf{r}_{rt}, f)$ est la phase non nulle et dépendante de la fréquence. Cela veut dire qu'ailleurs dans la cavité, les modes ne sont pas forcément sommés en phase et il n'y a pas a priori de pic. De plus, la forme de ces signaux est conditionnée par le produit $|H(f)| \cdot |H_s(\mathbf{r}, f)|$, qui dépend des propriétés de la cavité comme nous allons le voir dans la suite de cette partie.

Le contraste spatial du RT peut s'écrire [HLH14] ; [Alb+19] :

$$C_s = \frac{I_{pic}}{I_{bruit_{spatial}}} = \frac{s_{rt}^2(\mathbf{r}_{rt}, t_{rt})}{s_{rt}^2(\mathbf{r} \neq \mathbf{r}_{rt}, t_{rt})} = \frac{rt^2(t_{rt})}{s_{rt}^2(\mathbf{r} \neq \mathbf{r}_{rt}, t_{rt})} \quad (3.19)$$

avec I_{pic} l'intensité du pic au moment et à l'endroit de la focalisation (qui est la même que celle utilisée dans le calcul du contraste temporel de l'équation 3.14) et $I_{bruit_{spatial}}$ l'intensité du bruit spatial au moment de la focalisation t_{rt} . Comme nous allons le voir ce contraste est meilleur pour une cavité chaotique que régulière.

Généralement la géométrie de la cavité est prise chaotique pour obtenir une bonne qualité de la focalisation spatiale par RT [DF97] ; [Ros00]. Dans ce cas, le contraste spatial est souvent supposé égal au temporel [Qui04]. Récemment des auteurs ont étudié l'influence du degré de chaoticité de la cavité sur les contrastes temporel et spatial [Alb+19]. Ils ont montré que, dans certaines conditions, la chaoticité n'a pas beaucoup d'influence sur le contraste temporel. Par contre, ils remarquent que plus la chaoticité de la cavité augmente, plus le contraste spatial augmente et tend vers le temporel [Alb+19]. De plus, comme pour le contraste temporel, le régime de recouvrement modal influe aussi sur le contraste spatial, comme nous allons le voir dans la suite de cette partie.

Tout d'abord, notons comme nous l'avons déjà remarqué, que dans un régime de recouvrement modal fort, la réponse d'une cavité à une excitation "est alors statistiquement uniforme et isotrope quelle que soit [...] la configuration de la cavité réverbérante" [Gro14]. Dans ce cas la forme de $|H(f)|$ est à peu près "plate" en fréquence (comme représenté sur la Figure 3.3). Et cela est le cas pour toutes les fonctions $|H_s(\mathbf{r}, f)|$ quel que soit $\mathbf{r}$. Ainsi, on peut écrire

$|H_s(\mathbf{r}, f)| \approx |H(f)|$ pour tout $\mathbf{r}$. L'équation 3.18 devient alors :

$$S_{rt}(\mathbf{r}, f) \approx E^*(f) \cdot |H(f)|^2 \cdot e^{j(\phi_{H_s}(\mathbf{r}, f) - \phi_H(f))} \quad (3.20)$$

Au point de la focalisation $\mathbf{r} = \mathbf{r}_{rt}$, on retrouve l'équation 3.12 correspondant au signal de RT $RT(f)$. Comme expliqué dans la partie précédente, à l'endroit de la focalisation $\mathbf{r}_{rt}$ les modes sont donc sommés en phase. Un pic sortant du bruit en résulte comme représenté sur la Figure 3.5. Pour toutes les autres positions dans la cavité $\mathbf{r} \neq \mathbf{r}_{rt}$, le signal obtenu est décrit par la même équation 3.12 avec un terme de déphasage dépendant de la fréquence en plus correspondant à la différence de phase entre les arguments des fonctions de transfert $H(f)$ et $H_s(\mathbf{r}, f)$: $\phi_{H_s}(\mathbf{r}, f) - \phi_H(f)$. Donc, le pic de focalisation (issu d'une mise en phase des sinusoides) n'apparaît qu'au point et au moment de la focalisation et partout ailleurs les modes ne sont pas sommés en phase, donnant un certain niveau de bruit. Ce niveau de bruit sera le même que celui mesuré à l'endroit de la focalisation $\mathbf{r}_{rt}$ autour du pic de focalisation, puisque lui même est issu de la somme des modes possédant la même amplitude ($|H_s(\mathbf{r}, f)| \approx |H(f)|$). Cela illustre une propriété des systèmes chaotiques : l'*ergodicité* (dans une cavité chaotique, on parle de modes ergodiques [Gro14]). Elle stipule dans ce cas l'équivalence entre les moyennes spatiales et temporelles. Elle traduit le fait que la répartition spatiale de l'énergie électromagnétique est uniforme et isotrope dans la cavité et donc l'information mesurée en un point à un certain instant est une image de ce qui s'est déroulé dans la cavité partout ailleurs à des temps antérieurs. En conséquence, le contraste spatial sera égal au contraste temporel : $C_s = C_t$.

Dans le cas d'un recouvrement modal modéré (ou faible), la chaotité de la cavité est déterminante sur la qualité de la focalisation spatiale. En effet dans ce cas, la distribution spatiale du champ est proche de celle de la somme d'un certain nombre de modes propres $\mathbf{E}_n(\mathbf{r})$. C'est donc la façon dont ces modes répartissent spatialement l'énergie qui va conditionner la qualité de la focalisation spatiale au moment de la focalisation et c'est la géométrie de la cavité qui détermine cette répartition spatiale :

- Une cavité régulière (comme un parallélépipède par exemple, largement répandu en électromagnétisme grâce à sa simplicité de réalisation) va avoir tendance à concentrer le champ électromagnétique selon une certaine organisation spatiale due aux surfaces parallèles en regard. Dans ce cas, $|H_s(\mathbf{r}, f)| \neq |H(f)|$ pour tout $\mathbf{r}$. Au moment t_{rt} de la focalisation, le champ électromagnétique dans la cavité résulte de la somme de modes organisés et présente des lobes secondaires conséquents, comme illustré sur la Figure 3.11.
- Pour “casser” cette répartition organisée, il faut supprimer les symétries et les surfaces parallèles en regard. De cette façon l'énergie va se répartir spatialement de façon uniforme et isotrope, ne suivant donc pas une organisation précise et la cavité devient chaotique⁴. De plus, comme nous l'avons déjà vu, la chaotité d'une cavité se traduit par une répulsion des modes (disparition des modes dégénérés). Autrement dit, il y aura plus de fréquences permettant l'existence du champ électromagnétique dans la cavité et

4. Il peut exister des modes réguliers même dans une cavité chaotique mais ceux-ci sont minoritaires [Gro14].

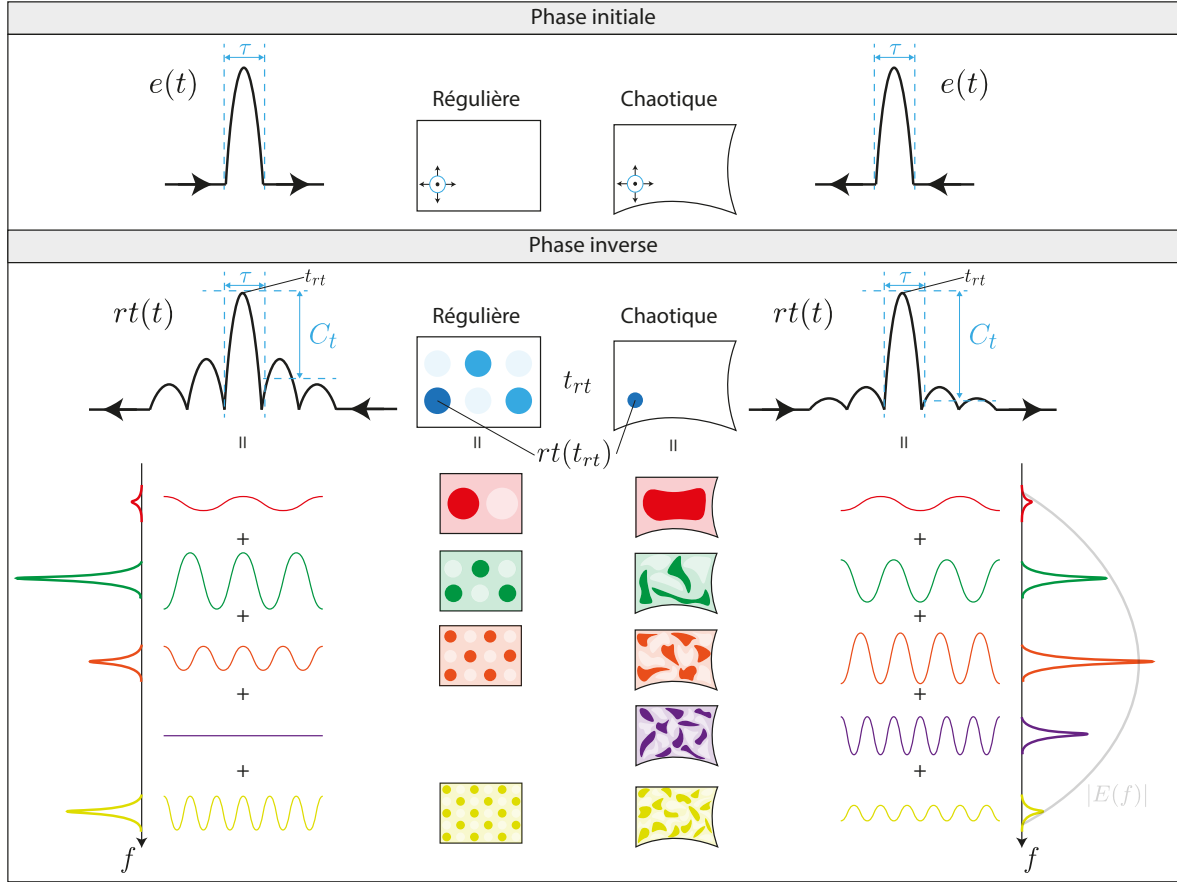


FIGURE 3.11 – Illustration de la répartition spatiale des modes lors d'un RT mono-voie en cavité réverbérante au moment t_{rt} de la focalisation pour une cavité régulière et chaotique en régime de recouvrement modal modéré (ou faible).

donc plus de répartitions spatiales différentes (la fréquence détermine en partie la façon dont l'énergie peut se répartir dans la cavité). Dans ce cas, au moment de la focalisation t_{rt} la somme de ces modes se traduit spatialement par un fort pic de focalisation comme illustré sur la Figure 3.11.

Finalement, on trouve que, comme pour le contraste temporel, la chaotité de la cavité permet d'augmenter aussi le contraste spatial. **La chaotité de la cavité est donc une propriété essentielle pour le contrôle des plasmas par RT.**

3.1.5 Synthèse

Nous avons montré dans cette partie la dépendance des contrastes temporel et spatial au recouvrement modal et à la chaotité de la cavité. Pour faire simple, pour une cavité donnée excitée avec une certaine bande passante autour d'une fréquence centrale également fixée, plus les pertes sont faibles plus le contraste temporel augmente jusqu'à ce que les modes

soient résolus (dans ce cas le contraste a atteint sa valeur maximale). Si, dans ce cas, la cavité est chaotique alors le contraste spatial est le même que le temporel. Sinon, si la cavité est régulière, alors le contraste spatial peut être inférieur au temporel [Alb+19].

Dans [Alb+19], il est introduit aussi deux autres contrastes, au sens de rapport entre le niveau du champ de la focalisation et le niveau du lobe secondaire maximal temporel (C_t^{max}) ou spatial (C_s^{max}) :

$$C_t^{max} = \frac{|s_{rt}(\mathbf{r}_{rt}, t_{rt})|}{\max(|s_{rt}(\mathbf{r}_{rt}, t \neq t_{rt})|)} = \frac{|rt(t_{rt})|}{\max(|s_{rt}(\mathbf{r}_{rt}, t \neq t_{rt})|)} \quad (3.21)$$

$$C_s^{max} = \frac{|s_{rt}(\mathbf{r}_{rt}, t_{rt})|}{\max(|s_{rt}(\mathbf{r} \neq \mathbf{r}_{rt}, t_{rt})|)} = \frac{|rt(t_{rt})|}{\max(|s_{rt}(\mathbf{r} \neq \mathbf{r}_{rt}, t_{rt})|)} \quad (3.22)$$

Ils montrent que pour ces contrastes aussi, dans une cavité chaotique le contraste spatial est le même que le temporel.

Pour amorcer un plasma par RT au moment et à l'endroit de la focalisation, il est nécessaire que le niveau du champ électrique **partout dans la cavité et durant toute la durée de l'expérience** soit inférieur au niveau du champ électrique de focalisation. Nous introduisons ici un contraste plus général que C_s^{max} , le contraste spatio-temporel C_{st}^{max} , défini comme :

$$C_{st}^{max} = \frac{|s_{rt}(\mathbf{r}_{rt}, t_{rt})|}{\max(|s_{rt}(\mathbf{r} \neq \mathbf{r}_{rt}, t \neq t_{rt})|)} = \frac{|rt(t_{rt})|}{\max(|s_{rt}(\mathbf{r} \neq \mathbf{r}_{rt}, t \neq t_{rt})|)} \quad (3.23)$$

Plus ce contraste sera grand, plus le contrôle des plasmas par RT sera bon. Si il est inférieur à 1, cela signifie qu'il y a un lobe secondaire à un moment dans la cavité dont le niveau est supérieur à celui du pic de focalisation. Et en conséquence un plasma sera généré à l'endroit du lobe secondaire avant de pouvoir l'être à l'endroit de la focalisation.

Nous allons dans la section suivante présenter le code de simulation FDTD 2D développé pendant cette thèse et permettant de simuler des expériences de RT en cavité selon des conditions similaires à celles du dispositif expérimental présenté au chapitre suivant. Nous vérifierons dans cette prochaine section les prédictions sur les contrastes temporel et spatial prévues par la théorie développée dans cette section. Ensuite, dans la section suivante, nous montrerons en simulation l'influence de la chaoticité sur l'amorçage de plasmas par RT.

3.2 Modélisation du retournement temporel en cavité réverbérante : Modèle FDTD 2D

Nous avons développé dans la première section de ce chapitre les modèles permettant de décrire et de comprendre le RT mono-voie en cavité réverbérante. Nous avons vu que ce dernier présente des lobes secondaires temporels et spatiaux, conséquences de l'émission dans toutes les directions lors de la phase inverse. Nous avons ensuite développé les équations qui expliquent l'influence de ces lobes à travers le concept de contraste. Ces équations nous disent que ces contrastes sont meilleurs en cavité chaotique. L'objectif de ce chapitre est de montrer l'importance de la chaoticité de la cavité sur le contrôle de plasmas par RT. Pour cela un code de simulation FDTD 2D a été développé. L'objectif de cette section est de présenter la partie électromagnétique du code (équations de Maxwell) et d'illustrer l'influence de la chaoticité de la cavité sur la qualité de la focalisation par RT. Pour cela deux cavités sont étudiées : une régulière et une chaotique. Nous verrons que la simulation présente les avantages de pouvoir contrôler facilement la géométrie des cavités étudiées et visualiser la répartition spatiale du champ électrique dans la cavité durant les simulations de RT. Enfin, l'influence de la chaoticité sur les contrastes temporel et spatial est vérifiée, en accord avec la théorie développée dans la partie précédente.

Les paramètres des simulations effectuées par la suite sont choisis de sorte à ce que le RT possède des propriétés spatio-temporelles similaires à celles obtenues avec le dispositif expérimental présenté au chapitre suivant. Les caractéristiques de la cavité expérimentale et du signal excitateur utilisés expérimentalement ainsi que les contrastes temporels C_t et C_t^{max} mesurés sont rassemblées dans le tableau suivant (les méthodes utilisées pour obtenir les valeurs de ces contrastes temporels expérimentaux sont présentées dans le chapitre suivant) :

Propriétés du dispositif expérimental	
Fréquence centrale f_c	≈ 2.4 GHz
Bande passante Δf	250 MHz
Volume de la cavité V	$0.6 \times 0.6 \times 0.3$ m ³
Contraste temporel C_t	168
Contraste temporel C_t^{max}	2.4

3.2.1 Présentation du modèle FDTD 2D

Nous allons dans cette partie présenter le code de simulation FDTD 2D développé durant cette thèse. Notre choix s'est porté sur un code basé sur la méthode FDTD pour sa facilité d'implémentation avec des géométries complexes. De plus, elle fournit une solution du champ total en implémentant directement les équations de Maxwell et peut être appliquée à des problèmes harmoniques ou larges bandes [ED15a]. Elle permet en outre, de suivre l'évolution temporelle du champ électromagnétique dans la cavité avec un seul "run". L'objectif de ce code de simulation étant de comprendre les mécanismes physiques à l'œuvre lors d'une expérience de RT, cette dernière propriété est très intéressante. Le choix de travailler seulement en deux

dimensions (2D) a été fait pour minimiser le temps de calcul et les ressources mémoires mobilisées.

Les simulations sont réalisées dans une cavité dont les parois possèdent une conductivité σ_{parois} et dont le milieu de propagation à l'intérieur est défini par les propriétés du vide ε_0 et μ_0 ($\sigma = 0$). On considère par ailleurs le cas transverse électrique (TE), soit un champ électrique orienté selon l'axe z et donc avec une composante selon cette direction. Il est donc perpendiculaire à la direction de propagation, contenue dans le plan xy . La grandeur qui nous intéresse pour l'amorçage de plasmas étant le champ électrique, ce choix semble le plus évident puisque dans ce cas ce champ sera facile à manipuler et à contrôler numériquement. Les équations de Maxwell 2.4 et 2.5 sans plasma se simplifient et deviennent alors :

$$\frac{\partial E_z(\mathbf{r}, t)}{\partial t} = \frac{1}{\varepsilon_0} \left(\frac{\partial H_y(\mathbf{r}, t)}{\partial x} - \frac{\partial H_x(\mathbf{r}, t)}{\partial y} \right) - \frac{1}{\varepsilon_0} J_{c_z}(\mathbf{r}, t) \quad (3.24)$$

$$\frac{\partial H_x(\mathbf{r}, t)}{\partial t} = -\frac{1}{\mu_0} \frac{\partial E_z(\mathbf{r}, t)}{\partial y} \quad (3.25)$$

$$\frac{\partial H_y(\mathbf{r}, t)}{\partial t} = \frac{1}{\mu_0} \frac{\partial E_z(\mathbf{r}, t)}{\partial x} \quad (3.26)$$

avec $J_{c_z}(\mathbf{r}, t) = \sigma(\mathbf{r}, t)E_z(\mathbf{r}, t)$ le courant de conduction, qui est nul à l'intérieur de la cavité et qui vaut $\sigma_{parois}(\mathbf{r}, t)E_z(\mathbf{r}, t)$ sur les parois.

Ces équations sont ensuite discrétisées en espace et en temps selon un schéma classique de différences finies centrées. Elles sont rappelées en Annexe A.1. Un enjeu important dans le développement d'un modèle numérique dans le domaine temporel est la condition de stabilité. En effet, comme nous l'avons dit, le développement d'un code FDTD nécessite d'échantillonner spatialement et temporellement les champs électriques et magnétiques. Le choix du pas spatial Δ_s (le maillage spatial retenu ici utilise des mailles carrées tel que $\Delta_x = \Delta_y = \Delta_s$) et du pas temporel Δ_t doit respecter certaines conditions pour assurer la stabilité de la solution. Dans le cas à deux dimensions, cette condition s'écrit [ED15a] :

$$c\Delta_t \leq \frac{\Delta_s}{\sqrt{2}} \text{ ou } \Delta_t \leq \frac{\Delta_s}{c\sqrt{2}} \quad (3.27)$$

avec c la vitesse des ondes dans le milieu. Cette relation permet d'assurer qu'entre deux itérations temporelles les ondes ne se sont pas propagées sur une distance supérieure à la dimension d'une cellule spatiale. Autrement dit, les ondes présentes sur une certaine cellule à un instant t ne peuvent pas s'être propagées plus loin que les cellules adjacentes au temps $t + \Delta_t$. Si elle n'est pas respectée, alors le modèle numérique perd de l'information et ne décrit plus le phénomène physique. Il en résulte potentiellement une instabilité lors du calcul itératif.

Un autre enjeu de la modélisation numérique par FDTD est la dispersion numérique. Elle est une conséquence de la discrétisation de fonctions continues (puisque représentant une

grandeur physique). Elle se traduit par la dépendance de la vitesse de propagation des ondes à leur fréquence et à leur direction de propagation. Cet effet est d'autant plus faible que le pas spatial est petit devant la plus petite longueur d'onde λ_{min} de la bande passante d'analyse. Cependant, cela implique aussi un pas temporel de plus en plus petit selon la relation de stabilité 3.27 et donc des temps de simulation de plus en plus grand. Généralement un bon compromis consiste à prendre un pas spatial [TH05] :

$$\Delta_s = \frac{\lambda_{min}}{10} \quad (3.28)$$

Puisque les expériences de RT en cavité réverbérante demandent des temps de simulation longs, la dispersion numérique ne sera pas négligeable [DF99]. Cependant, le RT fonctionne dans des milieux dispersifs, la dispersion étant compensée intrinsèquement par le processus [DF99] ; [ED15b] ; [PKS09].

3.2.2 Équivalence simulation - expérience

L'objectif du code de simulation FDTD 2D est de mieux comprendre la physique des plasmas amorcés par RT avec le dispositif expérimental développé pendant cette thèse. La présentation de ce dernier ainsi que des résultats expérimentaux font l'objet du chapitre suivant. Afin de se placer dans des conditions similaires aux conditions expérimentales, nous allons prendre les mêmes propriétés du signal d'excitation : une fréquence de porteuse $f_c = 2.4$ GHz avec une bande passante de $\Delta f = 250$ MHz. Comme nous l'avons expliqué précédemment dans ce chapitre, la fréquence de porteuse détermine la dimension de la focalisation ($\lambda_c/2$) et la bande passante le nombre de modes et donc les contrastes temporel et spatial.

Dans le problème numérique à deux dimensions, la formule qui donne une approximation du nombre de modes TE jusqu'à une fréquence f en fonction de la surface S de la cavité s'écrit (voir Annexe A.2) :

$$N_{2D}^{TE}(f) = \frac{\pi S}{c_0^2} f^2 \quad (3.29)$$

Sur une bande passante $\Delta f \ll f$ au voisinage de f , il y a donc statistiquement un nombre $N_{\Delta f_{2D}}^{TE}$ de modes TE :

$$N_{\Delta f_{2D}}^{TE}(f) = \frac{dN_{2D}^{TE}(f)}{df} \Delta f = \frac{2\pi S}{c_0^2} f \Delta f \quad (3.30)$$

Nous avons vu dans le début de ce chapitre qu'un paramètre déterminant était le nombre de modes $N_{\Delta f}(f)$ contenus dans la bande passante utile Δf (autour d'une fréquence porteuse f_c). Pour que les simulations se déroulent dans des conditions similaires aux conditions expérimentales du point de vue des propriétés de focalisation spatio-temporelle du RT, il faut donc que la relation suivante soit vérifiée :

$$N_{\Delta f_{2D}}^{TE}(f) = N_{\Delta f}(f) \quad (3.31)$$

Ce qui donne une relation entre la surface S de la cavité 2D par rapport au volume V de la cavité expérimentale (3D), dépendant de la fréquence d'excitation f :

$$S = \frac{4V}{c_0} f \quad (3.32)$$

Sachant que le volume de la cavité expérimentale est de $V = 0.6 \times 0.6 \times 0.3 \text{ m}^3$, et pour une fréquence centrale $f_c = 2.4 \text{ GHz}$, il faut considérer en simulation une surface de la cavité 2D de :

$$S = 3.5 \text{ m}^2 \quad (3.33)$$

En prenant une surface de cavité en simulation de 3.5 m^2 , on s'assure donc le même nombre de modes que pour la cavité expérimentale. Ce faisant, on retrouve le même temps de Heisenberg, qui s'exprime en cavité 2D pour le modes TE :

$$t_{H_{2D}}^{TE}(f) = \frac{N_{\Delta f_{2D}}^{TE}}{\Delta f} = \frac{2\pi S}{c_0^2} f \quad (3.34)$$

On retrouve bien, pour une fréquence centrale de $f_c = 2.4 \text{ GHz}$, le même temps de Heisenberg $t_{H_{2D}}^{TE}(f_c) = 600 \text{ ns}$ pour la cavité simulée 2D que pour la cavité expérimentale (relation 3.6 avec $V = 0.6 \times 0.6 \times 0.3 \text{ m}^3$).

Comme nous l'avons vu dans la section précédente, les deux autres paramètres qui influent sur les contrastes (temporel et spatial) d'une expérience de RT sont les pertes (qui déterminent le temps de réverbération t_{reverb}) et le degré de chaoticité de la cavité. Nous allons dans la partie suivante étudier en simulation l'influence de ces paramètres sur les contrastes. L'objectif est d'utiliser le code de simulation pour mieux comprendre le comportement des ondes dans la cavité lors d'un RT mono-voie. Nous vérifierons au passage que les propriétés

spatio-temporelles de la focalisation par RT obtenues suivent bien les tendances prévues par la théorie développée dans la section précédente.

3.2.3 Étude des propriétés spatio-temporelles du retournement temporel simulé

Nous allons dans cette partie utiliser le code de simulation pour illustrer l'influence des pertes et de la chaoticité de la cavité sur les propriétés de la focalisation par RT. Pour cela, nous allons tracer l'évolution des contrastes temporels C_t et C_t^{max} et spatio-temporel C_{st}^{max} en fonction de la durée ΔT de la fenêtre temporelle appliquée à la réponse impulsionnelle, et ce pour deux cavités : une régulière (carré) et l'autre chaotique (cercle tronqué⁵, comme reporté dans le chapitre 1), dont les géométries sont représentées sur la Figure 3.12. Les surfaces de ces deux cavités sont prises égales à 3.5 m^2 pour que leur temps de Heisenberg soient le même que celui de la cavité expérimentale.

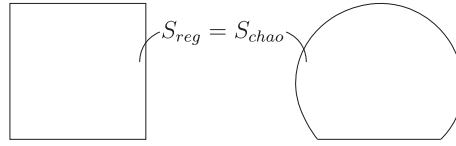
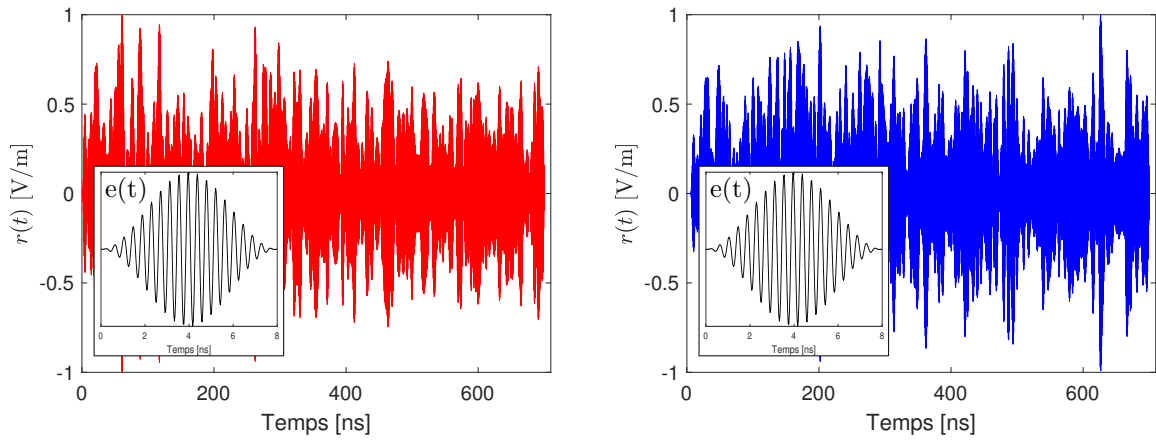


FIGURE 3.12 – Illustration des cavités utilisées en simulation. (a) Régulière (b) Chaotique.

La démarche est la suivante. La conductivité des parois des cavités (qui détermine les pertes) est choisie de sorte à ce que le temps de réverbération de ces cavités soit bien supérieur à leur temps de Heisenberg : $t_{reverb} \gg t_H$. Cela correspond au cas 2 de la Figure 3.7, les modes sont alors résolus. Quatre émetteurs et quatre récepteurs sont placés aléatoirement dans ces cavités. Lors de la première phase du RT, les réponses impulsionnelles entre chaque émetteur et chaque récepteur sont simulées. Puis, lors de la deuxième phase du RT, ces réponses sont retournées temporellement et transmises par chaque récepteur, jouant alors le rôle d'émetteur (cette opération est effectuée 16 fois, pour chaque couple émetteur-récepteur). Un exemple de signaux obtenus lors d'une expérience de RT (entre un récepteur et un émetteur) est présenté sur la Figure 3.13 pour les deux cavités de la Figure 3.12. Les pertes aux parois étant dans ce cas très faibles, sur les 700 ns de la simulation, on ne voit pas la décroissance exponentielle de l'amplitude de la réponse impulsionnelle (cas de gauche de la Figure 3.6). La durée de la fenêtre temporelle est alors de $\Delta T = 700 \text{ ns}$. Les contrastes sont obtenus en moyennant sur toutes les simulations entre chaque émetteur et chaque récepteur (il y en a donc 16). Les réponses impulsionnelles sont ensuite progressivement tronquées en diminuant ΔT : les contrastes correspondants sont alors calculés. Ces opérations sont répétées jusqu'à un ΔT bien inférieur au temps de Heisenberg : $\Delta T \ll t_H$. Dans ce cas les modes ne sont pas résolus, la situation est alors similaire au cas 1 de la Figure 3.7. Pour le calcul des contrastes spatiaux, les mesures sont faites en excluant une surface d'environ quelques $(\lambda/2)^2$ autour de la source.

5. Le cercle tronqué est une géométrie classiquement utilisée pour obtenir une cavité chaotique [Mic09]. Les premières expériences du RT ont été réalisées en acoustique avec un design en cercle tronqué [DAF99]; [Ros00]

(a) Réponses impulsionnelles



(b) Focalisation par RT

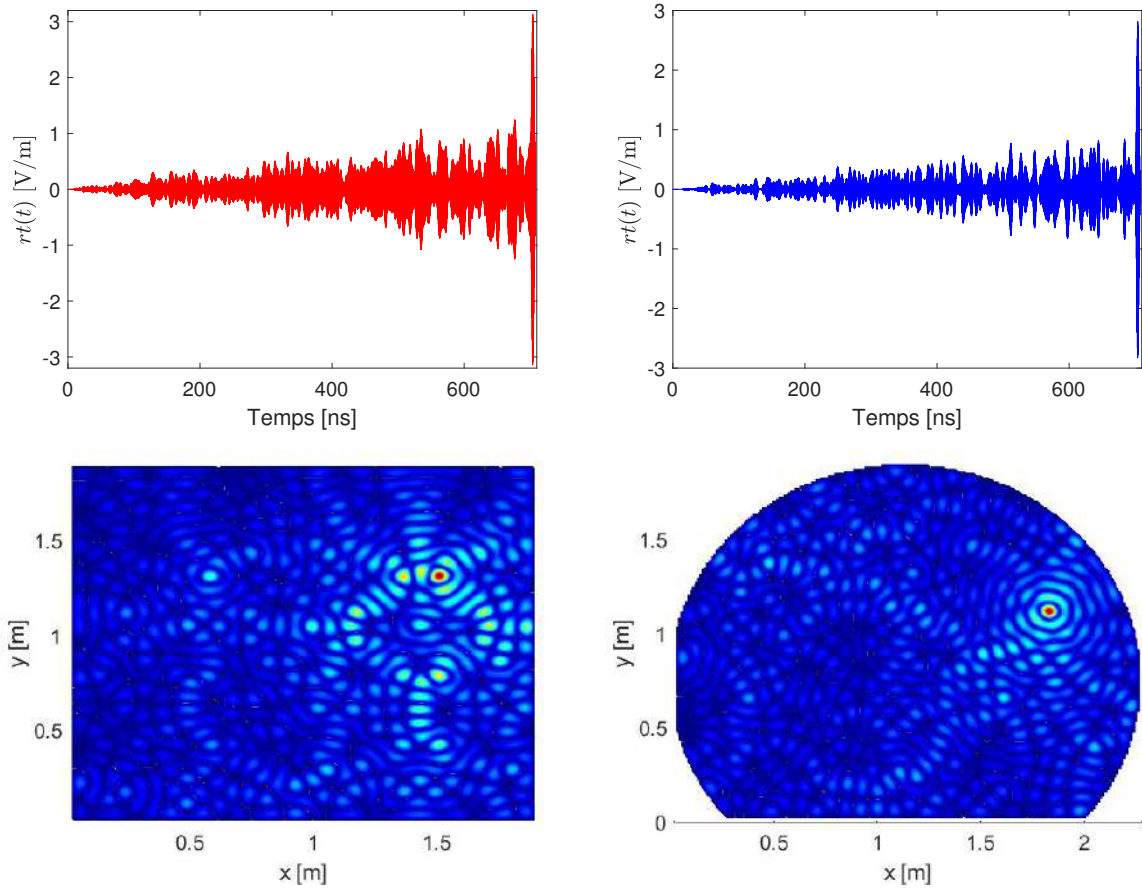


FIGURE 3.13 – Résultats des simulations avec $\Delta T = 700$ ns. À gauche les résultats obtenus avec la cavité régulière et à droite avec la chaotique. (a) Réponses impulsionnelles obtenues avec une forte conductivité des parois. Insert : Impulsion initiale de 8 ns modulée à 2.45 GHz. (b) Signaux obtenus à l'endroit de la focalisation ainsi que la répartition spatiale de la norme du champ électrique $|E_z(\mathbf{r}, t_{rt})|$ dans la cavité au moment de la focalisation.

Les contrastes ainsi obtenus en simulation sont tracés sur la Figure 3.14. Les valeurs des contrastes temporels $C_{t_{\text{experimental}}}$ et $C_{t_{\text{experimental}}}^{\text{max}}$ mesurés sur le dispositif expérimental sont aussi reportés sur cette Figure. Nous interpréterons ces derniers dans le chapitre suivant.

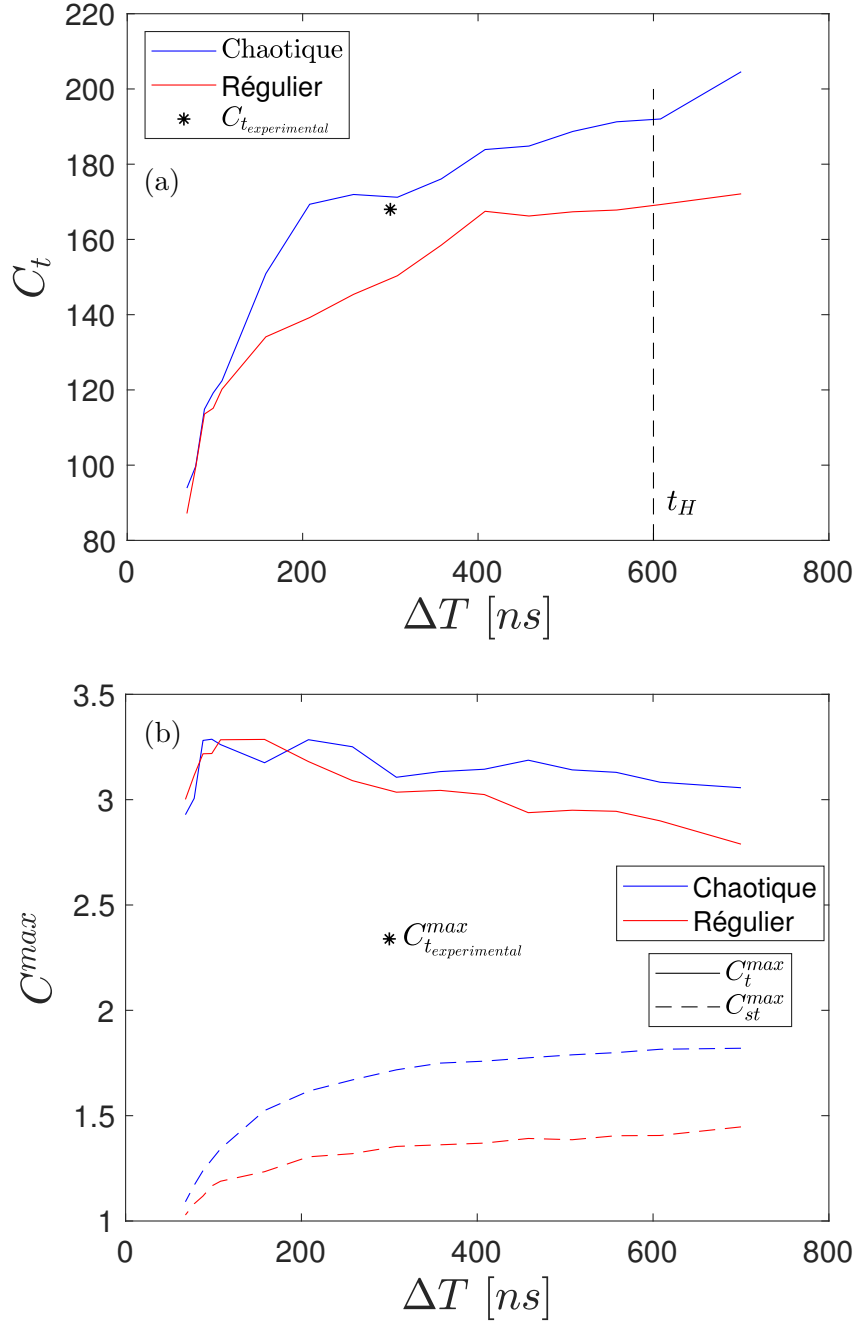


FIGURE 3.14 – Évolution des contrastes, en bleu pour la cavité régulière et en rouge pour la cavité chaotique, en fonction de ΔT . (a) Contraste temporel C_t . (b) Contrastes temporel C_t^{max} (trait plein) et spatio-temporel C_{st}^{max} (pointillé).

Pour le contraste temporel C_t , on retrouve le comportement prévu par la théorie. À faible temps de réverbération par rapport à t_H , la chaotité de la cavité n'a pas d'influence et le contraste augmente linéairement avec le temps de réverbération : $C_t \rightarrow 4\sqrt{\pi}\Delta T/\tau$ [Ros00]. Ensuite, plus le temps de réverbération augmente et plus le contraste augmente. Dans le domaine spectral, cela se traduit par le fait que plus en plus de modes sont résolus et que la quantité d'information contenue dans la bande passante augmente. Pour un certain temps de réverbération, ce nombre de modes résolus est plus important pour une cavité chaotique grâce au phénomène de répulsion des modes. Cela se traduit par un contraste supérieur pour la cavité chaotique. Puis, pour des temps de réverbération plus grand que le temps de Heisenberg, les modes sont tous résolus et il n'est pas possible d'avoir plus d'information décorrélée dans la bande passante. Le contraste sature alors : $C_t \rightarrow 4\sqrt{\pi}/k_u \cdot t_H/\tau$ [Ros00]. On trouve bien que ce contraste sature à une valeur plus élevée pour une cavité chaotique que régulière. Cela vient du kurtosis k_u , qui traduit la façon dont se répartissent les amplitudes des modes et qui est aux alentours de 3 pour une cavité chaotique. Dans sa thèse, Julien De Rosny trouve un kurtosis de 2.7 pour le cercle tronqué [Ros00]. En prenant cette valeur, on trouve un plateau de saturation de $C_t \rightarrow 197$, en bonne concordance avec la courbe du contraste obtenue en simulation pour le cercle tronqué (Figure 3.14). Pour la cavité régulière, l'amplitude des modes est répartie de façon moins concentrée et ce kurtosis sera plus élevé que 3, donnant un plateau de saturation inférieur à celui de la cavité chaotique.

Ensuite on remarque que l'on obtient de meilleurs contrastes temporel C_t^{max} et spatio-temporel C_{st}^{max} pour la cavité chaotique que régulière (comme reporté dans [Alb+19] pour C_t^{max}). Pour le contraste temporel C_t^{max} , celui-ci reste toujours supérieur à 3 pour la cavité chaotique et diminue avec ΔT pour la cavité régulière. L'évolution du contraste spatio-temporel C_{st}^{max} en fonction de ΔT semble suivre les mêmes tendances que le contraste temporel C_t . À faible ΔT , ce contraste spatio-temporel est proche de 1 pour les deux cavités. Cela signifie que durant la seconde phase du RT, il y a eu, à un endroit dans la cavité autre que celui de la focalisation, un niveau de champ similaire à celui de la focalisation. Puis, plus ΔT augmente, plus ce contraste C_{st}^{max} augmente jusqu'à saturer pour des valeurs supérieures au temps de Heisenberg (600 ns), autour de 1.8 pour la cavité chaotique et 1.5 pour la cavité régulière. Notons que, même dans la cavité chaotique, ces valeurs obtenues de contraste spatio-temporel ne sont pas très élevées. Elles signifient qu'il y a dans la cavité à un moment un (ou plusieurs) lobes secondaires dont le niveau est moins de deux fois inférieur à celui du pic de RT. Il est possible que ces lobes correspondent aux fronts d'onde circulaire se re-crétant autour du pic de RT au moment du RT. On pourrait alors peut-être considérer que ces lobes font partie de la focalisation et ne constituent alors pas des lobes secondaires limitant le processus de RT. La compréhension de ces faibles valeurs de contrastes spatio-temporel demanderait une étude plus approfondie. Nous verrons cependant dans la section suivante que ces contrastes suffisent à illustrer l'influence de la chaotité sur le contrôle des plasmas par RT et nous en resterons donc là pour cette analyse.

La différence de contraste entre les cavités régulière et chaotique est illustrée par exemple sur les résultats obtenus par RT avec un $\Delta T = 700$ ns présentés sur la Figure 3.13. Sur les signaux mesurés à l'endroit de la focalisation en fonction du temps, on remarque un contraste temporel maximal C_t^{max} similaire pour les cavités chaotique et régulière (on retrouve

bien aussi un contraste temporel C_t supérieur pour la cavité chaotique). Par contre, sur la répartition spatiale du module $|E_z(\mathbf{r}, t_{rt})|$ du champ électrique on remarque la présence de lobes secondaires beaucoup plus importants dans la cavité régulière que dans celle chaotique. Cela vient de la présence de surfaces parallèles, qui vont avoir tendance à répartir l'énergie électromagnétique dans la cavité de façon non-uniforme, avec des concentrations de champ à certains endroits. Pour le cercle tronqué, en dehors de l'endroit de la focalisation on retrouve une répartition typique de l'énergie dans une cavité chaotique, qui se fait de manière uniforme, sans concentration apparente de champ. Dans les deux cavités, on retrouve bien une dimension de la tache de focalisation de $\lambda_c/2 = 6$ cm [Ler+04].

Il faut aussi noter que cette répartition spatiale du module du champ électrique est tracée au moment du niveau maximal de champ électrique de focalisation, mais cette focalisation dure 8 ns et est modulée par une fréquence porteuse de $f_c = 2.4$ GHz. À l'endroit de la focalisation, le champ va donc osciller à 2.4 GHz pendant 8 ns sur une tache de 6 cm. C'est pour cela que l'on voit, autour du pic de focalisation, des fronts d'onde convergeant vers l'endroit de la focalisation (puisque la focalisation va encore durer 4 ns). Ces fronts d'onde sont beaucoup plus proches de fronts d'onde sphérique dans la cavité chaotique que ceux de la cavité régulière. Pour cette dernière, il semble qu'il y ait trois directions privilégiées par lesquelles les ondes convergent vers le point de focalisation. Cela vient encore de la structure régulière de la cavité. Pour une cavité chaotique, il est possible de recréer un front d'onde sphérique puisque les ondes vont arriver de toutes les directions [DAF99] ; [Qui04], conséquence de l'isotropie du champ.

Les répartitions spatiales du module du champ électrique au moment du pic de RT dans la cavité de la Figure 3.13 illustrent clairement les deux propriétés intéressantes pour le RT des cavités chaotiques : l'uniformité et l'isotropie de la distribution spatiale du champ. Finalement, le code simulation permet bien de simuler des processus de RT dont les tendances sont conformes aux prédictions théoriques. Nous avons illustré l'importance de la chaoticité sur la qualité de la focalisation temporelle et spatiale. Cela montre en quoi la chaoticité de la cavité peut être déterminante dans le contrôle des plasmas par RT. Nous allons dans la section suivante, coupler ce code de simulation à un modèle de plasma fluide, permettant de décrire l'amorçage du plasma. Nous montrerons le caractère essentiel de la chaoticité de la cavité sur le contrôle de plasmas par RT.

3.3 Modélisation du retournement temporel en cavité pour l'amorçage de plasmas : Couplage Maxwell-fluide

Nous avons présenté dans la première section de ce chapitre les équations permettant de décrire un RT mono-voie en cavité réverbérante. Nous avons vu que ce dernier présente des lobes secondaires temporels et spatiaux, conséquence de l'émission dans toutes les directions lors de la phase inverse du RT. Puis, dans la deuxième section, nous avons présenté la partie électromagnétique du code de simulation FDTD 2D développé durant cette thèse. Nous avons vérifié que les résultats de RT obtenus sont bien décrits par les équations développées dans la première section. Nous avons pu montrer l'intérêt d'une étude en simulation : le contrôle de la géométrie de la cavité et la possibilité de visualiser le champ électromagnétique partout dans la cavité. L'objectif de cette dernière section est d'étudier les plasmas amorcés par RT en simulation FDTD 2D dans des conditions similaires à ceux amorcés expérimentalement. À travers quelques exemples, nous pourrions visualiser ce qu'il se passe dans la cavité lors de l'amorçage des plasmas par RT et illustrer avec un modèle simple les mécanismes physiques ayant lieu lors de l'amorçage des plasmas par RT. Notons tout de même que si ce code constitue un outil intéressant pour extraire des tendances, il n'a pas la prétention d'être une description complètement fidèle des mécanismes ayant lieu lors d'une expérience de RT sur le dispositif expérimental. En effet, la différence majeure réside dans la représentation des mécanismes ayant lieu en 2D et en 3D (diffusion du plasma différente par exemple).

Nous présenterons d'abord le principe de l'amorçage de plasmas par RT utilisé en simulation. Ensuite le code fluide implémenté pour décrire l'amorçage de plasmas par RT en cavité est détaillé, avant de présenter les résultats.

Les paramètres de simulation sont choisis de sorte à ce que les plasmas soient amorcés dans des conditions similaires à celles du dispositif expérimental présenté au chapitre suivant, à savoir :

Propriétés des plasmas amorcés par RT expérimentalement	
Pression dans la cavité pour la phase inverse p_{inv}	$\in \{0.5; 4\}$ torr
Gaz	argon

3.3.1 Principe

Le principe de l'amorçage d'un plasma par RT consiste à utiliser l'énergie contenue dans le pic de focalisation à l'endroit et au moment de la focalisation pour claquer le gaz à cet endroit dans la cavité. La première phase du RT est la même que lors d'un RT sans plasma : elle consiste à enregistrer la réponse impulsionnelle $r(t)$ entre deux points de la cavité, comme illustré sur la première ligne de la Figure 3.15. Dans ce cas seulement, le modèle électromagnétique est utilisé avec du vide (ϵ_0, μ_0) et il n'y a donc pas de claquage. Dans un cas plus pratique cela correspond à mesurer la réponse impulsionnelle avec une pression p_{init} assez élevée pour éviter le claquage. Lors de la deuxième phase du RT, *i.e.* la réémission de la

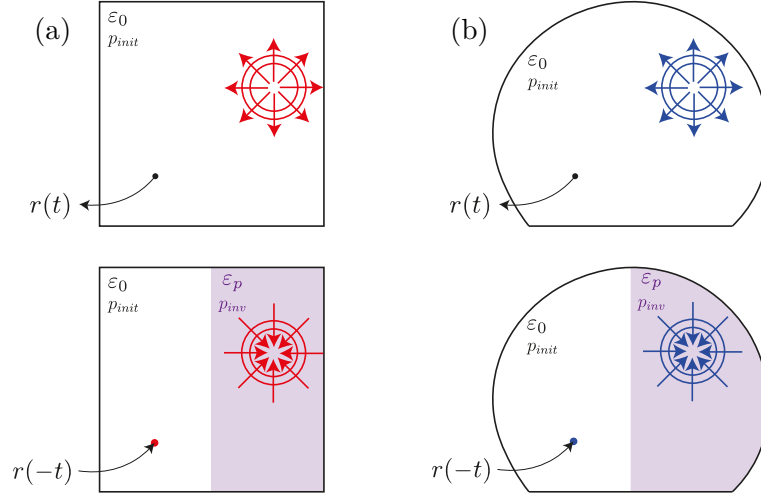


FIGURE 3.15 – Principe de l’amorçage de plasmas par RT en simulation. Ligne du haut : enregistrement de la réponse impulsionnelle $r(t)$ entre deux points dans la cavité. Ligne du bas : émission de la réponse retournée temporellement $r(-t)$ dans une cavité dans laquelle le modèle fluide est implémenté sur la moitié de la cavité contenant l’endroit de la focalisation et excluant la source. (a) Cavité régulière. (b) Cavité chaotique.

réponse impulsionnelle retournée temporellement, le modèle fluide est couplé avec les équations de Maxwell sur la moitié de la cavité contenant l’endroit de la focalisation et excluant la source, comme représenté sur la deuxième ligne de la Figure 3.15. Ce point est primordial pour éviter le claquage au niveau de la source d’émission, là où le champ électrique est très intense. Comme nous le verrons dans le chapitre suivant, sur le dispositif expérimental, l’antenne d’émission est placée dans une annexe de la cavité, laissée à la pression atmosphérique (p_{init}) et séparée du reste de la cavité par un hublot de verre. Dans la cavité, en revanche, de l’argon est injectée à la pression p_{inv} permettant le claquage par RT.

Comme nous l’avons vu dans ce chapitre, lors de la phase inverse, l’émission dans toutes les directions se traduit par la présence de lobes secondaires temporels et spatiaux. C’est à dire que pendant toute la durée de la seconde phase (émission de la réponse impulsionnelle retournée temporellement), les ondes vont se propager et interférer dans la cavité. De ces interférences résulte une certaine organisation spatiale de l’énergie électromagnétique dans la cavité qui évolue avec le temps. Afin d’amorcer un plasma au moment et à l’endroit de la focalisation par RT, le niveau du champ électrique du pic de RT obtenu sans le modèle fluide est réglé (en jouant sur l’amplitude de la réponse impulsionnelle retournée temporellement) de sorte à ce qu’il soit supérieur au champ de claquage pour des pressions autour du torr et une durée d’application du champ de $\tau = 8$ ns.

Cependant, puisque ce niveau du seuil de claquage dépend aussi de la durée d’application du champ, un contraste spatio-temporel C_{st}^{max} supérieur à 1 peut ne pas suffire pour assurer la génération du plasma seulement au moment et à l’endroit de la focalisation. Ainsi, la façon dont se répartit l’énergie dans la cavité (et donc sa chaoticité) est déterminante. Dans une

cavité régulière, le contrôle de la position du plasma sera rapidement limité par la présence de zones d'intensification du champ électrique où des plasmas peuvent être générés. De plus, une fois présents dans la cavité, ces plasmas vont changer le comportement des ondes. Les chemins que les ondes étaient censés emprunter entre les deux transducteurs afin de focaliser au moment de la focalisation peuvent alors être modifiés, résultant en une moins bonne focalisation. Dans une cavité chaotique, l'énergie se répartit plus uniformément comme nous l'avons vu dans ce chapitre. Dans ce cas le contrôle de la position du plasma par RT sera grandement facilité, comme illustré dans la suite de cette section.

3.3.2 Description du modèle fluide FDTD 2D

Le modèle de plasma fluide développé durant cette thèse est largement inspiré de celui utilisé dans [CB10]. Comme nous l'avons expliqué au chapitre 1, le plasma peut être modélisé par la densité de courant $\mathbf{J}_e$ des électrons dans le vide. Les équations du modèle fluide développées dans le chapitre 1 sont utilisées dans le code de simulation. Par souci de clarté, elles ne sont par ré-écrites ici mais sont rassemblées dans l'annexe A.3. Elles sont discrétisées en utilisant un schéma centré tel que présenté dans la section précédente, suivant le même développement que celui présenté dans [CB10] avec en plus la présence du courant de conduction $\mathbf{J}_c(\mathbf{r}, t)$ permettant d'imposer la conductivité des parois de la cavité. En omettant les dépendances spatiales et au champ électrique par souci de clarté, les équations sur la vitesse des électrons et le champ électrique deviennent alors :

$$v_e^{n+1} = \alpha v_e^n - \frac{e\Delta_t}{2m_e\zeta} (E_z^{n+1} + E_z^n) \quad (3.35)$$

$$E_z^{n+1} = \frac{1-\beta}{1+\beta} E_z^n + \frac{en_e\Delta_t}{2\varepsilon_0} \frac{1+\alpha}{1+\beta} v_e^n + \frac{\Delta_t}{\varepsilon_0(1+\beta)} \left(\frac{\partial H_y^{n+1/2}}{\partial x} - \frac{\partial H_x^{n+1/2}}{\partial y} \right) \quad (3.36)$$

avec :

$$\alpha = \frac{1-a}{1+a}, \quad \zeta = 1+a, \quad a = \frac{\nu_m\Delta_t}{2} \text{ et } \beta = \frac{\omega_p^2\Delta_t^2}{4\zeta} + \frac{\Delta_t}{2\varepsilon_0}\sigma \quad (3.37)$$

De plus, l'équation sur la densité électronique s'écrit :

$$\begin{aligned} n_e^{n+1}(i_x, i_y) = & (1 + \Delta_{t_F}\nu_i(i_x, i_y))n_e^n(i_x, i_y) + \frac{D_{eff}(i_x, i_y)\Delta_{t_F}}{\Delta_s^2} [n_e^n(i_x + 1, i_y) \\ & + n_e^n(i_x - 1, i_y) + n_e^n(i_x, i_y + 1) + n_e^n(i_x, i_y - 1) - 4n_e^n(i_x, i_y)] \end{aligned} \quad (3.38)$$

avec Δ_{t_F} le pas temporel du modèle fluide, correspondant à la période des ondes électromagnétiques $\Delta_{t_F} = 1/f_c$ [CB10]. En effet, les grandeurs fluides dépendent de la valeur rms du champ électromagnétique, *i.e.* la seule connaissance de la valeur rms du champ électromagnétique à

chaque période d'oscillation suffit à décrire l'évolution de ces grandeurs physiques.

Les modèles fluides des décharges plasma nécessitent en entrée des coefficients de transport et des taux de réaction qui dépendent de la fonction de distribution. Ces coefficients sont souvent calculés à partir des sections efficaces de collision en résolvant l'équation de Boltzmann pour l'électron (équation 1.13). Le modèle fluide présenté ci-dessus nécessite la connaissance de l'évolution des grandeurs $D_e(\mathbf{E})$, $\nu_i(\mathbf{E})$, $\nu_m(\mathbf{E})$ et $\mu_e(\mathbf{E})$ en fonction du champ électrique $\mathbf{E}$. Elles sont obtenues à l'aide du solveur Bolsig + [HP05]. Ce dernier propose une solution de résolution de l'équation de Boltzmann qui rend compte de phénomènes quasi-stationnaires et de champs oscillants.

Au final, ce modèle fluide permet de suivre l'évolution spatio-temporelle de la densité électronique n_e durant les simulations. La densité électronique initiale dans les cavités est prise égale à ⁶ $n_{e0} = 10^{10} \text{ m}^{-3}$. Généralement, on considère qu'il faut que la densité ait augmenté d'un facteur 10^8 pour pouvoir parler d'amorçage d'un plasma [GR56]. Nous prendrons ce critère pour déterminer si un plasma est généré par RT.

3.3.3 Équivalence simulation - expérience

Tout d'abord la conductivité des parois des cavités en simulation est choisie de sorte à obtenir le même temps de réverbération que l'expérience : environ 300 ns (selon l'équation 3.16). Les réponses impulsionnelles obtenues dans les cavités régulière et chaotique pour une configuration (entre un émetteur et un capteur) sont tracées sur la Figure 3.16(a-b). Dans ce cas, comme sur le dispositif expérimental l'amplitude de ces réponses suit une décroissance exponentielle à cause des pertes lors des réflexions des ondes dans la cavité et le temps de réverbération est défini comme :

$$t_{\text{reverb}} = \pi\tau_{\text{exp}} = \pi\sqrt{\frac{\langle \int t^2 \cdot r^2(t) dt \rangle}{\langle \int r^2(t) dt \rangle}} \quad (3.39)$$

avec $\langle . \rangle$ une moyenne sur 16 couples capteur - émetteur. La largeur spectrale des modes Γ_{2D}^{TE} des cavités simulées dans ces conditions, c'est à dire avec un temps de réverbération d'environ 300 ns, est donc égale à $2/(\pi\tau_{\text{exp}}) = 6 \text{ MHz}$ (comme c'est le cas sur le dispositif expérimental). Le facteur 2 est une conséquence de la modulation. On retrouve bien cette largeur spectrale sur les transformées de Fourier tracées sur la Figure 3.16(c-d). Au passage on remarque l'effet de la chaoticité sur le spectre : comme attendu le spectre de la cavité chaotique possède une réponse plus plate que celui de la cavité régulière, conséquence de la répulsion des modes et du kurtosis plus faible.

À partir de l'équation 3.30, donnant le nombre de modes TE d'une cavité 2D contenus dans une bande passante Δf , le recouvrement modal correspondant s'écrit donc :

⁶. Des simulations réalisées avec différentes valeurs de densité initiale donnent des comportements similaires. Les résultats du code ne sont donc pas sensibles à cette densité initiale.

$$d_{2D}^{TE}(f) = \frac{2\pi S}{c_0^2} f \Gamma_{2D}^{TE} \quad (3.40)$$

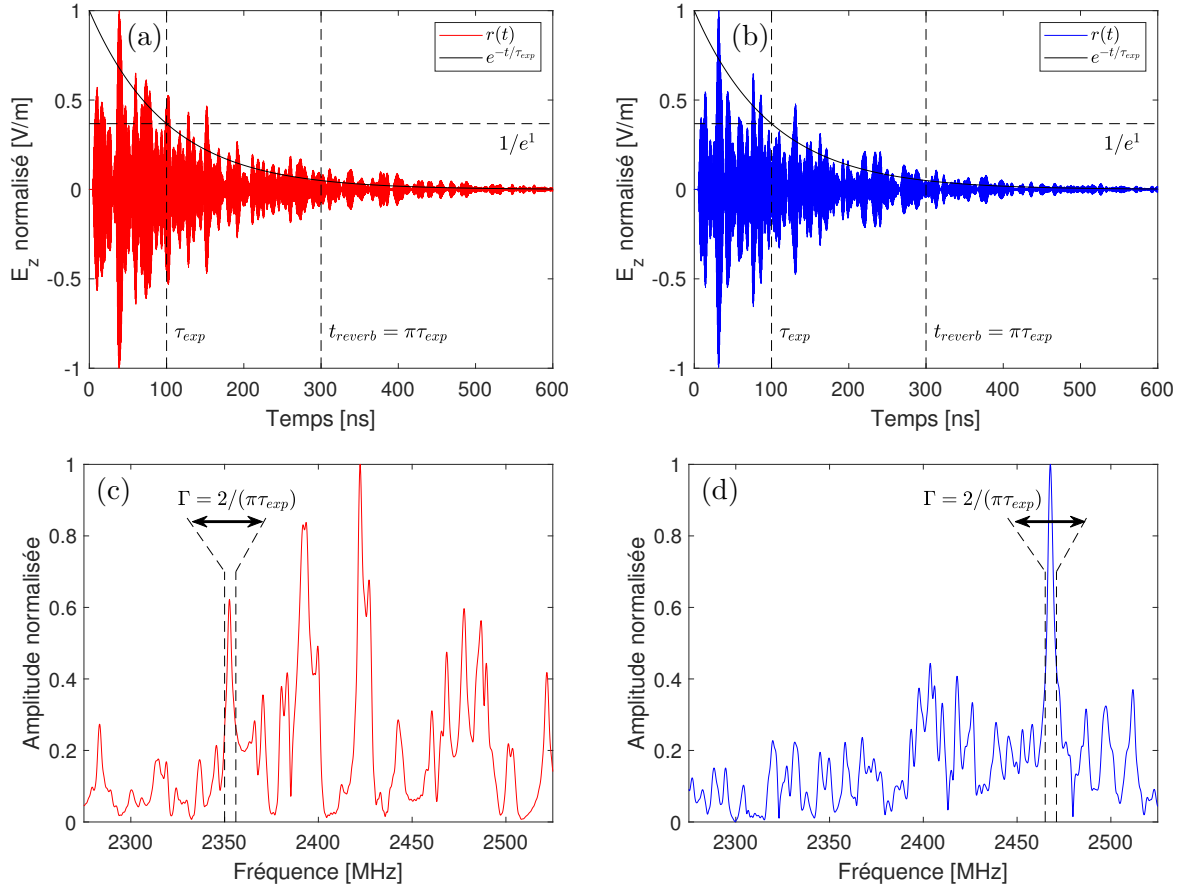


FIGURE 3.16 – Réponses impulsionnelles obtenues en simulation avec une constante de temps de l'exponentielle d'environ $\tau_{exp} = 100$ ns, et donc un temps de réverbération d'environ $t_{reverb} = 300$ ns. A gauche (rouge) : cavité régulière, (a) domaine temporel et (c) domaine fréquentiel. A droite (bleu) : cavité chaotique, (b) domaine temporel et (d) domaine fréquentiel.

Autour de la fréquence centrale $f_c = 2.4$ GHz le recouvrement modal en simulation est donc de $d_{2D}^{TE}(f_c) \approx 3.5$ (que l'on retrouve bien expérimentalement). D'après le critère de Schroeder, les cavités simulées se trouvent juste au dessus de la limite à partir de laquelle le régime de recouvrement modal devient fort [Sch87]. Si c'est le cas, alors les modes se recouvrent assez les uns les autres spectralement, de sorte à ce que le champ soit uniforme et isotrope dans la cavité indépendamment de la chaotité de la cavité. Les contrastes en cavité régulière ou chaotique sont alors censés être similaires. Cependant, on remarque sur la Figure 3.14 que pour une durée de la fenêtre temporelle de 300 ns, les contrastes temporels et spatiaux sont meilleurs dans la cavité chaotique. Cela signifie que le régime de recouvrement modal fort ($d \gg 1$) n'est pas tout à fait atteint dans ces conditions. La chaotité de la cavité

aura ainsi une influence importante sur la qualité de la focalisation spatio-temporelle par RT et donc sur le contrôle des plasmas. Ce point est illustré dans la partie suivante.

De plus, la pression avec laquelle les plasmas d'argon sont amorcés dans le dispositif expérimental est comprise sur la gamme : $p_{inv} \in \{0.5; 4\}$ torr. La pression à laquelle le transfert de puissance est trouvé maximal est $p_{inv} = 1.5$ torr, en accord avec la littérature [Fel66]. Cette pression dépend de la fréquence des ondes, du gaz et de la durée de l'impulsion [Fel66]. Nous effectuerons donc les simulations des plasmas par RT dans l'argon autour de cette pression.

3.3.4 Plasmas amorcés par retournement temporel en simulation

Nous allons présenter dans cette partie les plasmas amorcés par RT en simulation. Dans un premier temps, deux exemples d'amorçage de plasmas par RT en cavités régulière et chaotique sont présentés afin d'illustrer l'importance de la chaoticité sur le contrôle des plasmas par RT. Pour finir, une étude en pression et une en régime pulsé sont effectuées pour illustrer des tendances sur les mécanismes ayant lieu lors d'un amorçage sur le dispositif expérimental.

3.3.4.1 Influence de la chaoticité

Les simulations d'amorçage de plasmas par RT sont réalisées dans les deux cavités (régulière et chaotique). Dans les deux cas, le niveau maximal de l'amplitude de la réponse impulsionnelle retournée est choisi de sorte à obtenir un champ électrique au moment de la focalisation suffisant pour amorcer un claquage, soit d'environ 200 kV/m. Dans [Fel66], dans des conditions similaires à celles des simulations, *i.e.* pour un amorçage de plasma avec une durée d'application du champ microonde (3 GHz) de 10 ns dans l'air, le champ de claquage est de l'ordre de 100 kV/m. On retrouve bien cet ordre de grandeur en simulation. Les résultats présentés dans cette partie ont été obtenus dans l'argon avec une pression de $p = 1.5$ torr.

Un exemple de résultat d'amorçage de plasmas par RT en simulation est présenté sur la Figure 3.17, à gauche pour la cavité régulière et à droite pour la cavité chaotique. Sur la première ligne sont tracées les évolutions du champ électrique et de la densité électronique $n_{e_{rt}}(t)$ (en pointillé) mesurées à l'endroit de la focalisation $\mathbf{r}_{rt}$ (représentée en encart). En vert il s'agit des évolutions temporelles du champ électrique mesuré sans plasma, noté $|E_{rt}(t)|$. Le pic de RT est alors obtenu à 300 ns avec le même niveau pour les deux cavités. Nous remarquons tout d'abord la parfaite symétrie entre les lobes précédant le pic de RT et ceux lui succédant, et ce pour les deux cavités. Cela illustre parfaitement le fait qu'au moment de la focalisation les ondes convergent à l'endroit de la focalisation puis ne s'arrêtent pas et continuent de se propager dans la cavité. Au moment et à l'endroit de la focalisation se somment alors une onde convergente et une divergente, se traduisant par une limite de focalisation de $\lambda_c/2$. Après avoir focalisées, ces ondes vont emprunter les mêmes chemins qu'avant le pic de RT mais dans le sens inverse. En rouge et bleu on retrouve les évolutions temporelles du champ électrique obtenues en simulation avec le modèle fluide $|E_{rt}^{fluide}(t)|$

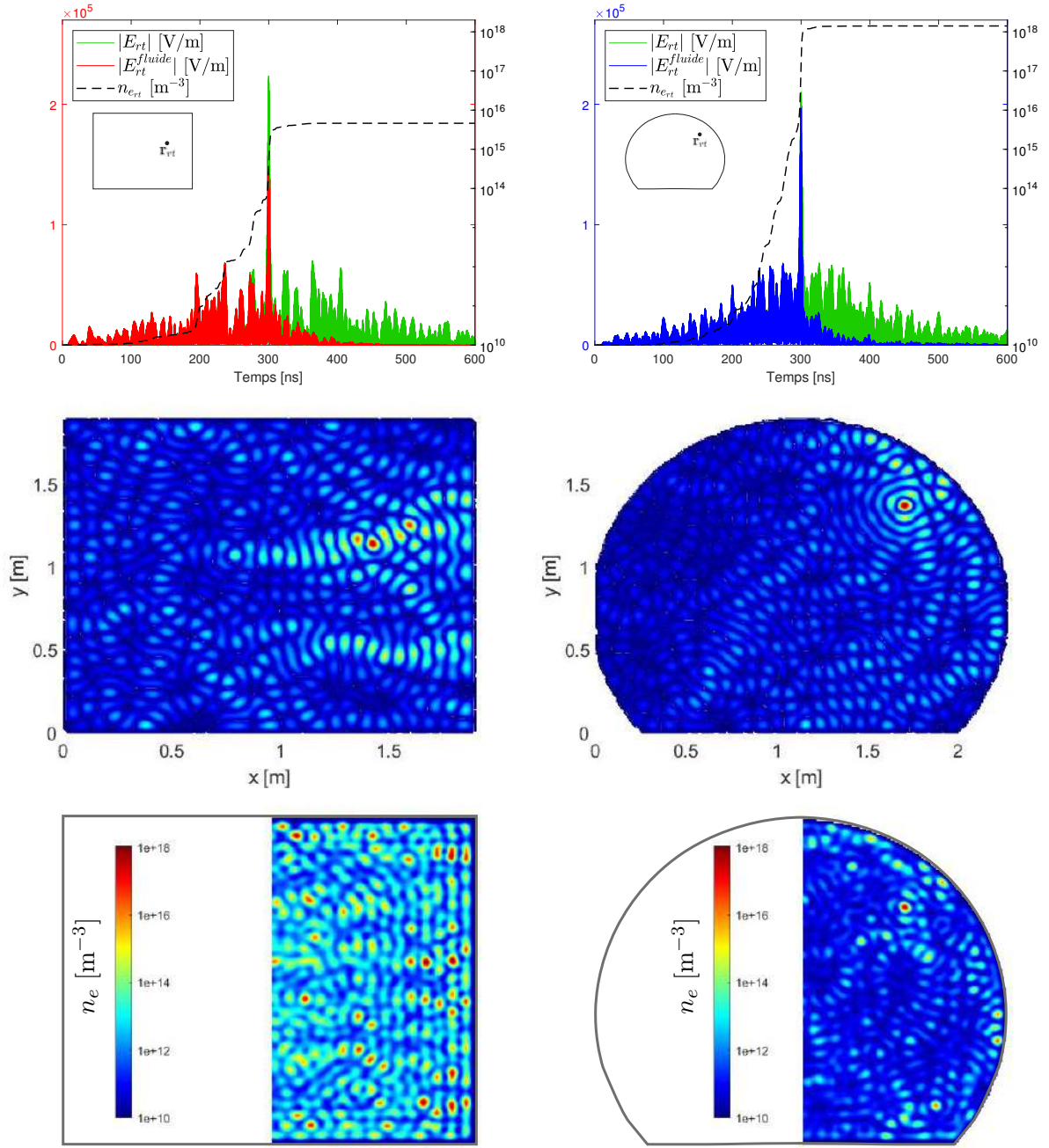


FIGURE 3.17 – Résultats de simulation de l'amorçage plasma par RT. À gauche les résultats obtenus avec la cavité régulière et à droite avec la chaotique. Sur la première ligne sont tracés les signaux microondes et la densité électronique mesurés au point de focalisation $\mathbf{r}_{rt}$. Sur la deuxième ligne est tracé la norme du champ électrique dans la cavité au moment de la focalisation $|E_z^{fluide}(\mathbf{r}, t_{rt})|$ et sur la troisième ligne est tracé la densité électronique juste après la focalisation $n_e(\mathbf{r}, t = 308 \text{ ns})$.

respectivement pour la cavité régulière et chaotique. Sur la seconde ligne est représenté la répartition spatiale du champ électrique au moment de la focalisation avec le modèle fluide $|E_z^{fluide}(\mathbf{r}, t_{rt})|$, et sur la troisième ligne la densité électronique obtenue avec le modèle fluide juste après le pic de RT (308 ns).

Pour la cavité régulière (à gauche sur la Figure 3.17), les champs électriques avec ($|E_{rt}^{fluide}(t)|$) et sans ($|E_{rt}(t)|$) plasma se superposent parfaitement jusqu'à environ 150 ns. Puis, le niveau avec le modèle fluide diminue par rapport à celui obtenu sans modèle fluide (différence de niveau que l'on ne voit pas sur la Figure 3.17). Cela provient de l'absorption d'énergie électromagnétique par le gaz, se traduisant par une augmentation de la densité électronique au point de focalisation $n_{e_{rt}}(t)$. Cette dernière augmente jusqu'au moment du pic de RT, durant lequel elle augmente de moins de deux ordres de grandeur jusqu'à atteindre une valeur maximale de $4 \times 10^{15} \text{ m}^{-3}$. On retrouve sur la seconde ligne une répartition spatiale du champ électrique $|E_z^{fluide}(\mathbf{r}, t_{rt})|$ typique d'une focalisation en cavité régulière. La focalisation n'est pas isotrope mais semble être la conséquence de deux directions privilégiées. De plus, l'organisation spatiale non uniforme du champ électrique se traduit par la présence de lobes secondaires importants durant le processus de RT. Au niveau de ces lobes secondaires, la densité électronique augmente significativement. Sur la répartition spatiale de la densité électronique juste après le pic de RT $n_e(\mathbf{r}, t = 308 \text{ ns})$, on voit clairement que la position du plasma n'est pas du tout contrôlée, puisque la densité électronique a significativement augmenté un peu partout, et notamment le long de la paroi parallèle à l'axe y . Certains de ces pics de densité atteignent des densités supérieures à 10^{18} m^{-3} , traduisant la présence de plasmas. Les ondes ne se comportent alors pas de la même façon que lors de la phase initiale. Elles peuvent être absorbées et/ou réfléchies par ces fortes densités électroniques. Le milieu de propagation diffère alors de la phase initiale et les chemins empruntés par les ondes seront différents. Une étude attentive des lobes précédant le pic de RT révèle la présence d'un lobe de niveau supérieur avec modèle fluide que sans. Cela traduit bien une réorganisation des ondes dans la cavité avec le modèle fluide. En conséquence, le niveau du pic de RT sera fortement diminué par rapport au cas sans modèle fluide. Dans ce cas la différence des niveaux de pic de RT sans et avec modèle fluide ne correspond pas, en tout cas pas entièrement, à un transfert de puissance au gaz. On pourrait parler dans ce cas de non-réciprocité du système, puisque l'augmentation significative de densité électronique en différents endroits de la cavité change significativement le milieu de propagation. Dans le cas de la cavité régulière, dans ces conditions (et pour cette position des récepteur-émetteur) il n'est donc pas possible de générer un plasma par RT (la densité électronique n'a pas augmenté d'un facteur 10^8), ou en tout cas pas sans la présence de plasmas parasites de densités supérieures à celle obtenue à l'endroit de la focalisation. Pour amorcer un plasma par RT, il faut que le système reste réciproque jusqu'au moment du pic de focalisation, comme nous l'avons vu au chapitre précédent.

Pour la cavité chaotique (à droite de la Figure 3.17), les champs électriques avec ($|E_{rt}^{fluide}(t)|$) et sans ($|E_{rt}(t)|$) plasma se superposent aussi parfaitement jusqu'à environ 150 ns. Puis, là aussi, le niveau avec modèle fluide diminue par rapport à celui obtenu sans modèle fluide (différence de niveau que l'on ne voit pas sur la Figure 3.17). Ainsi la densité électronique au point de focalisation $n_{e_{rt}}(t)$ augmente jusqu'au moment du pic de RT, durant lequel elle augmente d'environ trois ordres de grandeurs jusqu'à atteindre une valeur maximale de 1.4×10^{18}

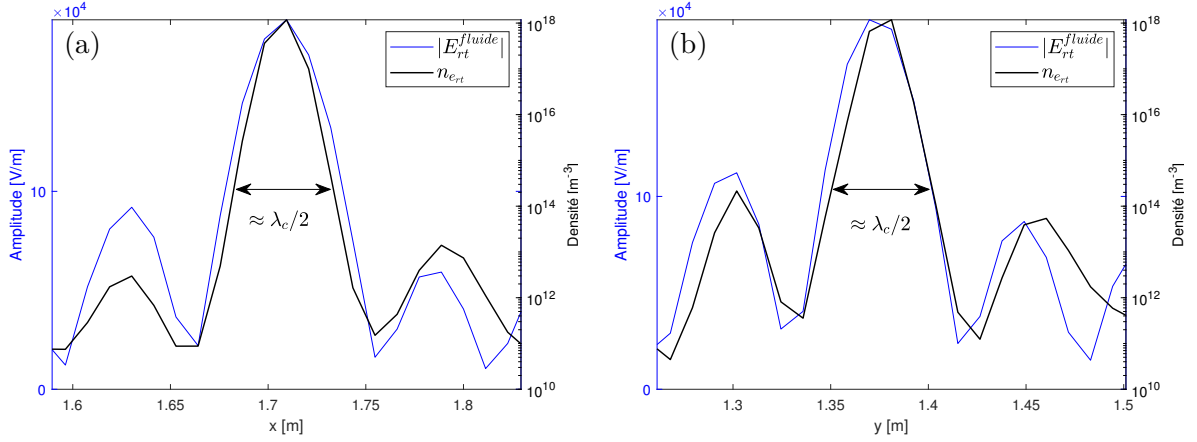


FIGURE 3.18 – Tracés du champ électrique au moment de la focalisation $|E_z^{fluide}(\mathbf{r}, t_{rt})|$ et de la densité électronique juste après $n_e(\mathbf{r}, t = 308ns)$ de la Figure 3.17 autour de l'endroit de la focalisation ($x = 1.7$ m et $y = 1.37$ m) le long (a) d'une coupe selon l'axe x ($y = 1.37$ m) et (b) d'une coupe selon l'axe y ($x = 1.7$ m).

m^{-3} , traduisant l'amorçage d'un plasma. Sur la répartition spatiale du champ au moment de la focalisation $|E_z^{fluide}(\mathbf{r}, t_{rt})|$ on retrouve les caractéristiques d'une focalisation en cavité chaotique. La focalisation est isotrope et le champ est réparti de façon uniforme dans la cavité. En conséquence, il n'y a pas d'endroit d'intensification du champ durant le processus de RT et on observe beaucoup moins de lobes de densité électronique élevée que dans le cas de la cavité régulière. De ce fait, les chemins empruntés par les ondes ne sont pas significativement modifiés avant le pic de RT. On observe alors, avec le modèle fluide, un niveau de pic de RT bien supérieur à celui obtenu dans la cavité régulière. Ce niveau du pic de RT obtenu avec le modèle fluide dans la cavité chaotique est légèrement inférieur à celui obtenu sans modèle fluide, correspondant ici à un transfert de puissance au gaz. Il est par ailleurs bien corrélé avec une forte augmentation de la densité électronique pendant le pic. Sur la répartition spatiale de la densité électronique juste après le pic de RT $n_e(\mathbf{r}, t = 308ns)$, on voit que la position du plasma est bien contrôlée par RT. Ce plasma ainsi généré est de forme circulaire. Le champ électrique et la densité électronique sont tracés autour de l'endroit de la focalisation le long d'une coupe selon les axes x et y sur la Figure 3.18. On observe une forte corrélation entre la dimension de la focalisation par RT et la dimension du plasma amorcé par RT, qui est donc de $\lambda_c/2$. On note aussi une corrélation entre les lobes secondaires de champ électrique et les lobes secondaires de densité électronique. Tout ceci illustre parfaitement la notion de réciprocité transitoire présentée au chapitre précédent. Dans ce cas les ondes se comportent de la même façon avec et sans modèle fluide jusqu'au moment du pic de RT, moment auquel un plasma est amorcé à l'endroit de la focalisation. On note toutefois la présence de lobes secondaires spatiaux de densité électronique relativement importants. Le lobe secondaire le plus important est de deux ordres de grandeur inférieur à celui obtenu au niveau du pic de RT, à savoir de $5 \times 10^{16} m^{-3}$.

Pour les deux cavités, le champ électrique avec le modèle fluide $|E_{rt}^{fluide}(t)|$ diminue fortement après le pic de RT comparé au cas sans modèle fluide. Pour la cavité régulière, cela est

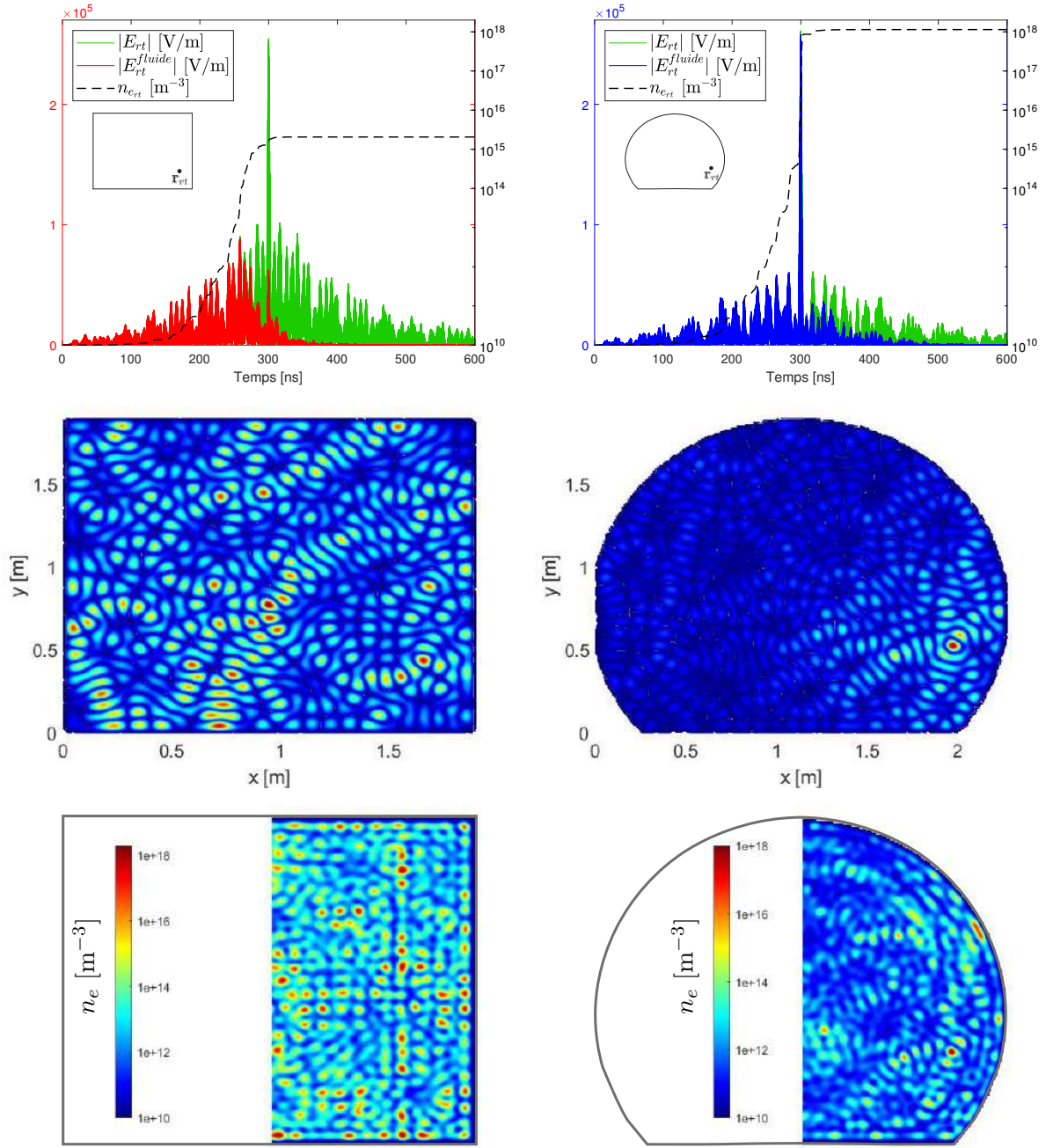


FIGURE 3.19 – Résultats de simulations de l’amorçage plasma par RT pour un autre point de focalisation que celui de la Figure 3.17. À gauche les résultats obtenus avec la cavité régulière et à droite avec la chaotique.

dû à l’absorption des ondes par les fortes densités un peu partout dans la cavité. Pour la cavité chaotique, la densité électronique obtenue par RT à l’endroit de la focalisation est supérieure à la densité critique $n_{e_{critique}} = 7 \times 10^{16} \text{ m}^{-3}$ (pour une fréquence $f_c = 2.4 \text{ GHz}$). Dans ce

cas le faible niveau de champ électrique $|E_{rt}^{fluide}(t)|$ mesuré après le pic est la conséquence de deux phénomènes locaux : la réflexion et l'absorption des ondes.

Pour finir sur l'illustration de l'importance de la chaotité sur le contrôle de plasmas par RT, un autre exemple d'amorçage de plasmas à un autre endroit dans les cavités est présenté sur la Figure 3.19. Pour la cavité régulière, nous notons dans ce cas que le pic de focalisation avec le modèle fluide se retrouve drastiquement diminué, devenant même inférieur aux lobes secondaires temporels le précédant. De la même façon on trouve sur la répartition spatiale du champ électrique dans la cavité des lobes secondaires spatiaux de niveau supérieur à celui du pic de RT. Cela est causé par un fort changement du milieu de propagation avec la génération d'endroits de fortes densités un peu partout dans la cavité, comme nous pouvons le voir sur la répartition spatiale de la densité électronique. On retrouve là aussi sur cette répartition une structure un peu similaire à celle obtenue lors du précédent exemple (Figure 3.17). Il y a des zones de forte densité s'organisant parallèlement aux parois de la cavité. Pour la cavité chaotique, la position du plasma est bien contrôlée par RT. Nous notons qu'au moment du pic de RT la densité électronique augmente de trois ordres de grandeur. On retrouve bien une bonne focalisation avec modèle fluide, comme le montre la présence du pic de RT de fort niveau sur l'évolution temporelle du champ $|E_{rt}^{fluide}(t)|$ et sur la répartition spatiale du champ électrique au moment du pic. Nous observons bien un maximum de densité électronique au niveau de la focalisation de dimension $\lambda_c/2$. Nous notons enfin la présence de lobes secondaires spatiaux de densité électronique relativement élevés, notamment un le long de la paroi de la cavité avoisinant les $5 \times 10^{16} \text{ m}^{-3}$.

Nous venons d'illustrer l'importance de la chaotité de la cavité sur le contrôle des plasmas par RT. Les propriétés de répartition spatiale du champ dans la cavité chaotique permettent ainsi un bon contrôle de la position des plasmas par RT, contrairement à une cavité régulière dans les mêmes conditions. Les plasmas amorcés dans la cavité chaotique sont de forme circulaire, conséquence de l'isotropie du champ électrique et de dimension égale à la dimension de la focalisation par RT, *i.e.* $\lambda_c/2$. Nous notons tout de même la présence de lobes secondaires spatiaux de densité électronique dans la cavité. Leur présence est inhérente à l'amorçage de plasmas par RT mono-voie en cavité réverbérante. En effet, la présence de lobes secondaires spatio-temporels donne forcément naissance à une augmentation locale de la densité électronique au niveau de ces lobes.

3.3.4.2 Vers le dispositif expérimental : Gamme de pression, régime pulsé et initiateur

Nous finirons ce chapitre en utilisant le code de simulation pour présenter les conditions d'amorçage des plasmas par RT présentés dans le chapitre suivant.

Commençons par nous intéresser à l'influence de la pression sur les plasmas amorcés par RT. La densité électronique maximale obtenue par RT n_{ert}^{max} à l'endroit de la focalisation (avec le même niveau maximal de l'amplitude de la réponse impulsionnelle retournée et les mêmes positions récepteur-émetteur que celles de la Figure 3.17) est tracée en fonction de la

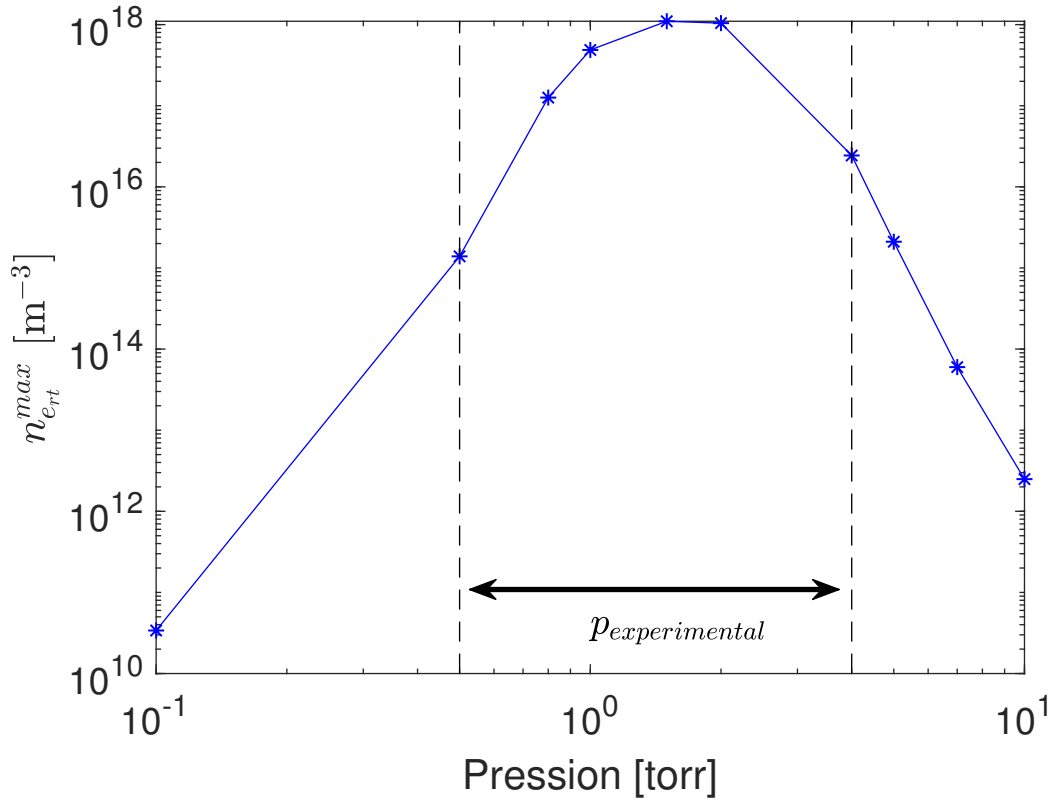


FIGURE 3.20 – Densité maximale $n_{e_{max}}$ mesurée à l’endroit de la focalisation par RT et juste après cette focalisation ($t = 308$ ns) dans la cavité chaotique en fonction de la pression pour le même niveau maximal de l’amplitude de la réponse impulsionnelle retournée que pour la Figure 3.17.

pression pour la cavité chaotique sur la Figure 3.20. On remarque que la densité électronique est maximale pour une pression de 1.5 torr. On retrouve cette valeur de pression optimale pour minimiser le champ électrique de claquage dans des conditions proches [Fel66]. À cette pression, le transfert de puissance de l’énergie microonde aux électrons est maximal. Cela se traduit par une densité électronique plus élevée avec le même niveau de champ électrique. Cela permet de valider la gamme de pression choisie expérimentalement ($\{0.5; 4\}$ torr). En effet, à ces pressions le niveau de champ requis pour amorcer un plasma est minimal. Nous verrons dans le chapitre suivant que nous retrouvons bien expérimentalement la même pression optimale (1.5 torr) à laquelle le transfert de puissance microonde avec le plasma est maximal.

Ensuite, sur le dispositif expérimental, les plasmas sont amorcés et contrôlés en régime pulsé. C’est à dire que le pic de RT est généré avec une certaine fréquence de répétition f_{rep} (typiquement : $f_{rep} = 6$ kHz). Dans ce cas, l’amorçage du plasma par RT peut s’étaler sur plusieurs pics de RT et le niveau du champ électrique nécessaire pour claquer le gaz est potentiellement diminué par rapport au cas mono-pulse. Pour faire du pulsé en simulation, la méthode que nous avons utilisée est la suivante. Un pic de RT est généré en couplant le modèle fluide aux équations de Maxwell. Puis, après 600 ns, seulement le modèle fluide

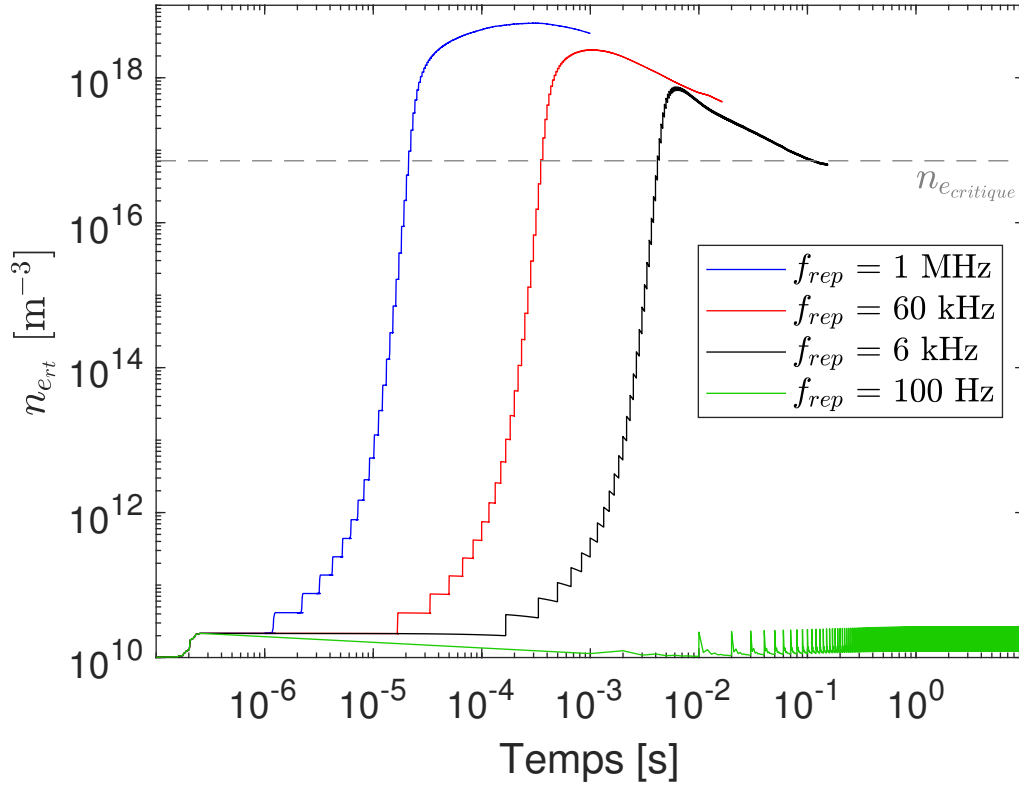


FIGURE 3.21 – Évolution de la densité électronique $n_{e,rt}(t)$ mesurée à l'endroit de la focalisation (avec une pression de 1.5 torr) en fonction du temps en régime pulsé pour plusieurs fréquences de répétitions f_{rep} .

est utilisé pendant un temps $T_{rep} = 1/f_{rep}$. La densité électronique issue du pic de RT va alors diffuser dans la cavité. Ensuite, un nouveau pic de RT est généré en couplant le modèle fluide aux équations de Maxwell. Et ainsi de suite. Par exemple, à une pression de 1.5 torr, les densités électroniques $n_{e,rt}(t)$ mesurées à l'endroit de la focalisation sont tracées pour différentes fréquences de répétition f_{rep} sur la Figure 3.21. L'amplitude du pic de RT est choisie de sorte à ce que la densité obtenue au moment du premier pic de RT soit largement inférieure à 10^{18} m^{-3} , soit égale à $2 \cdot 10^{10} \text{ m}^{-3}$, seulement le double de la densité initiale ($n_{e0} = 10^{10} \text{ m}^{-3}$). Pour une fréquence de répétition de 100 Hz, le second pic de RT est donc généré 10 ms après le premier. Entre ces deux impulsions, la densité électronique a diminué jusqu'à environ la même densité initiale n_{e0} (le temps de décroissance de la densité électronique dépend du gaz et de la pression) et on observe aucun plasma même pour un grand nombre d'impulsion. Pour des fréquences de répétition supérieures, la densité électronique au moment du second pic de RT est toujours supérieure à la densité initiale n_{e0} . Ainsi, la densité augmente encore et suit une courbe en forme d'escalier. Avec le temps, la densité créée à l'endroit du pic de focalisation diffuse et la densité alors présente spatialement autour de ce pic devient de plus en plus importante (sur une distance de l'ordre de la demi longueur d'onde). En augmentant, cette dernière se rapproche de la densité critique $n_{e,critique}$, à partir de laquelle les ondes sont

réfléchies par le plasma. Ainsi de moins en moins d'énergie électromagnétique pénètre jusqu'à l'endroit de la focalisation et la densité $n_{ert}(t)$ passe par un maximum et commence à diminuer. Il semble qu'elle tende ensuite vers la valeur de la densité critique $n_{ecritique}$. Un équilibre se fait entre la quantité d'énergie électromagnétique qui arrive à atteindre l'endroit de la focalisation et celle qui est réfléchi. Le régime permanent est alors atteint. Les diagnostics utilisés dans le chapitre suivant pour étudier le comportement spatio-temporel des plasmas générés par RT expérimentalement ne permettent pas d'étudier le régime transitoire de l'amorçage pulsé (correspondant à la Figure 3.21). L'étude expérimentale est donc concentrée sur l'analyse de ces plasmas pulsés en régime permanent.

Pour finir, et du fait du temps d'application du champ électrique très court (8 ns), il faut noter que le niveau du champ de claquage du gaz est très élevé, de l'ordre de la centaine de kV/m (bien que l'amorçage en régime pulsé permette de réduire ce niveau de champ de claquage comme en atteste la Figure 3.21). Pour contourner cette difficulté, une solution consisterait à utiliser un système de pré-ionisation local (comme cela est par exemple utilisé dans d'autres types de plasmas pour des dispositifs de limitation de puissance microonde avec pour pré-ionisation des décharges DC [Sim18]). Une fois le premier plasma amorcé par RT grâce à cette pré-ionisation, il serait alors possible de le déplacer en focalisant par RT proche de ce plasma. De cette façon le premier plasma joue le rôle de la pré-ionisation pour l'amorçage suivant et on pourrait, de proche en proche déplacer le plasma. Une autre alternative consiste à utiliser un initiateur à l'endroit de la focalisation. Il a la même fonction qu'un système de pré-ionisation. Le principe est cependant différent. Dans ce cas c'est l'intensification locale du champ électrique au niveau de cet initiateur qui permet un amorçage plus facile. Sur le dispositif expérimental, les monopoles métalliques utilisés lors de la phase initiale restent présents lors de la phase inverse. La focalisation a donc lieu sur ces monopoles, qui servent en fait d'initiateurs. Les résultats expérimentaux présentés dans cette thèse se concentrent sur l'étude de ces plasmas amorcés par RT sur les monopoles, comme exposé dans le chapitre suivant.

CHAPITRE 3 : Ce qu'il faut retenir

L'objectif de ce chapitre était d'identifier les paramètres influant sur la qualité du contrôle des plasmas par RT en cavité réverbérante. Pour cela nous avons montré, à travers un outil de simulation, comment le RT se met en place en pratique en environnement réverbérant. Cela nous a ensuite permis de déterminer les paramètres clés influant sur le contrôle des plasmas par RT.

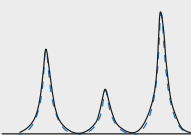
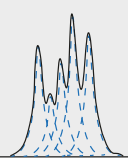
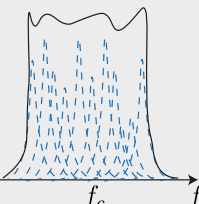
Contrôle des plasmas par RT			
	faible	modéré	fort
Recouvrement modal	 f_c $d \ll 1$	 f_c $d \sim 1$	 f_c $d \gg 1$
$\nearrow \Delta f$	+	+	+
$\nearrow t_H (\Leftrightarrow \nearrow V)$	+	=	=
$\nearrow t_{reverb} (\Leftrightarrow \searrow \text{Pertes})$	=	=	+
$\nearrow$ Chaoticité	+	+	=

FIGURE 3.22 – Synthèse du chapitre : influence des paramètres de la cavité sur le contrôle des plasmas par RT en fonction du régime de recouvrement modal. Le signe “+” signifie que le contrôle est meilleur et le signe “=” signifie que le contrôle est inchangé.

Les résultats de ce chapitre sont synthétisés dans le tableau de la Figure 3.22. Ils sont organisés en fonction du régime de recouvrement modal de la cavité. Tout d'abord, une augmentation de la bande passante se traduira toujours par un meilleur contrôle des plasmas par RT. Ensuite, par exemple pour une cavité excitée en régime de recouvrement faible avec une certaine fréquence centrale f_c , une augmentation du temps de Heisenberg (*i.e.* une augmentation du volume V) se traduira par une augmentation du nombre de modes dans la bande passante et donc un meilleur contrôle des plasmas par RT. Cela sera effectif jusqu'à ce que les modes se recouvrent les uns les autres, *i.e.* jusqu'à ce que d soit égal à 1. Au delà de $d = 1$, une augmentation du volume ne permettra pas d'améliorer le contrôle des plasmas par RT. De façon symétrique, pour une cavité excitée en régime de recouvrement fort avec une certaine fréquence centrale f_c , diminuer les pertes améliore le contrôle des plasmas par RT jusqu'à ce que les modes soient résolus, *i.e.* jusqu'à ce que $d = 1$. Enfin, augmenter la chaoticité de la cavité permet surtout une amélioration significative du contrôle des plasmas par RT en régimes de recouvrement modal faible et modéré.

Les plasmas amorcés par retournement temporel : Résultats expérimentaux

Sommaire

4.1	Le retournement temporel expérimental : Dispositif et méthode . . .	144
4.1.1	Dispositif expérimental	144
4.1.2	Le retournement temporel : Premier essai dans le domaine temporel . . .	148
4.1.3	Le retournement temporel : Développement dans le domaine fréquentiel .	155
4.2	Résultats : Premiers plasmas amorcés par retournement temporel . .	161
4.2.1	Le dispositif expérimental destiné à l'amorçage des plasmas par RT . . .	161
4.2.2	Les propriétés du retournement temporel utilisé pour l'amorçage de plasmas	166
4.2.3	Le contrôle des plasmas par retournement temporel	168
4.3	Caractérisation des plasmas amorcés par retournement temporel . .	175
4.3.1	Dispositif expérimental	175
4.3.2	Caractérisation électrique	177
4.3.3	Caractérisation par imagerie rapide	182
4.3.4	La physique des plasmas amorcés par retournement temporel	187

A completely predictable future is already the past, you've had it. That's not what you want. You want a surprise.

Alan Watts

L'objectif de ce chapitre est de présenter les résultats expérimentaux obtenus durant cette thèse. Nous montrerons ainsi les premiers plasmas amorcés et contrôlés par RT. Nous avons vu au chapitre précédent comment le RT se comporte en environnement réverbérant. Comme nous l'avons alors évoqué, l'utilisation d'une cavité réverbérante permet à la fois de n'utiliser qu'un seul transducteur dans le MRT (RT mono-voie) et de contrôler l'environnement (gaz et pression). Dans le domaine des microondes, les transducteurs utilisés en émission et en réception sont des antennes. Dans ce cas, la première étape du RT consiste à mesurer au niveau

des accès de ces antennes, disposées dans la cavité, la réponse impulsionnelle entre deux antennes. La seconde étape consiste à inverser temporellement cette réponse et à la transmettre à la cavité. Il en résulte une focalisation spatio-temporelle au niveau de l'antenne réceptrice, possédant des lobes secondaires et spatiaux comme expliqué dans le chapitre précédent. Si le niveau du champ électrique de la focalisation est assez élevé, alors un plasma peut être généré et maintenu en régime pulsé (*i.e.* avec un pic de RT répété périodiquement). Notons que la qualité de cette focalisation conditionne le contrôle que l'on va avoir sur les plasmas par RT. La première originalité de cette source plasma pulsée est ainsi l'indépendance de la position des plasmas par rapport au design de la cavité. De plus, les faibles temps de montée (< 1 ns) et les faibles rapport cycliques (typiquement 0.05 %) rendent ces décharges très atypiques.

Dans la première partie de ce chapitre, nous commencerons par décrire le dispositif expérimental utilisé ainsi que la cavité dans laquelle les plasmas sont amorcés par RT. Nous présenterons tout d'abord les premiers essais de RT qui ont été effectués sur ce dispositif au début de la thèse. Ils ont été réalisés en manipulant les signaux dans le domaine temporel comme cela est fait classiquement. Nous présenterons ensuite une méthode dans laquelle les opérations du RT sont réalisées dans le domaine fréquentiel à partir de la mesure des paramètres S de la cavité. Ces derniers sont mesurés à l'aide d'un analyseur de réseau vectoriel (VNA - pour vector network analyser), possédant une dynamique de mesure de 123 dB. Cette précision considérable sur la mesure de l'information de la propagation des ondes entre les antennes permet d'obtenir un meilleur RT. De plus, nous verrons que des études paramétriques sont rendues possibles, permettant notamment de trouver la fréquence centrale optimale (*i.e.* pour laquelle le RT possède les meilleurs contrastes temporels et spatiaux) pour chaque configuration.

Ensuite, nous commencerons la deuxième partie en présentant un exemple de RT obtenu numériquement par méthode fréquentielle. Nous poursuivrons en caractérisant, par méthode fréquentielle, les contrastes du RT expérimental autour de la bande de fréquence utile imposée par les appareils (quelques centaines de MHz autour de 2.4 GHz). La proximité entre le contraste temporel C_t obtenu expérimentalement et celui obtenu en simulation au chapitre précédent dans une cavité chaotique dans des conditions similaires permet de supposer la chaoticité de la cavité. Comme nous l'avons illustré au chapitre précédent, cette propriété est primordiale pour le contrôle des plasmas en cavité. Nous observerons que les résultats mesurés sont bien en accord avec les prédictions pour une cavité chaotique. Enfin, nous montrerons que les propriétés de la focalisation spatio-temporelle obtenue sur notre dispositif permettent bien de contrôler la position des plasmas dans la cavité. Ces plasmas sont amorcés en régime pulsé, c'est à dire que le pic de RT est répété avec une fréquence de répétition f_{rep} . Nous verrons que la position d'un ou même de plusieurs plasmas simultanés est entièrement déterminée par la forme d'onde transmise à la cavité.

Dans la dernière partie de ce chapitre, nous discuterons de la physique de ces plasmas amorcés par RT. Nous nous intéressons d'abord à l'influence de la présence d'un plasma sur le comportement des ondes dans la cavité. Nous verrons que le plasma modifie ce comportement, et que cette modification est dépendante de la pression. Afin de comprendre plus en détails la dynamique spatio-temporelle de ces plasmas, des mesures par imagerie rapide

sont présentées en fonction de la pression et de la fréquence de répétition du pic de RT. Nous verrons que la pression joue un rôle clé dans la dynamique de ces décharges. De plus, l'étude en fonction de la fréquence de répétition illustre l'effet d'une impulsion précédente sur la suivante, classiquement appelé "effet mémoire". Nous finirons par une discussion sur la physique de ces plasmas et à l'aide de calculs simples, nous décrirons les différents mécanismes à l'œuvre.

4.1 Le retournement temporel expérimental : Dispositif et méthode

Dans cette partie, le dispositif expérimental et la méthode fréquentielle développée durant cette thèse pour amorcer des plasmas par RT sont présentés. Les expériences ont lieu dans une cavité réverbérante, permettant la mise en place d'un RT mono-voie et le contrôle de l'environnement (gaz et pression). Les transducteurs utilisés sont alors des antennes métalliques. Les antennes utilisées durant cette thèse pour les expériences d'amorçage de plasmas par RT sont des antennes monopolaires de dimension centimétrique, sauf indication contraire. Le RT consiste alors à mesurer la réponse impulsionnelle entre deux antennes afin de l'inverser temporellement et de la re-émettre pour focaliser sur une antenne. Nous verrons dans cette partie comment le RT se traduit sur notre dispositif expérimental, avant de présenter la méthode fréquentielle qui possède des avantages par rapport à une méthode temporelle, plus classique.

4.1.1 Dispositif expérimental

Cette partie a pour objectif de présenter le dispositif expérimental qui a été mis en place durant cette thèse. Des photos de ce dispositif sont présentées sur la Figure 4.1. Il est schématisé sur la Figure 4.2.

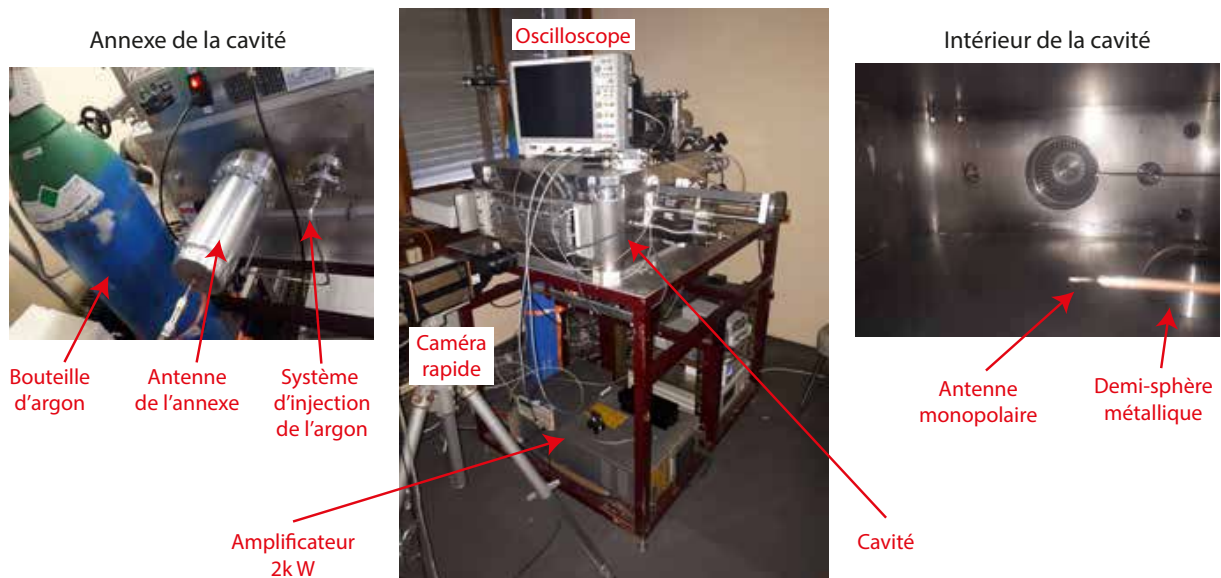


FIGURE 4.1 – Photos du dispositif expérimental.

Le principe permettant d'obtenir le RT en cavité réverbérante a été énoncé précédemment. Il nécessite la génération de signaux de forme adaptée durant la phase de réémission. Nous avons vu dans le chapitre précédent que le RT permet la focalisation sur une dimension de $\lambda_c/2$. On comprend alors qu'il est intéressant de manipuler des signaux hautes fréquences. Cependant les appareils permettant leur génération deviennent vite très onéreux. Nous ver-

rons dans la partie suivante, en nous appuyant sur les travaux de Geoffroy Lerosey [Ler06], comment la technique de modulation IQ (In phase and Quadrature) permet la manipulation de ces signaux haute fréquence avec des appareils courants (parce que largement développés pour les systèmes WiFi et Bluetooth) et donc à moindre coût. On parlera alors de signaux modulés à une fréquence centrale, prise autour de $f_c = 2.4$ GHz, donnant une dimension de focalisation de $\lambda_c/2 = 6$ cm. Cette modulation permet de générer des signaux haute fréquence avec la seule connaissance de l'information de ces signaux en bande de base (basse fréquence). Ainsi un générateur de signaux arbitraires basse fréquence permet de générer des signaux arbitraires haute fréquence par l'intermédiaire d'une porteuse.

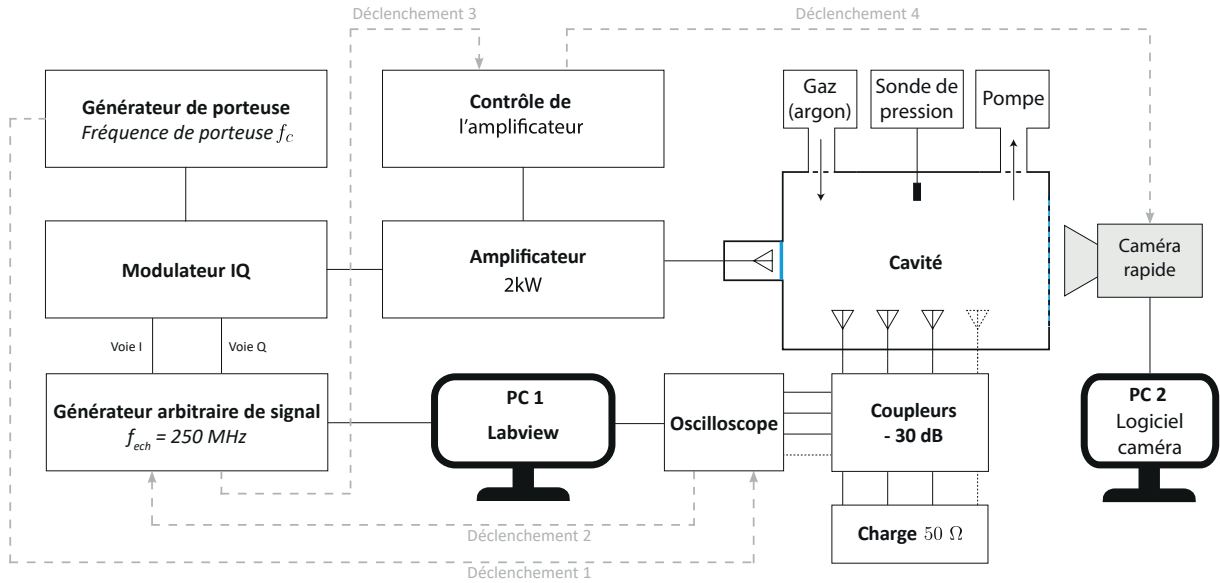


FIGURE 4.2 – Schéma du dispositif expérimental.

Le dispositif expérimental est composé des éléments suivants :

- Le générateur de signaux arbitraires Keysight 33600A qui possède deux voies synchronisées avec une fréquence d'échantillonnage de 250 MHz et une bande passante de 125 MHz. Ces deux voies sont utilisées pour la modulation IQ.
- Le modulateur QM2326A qui fonctionne entre 2.3 et 2.6 GHz. La modulation se fait sur une porteuse externe. Les signaux IQ doivent avoir une amplitude crête maximale de 320 mV, ce qui correspond à une puissance de 0 dBm.
- La porteuse qui est générée par un synthétiseur (IFR MARCONI 2041).
- Un amplificateur qui permet d'augmenter la puissance en entrée de cavité (avec un circulateur ou un système de sécurité permettant de protéger l'amplificateur de la puissance qui sera réfléchi par la cavité). Pour les études du RT à bas niveau, cet amplificateur est de 4 W (Khune KU PA 240270-4B). Pour l'amorçage des plasmas par RT, l'amplificateur utilisé est un amplificateur pulsé de 2 kW à tube à ondes progressives (TMD PTC7353). Il est contrôlé par un générateur de porte temporelle, pilotée par le générateur arbitraire de signal (déclenchement 3).
- L'acquisition des signaux est réalisée avec un oscilloscope Keysight MSO9254A (20

GS/s). Notons aussi que des coupleurs -30 dB sont utilisés entre les antennes et l'oscilloscope quand l'amplificateur utilisé est celui de 2 kW pour limiter la puissance en entrée de l'oscilloscope.

- La cavité dont la forme et les dimensions sont données sur la Figure 4.3.
- Un système de pompage permettant de faire le vide dans la cavité ainsi qu'un système permettant d'injecter de l'argon dans la cavité. Une sonde de pression permet d'assurer l'injection de sorte à obtenir la pression désirée dans la cavité (autour du torr).
- Une caméra rapide (PI-MAX-512) permettant de prendre des images des plasmas dans la cavité à travers un hublot faradisé.

On remarque sur le schéma de la Figure 4.2 la complexité de ce dispositif expérimental. Le générateur arbitraire de signaux et l'oscilloscope sont reliés à l'ordinateur par le réseau du laboratoire et sont contrôlés à l'aide d'un programme Labview. Celui-ci permet de gérer la partie émission et réception du dispositif directement par ordinateur. La reconstruction d'un signal par RT nécessitant un contrôle précis de la phase du signal, les horloges du générateur de porteuse, de l'oscilloscope et du générateur arbitraire de signaux sont donc synchronisées entre elles (déclenchements 1 et 2). De plus lors de l'amorçage des plasmas par RT, le système de contrôle de l'amplificateur 2 kW pulsé est déclenché par le trigger du générateur arbitraire de signal (déclenchement 3) et la caméra rapide est déclenchée par le système de contrôle de l'amplificateur (déclenchement 4).

Les expériences de RT sont réalisées entre les différentes antennes de la cavité, avec une impulsion initiale de $\tau = 8$ ns (durée la plus courte réalisable par le générateur arbitraire), modulée à la fréquence de porteuse f_c comprise dans la bande du modulateur, *i.e.* entre 2.3 et 2.6 GHz (typiquement $f_c = 2.4$ GHz). La bande passante (au sens de largeur à mi-hauteur) de l'impulsion initiale de $\tau = 8$ ns est de $\Delta f = 1/\tau = 125$ MHz et devient après modulation égale à $\Delta f = 2/\tau = 250$ MHz. L'opération de modulation permet ainsi de doubler la bande passante. Cela s'explique par le décalage du spectre du signal autour de la fréquence de porteuse qu'elle opère, comme nous le verrons dans cette partie.

La cavité est constituée d'une partie principale dans laquelle le gaz et sa pression sont contrôlés (permettant de se placer dans les conditions optimales de claquage) et une annexe qui reste à la pression atmosphérique (760 torr) dans l'air. L'antenne utilisée en émission lors de la deuxième phase du RT est placée dans l'annexe de la cavité pour éviter de claquer au niveau de l'antenne d'émission (où le champ électrique est élevé à la réémission). Cette antenne est d'une dimension voisine de $\lambda_c/4$ afin de maximiser la quantité de puissance transmise par cette antenne à la cavité (antenne adaptée à la fréquence centrale f_c en espace libre).

Dans la cavité principale, plusieurs antennes peuvent être positionnées en exploitant différentes traversées coaxiales pour vide. Ces antennes sont nécessaires lors de la première phase du RT pour mesurer l'information de la propagation des ondes entre ces antennes et celle de l'annexe (mesure des réponses impulsives). Lors de la deuxième phase du RT, ces antennes doivent rester présentes pour ne pas modifier le comportement des ondes dans la cavité par rapport à la première phase. Ces antennes ne servent alors qu'à mesurer des signaux, elles n'ont pas pour fonction la transmission de puissance dans la cavité. Elles jouent aussi le rôle

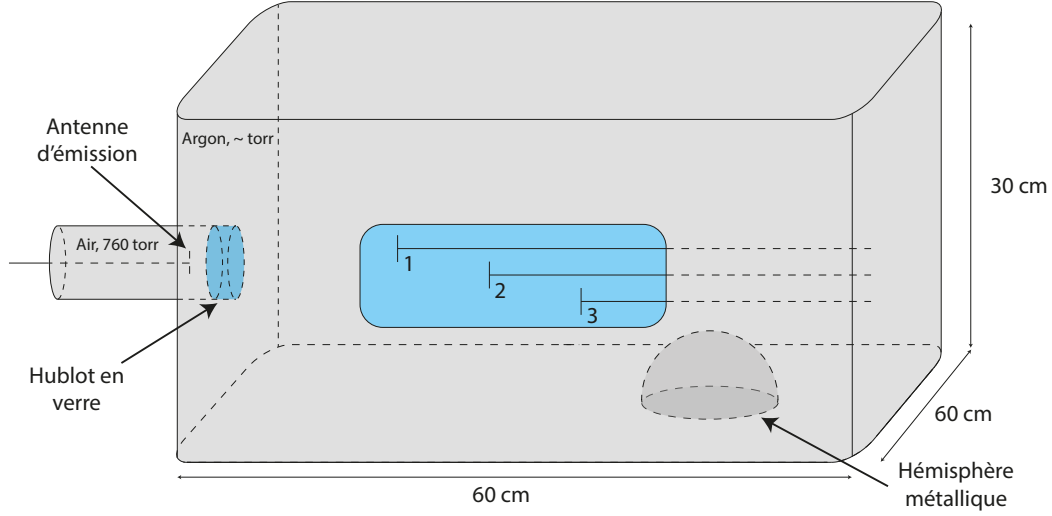


FIGURE 4.3 – Schéma de la cavité.

d'initiateur pour l'amorçage de plasmas par RT¹. Ainsi, il peut être intéressant de choisir des antennes complètement désadaptées, afin qu'elles ne puisent pas beaucoup d'énergie de la cavité. Dit autrement, la présence de ces antennes dans la cavité ne doit pas diminuer significativement le facteur de qualité $Q(f)$ de celle-ci, *i.e.* le rapport entre l'énergie stockée dans la cavité sur l'énergie dissipée. Ce dernier peut s'écrire (sans ouverture et sans absorbant dans la cavité) [Val16] ; [Hil98] :

$$\frac{1}{Q(f)} = \frac{1}{Q_{\text{parois}}(f)} + \frac{N_{\text{Aréception}}}{Q_{\text{ant}}(f)} \quad (4.1)$$

avec $Q_{\text{parois}}(f)$ le facteur de qualité correspondant aux pertes aux parois et $Q_{\text{ant}}(f)$ celui correspondant à la puissance collectée par les antennes réceptrices (au nombre de $N_{\text{Aréception}}$). D'après cette expression, le facteur de qualité d'une cavité résulte de la compétition entre deux phénomènes de pertes : les parois et les antennes réceptrices. Celui de ces deux phénomènes dont les pertes seront les plus importantes (et donc à plus faible facteur de qualité) dominera l'autre et le facteur de qualité de la cavité $Q(f)$ sera alors limité par lui. Le facteur de qualité des antennes réceptrices $Q_{\text{ant}}(f)$ est proportionnel à l'inverse du facteur de désadaptation d'impédance [Val16] ; [Hil98]. Plus les antennes réceptrices sont désadaptées, plus $Q_{\text{ant}}(f)$ est grand et donc moins il contribue au facteur de qualité $Q(f)$. La configuration typique avec laquelle les plasmas ont été amorcés par RT durant cette thèse est représentée sur la Figure 4.3. La dimension des trois antennes ($N_{\text{Aréception}} = 3$) de type monopole de la cavité principale est alors prise égale à $\lambda_c/10$, de sorte à ce que les antennes soient désadaptées à la fréquence f_c . Le facteur de qualité des antennes $Q_{\text{ant}}(f)$ étant proportionnel à la fréquence au cube,

1. Notons que pour la prochaine étape du développement du pinceau plasma ondulatoire, il faudra décrocher le plasma de ces initiateurs. L'utilisation d'une sonde non intrusive (type électro-optique) peut ainsi être envisagée pour l'acquisition des signaux utiles. Nous reviendrons sur ce point dans les perspectives de cette thèse.

à haute fréquence sa contribution au facteur de qualité $Q(f)$ est donc négligeable [Val16]; [Hil98]. Dans les cavités utilisées dans [Val16], c'est le cas aussi à basse fréquence. Cependant, utiliser des antennes désadaptées en réception ne peut avoir qu'un effet positif sur le facteur de qualité $Q(f)$ de la cavité.

De plus, notons aussi qu'une antenne se comporte différemment en émission et en réception [SHK97]; [Kun07]. Une discussion plus approfondie sur cette différence est proposée en annexe B. Pour les expériences réalisées durant cette thèse nous considérerons être dans le cas de signaux à bande étroite ($\Delta f/f_c \ll 1$) et négligerons donc cette différence de comportement. De cette façon nous pouvons nous attendre, en émission, à ce que le champ électromagnétique rayonné suive bien l'évolution du signal transmis au niveau du câble coaxial et, en réception, que le signal mesuré au niveau du câble coaxial soit une image fidèle du champ électromagnétique arrivant sur l'antenne. Notons quand même que ces réactions différentes des antennes en émission et en réception (même si cet effet est négligé ici) ne fait qu'éloigner encore plus la réalité de la vision du "film passé à l'envers" lors d'un RT. Nous avons déjà vu au chapitre précédant que l'émission dans toutes les directions lors de la phase inverse et la présence d'un environnement réverbérant ont pour conséquence la superposition du film passé à l'envers avec d'autres, non présents lors de la phase initiale.

Enfin, nous avons illustré au chapitre précédent l'importance de la chaotité sur le contrôle des plasmas par RT, particulièrement en régime modal faible et modéré. Nous verrons dans la deuxième section de ce chapitre que les expériences menées durant cette thèse sont faites avec un régime modal $d_{exp} = 3.5$. Dans ces conditions, nous avons montré au chapitre précédent que la chaotité de la cavité est indispensable pour avoir un bon contrôle des plasmas par RT. Il paraît donc important de rendre la cavité expérimentale chaotique. Notons tout d'abord que le simple fait d'introduire des antennes dans la cavité brise certaines symétries et rend la cavité plus chaotique [St9]. De plus, comme on peut le voir sur la photo de l'intérieur de la cavité de la Figure 4.1, la cavité expérimentale possède des aspérités réparties de façon irrégulière sur ses parois verticales. Afin de rendre la cavité plus chaotique, il reste seulement à supprimer les surfaces horizontales en regard. Pour ce faire, une demi-sphère métallique est placée sur la surface basse horizontale [Sel14]; [Gro14], comme illustré sur la Figure 4.3.

4.1.2 Le retournement temporel : Premier essai dans le domaine temporel

Nous allons dans cette partie réaliser le RT comme il est classiquement mis en œuvre expérimentalement dans la littérature [Ler06]. Les résultats présentés dans cette partie constituent les premiers résultats d'expériences de RT obtenus au début de la thèse. L'annexe de la cavité n'était alors pas présente et les expériences de RT étaient effectuées entre des antennes placées à différents endroits de la cavité (séparées d'une distance supérieure à $\lambda_c/2$). Ces antennes étaient des dipôles électriques de dimension $\lambda_c/2$.

L'objectif de cette partie est purement pédagogique. Tout d'abord, les résultats obtenus par RT classique, *i.e.* avec les développements dans le domaine temporel, sont présentés afin de suivre les signaux expérimentaux à chaque étape du RT. De plus, le tracé à chaque étape de

ces signaux dans le domaine fréquentiel permet d'enchaîner facilement sur les développements de la partie suivante : le retournement temporel dans le domaine fréquentiel. Nous verrons ainsi l'intérêt de ce dernier.

Le protocole mis en œuvre est le suivant, adapté de [Ler06] :

1 - Phase d'acquisition : la réponse impulsionnelle de la cavité mesurée sur l'antenne réceptrice est enregistrée.

2 - Phase de réémission : cette réponse impulsionnelle est retournée temporellement puis réémise au niveau de l'antenne émettrice.

Tout au long de cette étude, les signaux étudiés seront présentés dans les domaines temporel et fréquentiel. Une impulsion de 8 ns est envoyée sur la voie I (en phase) du modulateur, correspondant à l'impulsion la plus courte possible avec notre équipement (Figure 4.4(a)). Un signal nul est simultanément envoyé sur la voie Q (en quadrature).

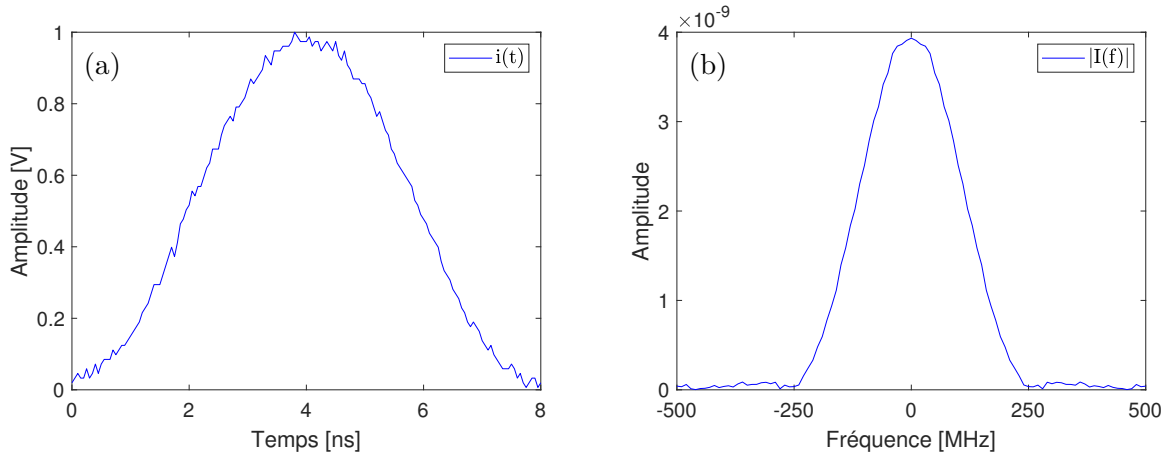


FIGURE 4.4 – (a) Impulsion de 8 ns mesurée à l'oscilloscope en sortie du générateur arbitraire de signaux et (b) son spectre obtenu par la fonction FFT de Matlab.

Cette impulsion $i(t)$ peut se modéliser de la façon suivante :

$$\begin{aligned} i(t) &= A_i(1 - \cos(2\pi f_i t)) \text{ pour } 0 \leq t \leq t_i \\ i(t) &= 0 \text{ pour } t > t_i \end{aligned} \quad (4.2)$$

t_i étant la durée de l'impulsion (8 ns ici). Dans ce cas-là : $A_i = 1/2$ et $f_i = 1/t_i = 125\text{MHz}$.

La norme de la transformée de Fourier de cette impulsion, correspondant au spectre du signal, est reportée sur la Figure 4.4(b). Cette transformée est obtenue en traitant le signal temporel sous Matlab avec la fonction FFT (Fast Fourier Transform). Toutes les représentations fréquentielles seront obtenues de cette manière sauf contre-indication. Le spectre obtenu a pour amplitude $A_i t_i$ et pour bande passante (largeur à mi-hauteur) f_i . Cette information se retrouve dans la transformée de Fourier du signal $i(t)$:

$$\mathcal{I}(f) = TF[i(t)] = At_i \left[\text{sinc}(t_i f) + \frac{1}{2} [\text{sinc}(t_i(f_i - f)) + \text{sinc}(t_i(f_i + f))] \right] \quad (4.3)$$

Le signal $i(t)$ est transféré dans la mémoire du générateur arbitraire de signal avant d'être envoyé sur la voie I. La modulation utilisée étant de type IQ (In phase and Quadrature), le signal de sortie du modulateur peut s'écrire sous la forme :

$$e(t) = i(t)\cos(2\pi f_c t) + q(t)\sin(2\pi f_c t) \quad (4.4)$$

avec $i(t)$ et $q(t)$ appelées signaux en bande de base. La transformée de Fourier $E(f)$ d'un signal obtenu par modulation IQ s'obtient simplement en connaissant les transformées de Fourier des signaux en bande de base. On note $\mathcal{I}(f)$ et $\mathcal{Q}(f)$ ces transformées :

$$i(t) \xrightarrow{TF} \mathcal{I}(f) \text{ et } q(t) \xrightarrow{TF} \mathcal{Q}(f) \quad (4.5)$$

D'après le théorème de Plancherel, la transformée de Fourier d'un produit simple est un produit de convolution [Cot]. Un signal obtenu par modulation IQ avec une fréquence de porteuse f_c , peut ainsi s'écrire dans le domaine fréquentiel de la façon suivante :

$$E(f) = \frac{\mathcal{I}(f - f_c) + \mathcal{I}(f + f_c)}{2} + j \frac{\mathcal{Q}(f - f_c) - \mathcal{Q}(f + f_c)}{2} \quad (4.6)$$

Dans notre cas $i(t)$ correspond à l'impulsion et $q(t)$ est nul. Le signal en sortie du modulateur s'écrit donc :

$$e(t) = i(t)\cos(2\pi f_c t) = A(1 - \cos(2\pi f_i t))\cos(2\pi f_p t) \quad (4.7)$$

Sa transformée de Fourier est ainsi donnée par :

$$E(f) = \frac{\mathcal{I}(f - f_c) + \mathcal{I}(f + f_c)}{2} \quad (4.8)$$

Ce signal $e(t)$ correspond à une impulsion de 8 ns modulée à la fréquence porteuse $f_c = 2.4$ GHz (Figure 4.5(a)).

D'après l'équation 4.8, la transformée de Fourier de l'impulsion modulée aura la même forme que la transformée de l'impulsion mais centrée autour de f_c et $-f_c$ avec l'amplitude divisée par deux. On obtient ainsi le spectre de la Figure 4.5(b). On remarque que celui-ci est non nul pour des fréquences nulles, ce qui est dû à un défaut du modulateur réel qui introduit une composante continue. De plus, la partie du spectre initial qui était dans les fréquences négatives se retrouve alors décalée du côté des fréquences positives. La modulation

a donc pour conséquence de doubler la bande passante du signal initial. La bande passante de l'impulsion modulée devient alors : $\Delta f = 2f_i = 250$ MHz.

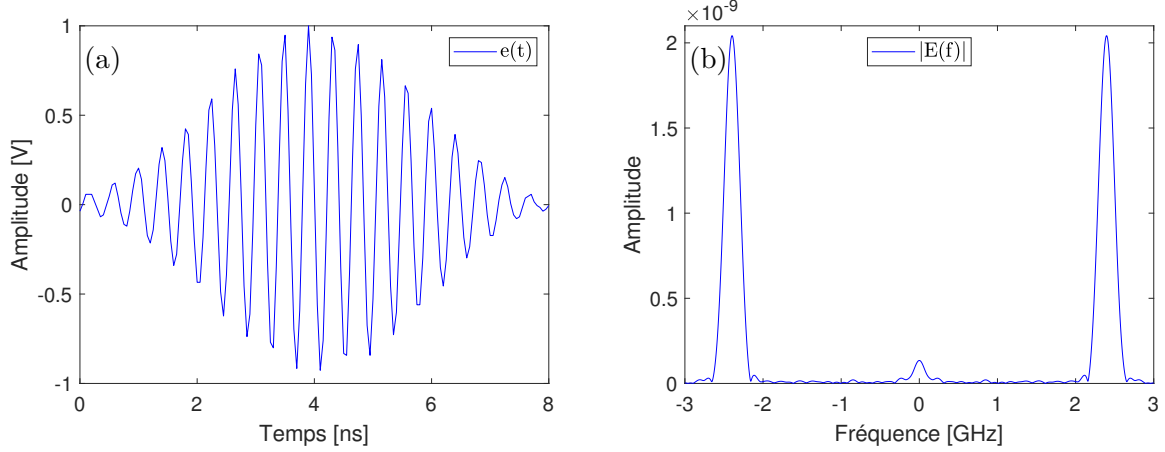


FIGURE 4.5 – (a) Impulsion de 8 ns modulée par une fréquence $f_c = 2.4$ GHz mesurée à l'oscilloscope en sortie du générateur arbitraire de signaux et (b) son spectre obtenu par la fonction FFT de Matlab.

Ce signal est ensuite amplifié avant d'être transmis dans la cavité par l'intermédiaire d'une antenne, comme décrit sur la Figure 4.2. L'objectif ici n'étant pas d'amorcer des plasmas mais seulement de présenter les signaux du RT mesurés sur le dispositif expérimental, un amplificateur de 4 W est utilisé. Nous verrons dans la section qui suit que pour atteindre le seuil de claquage du gaz, un amplificateur de 2 kW a été utilisé. Le champ ondulatoire ainsi généré va être réfléchi dans la cavité, et une partie de ce champ va être capté par une seconde antenne. Le signal reçu est enregistré par l'oscilloscope, puis récupéré sur l'ordinateur. Pour la suite du développement, toutes les amplitudes des signaux temporels et fréquentiels seront normalisées. En effet, le but de cette partie est de détailler, étape par étape, le principe du RT, et le niveau des signaux obtenus après être passé dans la cavité ne nous intéresse pas à ce stade.

La Figure 4.6(a) montre la réponse impulsionnelle mesurée avec l'antenne réceptrice. L'origine des temps a été choisie arbitrairement. Plusieurs choses sont à noter sur la forme de cette réponse impulsionnelle. Comme nous l'avons vu en simulation dans le chapitre précédent, la réponse impulsionnelle s'organise en lobes. Ces lobes sont le résultat des réflexions des ondes à l'intérieur de la cavité. Le premier lobe correspond à un trajet direct des ondes entre les deux antennes, et le signal émis initialement est donc reçu directement par l'antenne réceptrice (à un retard près de l'ordre de 2 ns). Les autres lobes sont des échos du signal d'origine, obtenus après un certain nombre de réflexions des ondes sur les parois de la cavité, ce qui allonge le retard. On remarque que l'amplitude maximale de ces lobes semble suivre une décroissance exponentielle. Si on s'intéresse de plus près à l'évolution temporelle du niveau de l'amplitude maximale des lobes, on remarque que certains lobes ont une amplitude maximale supérieure à des lobes qui les précèdent. Pour comprendre l'origine de ce phénomène, deux explications sont possibles.

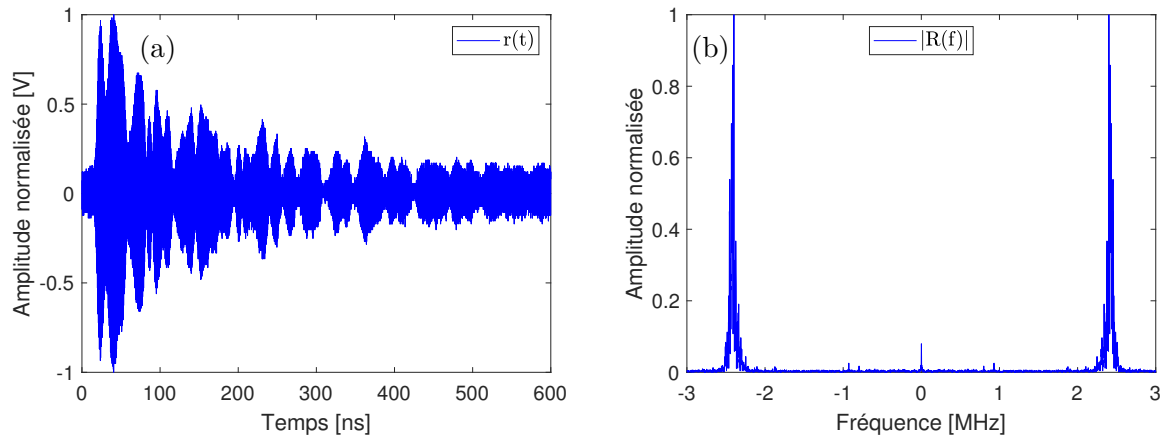


FIGURE 4.6 – (a) Réponse impulsionnelle mesurée sur l'antenne réceptrice à l'oscilloscope et (b) son spectre obtenu par la fonction FFT de Matlab.

Tout d'abord la durée de l'impulsion initiale (8 ns) peut être plus longue que la durée du trajet que met une onde entre deux antennes. En effet, les ondes mettent ici environ 2 ns pour arriver d'une antenne à l'autre par le trajet direct. Les ondes réfléchies par les parois vont forcément mettre un temps supérieur à 2 ns pour arriver d'une antenne à l'autre, mais il est possible qu'elles mettent un temps inférieur à 8 ns. Ainsi il est tout à fait possible que des échos de cette impulsion se chevauchent, c'est-à-dire que les lobes obtenus sur la réponse impulsionnelle ne soient pas discriminés. Ce chevauchement peut tout aussi bien se traduire par une interférence constructive ou destructive, impliquant dans ce dernier cas une diminution de l'amplitude apparente des lobes.

Il existe une autre façon d'expliquer ce phénomène. Sans parler de chevauchement de deux lobes voisins, l'amplitude d'un lobe dépend du nombre de réflexions que font les ondes correspondantes dans la cavité. Et il est tout à fait possible (géométriquement parlant) que certaines ondes arrivent plus rapidement au niveau de l'antenne réceptrice en ayant fait plus de réflexions que d'autres. Ce qui résulte en l'obtention d'un lobe d'amplitude maximale inférieure à d'autres lobes qui arrivent plus tardivement.

Ces lobes sont obtenus sur une durée maximale d'environ 600 ns, soit 70 fois la durée de l'impulsion initiale. La durée de cette réponse impulsionnelle permet de qualifier le temps pendant lequel les ondes sont réfléchies dans la cavité, et donc le temps durant lequel on a de l'information sur cette réponse. L'amplitude de la réponse impulsionnelle suit une décroissance exponentielle. Afin de caractériser la cavité, il est intéressant d'utiliser le temps de réverbération t_{reverb} dont nous avons déjà discuté au chapitre précédent. Nous reviendrons sur l'estimation de ce temps de réverbération dans la section suivante.

Au-delà de 600 ns, le signal obtenu n'a pas pour origine la réflexion des ondes générées par l'impulsion initiale, mais il correspond plutôt à du bruit généré par le modulateur, dont le signal de sortie n'est pas nul lorsqu'il reçoit des signaux nuls en entrée sur les voies I et Q. On remarque un rapport signal sur bruit (RSB) de 15 dB. Ce rapport signal sur bruit

est dépendant de l'étage d'amplification choisi (ici 4 W) et de la sensibilité de l'oscilloscope utilisé. Nous verrons dans la partie suivante que dans le domaine fréquentiel, l'utilisation d'un VNA permet la mesure de l'information de la propagation des ondes entre deux antennes avec un RSB de 123 dB dans notre cas.

On s'intéresse maintenant au spectre de cette réponse impulsionnelle, tracé sur la Figure 4.6(b). Comme nous l'avons vu dans le chapitre précédent, en notant $h(t)$ la réponse de notre système à une impulsion de Dirac $\delta(t)$ et $e(t)$ le signal en entrée, le signal de sortie sera obtenu par l'équation [Cot] :

$$r(t) = e(t) * h(t) \quad (4.9)$$

et donc, par le théorème de Plancherel, dans le domaine spectral on aura :

$$R(f) = E(f).H(f) \quad (4.10)$$

avec $r(t) \xrightarrow{TF} R(f)$ et $h(t) \xrightarrow{TF} H(f)$. $H(f)$ est la fonction de transfert de notre système antennes/cavité. Sur le spectre de la réponse impulsionnelle de la Figure 4.6(b), son amplitude a des valeurs non nulles pour des fréquences auxquelles le spectre du signal d'entrée (Figure 4.5(b)) avait lui aussi des valeurs non nulles. Ceci est cohérent puisque le spectre de la réponse impulsionnelle résulte du produit des spectres du signal d'entrée et de la fonction de transfert du système. Les pics observés sur le spectre de la réponse impulsionnelle correspondent aux modes de résonance de l'ensemble antennes/cavité, c'est-à-dire aux fréquences auxquelles le système va pouvoir être excité.

Si l'on transmet dans un milieu une onde électromagnétique modulée en phase et en quadrature de phase, on peut toujours exprimer en n'importe quel point du milieu, et en particulier au niveau de l'antenne réceptrice, le signal reçu $r(t)$ sous la forme [Ler06] :

$$r(t) = r_I(t)\cos(2\pi f_c t) + r_Q(t)\sin(2\pi f_c t) \quad (4.11)$$

Pour faire du RT, il faut réémettre le signal en changeant les sens du temps, ce qui veut dire changer le signe de t . Le signal à émettre est donc :

$$r(-t) = r_I(-t)\cos(2\pi f_c t) - r_Q(-t)\sin(2\pi f_c t) \quad (4.12)$$

L'équation 4.12 (empruntée à [Ler06]) montre que pour retourner temporellement une réponse impulsionnelle $r(t)$, la seule connaissance des signaux en bande de base permet de réaliser l'opération. Ainsi, les signaux reçus en réponse à une impulsion contiennent toute l'information de la propagation de l'onde dans le milieu. Pour cela : 1. On démodule numériquement sous Matlab la réponse impulsionnelle en phase et quadrature de phase. Le signal est d'abord multiplié par un cosinus et un sinus à la fréquence de la porteuse f_c et un filtre passe bas numérique est ensuite appliqué aux deux signaux issus de la précédente opération afin d'obtenir l'information de la réponse impulsionnelle en bande de base. 2. Ces signaux sont ensuite inversés temporellement et on multiplie par -1 celui en quadrature de phase. Un

exemple de résultat est reporté sur la Figure 4.7. La durée de ces signaux en bande de base est bien la même que la durée sur laquelle on a enregistré la réponse impulsionnelle (600 ns).

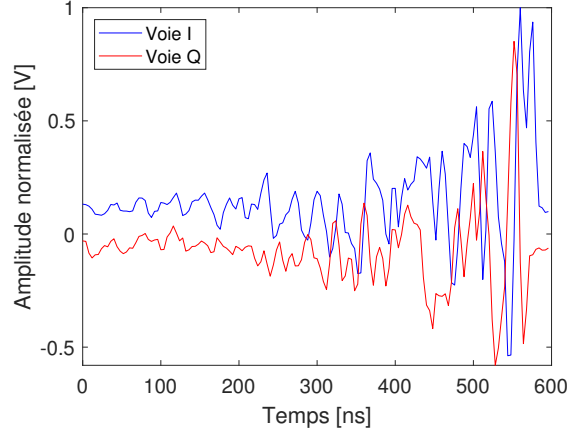


FIGURE 4.7 – Signaux numériques obtenus par démodulation numérique de la réponse impulsionnelle retournée temporellement.

On échantillonne ensuite ces signaux à 250 MHz, à savoir la fréquence d'échantillonnage du générateur arbitraire de signaux. Les signaux échantillonnés sont alors normalisés par le maximum d'amplitude en valeur absolue et chargés dans la mémoire du générateur. Le signal issu de la démodulation en phase est envoyé sur la voie I et celui de la démodulation en quadrature de phase sur la voie Q. La modulation de type IQ est ainsi réalisée à la fréquence de porteuse f_c . Le signal est ensuite amplifié et transmis dans la cavité par l'intermédiaire de la même antenne que celle qui avait émis l'impulsion initiale (exploitation de la réciprocité).

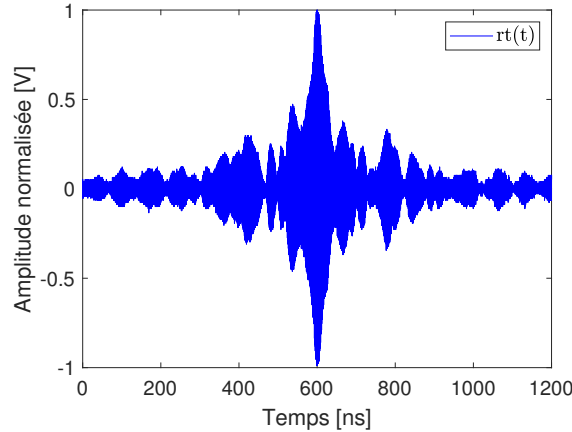


FIGURE 4.8 – Signal de RT obtenu sur l'antenne réceptrice mesuré à l'oscilloscope.

Le signal de RT $rt(t)$ obtenu au niveau de la deuxième antenne est tracé sur la Figure 4.8. Ce signal peut être défini comme le produit de convolution suivant :

$$rt(t) = r(-t) * h(t) \quad (4.13)$$

On observe bien une croissance exponentielle de l'enveloppe du signal avant le pic de RT, qui correspond à la phase où les ondes se réorganisent de façon à recréer l'impulsion suivant les mécanismes décrits dans le chapitre précédent. Puis, une fois que l'on arrête d'émettre (aux alentours de 600 ns), une décroissance exponentielle qui correspond à la phase où les ondes qui ont généré le pic de RT continuent de se propager et de s'amortir dans la cavité.

Dans cette partie, nous avons décrit le protocole permettant de faire du RT classique, *i.e.* avec les développements dans le domaine temporel. Il a ainsi été montré que la modulation IQ permettait l'utilisation d'ondes très hautes fréquences à moindre coût. Nous avons donc été capables d'obtenir une focalisation dans l'espace et dans le temps en utilisant cette méthode de RT.

La suite de l'étude va consister à optimiser la puissance injectée localement par RT. En effet, plusieurs limitations sont présentes tout au long de cette démarche et il convient de les identifier afin de connaître les limites de notre dispositif expérimental mais aussi, si possible, de les contourner. Premièrement, la durée de l'impulsion initiale est limitée à 8 ns par la fréquence d'échantillonnage du générateur arbitraire de signaux. Le principe de RT consistant à recréer cette impulsion, il n'est donc pas possible d'obtenir un pic principal de retournement plus court. Ensuite, le bruit généré par le modulateur peut altérer la qualité de la réponse impulsionnelle (voir Figure 4.6) ce qui va influencer sur la qualité de l'information contenue dans cette dernière et donc sur la qualité du RT. De plus, le RSB obtenu est conditionné par le bruit de tous les appareils du dispositif expérimental et il est trouvé égal à 15 dB avec l'amplificateur 4 W. Lors de l'acquisition de la réponse impulsionnelle à l'oscilloscope, la qualité de l'information recueillie peut donc être légèrement dégradée. Toutes ces contraintes nous ont amené à imaginer une autre méthode originale pour faire du RT. Tous les calculs utilisés dans cette partie ont été fait dans le domaine temporel, et pourtant, la qualité des mesures dans le domaine fréquentiel permet de pressentir que des choses sont à explorer dans ce domaine pour lequel les équations sont déjà disponibles. C'est pourquoi, dans une optique d'optimisation et de meilleur contrôle du RT, il a été décidé de travailler sur l'analyse fréquentielle.

4.1.3 Le retournement temporel : Développement dans le domaine fréquentiel

Nous allons dans cette partie présenter la méthode nécessaire à la mise en place du RT par analyse fréquentielle. Ici, seulement les équations sont présentées. Nous verrons dans la partie suivante les signaux ainsi obtenus, illustrant le gain apporté par cette méthode.

4.1.3.1 Principe

L'objectif de cette sous-partie consiste à décrire analytiquement et à mettre en place expérimentalement cette méthode alternative basée sur l'analyse fréquentielle. Généralement, les expériences de RT sont menées dans le domaine temporel [Ler06]. Cette partie a pour but

d'expliquer analytiquement cette méthode avant d'en démontrer les avantages.

Comme expliqué dans la partie précédente, l'idée de l'utilisation de l'analyse fréquentielle est venue de l'observation des spectres des signaux à chaque étape du RT et en particulier de l'équation 4.10. Cette équation indique que le signal fréquentiel correspondant à une réponse de la cavité s'obtient par le produit du spectre en entrée de cavité et de la fonction de transfert de la cavité $H(f)$. La cavité étant un système linéaire, son comportement peut être étudié en utilisant sa matrice $[S]$, ou matrice de répartition, dépendant d'un choix d'accès. Le RT nécessite seulement l'utilisation de deux accès de la cavité, cette matrice $[S]$ sera donc d'ordre 2. Cette matrice est composée d'un paramètre S_{11} , traduisant la réflexion de l'antenne en entrée de cavité, de deux paramètres S_{21} et S_{12} égaux, caractérisant la transmission réciproque entre les deux antennes, et d'un paramètre S_{22} traduisant la réflexion de l'antenne à la sortie de la cavité. Ainsi les modes de l'ensemble antennes/cavité correspondent aux pics relevés sur le paramètre $|S_{21}(f)|$. Ce paramètre, comme l'ensemble des paramètres S , peut être mesuré à l'aide d'un VNA. Il permet de qualifier la capacité que va avoir l'énergie à être transmise à travers la cavité. Dans ces conditions, la fonction de transfert de la cavité $H(f)$ est directement donnée par $H(f) = S_{21}(f)$.

Afin d'obtenir la réponse impulsionnelle de la cavité en se basant sur la mesure de ses paramètres S , il suffira donc d'effectuer le produit de la transformée de Fourier de l'impulsion par le paramètre $S_{21}(f)$. Pour cela, il y a tout d'abord une petite subtilité d'ordre mathématique à relever. En effet, tous les signaux étudiés dans la partie 4.1.2 possèdent un spectre pair (la phase de ces signaux est quant à elle impaire). L'apparition de fréquences négatives dans la transformée de Fourier des signaux étudiés possède une origine purement mathématique. Ces fréquences négatives ne correspondent pas à un phénomène physique, mais elles sont juste un moyen de représenter les signaux dans le domaine fréquentiel. En utilisant les propriétés de la transformée de Fourier et du produit de convolution, il est facile de montrer que la transformée de Fourier d'un signal réel $x(t)$ possède la symétrie hermitienne, c'est-à-dire : $X^*(-f) = X(f)$, avec $x(t) \xrightarrow{TF} X(f)$ [Cot], d'où les propriétés de parité des spectres et d'imparité des phases. Ainsi, pour obtenir par analyse fréquentielle le signal fréquentiel correspondant à la réponse à une impulsion, il faut réaliser les opérations suivantes :

1. Mesurer au VNA le paramètre S_{21} de l'ensemble antennes/cavité sur l'intervalle fréquentiel utile (c'est-à-dire entre 2.2 et 2.6 GHz).
2. Multiplier la transformée de Fourier de l'impulsion par $S_{21}(f)$ sur les fréquences positives et par $S_{21}^*(f)$ sur les fréquences négatives.

Le signal fréquentiel obtenu possède la symétrie hermitienne et correspond donc bien à un signal temporel réel. En pratique, le paramètre $S_{21}(f)$ est mesuré au VNA, et on crée une fonction de transfert $S_{21_M}(f)$ sous Matlab en appliquant les opérations suivantes :

$$\begin{aligned} S_{21_M}(f) &= S_{21}^*(f) \text{ pour } f < 0 \\ S_{21_M}(f) &= S_{21}(f) \text{ pour } f \geq 0 \end{aligned} \tag{4.14}$$

De cette façon, le signal fréquentiel $R(f)$ résultant du produit de $S_{21_M}(f)$ avec la transformée de Fourier du signal en entrée $E(f)$ de la cavité correspond bien à un signal temporel réel.

L'analyse fréquentielle permet au final d'obtenir la réponse de la cavité à partir de la mesure du paramètre $S_{21}(f)$ de la cavité. On peut donc utiliser ce principe pour obtenir du RT alternativement au RT classique décrit dans la partie 4.1.2.

4.1.3.2 Mise en équation du RT par analyse fréquentielle

Dans cette sous-partie nous allons pouvoir utiliser l'analyse fréquentielle pour obtenir le signal de RT numériquement uniquement à partir de la mesure des paramètres S . Comme nous le verrons par la suite, cette technique permet d'étudier facilement l'influence de certains paramètres sur le RT. La réponse fréquentielle $R(f)$ de la cavité à un signal en entrée qui a pour transformée de Fourier $E(f)$ s'obtient de la façon suivante :

$$R(f) = E(f) \cdot S_{21_M}(f) = |E(f)| \cdot |S_{21_M}(f)| \cdot e^{j(\phi_E(f) + \phi_{S_{21_M}}(f))} = |R(f)| \cdot e^{j\phi_R(f)} \quad (4.15)$$

avec $\phi_E(f)$ et $\phi_{S_{21_M}}(f)$ respectivement la phase du signal d'entrée et la phase du paramètre $S_{21_M}(f)$, $\phi_R(f) = \phi_E(f) + \phi_{S_{21_M}}(f)$ la phase de la réponse de la cavité et $|R(f)| = |E(f)| \cdot |S_{21_M}(f)|$ la norme de la réponse de la cavité. Pour obtenir le RT, la réponse de la cavité est inversée temporellement, ce qui revient à inverser le signe de la phase du signal fréquentiel. Le signal à envoyer dans la cavité est donc :

$$S_{RT}(f) = |R(f)| \cdot e^{j(-\phi_R(f))} = |E(f)| \cdot |S_{21_M}(f)| \cdot e^{j(-\phi_E(f) - \phi_{S_{21_M}}(f))} \quad (4.16)$$

Ce signal est ensuite envoyé dans la cavité. Le signal fréquentiel du RT est donc :

$$\begin{aligned} RT(f) &= S_{RT}(f) \cdot S_{21_M}(f) = |R(f)| \cdot |S_{21_M}(f)| \cdot e^{j(-\phi_R(f) + \phi_{S_{21_M}}(f))} \\ &= |E(f)| \cdot |S_{21_M}(f)|^2 \cdot e^{j(-\phi_E(f))} \end{aligned} \quad (4.17)$$

Nous avons donc à présent notre signal de RT dans le domaine fréquentiel. Il suffit maintenant de lui appliquer la transformée de Fourier inverse pour obtenir ce signal en temporel :

$$RT(f) \xrightarrow{TF^{-1}} rt(t) \quad (4.18)$$

Ainsi, juste avec la mesure du paramètre $S_{21}(f)$, on peut effectuer les différentes étapes du RT et en particulier connaître les caractéristiques du pic de RT par le calcul numérique.

4.1.3.3 Génération des signaux s_{ri} et s_{rq} par analyse fréquentielle

Nous allons dans cette sous-partie décrire les développements permettant à partir de l'analyse fréquentielle de générer le RT dans la cavité. Il faut pour cela construire les signaux que l'on doit envoyer sur les voies I et Q du modulateur. La méthode consiste à démoduler en phase (pour la voie I) et en quadrature de phase (pour la Q), la réponse de la cavité inversée temporellement $S_{RT}(f)$ étant :

$$S_{RT}(f) = |R(f)| \cdot e^{j(-\phi_R(f))} \quad (4.19)$$

Démoduler un signal dans le domaine temporel revient, en phase à multiplier celui-ci par un cosinus à la fréquence de la porteuse et en quadrature de phase à multiplier celui-ci par un sinus à la fréquence de la porteuse. Ainsi, dans le domaine fréquentiel, les signaux démodulés en phase $S_{RI_D}(f)$ et en quadrature de phase $S_{RQ_D}(f)$ sont données par les équations :

$$S_{RI_D}(f) = S_{RT}(f) * \frac{\delta(f + f_c) + \delta(f - f_c)}{2} = \frac{S_{RT}(f + f_c) + S_{RT}(f - f_c)}{2} \quad (4.20)$$

$$S_{RQ_D}(f) = S_{RT}(f) * j \frac{\delta(f + f_c) - \delta(f - f_c)}{2} = j \frac{S_{RT}(f + f_c) - S_{RT}(f - f_c)}{2} \quad (4.21)$$

Afin de récupérer l'information en bande de base, un filtre passe bande idéal d'une valeur de bande passante $2\Delta f$ est appliqué à ces signaux. En traitement numérique, ceci se traduit par l'application de la fonction indicatrice sur les signaux précédents :

$$S_{RI_{BB}}(f) = 1_{[-\Delta f; \Delta f]}(f) \cdot (S_{RI_D}(f)) \quad (4.22)$$

$$S_{RQ_{BB}}(f) = 1_{[-\Delta f; \Delta f]}(f) \cdot (S_{RQ_D}(f)) \quad (4.23)$$

Avec :

$$1_U(w) = \begin{cases} 1, & \text{si } w \in U. \\ 0, & \text{si } w \notin U. \end{cases} \quad (4.24)$$

Nous avons donc à présent nos deux signaux dans le domaine fréquentiel en bande de base, l'un en phase $S_{RI_{BB}}(f)$ et l'autre en quadrature de phase $S_{RQ_{BB}}(f)$. Il suffit maintenant de leur appliquer la transformée de Fourier inverse pour obtenir ces signaux en temporel :

$$S_{RI_{BB}}(f) \xrightarrow{TF^{-1}} s_{ri}(t) \quad (4.25)$$

$$S_{RQ_{BB}}(f) \xrightarrow{TF^{-1}} s_{rq}(t) \quad (4.26)$$

Ces signaux sont ensuite échantillonnés à la fréquence d'échantillonnage du générateur arbitraire de signal, à savoir 250 MS/s et envoyés sur les voies I et Q. Ainsi la méthode

fréquentielle permet aussi d'obtenir du RT expérimentalement dans la cavité.

Notons qu'il est aussi possible d'obtenir du RT numériquement à partir de ces signaux $s_{ri}(t)$ et $s_{rq}(t)$. Pour cela, il faut tout d'abord les moduler numériquement puis sommer les deux signaux modulés ainsi obtenus (cette étape faite numériquement équivaut à envoyer les signaux $s_{ri}(t)$ et $s_{rq}(t)$ sur le modulateur en expérimental). Il suffit ensuite de calculer la transformée de Fourier du signal ainsi obtenu et de la multiplier à la fonction $S_{21_M}(f)$. Il ne reste plus qu'à appliquer la transformée de Fourier inverse pour obtenir le signal de RT.

Les intérêts de cette méthode, par rapport à la méthode classique décrite partie 4.1.2 sont les suivants :

- Tout d'abord cette méthode permet la mesure de l'information de la propagation des ondes entre deux antennes avec une précision bien supérieure à celle obtenue dans le domaine temporel. En effet, le RSB sur la réponse impulsionnelle mesurée à l'oscilloscope est de 15 dB avec un amplificateur 4 W. Le VNA, lui, mesure l'information dans le domaine fréquentiel avec un RSB de 123 dB. Même si l'amplificateur 2 kW (nécessaire pour l'amorçage de plasmas par RT sur notre dispositif) est utilisé pour amplifier l'impulsion initiale et que l'on suppose un cas idéalisé dans lequel il n'apporte qu'une augmentation du niveau du signal par rapport au bruit, le RSB serait alors de $15 + 10 \log_{10}(2 \text{ kW}/4 \text{ W}) = 42 \text{ dB}$. En pratique le RSB avec cet amplificateur sera bien inférieur à cette valeur puisque le bruit des appareils en amont de l'amplificateur sera lui aussi amplifié. Finalement la méthode fréquentielle permet une mesure de l'information simple et très précise.
- Cette méthode permet un gain de temps pour générer les signaux $s_{ri}(t)$ et $s_{rq}(t)$ à envoyer sur les voies I et Q du modulateur pour obtenir le RT. En effet il n'y a pas besoin d'enregistrer à l'oscilloscope la réponse impulsionnelle avant de la retourner et de la démoduler, il suffit de lancer les codes Matlab en donnant en entrée la fréquence de la porteuse et la durée de l'impulsion initiale. Il est tout à fait possible de se passer de l'utilisation du dispositif expérimental pour étudier le RT (tant qu'on a la mesure du paramètre S_{21}), puisqu'il est possible de l'étudier numériquement. Ainsi l'influence de certains paramètres sur la qualité du RT peut être évaluée uniquement numériquement, et cela permet un gain de temps énorme et une bien plus grande facilité de mise en œuvre.
- Le RT obtenu numériquement permet de s'affranchir des limitations dues au dispositif expérimental et permet ainsi d'obtenir des résultats pour une durée de l'impulsion initiale inférieure à 8 ns et pour une fréquence de porteuse non limitée par la bande passante du modulateur. Cela permet ainsi, dans l'optique d'éventuelles optimisations, de prévoir l'influence de dispositifs plus performants ou avec des bandes passantes différentes sur le RT.
- Afin de caractériser la cavité (largeur spectrale des modes, régime de recouvrement modal...) il est intéressant de pouvoir tracer sa fonction de transfert $H(f)$. Dans la méthode du RT développée dans le domaine temporel, l'information de la fonction de transfert de la cavité est contenue dans la réponse impulsionnelle à travers un produit de convolution : $r(t) = e(t) * h(t)$. Pour estimer cette fonction de transfert, il est alors plus

simple de passer dans le domaine fréquentiel : $H(f) = R(f)/E(f)$. Cependant, cette réponse impulsionnelle $R(f)$ est soumise aux performances des appareils de génération et de mesure des signaux microondes et va donc influencer sur la qualité de la fonction de transfert estimée. Avec la méthode de RT dans le domaine fréquentiel, la fonction de transfert est directement accessible (mesure des paramètres S) avec une très grande précision.

Tout cela prouve la pertinence de la méthode fréquentielle pour obtenir le RT.

La méthode fréquentielle de RT ayant désormais été mise en place, nous allons maintenant nous en servir pour caractériser la cavité et les propriétés spatio-temporelle du pic de RT. Nous verrons ensuite que ces propriétés permettent l'amorçage et le contrôle de plasmas par RT.

4.2 Résultats : Premiers plasmas amorcés par retournement temporel

L'objectif de cette partie est de présenter les plasmas amorcés par RT sur notre dispositif expérimental. Tout d'abord, la cavité dans laquelle les expériences d'amorçage de plasma ont été réalisées est présentée. L'analyse fréquentielle détaillée dans la partie précédente est utilisée ici pour tracer les signaux sur un exemple de RT dans cette cavité. Ensuite, cette méthode est utilisée pour caractériser les propriétés spatio-temporelles du RT sur le dispositif expérimental. Nous verrons enfin les résultats montrant les premiers plasmas amorcés et contrôlés par RT.

4.2.1 Le dispositif expérimental destiné à l'amorçage des plasmas par RT

La cavité utilisée pour les expériences d'amorçage de plasmas par RT est schématisée sur la Figure 4.9. Une antenne monopolaire (numérotée 4) de dimension $\lambda_c/4$ est localisée dans une annexe de la cavité, séparée par un hublot de la cavité principale. Dans cette dernière, trois antennes monopoles (numérotées 1, 2 et 3) de dimension $\lambda_c/10$ et distantes de $\lambda_c/2$ sont positionnées. Le volume de la cavité est égal à $V = 0.6 \times 0.6 \times 0.3 \text{ m}^3$. Le temps de Heisenberg correspondant est alors $t_{H_{exp}} = 8\pi V f^2 / c^3 \approx 600 \text{ ns}$ (équation 3.6) pour une fréquence de 2.4 GHz.

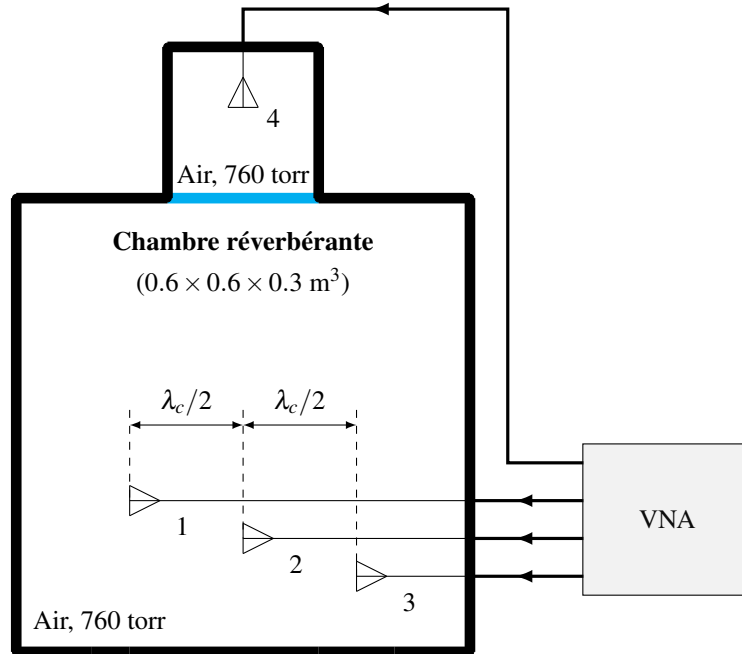


FIGURE 4.9 – Schéma de la cavité avec trois antennes dans la cavité principale distantes de $\lambda_c/2$. Les paramètres S sont mesurés au VNA entre les quatre antennes (celle de l'annexe et celles de la cavité principale).

Les paramètres S sont mesurés entre ces quatre antennes simultanément en reliant chaque antenne à un port du VNA (Keysight E5071C), qui en possède quatre, comme représenté sur la Figure 4.9 (la cavité principale est laissée à la pression atmosphérique, *i.e.* 760 torr, dans l'air). Comme expliqué dans la partie précédente, cette opération est nécessaire pour faire du RT par analyse fréquentielle. Elle permet ensuite d'obtenir numériquement la réponse impulsionnelle de la cavité (première étape du RT) et de réaliser des études paramétriques.

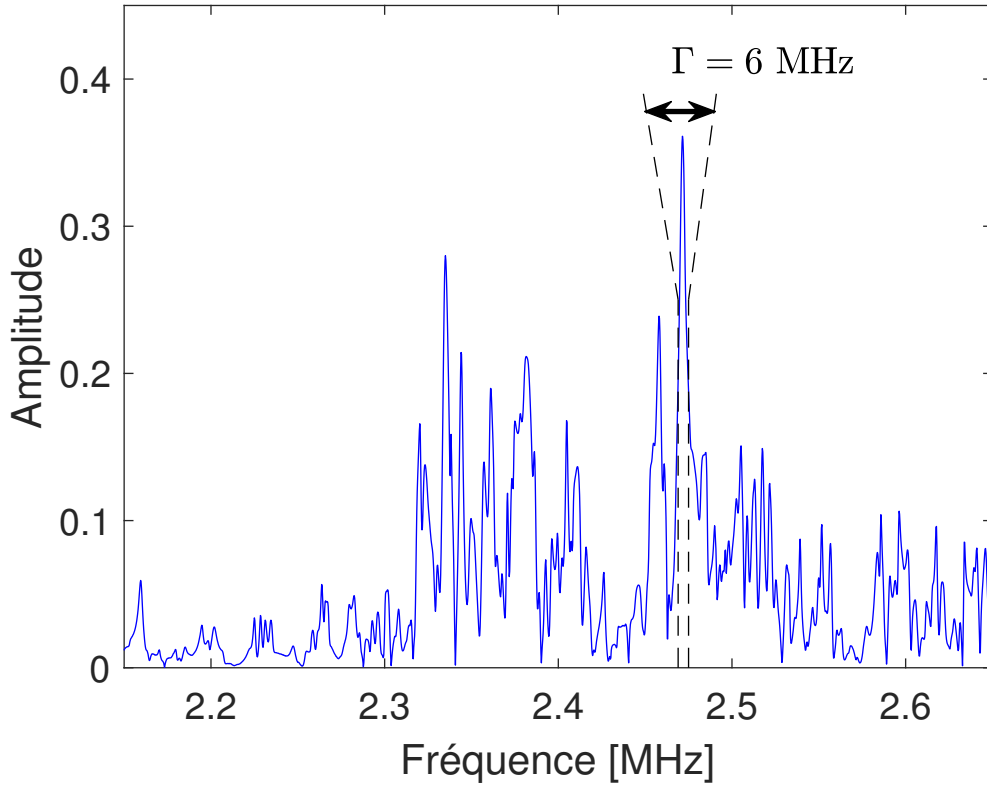


FIGURE 4.10 – $|S_{41}(f)|$ mesuré au VNA.

L'amplitude du paramètre $S_{41}(f)$ mesuré entre l'antenne de l'annexe et l'antenne n°1 dans la cavité est tracé à titre d'exemple sur la Figure 4.10 dans la bande passante sur laquelle les expériences ont été réalisées (500 MHz autour de 2.4 GHz). Le comportement en fréquence de la cavité s'organise en modes, autour desquels sont présents des pics d'une largeur spectrale d'environ $\Gamma = 6$ MHz (les paramètres des simulations du chapitre précédent ont été réglés de sorte à ce que les modes des cavités simulées possèdent justement la même largeur spectrale).

À partir des paramètres S mesurés au VNA, comme expliqué dans la partie précédente, il est possible d'obtenir numériquement par analyse fréquentielle la réponse impulsionnelle de la cavité à n'importe quelle impulsion (dont la bande passante est contenue dans la bande de fréquences mesurée des paramètres S). Par exemple, entre l'antenne de l'annexe de la cavité (n°4) et l'antenne n°1, il suffit de multiplier la transformée de Fourier de l'impulsion avec le paramètre S_{41_M} construit à partir du paramètre S_{41} (équations 4.15 et 4.14). Pour

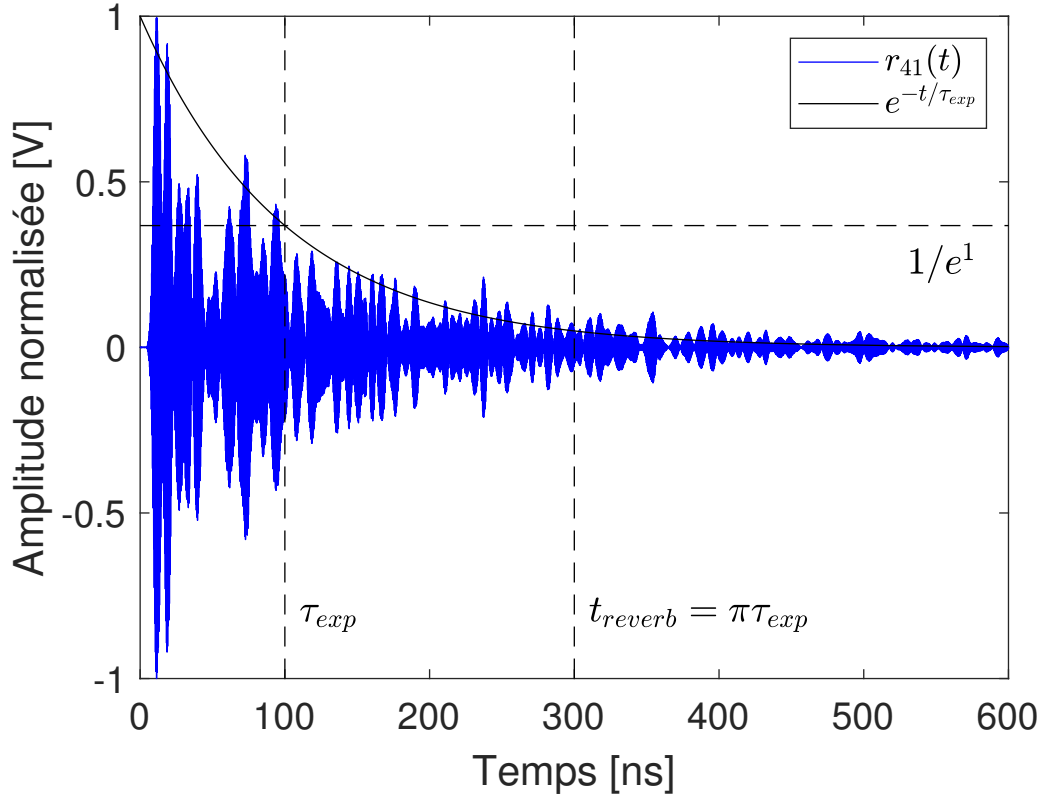


FIGURE 4.11 – Exemple de réponse impulsionnelle obtenue par analyse fréquentielle.

une impulsion de 8 ns modulée à $f_c = 2.4$ GHz, la réponse impulsionnelle $r_{41}(t)$ entre ces antennes est tracée sur la Figure 4.11. On retrouve bien une organisation en lobes, dont l'amplitude suit une décroissance exponentielle, traduisant l'amortissement des ondes dans la cavité. Nous pouvons remarquer ici que le niveau de bruit a considérablement diminué comparé aux mesures faites dans le domaine temporel et présentées dans la section précédente (Figure 4.6(a)). Cela illustre bien la meilleure précision avec laquelle la méthode fréquentielle mesure l'information. Le temps caractéristique de cette décroissance est obtenu en utilisant la formule 3.16 en moyennant sur les trois réponses obtenues entre l'antenne n°4 et les trois antennes de la cavité principale. Il est trouvé égal à $\tau_{exp} = 100$ ns. Le temps de réverbération, calculé par l'équation 3.16, correspond alors bien à l'inverse de la largeur spectrale des modes du paramètre $S_{41}(f)$: $t_{reverb} = \pi\tau_{exp} = 2/\Gamma = 300$ ns (le facteur 2 est une conséquence de la modulation). La largeur spectrale des modes et donc la durée de la réponse impulsionnelle des cavités simulées présentées au chapitre précédent sont bien les mêmes que celles mesurées sur la cavité expérimentale.

Nous avons discuté dans le chapitre précédent de l'importance de la chaotité de la cavité sur la qualité de la focalisation par RT. Nous avons expliqué pourquoi la chaotité est surtout importante à recouvrement modal faible ou modéré. En suivant l'équation 3.7 du recouvrement modal en fonction du volume de la cavité et de la fréquence d'excitation, le recouvrement

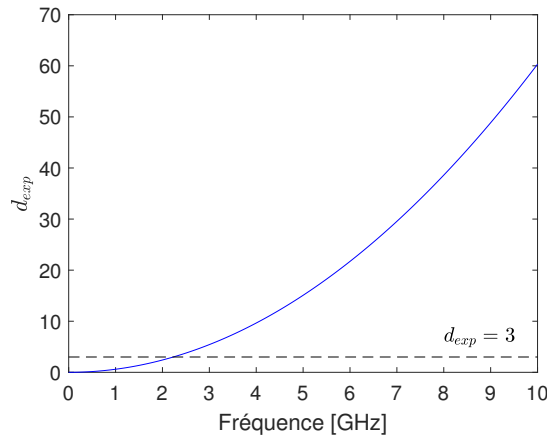


FIGURE 4.12 – d_{exp} en fonction de la fréquence.

modal expérimental $d_{exp}(f)$ est tracé en fonction de la fréquence sur la Figure 4.12. Selon le critère de Schroeder, le régime de recouvrement fort est atteint pour une fréquence de 2.2 GHz ($d_{exp} = 3$ [Sch87]). Pour une fréquence de 2.4 GHz, le recouvrement modal est de $d_{exp} = 3.5$ (information que l'on retrouve bien en comparant le temps de réverbération et le temps de Heisenberg : $t_{reverb} < t_H$). Nous avons montré dans le chapitre précédent que dans les cavités 2D simulées avec les mêmes régimes de recouvrement la chaotité avait encore une influence significative sur les propriétés d'uniformité et d'isotropie du champ dans la cavité et donc sur les propriétés du RT. Nous pouvons alors penser que cela sera aussi le cas pour la cavité expérimentale (en restant prudent sur la validité de cette analogie 2D - 3D). La chaotité de la cavité paraît alors importante. De toute façon, même en régime de recouvrement fort, la chaotité de la cavité ne peut qu'avoir un effet positif sur la qualité de la focalisation par RT. Pour qu'une cavité soit chaotique, il faut supprimer les surfaces parallèles en regard [Sel14] ; [Gro14]. La cavité expérimentale possède des aspérités réparties de façon irrégulière sur ses parois verticales (voir Figure 4.1). Pour supprimer la symétrie entre les surfaces horizontales, une demi-sphère métallique est placée sur la surface basse horizontale [Sel14] ; [Gro14], comme montré sur la photo de la Figure 4.1 et illustré sur la Figure 4.3. De plus, comme nous l'avons déjà évoqué, la présence d'antennes dans la cavité supprime déjà certaines symétries [St9]. Enfin, nous pouvons noter par exemple qu'en bande X (autour de 10 GHz), le régime de recouvrement modal atteint des valeurs de l'ordre de 60. À ces fréquences là, la chaotité de la cavité expérimentale a sûrement peu d'influence sur la qualité de la focalisation par RT.

La réponse impulsionnelle $r_{41}(t)$ de la Figure 4.11 est ensuite retournée temporellement et s'écrit donc $r_{41}(-t)$. Les réponses de la cavité entre l'émission de cette dernière au niveau de l'antenne n°4 et les trois antennes de la cavité principale sont finalement obtenues par analyse fréquentielle. Pour cela, la transformée de Fourier de la réponse impulsionnelle retournée $r_{41}(-t)$ est multipliée avec les paramètres $S_{41_M}(f)$, $S_{42_M}(f)$ et $S_{43_M}(f)$ (construits respectivement à partir des paramètres $S_{41}(f)$, $S_{42}(f)$ et $S_{43}(f)$). Les signaux ainsi obtenus sont tracés sur la Figure 4.13 dans le domaine temporel. On retrouve sur ces trois signaux la croissance exponentielle de l'amplitude des lobes pendant les 600 premières nanosecondes, correspondant à la durée de l'émission de la réponse impulsionnelle retournée. Le signal me-

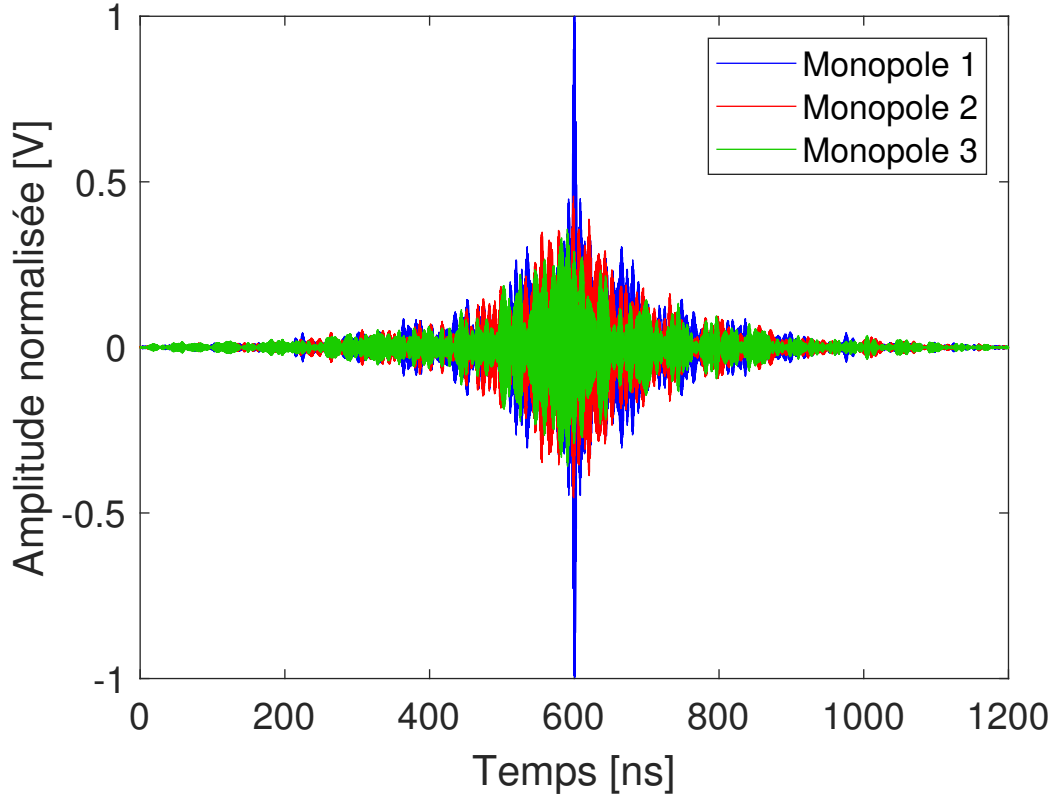


FIGURE 4.13 – Exemple de signaux obtenus lors d’une expérience de RT par analyse fréquentielle.

suré sur le monopole 1 (en bleu) correspond au signal de RT $rt_{41}(t)$. Au bout de 600 ns, les ondes se somment en phase au temps 600 ns au niveau du monopole 1. Ainsi, un pic de focalisation de durée égale à la durée de l’impulsion initiale (8 ns) apparaît. Sur les deux autres monopoles, les modes ne sont pas sommés en phase et il n’y a pas de pic de focalisation. Après ce pic, les ondes continuent de se propager dans la cavité et leur amplitude suit une décroissance exponentielle à cause des réflexions sur les parois de la cavité.

On observe sur la Figure 4.13 un contraste temporel C_t^{max} et un contraste spatial C_{st}^{max} (défini ici comme le rapport entre le niveau du pic de RT sur le monopole 1 et le niveau maximum du lobe obtenu sur les deux autres monopoles) similaires. Ces contrastes sont obtenus ici en considérant seulement trois positions dans la cavité (les trois antennes) et uniquement une seule opération de RT (sur l’antenne n°1 avec une fréquence centrale $f_c = 2.4$ GHz). Ils ne sont donc pas forcément caractéristiques de la cavité. La partie suivante se concentre sur la caractérisation de ces contrastes temporels et spatiaux expérimentaux.

4.2.2 Les propriétés du retournement temporel utilisé pour l'amorçage de plasmas

Comme nous l'avons évoqué précédemment, l'un des avantages que possède la méthode fréquentielle par rapport à la méthode temporelle est la possibilité de faire des études paramétriques en s'affranchissant des limites inhérentes aux appareils du dispositif expérimental. Nous allons ici tout d'abord tracer l'évolution du contraste temporel C_t en fonction de la bande passante $C_t(\Delta f)$ et en fonction de la fréquence de porteuse $C_t(f_c)$. À partir des paramètres $S_{41_M}(f)$, $S_{42_M}(f)$ et $S_{43_M}(f)$ il est possible d'étudier numériquement par analyse fréquentielle les propriétés de la focalisation par RT en fonction de certains paramètres. Dans l'étude qui suit, les contrastes temporels C_t et C_t^{max} sont obtenus en moyennant les trois contrastes temporels obtenus sur les trois antennes de la cavité principale (lors d'une focalisation par RT entre l'antenne n°4 et chaque antenne de la cavité).

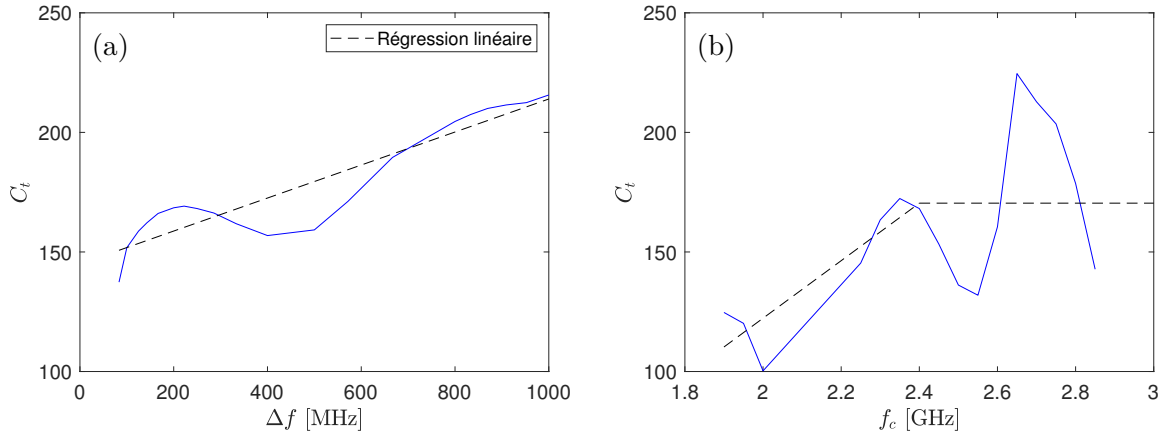


FIGURE 4.14 – Évolution du contraste temporel C_t (a) en fonction de la bande passante Δf pour une fréquence centrale $f_c = 2.4$ GHz (b) en fonction de la fréquence centrale f_c pour une bande passante $\Delta f = 250$ MHz.

L'évolution du contraste temporel C_t pour un RT avec une impulsion initiale modulée par une fréquence centrale $f_c = 2.4$ GHz obtenu par analyse fréquentielle est tracée en fonction de la bande passante Δf sur la Figure 4.14(a). La durée de l'impulsion initiale est donnée par $2/\Delta f$. Ce contraste temporel évolue bien de façon proportionnelle avec la bande passante, comme attendu par l'expression du contraste temporel développée dans le chapitre précédent (équation 3.14) et reporté par exemple dans [Ler06]. En effet, une augmentation de la bande passante se traduit directement par une augmentation de la quantité d'information décorrélée, que les modes soient résolus ou pas.

Ensuite l'évolution du contraste temporel C_t pour un RT avec une impulsion de 8 ns modulée (bande passante $\Delta f = 250$ MHz) obtenu par analyse fréquentielle est tracée en fonction de la fréquence centrale f_c sur la Figure 4.14(b). Ce contraste semble suivre une augmentation avec la fréquence centrale f_c jusqu'à une fréquence de $f_c = 2.4$ GHz avant d'osciller autour d'un contraste constant d'environ 170. Ce comportement est cohérent avec

la tendance prévue par la théorie développée dans le chapitre précédent (voir Figure 3.8) : le contraste temporel C_t augmente avec la fréquence centrale puisque le recouvrement modal augmente, jusqu'à ce que ce dernier soit très élevé ($d \gg 1$) correspondant à un recouvrement modal fort et donc une saturation du contraste temporel. Nous pouvons cependant remarquer des forts écarts autour de la valeur de 170, avec un minimum de 125 vers 2.6 GHz et un maximum de plus de 200 vers 2.7 GHz. Ceci vient sûrement du faible nombre de positions (les trois antennes n°1, 2 et 3) sur lesquelles est fait la moyenne pour évaluer C_t . Une évaluation de ce contraste en connaissant l'évolution du champ électromagnétique partout dans la cavité permettrait peut être de diminuer ces oscillations. Néanmoins, dans le cas qui nous intéresse dans cette partie, *i.e.* l'amorçage et le contrôle des plasmas sur les initiateurs (les trois antennes n°1, 2 et 3) nous pouvons nous rendre compte que la méthode fréquentielle permet, par un réglage fin de la fréquence centrale, d'optimiser le contraste.

Ces résultats illustrent l'intérêt d'une étude paramétrique par analyse fréquentielle. Pour un système antennes/cavité donné, il suffit de mesurer les paramètres S (sur une large bande de fréquence) pour obtenir ensuite numériquement les évolutions des contrastes en fonction de certains paramètres tels que la bande passante ou la fréquence d'excitation. On peut ensuite en déduire quelle serait la meilleure option pour améliorer la focalisation par RT. Par exemple, à partir de la Figure 4.14(a), on peut quantifier l'augmentation du contraste temporel avec la bande passante. Sur la 4.14(b), si le comportement du contraste temporel suit la tendance pour les plus hautes fréquences, on pourrait par exemple en conclure que travailler en bande X n'apporte pas grand chose sur le contraste temporel (pour s'en assurer il faudrait tout de même faire des mesures complémentaires des paramètres S jusqu'en bande X).

Dans les conditions expérimentales de ces travaux, *i.e.* pour une bande passante $\Delta f = 250$ MHz et une fréquence de porteuse $f_c = 2.4$ GHz, le contraste temporel est de 168. Dans des conditions similaires (même bande passante et durée de la fenêtre temporelle simulée égale au temps de réverbération expérimental), en simulation FDTD 2D le contraste temporel C_t obtenu est de 150 pour une cavité régulière et de 171 pour une cavité chaotique (voir Figure 3.14(a)). Le contraste temporel C_t expérimental est donc très proche de celui de la cavité chaotique simulée. Sachant que la formule du contraste temporel utilisée dans le chapitre précédent (équation 3.14) est valable pour les cavités 2D et 3D, mais en restant prudent sur la validité de l'analogie 2D (simulation) - 3D (expérimental), nous pourrions en conclure que la cavité expérimentale semble être chaotique (plus exactement de degré de chaoticité similaire à celui de la cavité chaotique simulée du chapitre précédent). De plus, pour les conditions typiques d'amorçage de plasmas par RT, le contraste temporel C_t^{max} obtenu expérimentalement est de 2.4. Pour les cavités simulées (régulière et chaotique) ce dernier était de 3 (voir Figure 3.14(b)). Pour ce contraste temporel C_t^{max} il n'y a pas de tendance prévue par la théorie. Il n'y a donc pas de raison à priori d'obtenir le même contraste expérimentalement et en simulation (dans les mêmes conditions). On remarque que le contraste temporel C_t^{max} est un peu inférieur expérimentalement par rapport à celui obtenu en simulation. La raison est peut être l'organisation différente de l'énergie électromagnétique dans une cavité 3D par rapport à une 2D. De toute façon, nous verrons dans la partie suivante que ce contraste est suffisant pour contrôler les plasmas par RT. Pour finir, le contraste spatio-temporel expérimental C_{st}^{max} est obtenu (avec une bande passante $\Delta f = 250$ MHz et une fréquence de porteuse $f_c = 2.4$

GHz) en moyennant les rapports entre le niveau du pic de RT obtenu sur une antenne et le niveau maximal mesuré sur chacune des deux autres. Il est trouvé égal à $C_{st}^{max} = 2.6$, soit très proche du contraste temporel C_t^{max} . Cela semble cohérent avec la propriété de chaoticité supposée de la cavité.

Propriétés du dispositif expérimental	
Fréquence centrale f_c	≈ 2.4 GHz
Bande passante Δf	250 MHz
Volume de la cavité V	$0.6 \times 0.6 \times 0.3$ m ³
Contraste temporel C_t	168
Contraste temporel C_t^{max}	2.4
Contraste spatio-temporel C_{st}^{max}	2.6

Les propriétés de la cavité expérimentale et celles de la focalisation par RT obtenues dans les conditions typiques d’amorçage de plasmas par RT sont reportées dans le tableau ci-dessus. Ces résultats paraissent cohérents avec ceux obtenus dans le chapitre précédent en simulation dans la cavité chaotique 2D. Notons tout de même que les contrastes présentés dans cette sous-partie, issus des paramètres S , correspondent uniquement à ceux du système antennes/cavité. Cela permet de caractériser ce système avec une très grande précision (123 dB de dynamique). En pratique chaque appareil du dispositif expérimental utilisé lors des expériences d’amorçage de plasmas par RT augmente le niveau de bruit et induit potentiellement une altération des signaux, ce qui se traduit donc par des contrastes mesurés sur le dispositif expérimental inférieurs à ceux issus uniquement des paramètres S . Cependant, nous allons voir dans la partie suivante que cette diminution n’est pas significative et que les propriétés spatio-temporelles de la focalisation par RT permettent bien d’amorcer et de contrôler des plasmas par RT dans la cavité.

Dans les résultats présentés dans la suite de ce chapitre, pour chaque nouvelle configuration (nombre et positions des antennes dans la cavité principale, position de la demi-sphère métallique...), une pré-étude par analyse fréquentielle a été effectuée sur la bande de fréquence du dispositif expérimental (entre 2.3 et 2.6 GHz) afin de déterminer la fréquence centrale f_c pour laquelle le RT possède les meilleurs contrastes (temporels et spatiaux). Afin de ne pas surcharger ce manuscrit, ces études ne sont pas montrées. Ainsi les expériences de RT sont effectuées avec cette fréquence centrale optimale, issue de cette pré-étude en fréquence et différente pour chaque configuration.

4.2.3 Le contrôle des plasmas par retournement temporel

L’objectif de cette partie est de présenter les premiers plasmas amorcés par RT expérimentalement. Sur ces plasmas, nous resterons descriptif, le but étant ici d’illustrer le contrôle de leur position par RT. Nous discuterons plus en détail dans la section suivante de l’explication physique des phénomènes ayant lieu lors de ces amorçages par RT. Les résultats présentés dans cette partie correspondent aux premiers plasmas amorcés par RT et font

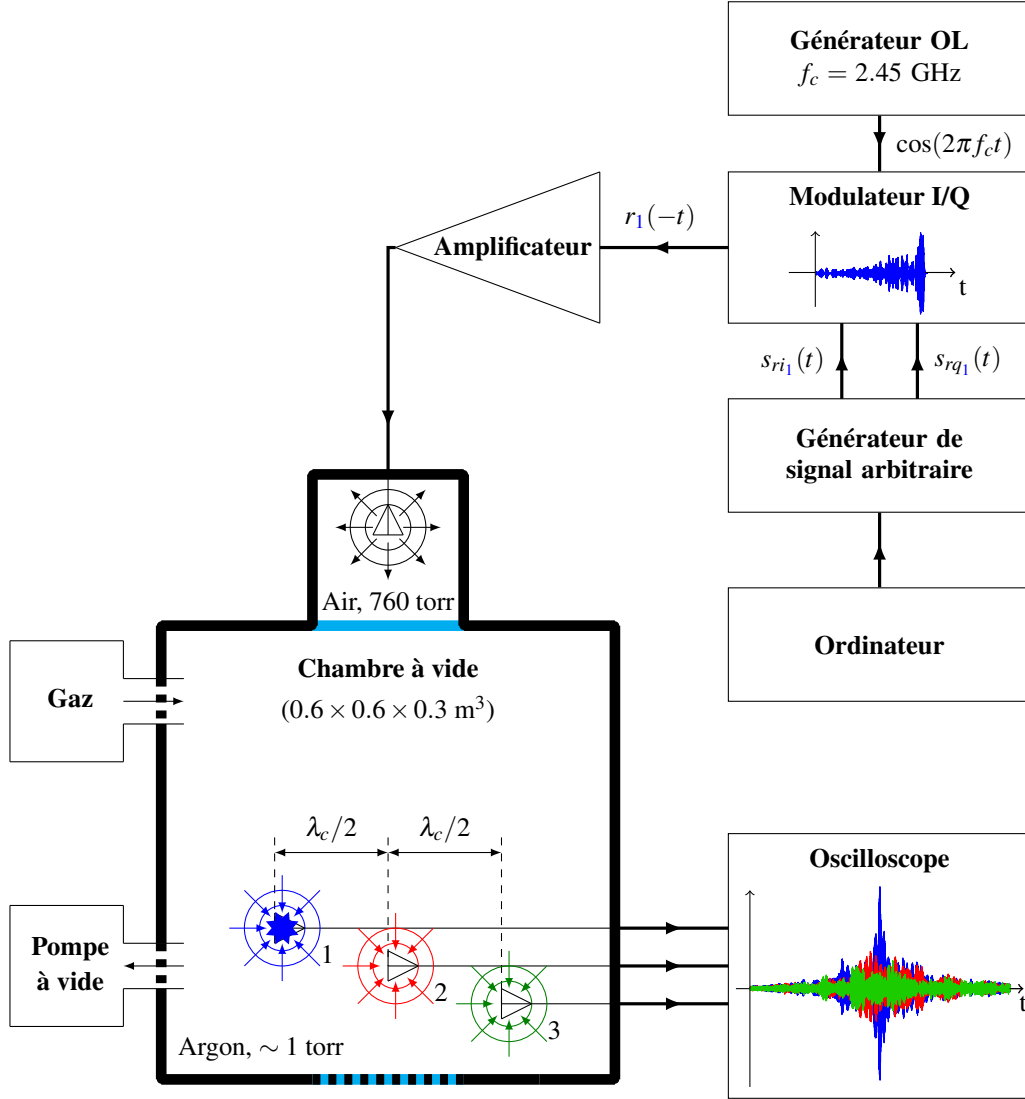


FIGURE 4.15 – Schéma d'une expérience d'amorçage de plasma par RT sur les monopoles n°1, 2 ou 3. Les signaux $s_{ri}(t)$ et $s_{rq}(t)$ obtenus par analyse fréquentielle sont transmis au modulateur I/Q. Le signal ainsi modulé est amplifié et transmis dans la cavité par l'antenne de l'annexe avec une fréquence de répétition $f_{rep} = 6$ kHz. Suivant le signal transmis $r_1(-t)$, $r_2(-t)$ ou $r_3(-t)$ la focalisation a lieu sur l'antenne correspondante (respectivement 1, 2 ou 3). Par exemple le signal $r_1(-t)$ permet d'obtenir un claquage seulement proche du monopole 1 (représenté par une étoile bleue) et les autres monopoles reçoivent un niveau microonde plus faible.

l'objet de la publication [Maz+19]. La méthode utilisée pour cela est la suivante. Les paramètres S entre les quatre antennes de la cavité sont mesurés avec le VNA. Puis les signaux en bande de base $s_{ri}(t)$ et $s_{rq}(t)$, issus d'une impulsion initiale de 8 ns (impulsion la plus courte possible avec le dispositif expérimental) modulée par une fréquence centrale de $f_c = 2.45$ GHz (correspondant à la fréquence centrale optimale) sont obtenus par analyse fréquentielle. Ils sont générés par le générateur arbitraire de signal et envoyés respectivement

sur les voies I et Q du modulateur. Le signal ainsi modulé correspond à la réponse impulsionnelle retournée entre une antenne de la cavité principale et l'antenne de l'annexe, noté : $r_{1,2,3}(-t) = s_{ri_{1,2,3}}(t)\cos(2\pi f_c t) + s_{rq_{1,2,3}}(t)\sin(2\pi f_c t)$, avec les indices 1, 2 et 3 dénotant les trois antennes de la cavité. Ce dernier est ensuite amplifié par un amplificateur 2 kW pulsé à tube à ondes progressives et transmis dans la cavité par l'antenne de l'annexe avec une fréquence de répétition $f_{rep} = 6$ kHz. Ainsi le pic de RT est obtenu de façon périodique sur l'antenne correspondante dans la cavité et le niveau microonde reste inférieur sur les autres antennes. De cette façon un plasma est amorcé sur l'antenne sur laquelle le pic de focalisation par RT est généré et maintenu par la répétition des impulsions. Par exemple, le schéma d'un amorçage de plasma par RT sur le monopole n°1 est représenté sur la Figure 4.15.

Les résultats mesurés durant les expériences d'amorçage de plasmas par RT dans ces conditions sont montrés sur la Figure 4.16. Les conditions gazeuses dans la cavité (argon à 1 torr) ont été choisies de sorte à minimiser le seuil de claquage associé à notre impulsion de 8 ns oscillant à 2.45 GHz [Fel66]. Une étude du comportement en pression de ces plasmas est proposée dans la partie suivante, dans laquelle nous illustrerons bien que le transfert de puissance est maximal autour de 1 torr. Les Figures 4.16(b-d) montrent les résultats obtenus respectivement lors de l'émission de $r_1(-t)$, $r_2(-t)$ et $r_3(-t)$. Pour chaque cas, les signaux mesurés sur les trois antennes sont tracés (par soucis de clarté, seulement les 300 ns avant et après le pic de RT sont montrées ici) ainsi que la photo de l'intérieur de la cavité prise à travers le hublot faradisé. Ces photos ont été prises avec le même angle de vue et un temps d'exposition de 40 ms. Nous remarquons tout d'abord l'efficacité de l'analyse fréquentielle développée dans la partie précédente. À partir de la seule mesure des paramètres S entre les quatre antennes, les signaux construits par analyse fréquentielle, modulés et transmis à la cavité (signaux $r_{1,2,3}(-t)$) permettent bien d'obtenir le pic de RT correspondant, avec des propriétés spatio-temporelles de focalisation similaires à celles obtenues par analyse fréquentielle dans la partie précédente.

Nous remarquons que l'opération de RT permet de générer des décharges plasmas localisées, dont la dimension est centimétrique à 1 torr. La position du plasma est bien contrôlée par la forme d'onde transmise à la cavité, illustrant le concept de "pinceau plasma ondulatoire".

Il est aussi important de noter que les monopoles, nécessaires lors de la phase initiale du RT (acquisition des réponses impulsionnelles par mesure des paramètres S au VNA) doivent rester présents lors de la seconde phase du RT pour ne pas modifier la fonction de transfert de la cavité. En effet, retirer un monopole change significativement la façon dont les ondes se comportent dans la cavité entre les deux phases du RT. Cela empêche donc d'obtenir une bonne focalisation par RT. De plus, leur présence est indispensable pour amorcer un plasma dans nos conditions expérimentales. En effet, au bout du monopole le champ électrique est localement augmenté d'environ 20 dB [Mil67] ; [Pod+06]. Cette intensification du champ électrique au niveau d'objets métalliques pointus est souvent utilisée pour amorcer des plasmas à haute pression, pour lesquelles le champ de claquage est très élevé. On parle dans ce cas d'initiateurs. L'initiateur le plus simple auquel nous puissions penser consiste juste en une pointe métallique placée dans le réacteur plasma [Bab+00] ; [Rus+01]. L'utilisation d'un initiateur est largement répandue pour amorcer une décharge dans des dispositifs à torches

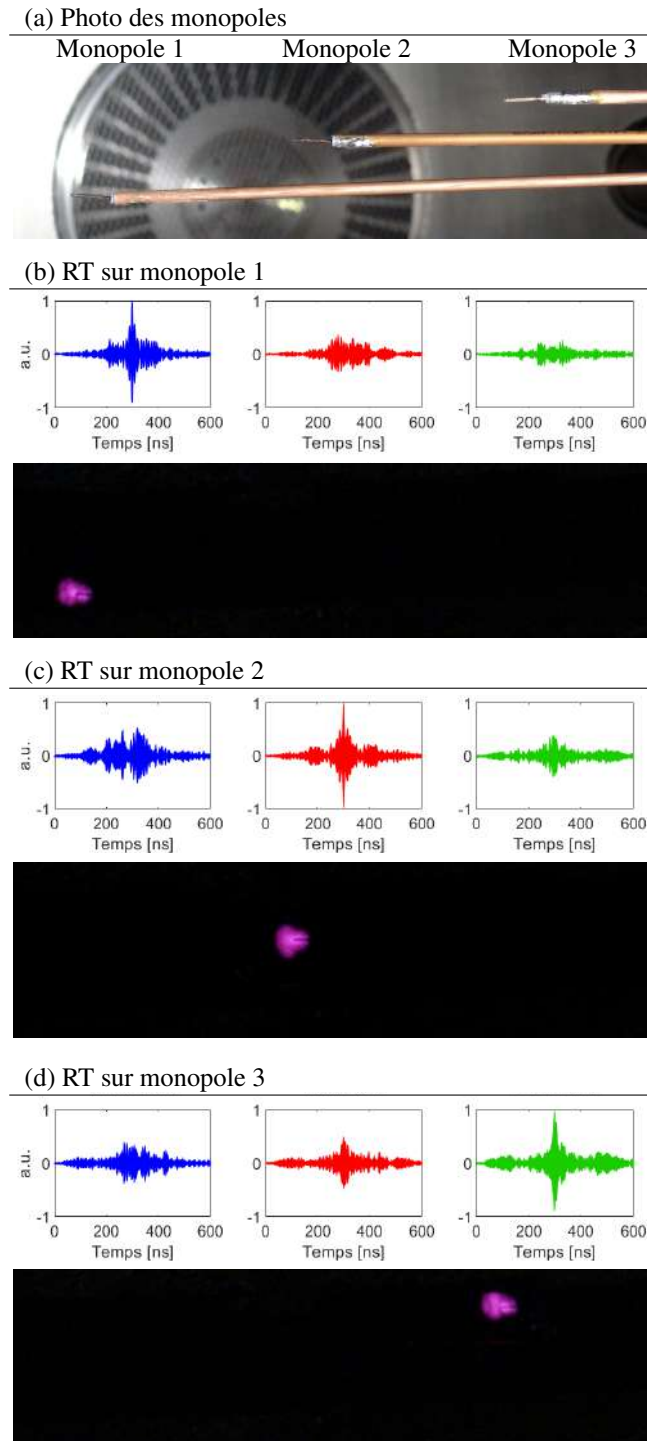


FIGURE 4.16 – (a) Photos des trois monopoles de la cavité principale (1, 2 et 3). (b),(c),(d) Pour chaque expérience de RT (respectivement sur les monopoles 1, 2 et 3) : signaux mesurés (en bleu, rouge et vert respectivement sur 1, 2 et 3) avec la photo du plasma correspondant. Les signaux sont normalisés par rapport au niveau du pic correspondant.

plasma [WBW15]; [AS+01]. Les monopoles sur lesquels nous amorçons des plasmas dans nos expériences peuvent être vus comme des initiateurs de structure simple que nous venons d'évoquer : des pointes métalliques. Sur notre dispositif expérimental à basse pression, suivant le même principe que sur les dispositifs à haute pression [Bab+00]; [Rus+01], le monopole facilite ainsi l'amorçage du plasma en diminuant localement la puissance microonde nécessaire pour claquer le gaz. Il est aussi important de noter que la présence d'initiateurs dans un réacteur plasma joue un rôle significatif sur la forme et la dimension des plasmas générés au niveau de ces initiateurs. À haute pression, cet effet est souligné par exemple dans [Bab+00] : "The discharge is seen as a bunch of thin plasma filaments emerging from the needle point". Dans les conditions de basse pression de nos expériences on ne verra pas la formation de filaments mais la présence de ces initiateurs joue sûrement un rôle significatif sur la forme et la dimension des plasmas générés par RT dans les expériences. Cet effet sera illustré plus clairement à l'aide des mesures obtenues avec une caméra rapide dans la partie suivante.

Pour finir avec la présentation des plasmas amorcés par RT, il est intéressant de montrer les expériences d'amorçage simultané de plusieurs plasmas. Les résultats de ces expériences sont présentés sur la Figure 4.17. Dans ce cas, les pics de focalisation de RT sont obtenus simultanément sur deux monopoles. En se basant sur le principe de superposition, il suffit pour cela d'envoyer sur l'antenne d'émission (antenne n°4) la somme des deux réponses impulsionnelles retournées correspondantes. Par exemple, les résultats présentés sur la Figure 4.17(b) sont obtenus en envoyant à l'antenne d'émission le signal $r_1(-t) + r_2(-t)$. On remarque tout d'abord sur les signaux microondes que les pics de RT sont bien obtenus sur les monopoles désirés. Les photos des plasmas montrent que la position des plasmas est bien contrôlée par la position des pics de RT. Ainsi l'amorçage de plasmas simultanés, même distants de 6 cm (amorçages simultanés sur les monopoles 1 et 2, et 2 et 3) ne modifie pas significativement le comportement des ondes au point d'empêcher la focalisation par RT sur l'un ou l'autre des deux monopoles concernés.

D'un point de vue plus applicatif, ces résultats sont intéressants tout d'abord dans le domaine du traitement de surface. En effet, il serait possible de traiter plusieurs endroits simultanément. Cela signifierait pouvoir contrôler simultanément plusieurs "pinceaux plasmas". Une autre application pourrait bénéficier de ce contrôle simultané des plasmas par RT : les "space-time-modulated (STM) media" [TK20]. Ce sont des milieux dont les paramètres constitutifs varient à la fois dans le temps et l'espace. Les propriétés uniques et exotiques des STM ont mené au développement de concepts physiques originaux et de nouveaux dispositifs, notamment dans le domaine des microondes [TK20]. Dans le cas des plasmas amorcés simultanément par RT, le contrôle de la position des plasmas par RT permettrait une modulation spatio-temporelle des propriétés (permittivité) du milieu de propagation. Pour ces deux applications, le "décrochage" des plasmas des initiateurs (monopoles) est nécessaire. Nous reparlerons de ce point dans les conclusions et perspectives de cette thèse.

Nous sommes restés volontairement descriptif dans cette partie, dont l'objectif est de montrer le contrôle des plasmas par RT. Cependant, comme nous l'avons vu dans le chapitre précédent, la présence d'un plasma dans la cavité lors de la phase inverse affecte le comportement des ondes. De plus, le claquage du gaz est un phénomène non-linéaire. Dans le domaine

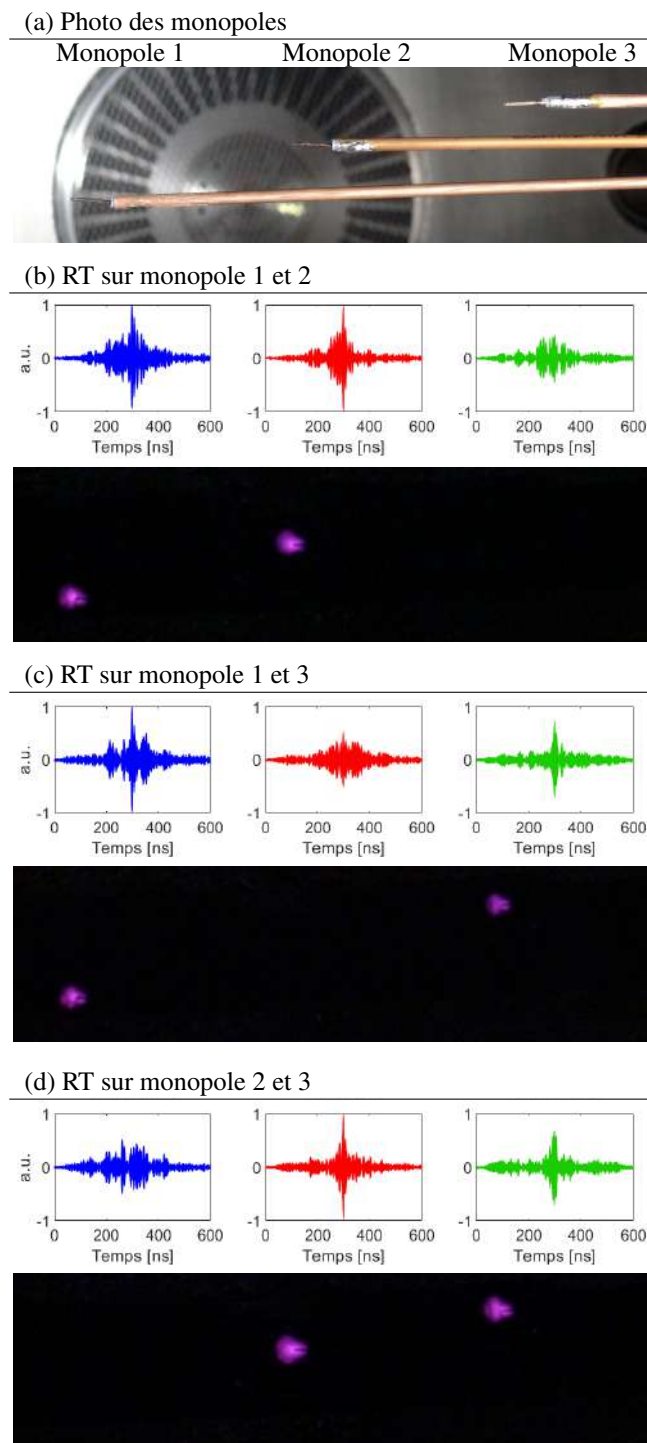


FIGURE 4.17 – Expériences de RT simultanément sur deux monopoles (1 et 2, 1 et 3, et 2 et 3) : signaux mesurés sur les trois monopoles 1, 2, 3 (respectivement en bleu, rouge et vert) avec la photo des plasmas correspondant.

des microondes, le RT a été appliqué à des milieux possédant une non-linéarité [Fra+13a]; [Fra+13b] et à des milieux qui changent entre la phase d'apprentissage et la phase de ré-émission du RT [NBEZ09]. Dans [Fra+13a], la non-linéarité est générée par la présence d'une diode sur une antenne. Cette non-linéarité est donc aussi localisée dans la cavité mais elle est présente lors des deux phases du RT, contrairement à notre application. Elle a pour conséquence la génération d'ondes de fréquences plus élevées que celles de la source initiale. Elle joue alors le rôle d'une seconde source dans la cavité, et en retournant l'information contenue dans la réponse impulsionnelle autour de ces fréquences élevées, il est possible de focaliser sur la diode [Fra+13a]. Dans [NBEZ09], la robustesse du principe du RT est étudiée dans une chambre réverbérante à brasseur de modes, le brasseur étant utilisé pour changer le milieu de propagation entre les deux phases du RT. Les résultats montrent que les performances du processus de RT dépendent du degré de corrélation entre les réponses impulsionnelles correspondant aux différents milieux de propagation. Pour les plasmas amorcés par RT, le claquage du gaz est un phénomène non-linéaire momentané (par opposition à la non-linéarité permanente de [Fra+13a]) et la présence du plasma dans la cavité change le milieu dans lequel les ondes se propagent. Ce changement engendre ainsi une décorrélation entre les réponses impulsionnelles sans et avec plasma et donc une perte de performance du processus de RT. Il paraît alors nécessaire de se pencher sur ces propriétés du milieu de propagation, afin de pouvoir déterminer comment la présence du plasma va affecter, et probablement limiter, le contrôle de la focalisation de l'énergie électromagnétique par RT et donc le contrôle des plasmas amorcés par RT. Pour cela, une analyse temps-fréquence est intéressante, comme présentée dans la section suivante. De plus, la dynamique spatio-temporelle de ces plasmas est étudiée par imagerie rapide afin de mieux comprendre le comportement et la physique de ces plasmas.

4.3 Caractérisation des plasmas amorcés par retournement temporel

Nous avons vu dans le premier chapitre l'intérêt d'utiliser une source plasma pulsée. Elle permet de limiter les phénomènes de chauffage du gaz et des parois. Dans ce cas, l'intensité du champ électrique, la durée du pulse τ , la fréquence de répétition f_{rep} , le temps de montée t_r et de descente, ainsi que le rapport cyclique D défini comme $D = \tau \times f_{rep}$ sont les paramètres principaux qui contrôlent la production et le maintien du plasma. Des études détaillées des plasmas microondes pulsés ont été menées sur les plasmas d'onde de surface. Par exemple, Hubner *et al.* et Carbone *et al.* ont étudié des impulsions avec τ descendant jusqu'à 1 μ s, t_r jusqu'à quelques dizaines de nanosecondes et $D > 10$ % [Hü+12] ; [Car+14].

Une caractéristique commune aux sources plasmas microondes habituellement utilisées est le fait que la position du plasma est déterminée et fixée par la géométrie de la cavité. Dans le cas des plasmas amorcés et contrôlés par RT, la position et la dimension du plasma est complètement décorrélée de la cavité, comme nous l'avons vu dans la partie précédente. Dans ce cas, la maîtrise de la forme d'onde transmise à la cavité permet le contrôle de l'organisation spatio-temporelle de l'énergie dans la cavité. Le plasma n'occupe alors pas la totalité du volume de la cavité, mais il est localisé à l'endroit de la focalisation par RT. En plus de ses propriétés spatiales innovantes, les plasmas obtenus par RT possèdent aussi des caractéristiques temporelles originales : ils sont amorcés en régime pulsé (typiquement $f_{rep} = 6$ kHz) avec des impulsions de $\tau = 8$ ns et des temps de montée $t_r < 1$ ns. Cela correspond à des rapport cycliques inférieurs à 0.05 %. À notre connaissance, aucune source plasma microonde avec des temps de montée et des rapports cycliques si faibles n'a jamais été étudiée.

Cette partie a pour objectif de décrire la dynamique spatio-temporelle des plasmas amorcés par RT. La plupart des résultats présentés dans cette partie ont fait l'objet d'une publication en collaboration avec Luc Stafford de l'université de Montréal [Maz+20b]. Une caractérisation électrique est d'abord effectuée dans un premier temps, en comparant les signaux mesurés avec et sans plasma. Nous verrons alors que la pression joue un rôle fondamental sur le comportement des plasmas amorcés par RT. Afin de comprendre plus en détail l'influence de la pression sur le comportement du plasma, des mesures en imagerie rapide sont présentées. Elles permettent d'illustrer l'influence de la pression sur la dynamique temporelle et sur la dimension du plasma. Les effets des lobes secondaires temporels précédant le pic de RT sont clairement illustrés. On les attribue à un "effet mémoire". Une étude sur l'influence de la fréquence de répétition permet ensuite de valider cette interprétation. Pour finir, quelques éléments d'interprétation de la physique des phénomènes prenant place durant ces plasmas amorcés par RT sont discutés.

4.3.1 Dispositif expérimental

La configuration expérimentale utilisée est représentée sur la Figure 4.18. Deux antennes monopoles (de dimension $\lambda_c/10$) numérotées 1 et 2 sont présentes dans la cavité principale.

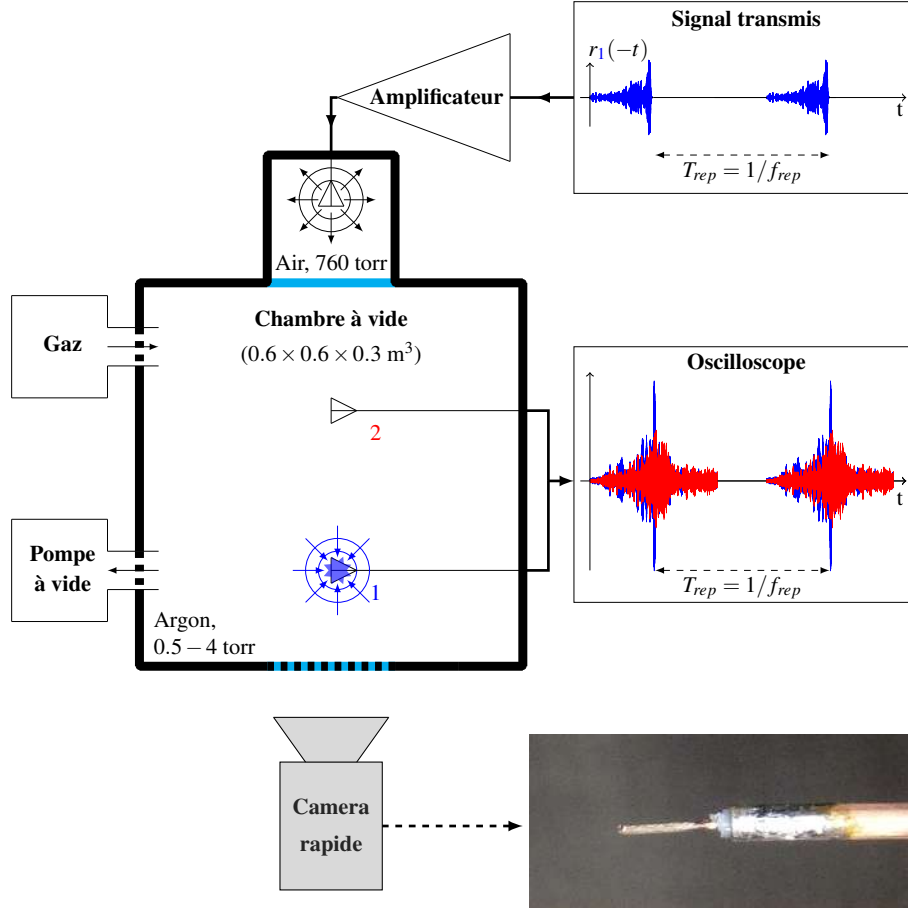


FIGURE 4.18 – Schéma d'une expérience d'amorçage de plasmas par RT pour l'étude spatio-temporelle de leur dynamique. Le signal $r_1(-t)$, *i.e.* la réponse impulsionnelle retournée temporellement entre l'antenne de l'annexe et l'antenne 1 de la cavité, permet la focalisation d'un pic de 8 ns sur ce monopole et l'amorçage d'un plasma [Maz+19].

Les plasmas sont amorcés uniquement sur l'antenne 1 et l'antenne 2 est présente seulement pour mesurer le champ à un autre endroit dans la cavité. L'information de la propagation des ondes entre l'antenne 1 et celle de l'annexe est obtenue en mesurant les paramètres S au VNA. L'impulsion initiale est prise de 8 ns modulée à $f_c = 2.55 \text{ GHz}$, soit la fréquence centrale optimale obtenue par analyse fréquentielle. Les signaux en bande de base à transmettre au modulateur I/Q pour obtenir la réponse impulsionnelle retournée temporellement $r_1(-t)$ (à une impulsion de 8 ns modulée à $f_c = 2.55 \text{ GHz}$) sont construits numériquement par analyse fréquentielle. Ce signal $r_1(-t)$ est transmis périodiquement à la cavité avec une fréquence de répétition f_{rep} . Un pic de RT est alors généré périodiquement à cette fréquence f_{rep} sur le monopole 1. La pression p dans la cavité est comprise entre 0.5 et 4 torr dans l'argon.

La puissance du pic de focalisation, définie comme le carré de la tension mesurée divisée par l'impédance d'entrée de l'oscilloscope (50Ω), est proportionnelle au niveau du champ électrique appliqué au monopole 1 au moment de la focalisation. Cependant, à cause de la

présence d'un milieu très diffusant entre les monopoles (cavité réverbérante) et l'utilisation de signaux large bande (250 MHz), la relation entre cette puissance mesurée et la valeur du champ électrique est inconnue et ne peut pas être facilement estimée. En conséquence, seules des mesures relatives peuvent être réalisées et la puissance est donnée en unités arbitraires.

Afin d'atteindre les conditions de claquage, le signal $s_1(-t)$ est amplifié par l'amplificateur 2 kW pulsé avant d'être transmis à la cavité, comme nous l'avons vu dans la partie précédente [Maz+19]. Notons que dans les mesures présentées dans cette partie, le signal d'entrée $s_1(-t)$ est transmis à la cavité avec le même niveau de puissance maximale d'environ 1 kW, sauf indication contraire. Le signal de RT alors mesuré à l'oscilloscope est tracé sur la Figure 4.18. La puissance du pic de focalisation correspondante dans cette condition est considéré égale à $P_{rt} = 7$ u.a..

4.3.2 Caractérisation électrique

4.3.2.1 Identification des limites d'entretien et d'extinction du plasma

Le monopole sur lequel la focalisation par RT a lieu sert d'initiateur pour la décharge plasma et il permet aussi d'étudier le comportement du champ électrique. On pourrait alors penser à utiliser ce signal mesuré en présence de plasma pour en déduire des informations sur ce dernier, en comparant notamment le signal de RT (issu du même signal $s_1(-t)$) mesuré dans un cas sans plasma (à pression atmosphérique) et dans un cas avec plasma (autour du torr). Par exemple, ces signaux de RT avec et sans plasma sont tracés sur la Figure 4.19 pour une pression $p = 1.5$ torr et une fréquence de répétition de $f_{rep} = 6$ kHz. Le pic de focalisation a lieu à 300 ns. En comparant ces deux signaux, on pourrait espérer obtenir une information sur les propriétés du plasma, comme par exemple l'énergie absorbée par le plasma. Cependant, la présence du plasma proche du monopole change sa sensibilité, *i.e.* la relation entre le champ électrique appliqué au monopole et la tension résultante alors mesurée. De plus, en fonction de la densité du plasma, les ondes peuvent être plus ou moins réfléchies par le plasma, donnant un niveau de signal mesuré plus faible. De ce fait, l'analyse de ces signaux ne peut pas aller plus loin et des données absolues ne peuvent pas être obtenues de ces mesures. Seule une tendance de ces signaux peut être remarquée, comme la diminution du niveau microonde et le décalage en phase induit par l'amorçage du plasma. C'est pourquoi une seconde antenne est placée ailleurs dans la cavité (antenne 2), sur laquelle il n'y a pas de décharge lors de l'amorçage des plasmas par RT sur l'antenne 1. Sa sensibilité est alors inchangée dans les cas avec et sans plasma et on peut alors espérer comparer les signaux mesurés dans ces deux cas. Nous montrerons par exemple dans la suite de cette partie comment cette caractérisation électrique permet de déterminer le mode le plus affecté par la présence du plasma dans la cavité.

Pour les sources plasmas pulsées, les courbes de maintien classiques, *i.e.* la puissance minimale requise pour entretenir le plasma en fonction de la pression, ne sont pas vraiment pertinentes. En effet, la décharge plasma obtenue au moment d'une impulsion peut dépendre de la précédente. Cela permet néanmoins de définir le domaine opérationnel. La méthode pour

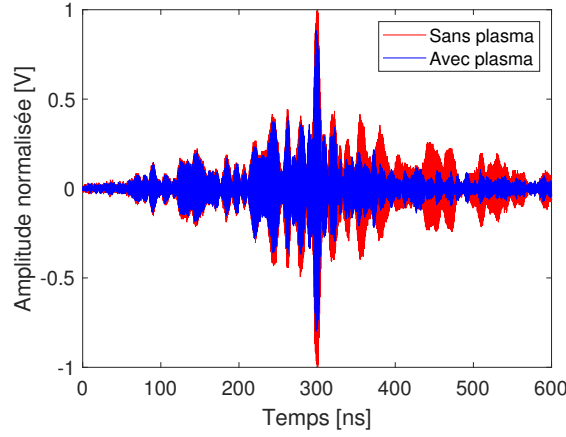


FIGURE 4.19 – Signaux microondes mesurés sur le monopole 1 durant une expérience de RT en rouge sans plasma (pression atmosphérique) et en bleu avec plasma ($p = 1.5$ torr) pour une fréquence de répétition de $f_{rep} = 6$ kHz.

déterminer la puissance de maintien est la suivante. Un plasma est amorcé sur le monopole 1 par RT avec un fort niveau de puissance du signal incident et maintenu par la répétition des impulsions à la fréquence f_{rep} . Ensuite ce niveau est diminué petit à petit jusqu'au niveau minimum permettant de maintenir la décharge. En dessous de ce niveau, la décharge s'éteint (visuellement). Cette opération est répétée en fonction de la pression pour une fréquences de répétition donnée. Les courbes de maintien correspondantes sont tracées sur la Figure 4.20(a) pour deux valeurs de f_{rep} . En considérant la théorie classique du dépôt de puissance dans les plasmas hautes fréquences, la densité de puissance absorbée par le plasma est donnée par (équation 1.32) [MP12] :

$$P_a = \theta_e n_e \quad \text{où} \quad \theta_e = \frac{e^2}{m_e} \frac{\nu_m}{\nu_m^2 + \omega_0^2} E_{0rms}^2 \quad (4.27)$$

avec θ_e la puissance absorbée par un électron, ω_0 la pulsation de l'onde et E_{0rms} le champ électrique rms requis pour maintenir la décharge. Les courbes expérimentales de la Figure 4.20(a) présentent une puissance minimale de maintien pour une pression d'argon proche de 1.5 torr. Cela signifie que le transfert d'énergie est maximal autour de cette pression. Mathématiquement l'optimum est obtenu lorsque le rapport $\frac{\nu_m}{\nu_m^2 + \omega_0^2}$ est maximal, soit $\nu_m = \omega_0$. Sachant que la bande passante (250 MHz) est faible comparée à la fréquence centrale f_c , ω_0 peut être approximé par $2\pi f_c \approx 16 \times 10^9$ rad/s. Cela donne un ν_m d'environ 16 GHz à 1.5 torr dans nos expériences, en bon accord avec les données trouvées dans la littérature dans ces conditions [LL05]. Il convient de noter par ailleurs que pour un même niveau de signal incident la densité électronique obtenue par RT en simulation au chapitre précédent se trouve aussi maximale autour de cette valeur de 1.5 torr (voir Figure 3.20).

En ce qui concerne l'influence de la fréquence de répétition, on pourrait s'attendre à ce qu'avec des rapports cycliques $D < 0.05$ %, l'effet de l'impulsion précédente disparaisse et que la puissance de maintien soit indépendante de la fréquence de répétition. Cependant, la Figure

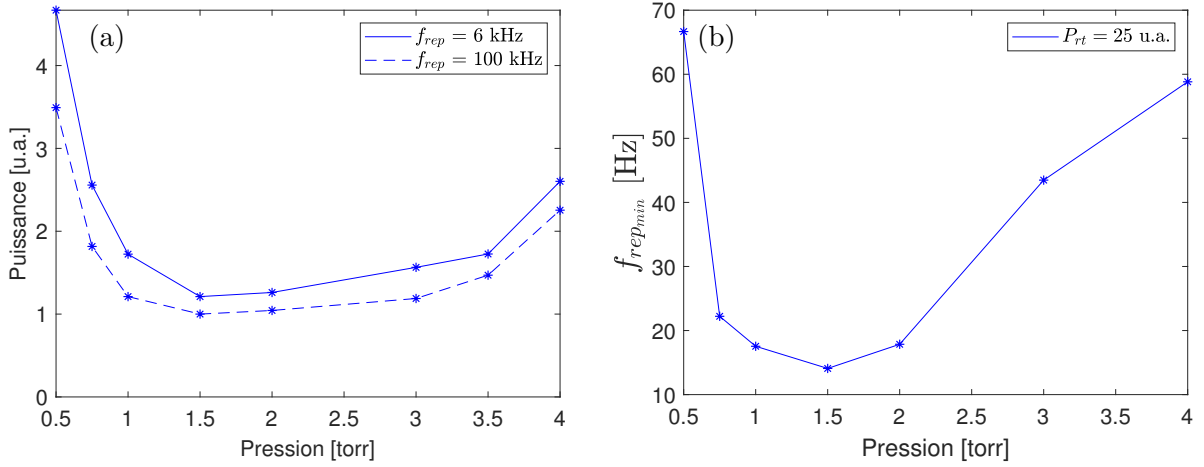


FIGURE 4.20 – (a) Puissance microonde minimum et (b) fréquence de répétition minimum du pic de RT nécessaire pour maintenir un plasma en fonction de la pression. Dans les deux cas, ces valeurs sont déterminées visuellement à partir de la disparition de la décharge.

4.20(a) montre que plus la fréquence de répétition augmente, plus la puissance nécessaire pour maintenir la décharge diminue, comme reporté dans [DL73] ; [VIS85]. Ces mesures attestent clairement de l’influence de l’impulsion précédente sur la suivante, soit un “effet mémoire”.

Afin d’approfondir l’investigation de ce phénomène, la fréquence de répétition minimale permettant de maintenir la décharge est aussi déterminée. Une procédure similaire à celle utilisée pour les courbes de maintien est utilisée. Le plasma est amorcé avec un fort niveau de puissance du pic de RT et une fréquence de répétition élevée ($P_{rt} = 25$ u.a. et $f_{rep} = 6$ kHz). Puis la fréquence de répétition est lentement diminuée jusqu’à la fréquence minimale $f_{rep_{min}}$ permettant de maintenir la décharge. Cette opération est répétée pour chaque pression et les résultats sont tracés sur la Figure 4.20(b). On note que la décharge peut être maintenue pour une fréquence de répétition aussi basse que 12 Hz pour la condition de transfert maximal trouvé autour de 1.5 torr. Cela signifie que l’effet de l’impulsion précédente est toujours présent après un intervalle d’environ 100 ms. Afin de mieux décrire et comprendre ce mécanisme, nous exploiterons dans la partie suivante des mesures obtenues par imagerie rapide.

4.3.2.2 Effet du plasma sur la cavité

Intéressons nous maintenant au changement de milieu de propagation induit par l’amorçage d’un plasma par RT. La nature de ce changement dépend des caractéristiques du plasma (densité, chimie du plasma, dimensions). L’interaction entre les ondes et le plasma peut influencer sur ces caractéristiques. Afin de mieux comprendre comment le plasma joue sur le comportement des ondes - et vice-versa - lors d’une expérience de RT, il paraît intéressant de pouvoir suivre dans le temps le comportement des ondes dans la cavité. La Short-Time Fourier Transform (STFT) se prête bien à cette étude. Cet outil, qui consiste à calculer la transformée de Fourier sur une fenêtre temporelle glissante, permet de tracer le spectrogramme de ce signal (diagramme temps-fréquence). De cette façon, il est possible de suivre l’évolution temporelle

des fréquences auxquelles la cavité est excitée entre deux antennes lors d’une expérience de RT. Nous allons finir cette partie sur la caractérisation électrique des plasmas en traçant les spectrogrammes des signaux mesurés lors d’une expérience de RT pour des cas avec et sans plasma pour différentes pressions.

Du point de vue expérimental, plusieurs paramètres influent sur le “type” de plasma généré par RT : la fréquence de répétition des pics de RT f_{rep} , l’amplitude de ce pic et la pression. On se concentrera ici sur l’étude de l’influence de la pression sur le processus de RT. La fréquence de répétition des pics de RT est fixée ici à $f_{rep} = 6$ kHz. Les signaux sont mesurés sur l’antenne 2 pour les cas avec et sans plasma généré sur l’antenne 1 pour différentes pressions (entre 0.5 et 4 torr). La présence du plasma sur l’antenne 1 peut modifier le comportement des ondes dans la cavité et cette éventuelle modification sera mesurée sur l’antenne 2.

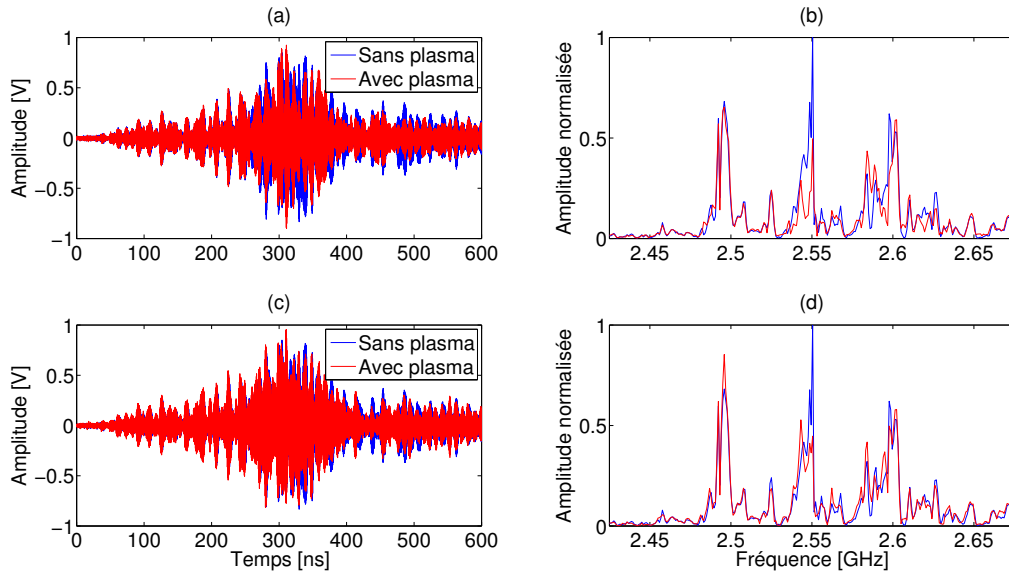


FIGURE 4.21 – Comparaison des signaux mesurés sur l’antenne 2 sans et avec plasma sur l’antenne 1 (en bleu et rouge respectivement) dans les domaines temporel et fréquentiel (colonne de gauche et droite respectivement). (a)(b) Pour une pression de 0.5 torr. (c)(d) Pour une pression de 4 torr.

Les signaux mesurés sur l’antenne 2 sans et avec plasma pour deux pressions 0.5 et 4 torr sont tracés sur la Figure 4.21 dans les domaines temporel et fréquentiel. Seule la “bande utile” de ces derniers est tracée, correspondant à une bande de 250 MHz autour de la fréquence porteuse ($f_c = 2.55$ GHz). Les ondes électromagnétiques obtenues au niveau de l’antenne 1 et notamment au moment du pic de RT vont augmenter la densité électronique. Cela change l’influence du plasma sur le comportement des ondes dans la cavité. Les signaux avec et sans plasma se superposent exactement jusqu’à environ 250 ns pour une pression de 0.5 torr (Figure 4.21(a)) et jusqu’à environ 300 ns pour une pression de 4 torr (Figure 4.21(c)). Ainsi le niveau de densité électronique au niveau de l’antenne 1 durant ces périodes n’est pas assez élevé pour modifier le comportement des ondes dans la cavité, ou du moins cette modification n’est pas perceptible sur la mesure du signal sur l’antenne 2. Ensuite les signaux mesurés avec

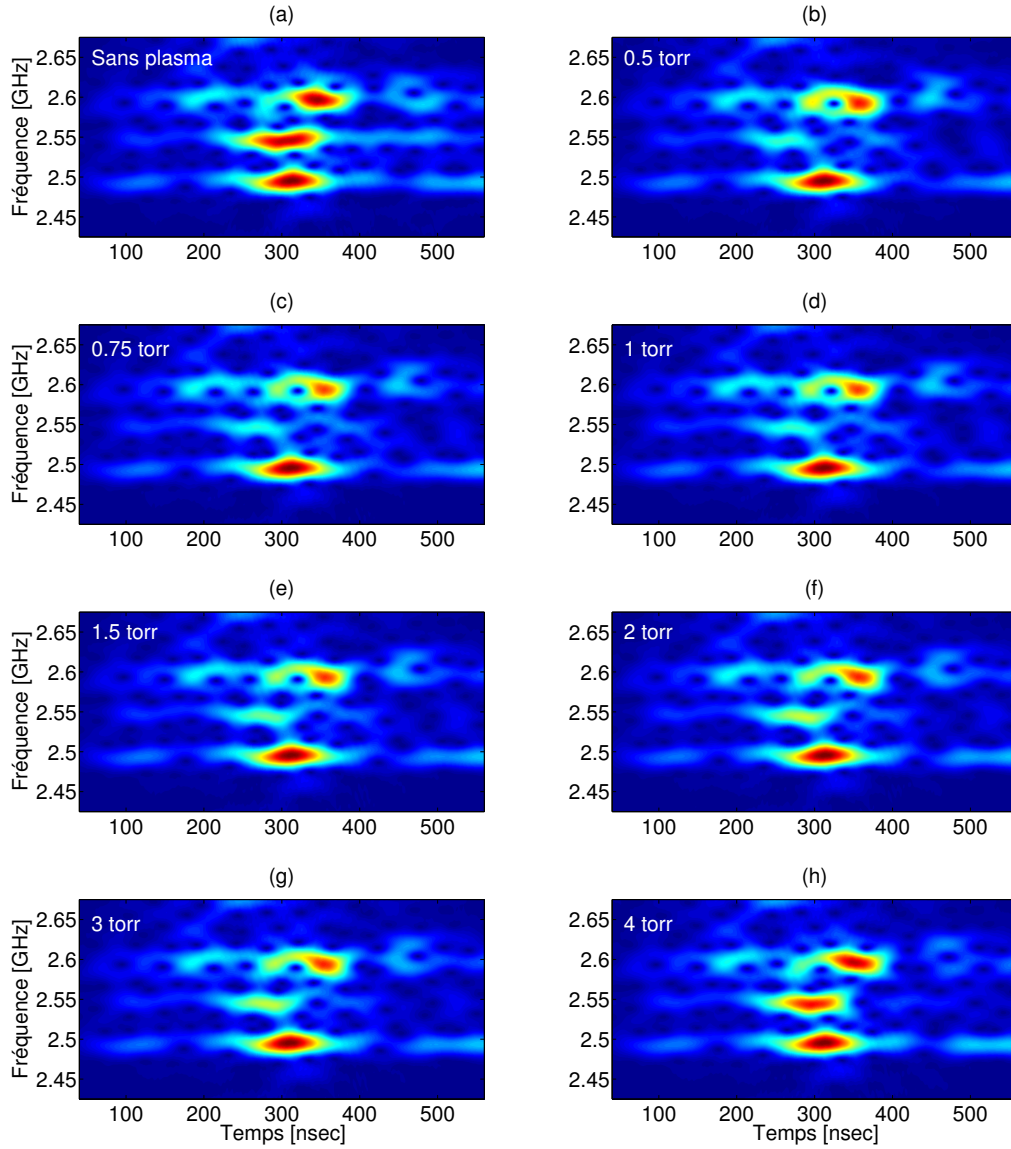


FIGURE 4.22 – Spectrogrammes des signaux mesurés sur l’antenne 2 lors d’une focalisation par RT sur l’antenne 1 : (a) Sans plasma sur l’antenne 1. (b)(c)(d)(e)(f)(g)(h) Avec plasma sur l’antenne 1 à des pressions de 0.5, 0.75, 1, 1.5, 2, 3 et 4 torr respectivement.

et sans plasma différent (Figure 4.21(a)(c)). La densité électronique est alors assez élevée pour modifier le comportement des ondes dans la cavité. La dimension du plasma joue aussi un rôle sur la différence observée. Nous verrons dans la partie suivante que la pression joue un rôle essentiel sur l’émission lumineuse et la dimension du plasma. Les Figures de ces signaux dans le domaine fréquentiel (Figure 4.21(c)(d)) nous indiquent une partie des fréquences auxquelles la cavité est excitée durant une expérience de RT sur l’antenne 1. La présence du plasma sur l’antenne 1 va potentiellement affecter la façon dont les différents modes associés vont être

excités. Les fréquences qui semblent les plus “affectées” par la présence du plasma sont celles autour de 2.55 GHz (Figure 4.21(b)(d)) et plus à 0.5 torr qu’à 4 torr. Afin d’analyser plus finement ces signaux, il paraît intéressant de tracer leurs spectrogrammes. Ils permettent de suivre dans le temps les fréquences contenues dans le signal reçu sur l’antenne 2 la cavité est excitée et de comparer les cas sans et avec plasma.

Les spectrogrammes des signaux mesurés sur l’antenne 2 lors d’une expérience de RT sur l’antenne 1 pour le cas sans plasma (Figure 4.22(a)) et avec plasma sur l’antenne 1 à différentes pressions (Figure 4.22(b-h)) sont présentés sur la Figure 4.22. Ces spectrogrammes ont été obtenus avec une fenêtre de Hamming de 40 ns, ce qui correspond à cinq fois la durée de l’impulsion initiale. Le choix du type de fenêtre et de sa durée est fait de sorte à avoir un bon compromis entre les résolutions temporelle et fréquentielle. Pour le cas sans plasma (Figure 4.22(a)), on retrouve bien les trois principaux pics obtenus par transformée de Fourier, autour de 2.5 GHz, 2.55 GHz et 2.6 GHz (courbes bleues des Figure 4.21(b) et Figure 4.21(d)). Les niveaux de ces trois pics sont plus intenses autour de 300 ns (Figure 4.22(a)), correspondant bien à l’instant auquel les niveaux des signaux temporels (Figure 4.21(a) et Figure 4.21(c)) sont les plus élevés.

Pour les cas avec plasma (Figure 4.22(b-h)) on retrouve les pics autour de 2.5 GHz et 2.6 GHz (courbes rouges des Figure 4.21(b) et Figure 4.21(d)). Le pic autour de 2.5 GHz semble être similaire dans les cas avec et sans plasma (pour toutes les pressions) et on peut noter quelques différences sur le comportement temporel du pic autour de 2.6 GHz avec la pression. Plus la pression augmente et plus les spectrogrammes obtenus avec plasma semblent se rapprocher de celui obtenu sans plasma. Les différences les plus importantes entre tous ces spectrogrammes se trouvent autour de 2.55 GHz (avec un comportement similaire pendant les 250 premières nanosecondes). Plus la pression est faible et plus le plasma semble avoir une influence significative sur le comportement des ondes dans la cavité à cette fréquence. Une hypothèse pour expliquer cela pourrait être la dimension du plasma, qui diminue avec la pression comme nous allons le voir dans la partie suivante. Afin d’étudier plus en détail la dynamique spatio-temporelle du plasma amorcé sur le monopole 1, nous allons dans la suite de cette partie utiliser un diagnostic optique pour observer le comportement des plasmas.

4.3.3 Caractérisation par imagerie rapide

Les mesures d’imagerie rapide présentées dans cette partie ont été obtenues avec une caméra rapide PI-MAX-512. Toutes les mesures caméra présentées dans cette partie ont été effectuées sans filtre optique et avec un temps d’exposition de 2 s et une durée de porte de 4 ns.

4.3.3.1 Dynamique de la décharge

L’évolution spatio-temporelle d’une décharge plasma par RT typique à $p = 1.5$ torr et pour une fréquence de répétition de $f_{rep} = 6$ kHz a été observée par imagerie rapide. La

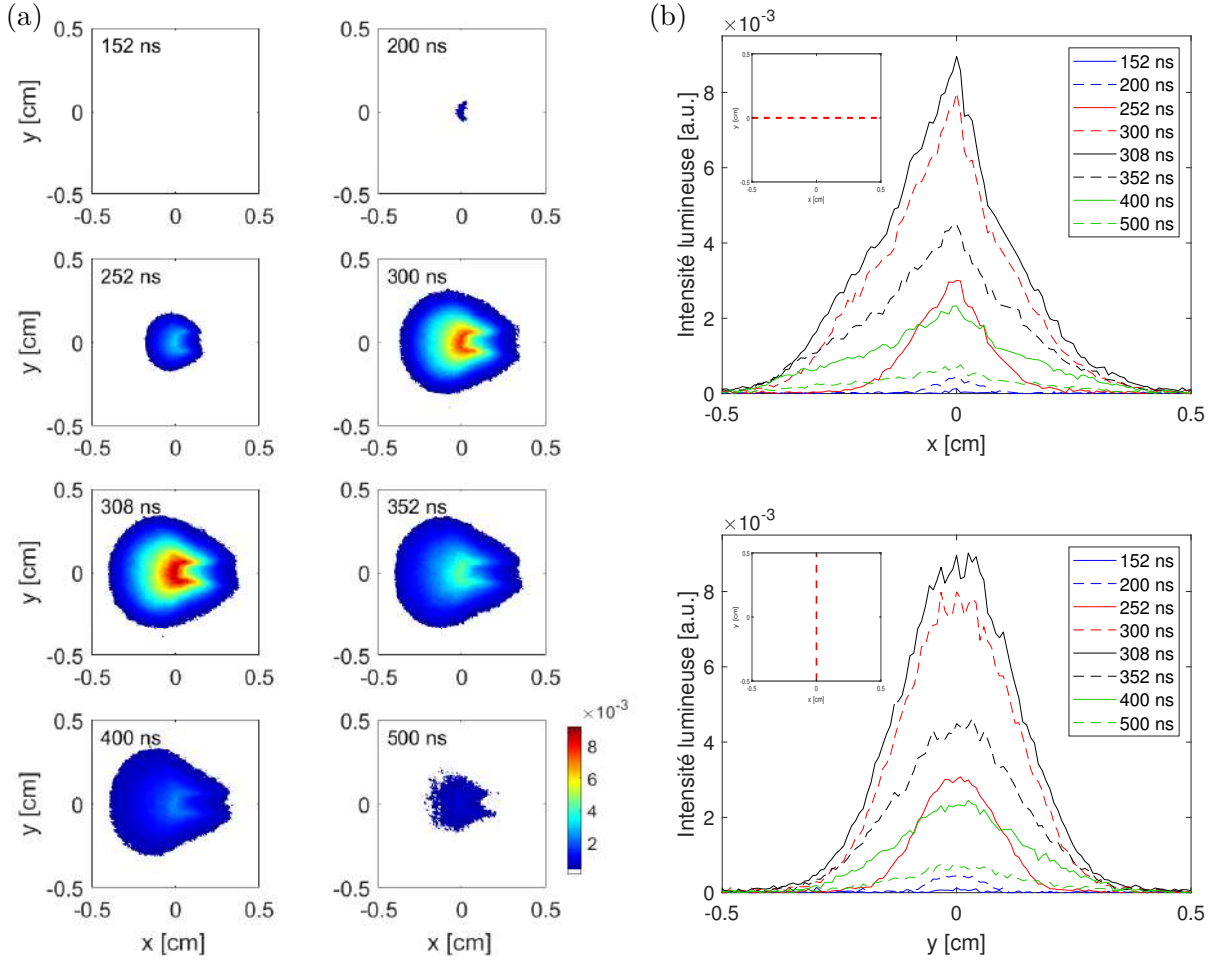


FIGURE 4.23 – (a) Images de l’intensité lumineuse du plasma mesurée par imagerie rapide pour une pression de $p = 1.5$ torr et une fréquence de répétition de $f_{rep} = 6$ kHz pour des temps différents autour du temps du pic de RT ($t = 300$ ns). L’intensité lumineuse est représentée par une carte de couleur. Les parties blanches correspondent aux régions en dessous d’un niveau seuil correspondant au niveau de bruit des mesures caméra et égal à $0.3 \cdot 10^{-3}$ a.u.. (b) Coupe horizontale (haut), *i.e.* le long de l’axe x pour $y = 0$ cm et coupe verticale (bas), *i.e.* le long de l’axe y pour $x = 0$ cm, pour les différents instants présentés en (a).

Figure 4.23(a) montre différents clichés pris autour du moment de la focalisation. On voit que la décharge est initiée proche de la pointe du monopole métallique. Elle reste ensuite à cet endroit et s’étend sur des dimensions centimétriques autour du monopole avec une forme quasi-sphérique. Afin de décrire plus finement la forme du plasma, les profils d’intensité lumineuse le long des axes x et y sont tracés pour différents instants sur la Figure 4.23(b). On remarque que, à l’inverse de l’intensité lumineuse qui possède une forme symétrique le long de la coupe verticale (pour $x = 0$ cm), le long de la coupe horizontale (pour $y = 0$ cm) nous pouvons observer une dissymétrie induite par la présence du monopole (du côté des x positifs). Par ailleurs, le rayon moyen du plasma, d’environ $R_{plasma} = 0.5$ cm, est bien infé-

rieur à la demi longueur d'onde ($\lambda_c/2 = 6$ cm), soit la dimension spatiale de la focalisation par RT [Ler+04]. Pourtant, sur les résultats des simulations en 2D présentés dans le chapitre précédent, la dimension du pic de densité électronique est en accord avec la dimension de la focalisation de l'énergie microonde $\lambda_c/2$ (voir Figure 3.18). Plusieurs hypothèses peuvent être avancées pour expliquer cette différence. Tout d'abord notons que les résultats des simulations présentés au chapitre précédent (Figure 3.18) ont été obtenus en 2D et en mono-pulse, contrairement aux plasmas expérimentaux qui sont générés en 3D en régime pulsé. Notons aussi que les grandeurs que l'on observe en simulation (densité) et expérimentalement (intensité lumineuse) sont différentes. Il se pourrait que la densité électronique expérimentale se trouve bien concentrée sur une dimension de $\lambda_c/2$. Mais il est surtout fort possible que cette dimension soit liée à la présence du monopole, au bout duquel le champ électrique est considérablement augmenté par un effet de pointe. Le plasma se créerait alors sur le bout du monopole, et son expansion dépendrait juste de la puissance injectée et de la pression. Nous illustrerons d'ailleurs l'influence de la pression sur la dimension dans la suite de cette partie.

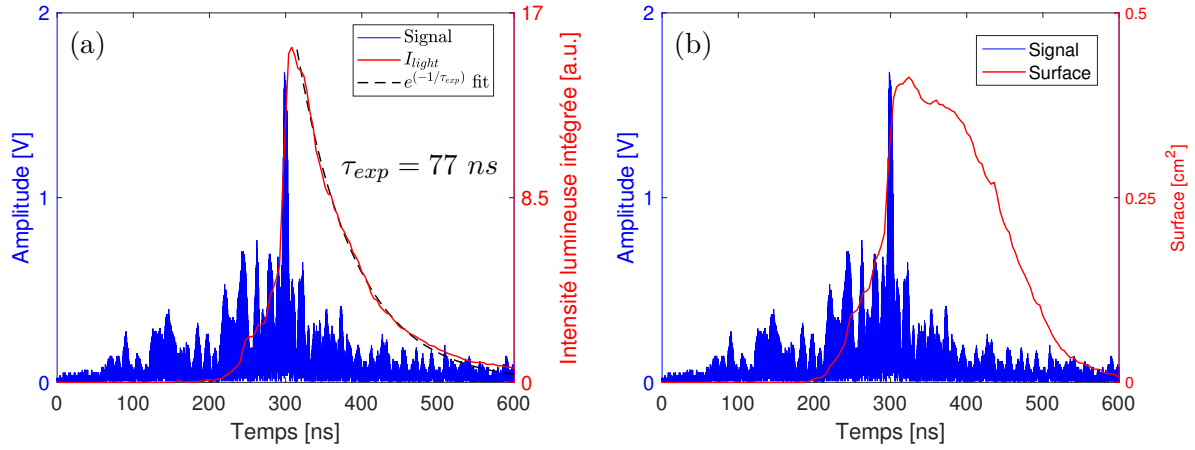


FIGURE 4.24 – Évolution temporelle de l'intensité lumineuse intégrée (a) et de la surface occupée par le plasma (b) avec la valeur absolue de la tension. La surface est définie comme la somme des pixels avec une intensité lumineuse supérieure au niveau seuil ($0.3 \cdot 10^{-3}$ u.a.).

Les évolutions temporelles de l'intensité lumineuse intégrée sur tous les pixels et la surface apparente plasma sont observées autour du pic de RT. L'intensité lumineuse intégrée sur tous les pixels, tracée avec la valeur absolue de la tension de la Figure 4.24(a), montre que l'intensité lumineuse augmente brutalement au moment du pic de RT jusqu'à un maximum trouvé à 308 ns, *i.e.* à la fin du pic de RT. Après ce pic, l'intensité lumineuse décroît en suivant une exponentielle décroissante (légèrement modulée par la présence des lobes secondaires temporels succédant le pic de RT), comme le montre le fit exponentiel tracé en pointillés sur la Figure 4.24(a) et nommé dans la suite la “post-décharge immédiate”. Pour une pression de $p = 1.5$ torr et une fréquence de répétition de $f_{rep} = 6$ kHz, ce temps de décroissance est estimé à 77 ns. Finalement, une observation plus attentive montre que le plasma est initié avant le pic de RT, par exemple autour de $t = 252$ ns, *i.e.* 50 ns avant, période dite “post-décharge tardive”. En effet, entre 220 et 300 ns, l'intensité lumineuse est clairement corrélée avec les lobes secondaires temporels précédant le pic de RT. Le même comportement est relevé sur la

surface plasma tracée sur la 4.24(b) : l'étalement spatial est fortement corrélé à la forme du signal mesuré sur le monopole et est principalement causé par le pic de RT, avec une légère augmentation avant ce pic due aux lobes secondaires temporels. Néanmoins, la surface ne décroît pas directement à la fin du pic de RT (308 ns) mais quelques dizaines de ns après avec quelques “rebonds” dus aux lobes secondaires. Ces deux périodes (“post-décharge immédiate” et “post-décharge tardive”) sont illustrées sur la Figure 4.25. Notons que la “post-décharge tardive” aurait aussi être pu nommée “pré-décharge” par rapport à la décharge suivante. Les phénomènes ayant lieu durant ces deux périodes sont discutés dans la partie suivante.

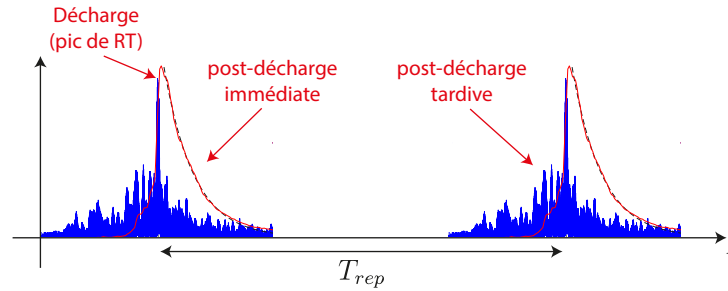


FIGURE 4.25 – Schéma illustrant les périodes nommées “post-décharge immédiate” et “post-décharge tardive”.

4.3.3.2 Étude en pression

La Figure 4.26(a) montre les images d'intensité lumineuse du plasma à la fin du pic de RT (308 ns) pour une fréquence de répétition de $f_{rep} = 6$ kHz pour différentes pressions.

La pression a clairement un effet sur la décharge : quand la pression augmente, le plasma devient de plus en plus contracté au bout du monopole. Ce comportement est typique d'un plasma contrôlé par la diffusion. Dans ce cas c'est l'augmentation de la fréquence de collision électron - neutre avec la pression qui explique ce phénomène. Il est donc raisonnable de supposer qu'un processus similaire a lieu durant les plasmas transitoires amorcés par RT. La dimension du plasma décroît ainsi en augmentant la pression. Puisque le plasma est localisé au niveau du bout du monopole où le champ électrique est le plus élevé, l'intensité lumineuse maximale émise par le plasma augmente aussi.

L'évolution temporelle de l'intensité lumineuse intégrée tracée sur la Figure 4.26(b) démontre que le pression joue un rôle clé. Tout d'abord, cela montre qu'en plus d'une contraction du plasma avec la pression, le maximum d'intensité lumineuse intégrée décroît avec la pression. La post-décharge est toujours caractérisée par une décroissance exponentielle avec une temps de décroissance de l'ordre de quelques dizaines de ns. Ce temps de décroissance décroît avec la pression de 100 ns à 0.5 torr jusqu'à 65 ns à 4 torr, suivant une évolution en $1/p$ (Figure 4.26(c)). Pour finir, avant le pic de RT, l'intensité lumineuse est corrélée avec les lobes secondaires dans toutes les conditions.

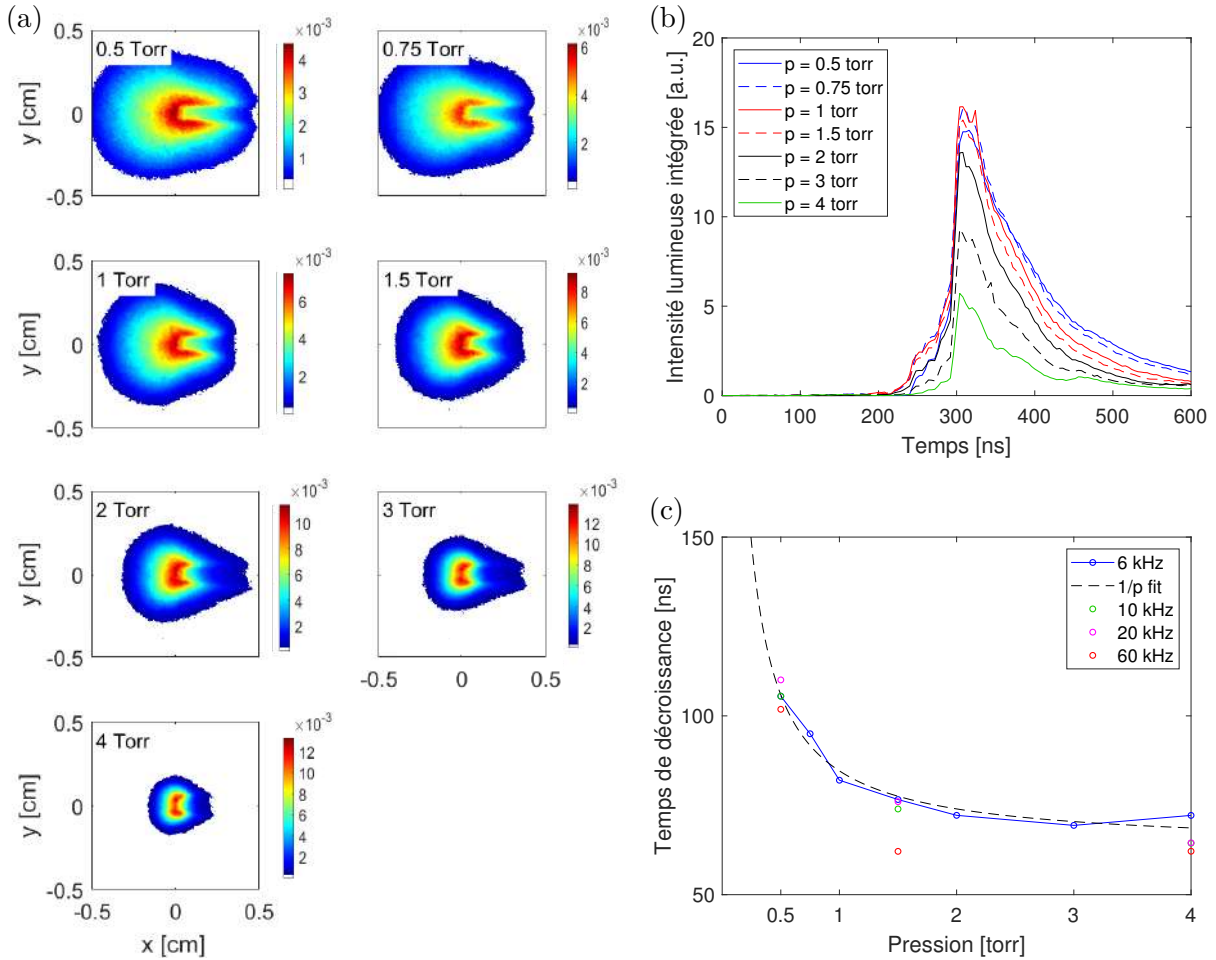


FIGURE 4.26 – (a) Images de l'intensité lumineuse du plasma par imagerie rapide juste après l'instant du pic de RT (308 ns) pour différentes pressions (avec une fréquence de répétition de $f_{rep} = 6$ kHz). (b) Évolution temporelle de l'intensité lumineuse intégrée pour différentes pressions (avec une fréquence de répétition de $f_{rep} = 6$ kHz). (c) Les temps de décroissances correspondant (ainsi que ceux obtenus pour différentes fréquences de répétitions).

4.3.3.3 Influence de la fréquence de répétition

Dans cette partie l'influence de la fréquence de répétition sur la dynamique de la décharge est étudiée. Son influence sur la décharge est particulièrement significative comme illustré sur la Figure 4.27 pour deux fréquences de répétition et trois pressions. Les courbes en traits pointillés correspondent au cas précédent avec une fréquence de répétition de 6 kHz alors que celles en traits pleins correspondent à une fréquence de répétition de 60 kHz. Pour une certaine pression, l'intensité lumineuse intégrée augmente avec la fréquence de répétition. C'est encore plus marqué à 4 torr. En outre, le temps de décroissance de l'intensité lumineuse avec le pic de RT reste presque identique dans la gamme de fréquence 6 - 60 kHz comme récapitulé sur la Figure 4.26(b) (Figure du bas). Ce temps de décroissance semble être principalement

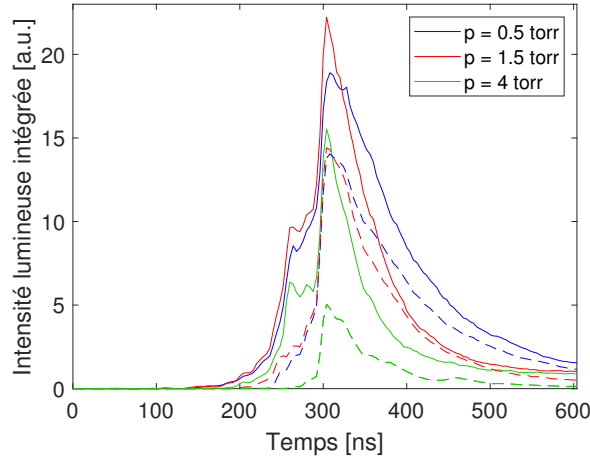


FIGURE 4.27 – Évolution de l'intensité lumineuse intégrée pour deux fréquences de répétition, $f_{rep} = 6$ kHz en traits pointillés et $f_{rep} = 60$ kHz en traits pleins, et pour trois pressions $p = 0.5$ torr (bleu), $p = 1.5$ torr (rouge) et $p = 4$ torr (vert).

conditionné par la pression. Pour finir, l'émission lumineuse précédant le pic de RT est aussi considérablement affectée par une augmentation de fréquence de répétition. À 60 kHz, son niveau peut atteindre jusqu'à 40 % du niveau maximal d'intensité, suggérant un effet mémoire entre deux pics successifs de RT. Nous discuterons dans la partie qui suit de la physique de ces décharges et notamment de cet effet mémoire.

4.3.4 La physique des plasmas amorcés par retournement temporel

Afin de décrire et d'analyser la production de plasmas par RT, la décharge pulsée sera étudiée à partir d'un pic de RT (300 ns) jusqu'au suivant.

4.3.4.1 Pic de RT

Comme cela a déjà été mentionné, la durée du pic de RT est assez inhabituelle pour les sources plasmas microondes. En effet, la durée du dépôt de puissance microonde au niveau du monopole est de $\tau = 8$ ns avec un temps de montée inférieur à $t_r = 1$ ns. Avec la technique du RT l'énergie électromagnétique injectée dans la cavité est focalisée en espace et en temps au moment du pic de RT. Ce pic se situe à 300 ns sur les Figures de cette partie. Il est spatialement concentré sur le monopole. Il déclenche une importante émission lumineuse due à la formation d'une décharge plasma. Pour tous les résultats présentés ici, la densité électronique reste inconnue. La seule information expérimentalement disponible est la différence observée sur la tension mesurée en sortie du monopole après le pic de RT avec et sans plasma. Cela suggère que la densité électronique change soit la sensibilité du monopole, *i.e.* la relation entre le niveau du champ électrique au niveau du monopole et la tension mesurée, soit la propagation des ondes électromagnétiques proches du monopole ou bien les deux.

4.3.4.2 Post-décharge immédiate

À la fin du pic de RT, l'émission de l'intensité lumineuse décroît exponentiellement, seulement légèrement modulée par les lobes secondaires. Sur la Figure 4.26(c), ce temps de décroissance est estimé inversement proportionnel à la pression. En supposant que l'émission lumineuse vient principalement de la décroissance spontanée des états $Ar\ 2p_i$ et que ces niveaux sont majoritairement excités par impact d'électrons sur le niveau fondamental des atomes Ar , l'émission totale peut s'écrire :

$$I_{Total} = \sum_i \frac{hc}{\lambda_i} f(\lambda_i) k_i(T_e) n_e n_{Ar} \quad (4.28)$$

où h est la constante de Planck, c la vitesse de la lumière, λ_i la longueur d'onde de la transition $Ar\ 2p_i - 1s_j$, $f(\lambda_i)$ la réponse optique du système de détection à la longueur d'onde λ_i , $k_i(T_e)$ le taux de réaction de l'excitation par impact d'électron des états $Ar\ 2p_i$ à partir de l'état fondamental des atomes Ar (dépendant de la température électronique T_e), n_e la densité électronique, et n_{Ar} la densité d'atomes dans l'état fondamental Ar . À partir de l'équation 4.28, la décroissance de l'émission lumineuse totale durant la post-décharge immédiate pour une certaine pression pourrait être due à une diminution soit de la densité électronique soit de la température électronique. Nous regardons maintenant en détails ces deux possibles causes.

Dans un plasma contrôlé par la diffusion, les pertes de particules chargées sont gouvernées par la diffusion libre ou ambipolaire en fonction de la densité électronique (diffusion libre pour des décharges faiblement ionisées et ambipolaire lorsqu'elles sont fortement ionisées, comme nous l'avons vu dans le chapitre 1). Dans ces conditions, la décroissance de la densité électronique est liée au temps de diffusion caractéristique qui peut être estimé par :

$$\tau_{diff} = \frac{\Lambda}{D_{diff}^2} \quad (4.29)$$

où Λ est la longueur de diffusion et D_{diff} est le coefficient de diffusion libre ou ambipolaire. Dans des géométries sphériques, la longueur de diffusion s'écrit $\Lambda = L/\pi$, où L est la taille caractéristique du réacteur plasma [Rai91]. Dans notre cas, il n'y pas de corrélation entre la dimension du plasma et celle du réacteur. En conséquence, L est pris égal à un rayon typique du plasma, *i.e.* $L = 5$ mm. En considérant une température électronique T_e d'environ 2 eV dans un plasma d'argon à 1.5 torr, le temps caractéristique est de quelques dizaines de nanosecondes en régime de diffusion libre (plasma à faible densité) et de quelques dizaines de millisecondes en régime de diffusion ambipolaire (plasma à forte densité). Cependant, dans ces deux cas, en augmentant la pression le coefficient de diffusion D_{diff} décroît se traduisant par un temps de décroissance de la densité électronique qui augmente (équation 4.29), ce qui est incompatible avec les données présentées sur la Figure 4.26(c).

La relaxation de l'énergie des électrons après la fin du pic de RT peut aussi induire une

décroissance de l'émission lumineuse totale durant la post-décharge immédiate. En négligeant les phénomènes de dépôt de puissance et de transport de chaleur, l'équation du bilan d'énergie pour les électrons peut s'écrire comme [LL05] :

$$\frac{dT_e}{dt} = -n_n \left(\sum_j k_j(T_e) \epsilon_j \right) \quad (4.30)$$

où n_n est la densité de neutre, $k_j(T_e)$ est le taux de réaction pour les collisions électron-atome (collisions élastiques et inélastiques) et ϵ_j est la perte d'énergie des électrons par collision électron-atome. Ici, la somme prend en compte le moment du transfert des collisions des électrons ainsi que l'excitation par impact d'électron et les processus d'ionisation [LL05]. À partir des taux de réaction correspondant, cette équation différentielle non-linéaire peut être résolue numériquement en utilisant la densité d'atomes d'argon dans l'état fondamental et la valeur initiale de la température électronique T_{e0} comme paramètres d'ajustement. Il est utile de souligner le fait que dans les plasmas d'argon avec des températures électroniques de quelques eV, la relaxation de l'énergie des électrons est dominée par les processus d'excitation et d'ionisation. En considérant une température électronique de 2 eV d'argon à 1.5 torr, $T_e(t)$ et donc $I_{Total}(t)$ peuvent être calculées en utilisant respectivement les équations 4.30 et 4.28. Dans ces conditions, en utilisant le jeu de sections efficaces décrit dans [DJDS19], le temps de décroissance caractéristique de l'intensité de l'émission totale est autour de 100 ns, comparable aux données expérimentales. De plus, l'équation 4.30 indique que le temps de relaxation de l'énergie des électrons est inversement proportionnel à la densité d'atomes d'argon dans l'état fondamental, en très bon accord avec les données tracées sur la Figure 4.26(c). Par conséquent, dans les conditions expérimentales étudiées durant cette thèse, on peut en déduire que les décroissances exponentielles de l'intensité totale d'émission observées durant la post-décharge immédiate sont principalement dues à la relaxation de l'énergie des électrons.

4.3.4.3 Post-décharge tardive

Entre deux pics de RT, une fois que T_e a relaxé comme discuté dans le paragraphe précédent, nous pouvons supposer une évolution post-décharge classique, avec des densités d'espèces chargées et de métastables évoluant ensemble. Les mesures montrent clairement l'existence d'un effet mémoire résultant des décharges précédentes et influençant le développement de la suivante. Cela peut être dû à la présence de charges résiduelles et/ou d'atomes métastables. En supposant une température électronique de 300 K, une nouvelle estimation des temps caractéristiques de diffusion des particules chargées peut être réalisée. Pour ces températures électroniques très basses, on considère un taux de collision de $2 \times 10^5 \text{ m}^3 \text{ s}^{-1}$ [LL05]. Les temps de diffusion caractéristiques résultant pour une pression de 1 torr sont de 20 ns et 500 μs respectivement pour le temps de diffusion libre et ambipolaire. Ainsi, si le plasma est suffisamment dense de sorte à ce que le processus ambipolaire gouverne la diffusion des charges libres, une quantité non négligeable de charges peut rester dans le milieu jusqu'à la

prochaine impulsion. La recombinaison électron-ion peut aussi avoir un effet. Le processus de recombinaison dissociative (RD) $Ar_2^+ + e \rightarrow Ar^* + Ar$ joue généralement un rôle majeur au tout début de la post-décharge alors que le processus de recombinaison collisionnel-radiatif (CR) $Ar^+ + e + e \rightarrow Ar^* + e$ domine sur des temps plus long quand la température électronique devient très basse [CDK15]. Leurs coefficients respectifs peuvent être trouvés dans [MB68] pour les recombinaisons (RD) et [BVY87] pour les recombinaisons (CR). En fonction de la densité plasma, des calculs simples indiquent que les électrons restent dans le milieu pendant des douzaines de millisecondes après le pic de RT. Néanmoins, il est intéressant de noter que ces processus ne résultent pas seulement des pertes des particules chargées mais aussi de la production d'atomes d'argon métastables. Les atomes métastables provenant soit de la décharge soit du processus de recombinaison peuvent jouer un rôle important dans l'effet mémoire comme montré par exemple par Carbone *et al.* [Car+14]. En se basant sur [Ste+14], le temps de vie de l'état métastable $Ar(^3P_2, 1s_5)$ est essentiellement déterminé par le taux de la désexcitation de deux et trois corps avec les neutres d'argon. En effet, la faible température électronique rend peu probable les processus impliquant des électrons énergétiques comme l'excitation vers des niveaux supérieurs ou le mélange des métastables et des états résonants $1s$ de l'argon [CDK15]. Il en résulte un temps de vie de $Ar(^3P_2, 1s_5)$ à 1.5 torr autour de 10 ms. Ainsi, pour $f_{rep} \geq 100$ Hz, les métastables peuvent grandement faciliter le claquage avant le pic suivant par des processus d'ionisation par étapes.

La présence de charges résiduelles et d'atomes métastables dans la post-décharge tardive est donc probablement responsable de l'effet mémoire observé. Cela explique l'émission lumineuse observée avant le pic de RT durant les lobes secondaires et la forte dépendance du comportement du claquage par rapport aux conditions opérationnelles telle que la fréquence de répétition. Ce comportement peut être qualifié d'ionisation à faible champ électrique.

CHAPITRE 4 : Ce qu'il faut retenir

L'objectif de ce chapitre était la présentation des résultats expérimentaux obtenus durant cette thèse. Ces derniers constituent les premiers plasmas amorcés et contrôlés par RT. Les propriétés de cette source plasma novatrice sont rappelées dans la liste suivante :

- Le RT permet un contrôle efficace de la position des plasmas dans la cavité.
- La position des plasmas est **indépendante du design de la cavité** et seulement contrôlée par la forme d'onde transmise à la cavité.
- Les temps de montée et les rapports cycliques sont très faibles (respectivement < 1 ns et < 0.05 %).
- Le contrôle des plasmas est aussi temporel : **contrôle spatio-temporel des plasmas**.
- Même à ces faibles rapports cycliques, un effet mémoire existe entre une impulsion et la suivante.
- La pression joue un rôle important sur la dimension des plasmas et sur l'influence qu'ils ont sur le comportement des ondes dans la cavité.
- La possibilité d'amorcer simultanément plusieurs plasmas à différentes positions pourrait être intéressante pour le traitement de surface et pour le développement de milieux modulés spatio-temporellement.
- La présence des initiateurs joue certainement un rôle important dans la structure spatiale de ces plasmas, nous discuterons dans les perspectives de cette thèse de stratégies possibles pour le décrochage du plasma de ces initiateurs.

Conclusion et perspectives

Conclusion

Dans le premier chapitre de cette thèse, nous avons présenté le contexte et les enjeux autour des sources plasma microonde. Nous avons commencé par décrire la physique des plasmas froids hors équilibre d'argon. Ensuite, les techniques standards de génération de plasmas microondes en cavité ont été présentées. Les limitations de ces techniques ont ainsi été identifiées. La limitation la plus restrictive est la dimension des objets à traiter. L'objectif de cette thèse est alors le développement d'une source plasma innovante et originale qui permet le contrôle des plasmas dans les cavités de grandes dimensions, contrôle actuellement impossible. Ainsi, cette nouvelle source pourrait révolutionner certaines applications dans le domaine des plasmas microondes et des procédés associés. Dans le premier chapitre de cette thèse, nous avons identifié et rassemblé selon trois axes principaux les défis que constituent le développement de cette source innovante. Ils sont rappelés ici, en mettant en évidence les défis auxquels nous avons répondu :

- ✓ 1) Études théoriques.
 - Chap.2 (a) Comprendre le RT en présence de milieux linéaires ou non-linéaires (plasmas).
 - Chap.3 (b) Identifier les propriétés du système qui fixent les performances en amorçage et en contrôle des plasmas.
- ✓ 2) Premières études expérimentales de la source avec du matériel générique.
 - (a) Proposer une architecture de banc expérimental permettant, à la fois la manipulation de formes d'ondes complexes pour le RT et le contrôle de l'environnement (gaz, pression) pour l'amorçage des plasmas.
 - Chap.4 (b) Choisir des initiateurs pour abaisser les niveaux de puissance requis et mesurer les signaux nécessaires.
 - (c) Contrôler l'amorçage et l'entretien (pulsé) des plasmas par RT sur des initiateurs.
 - (d) Étudier la physique (ns) de ce nouveau type de décharge.
- 3) Développement de la source (pinceau plasma ondulatoire) : vers du matériel dédié.
 - (a) Développer un prototype de source autorisant une étude expérimentale dans des conditions optimales.
 - (b) Décrocher le plasma de l'initiateur puis contrôler le plasma partout dans la cavité.
 - Chap.4 (c) Étudier l'influence de la présence d'un plasma dans la cavité sur le RT (non linéaire).
 - (d) Étudier les procédés qui pourront exploiter cette nouvelle source.

Le premier défi que nous avons identifié consistait en une étude théorique. Il était bien sûr nécessaire de comprendre la théorie avant de développer la source plasma à RT. Nous avons

proposé dans le second chapitre de cette thèse une description de la théorie autour du RT. Nous avons alors discuté de trois concepts fondamentaux dans le domaine de l'électromagnétisme : la causalité, la réversibilité et la réciprocité. Nous avons vu que le RT ne viole pas la causalité, qu'il est basé sur la réversibilité et, qu'en pratique, il est plus proche d'une expérience de réciprocité. Puis, nous avons discuté de la théorie autour du retournement temporel appliqué au contrôle plasma. Le point central est l'introduction d'un phénomène non-linéaire (claquage plasma) dans le processus de RT, qui lui est basé sur une physique linéaire. Ensuite, nous avons développé dans le troisième chapitre les équations permettant de décrire le RT en cavité réverbérante selon son régime de recouvrement modal. Pour illustrer ces développements et conclure sur les propriétés du système qui fixent les performances en amorçage et en contrôle des plasmas, nous avons utilisé un outil de simulation FDTD 2D spécialement développé à cet effet. Ensuite, le caractère essentiel de la chaotité sur le contrôle des plasmas par RT a été démontré en couplant les équations de Maxwell à un modèle plasma fluide. Finalement, les propriétés du système qui fixent les performances en amorçage et en contrôle des plasmas sont clairement identifiées et rassemblées dans un tableau récapitulatif en fin de chapitre.

Le second défi que nous avons identifié portait sur la première démonstration expérimentale de la source plasma avec du matériel générique. Afin de répondre à ce défi, nous avons développé dans le chapitre 4 un dispositif expérimental complexe, permettant à la fois la manipulation de formes d'ondes complexes pour le RT et le contrôle de l'environnement (gaz, pression) pour l'amorçage des plasmas. De plus, une méthode fréquentielle a été développée durant cette thèse. Elle permet une optimisation du RT sur le dispositif en rendant possible des études paramétriques numériques. De cette façon, le RT obtenu permet l'amorçage et le contrôle de la position des plasmas dans la cavité. Nous avons vu que leurs positions sur les initiateurs sont contrôlées par les formes d'ondes transmises à la cavité. Nous avons remarqué au passage la possibilité d'amorcer simultanément plusieurs plasmas sur les différents initiateurs. Puis nous avons décrit la physique de ces plasmas. Nous avons vu que l'originalité de cette nouvelle source plasma pulsée ne vient pas seulement de l'indépendance entre la position du plasma et le design de la cavité, mais aussi de ses faibles temps de montée (< 1 ns) et faibles rapports cycliques (typiquement 0.05 %). Nous avons ensuite étudié l'influence de la présence d'un plasma sur le comportement des ondes dans la cavité. Puis, lors des dernières discussions du chapitre, nous avons illustré un effet mémoire, dû à la présence de charges résiduelles (particules chargées et/ou métastables).

Le troisième défi que nous avons identifié consiste à développer le pinceau plasma ondulatoire. Pour cela, les études théoriques réalisées permettent le développement d'un prototype de source autorisant une étude expérimentale dans des conditions optimales. Nous avons, entre autres, identifié dans ce troisième défi l'étude de l'influence de la présence d'un plasma dans la cavité sur le RT. Dans le chapitre 4 de cette thèse, nous avons présenté un dispositif expérimental et une méthode permettant cette étude et nous avons proposé des pistes d'interprétation. Des études plus poussées seraient cependant souhaitables pour vraiment appréhender et interpréter les mécanismes à l'œuvre. Pour finir, le point le plus important est le décrochage du plasma de l'initiateur, puis le contrôle du plasma partout dans la cavité. Nous proposons dans les perspectives de cette thèse des pistes pour atteindre cet objectif.

Perspectives

I - La physique des plasmas amorcés par retournement temporel : Études complémentaires

Nous avons discuté, tout au long de cette thèse et notamment lors des résultats expérimentaux du chapitre 4, de la physique des plasmas amorcés par RT. Nous donnons ici deux pistes pour aller plus loin dans la compréhension des mécanismes physiques.

a - Le plasma : Un puits à retournement temporel électromagnétique ?

Lors d'une expérience pratique de RT, pour se rapprocher d'une vraie expérience de réversibilité restreinte, il faut aussi inverser la source de la phase initiale lors de la phase inverse (comme illustré sur la Figure 2.9). En effet, si cette opération n'est pas réalisée, les ondes convergentes continuent à se propager après la focalisation. Au moment de la focalisation, il y a donc des interférences entre une onde convergente et une divergente. Il en résulte une dimension de focalisation de $\lambda/2$. Cette opération de réversibilité sans inversion de la source est présentée sur la Figure P.1. Les ondes convergent vers la source puis divergent, la partie avant la focalisation étant symétrique par rapport à celle qui lui succède.

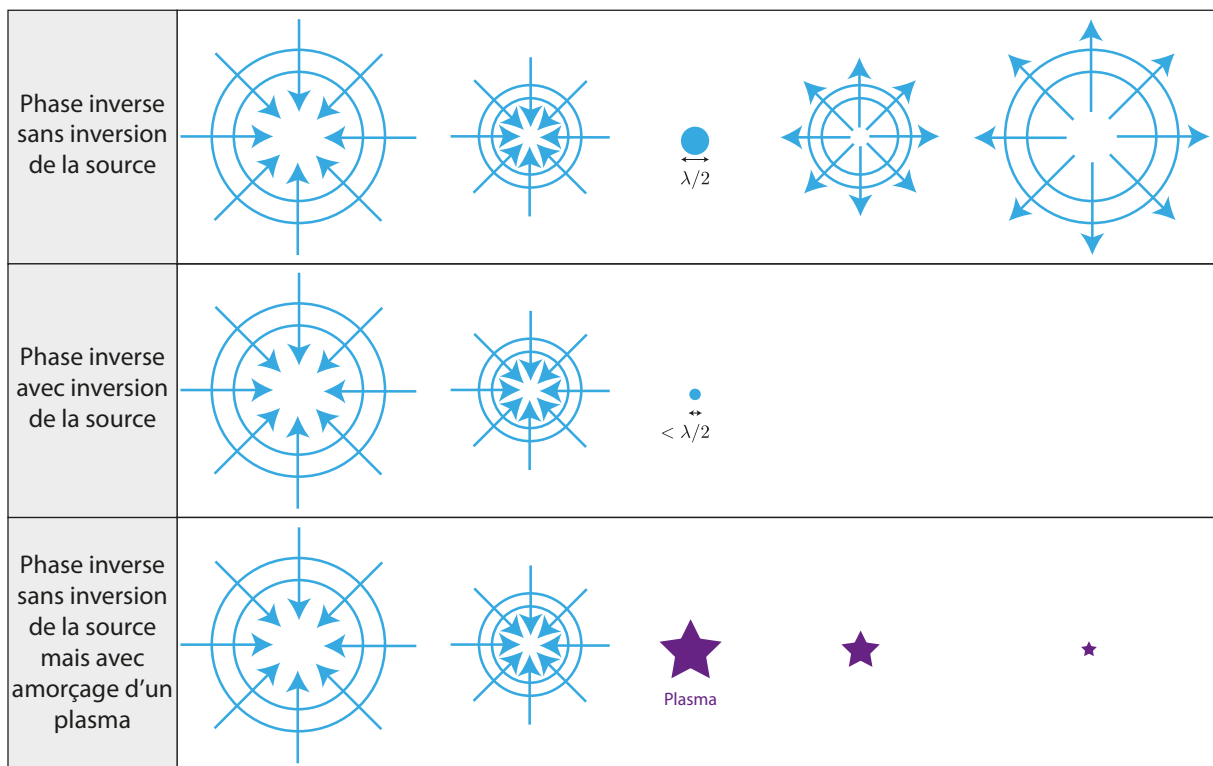


FIGURE P.1 – Illustration du puits à RT électromagnétique par plasma.

Si la source est inversée au moment de la focalisation comme illustré sur la Figure P.1, la

dimension de la tâche focale correspond alors à la dimension spatiale de la source de la phase initiale. On parle alors de puits actif à RT. Cela a été démontré en acoustique par Julien de Rosny [RF01]. Dans ce cas l'onde divergente peut être annulée. Le point source peut être vu comme parfaitement adapté avec l'onde convergente, c'est comme si à l'endroit de la source l'énergie était aspirée par la source, d'où l'idée de puits.

Une autre option pour diminuer la dimension de la focalisation consiste à utiliser, dans le domaine microonde, des métamatériaux placés au niveau de la source (et à l'endroit de la focalisation) [Ler+07]. Ce sont des structures périodiques dont la périodicité est inférieure à la longueur d'onde. Avec ce genre de structure, lors de la phase initiale une partie des ondes qui étaient évanescentes (et donc non propagatives) deviennent des ondes propagatives. Ces ondes possèdent des variations spatiales plus brutales que celles qui se propagent dans le cas classique, et donc lors de la phase inverse, la focalisation peut se faire sur des dimensions plus fines. On trouve aussi des travaux proposant des puits passifs. Par exemple dans [NPC13], un puits passif large bande est proposé dans le domaine optique. Le principe consiste à façonner la fonction de dispersion diélectrique d'une nanoparticule de sorte à ce qu'elle soit totalement absorbante pour une certaine impulsion incidente. Ces travaux sont proches des études en imagerie optique autour des "lentilles fish eye", dont le but est de focaliser sur des dimensions sub-longueur d'onde en façonnant l'indice de réfraction [Leo09] ; [QTMTH12]. On trouve des travaux similaires dans le domaine microonde [Loo+12].

Cependant, ces deux options pour réduire la dimension de la focalisation sont restrictives. En effet, la première nécessite la présence d'un élément actif à l'endroit de la focalisation, en plus de sa synchronisation avec le dispositif. La seconde nécessite la présence d'une structure plus ou moins encombrante à l'endroit de la focalisation, agissant comme un puits (on parle aussi de façon équivalente de syphon, "drain" en anglais). L'amorçage d'un plasma au moment de la focalisation pourrait être une solution pour supprimer les ondes divergentes post-focalisation et ainsi obtenir des dimensions de focalisation inférieures à la demi longueur d'onde sans structure encombrante ou élément actif au niveau de la focalisation. Dans ce cas, selon le même principe que les puits passifs mentionnés précédemment, les ondes convergentes sont absorbées par le plasma au moment et à l'endroit de la focalisation, faisant penser à un puits. Des études supplémentaires sont nécessaires pour vérifier ce comportement.

b - Spectrométrie résolue en temps

Nous avons vu que les propriétés des électrons (densité et température) jouent un rôle clé dans la physique des plasmas amorcés par RT. Il serait donc intéressant de pouvoir mesurer ces propriétés. La mesure de ces propriétés avec des diagnostics non intrusifs a toujours été un grand enjeu. Les deux plus grandes techniques utilisées sont celles utilisant la diffusion Thomson [MUB98] et celles basées sur la spectrométrie d'émission optique [ZP08]. Les premières sont plus compliquées à mettre en place puisqu'elles nécessitent un laser éclairant le plasma. Les secondes ne nécessitent que l'utilisation d'un spectromètre mesurant la lumière émise par le plasma. En effet, chaque longueur d'onde émise correspond à une transition précise et il est alors possible, à partir des spectres d'émissions mesurés, d'évaluer les propriétés du plasma en

utilisant un modèle collisionnel-radiatif [DJDS19]. La technique consiste à faire correspondre les spectres calculés par le modèle avec les spectres mesurés.

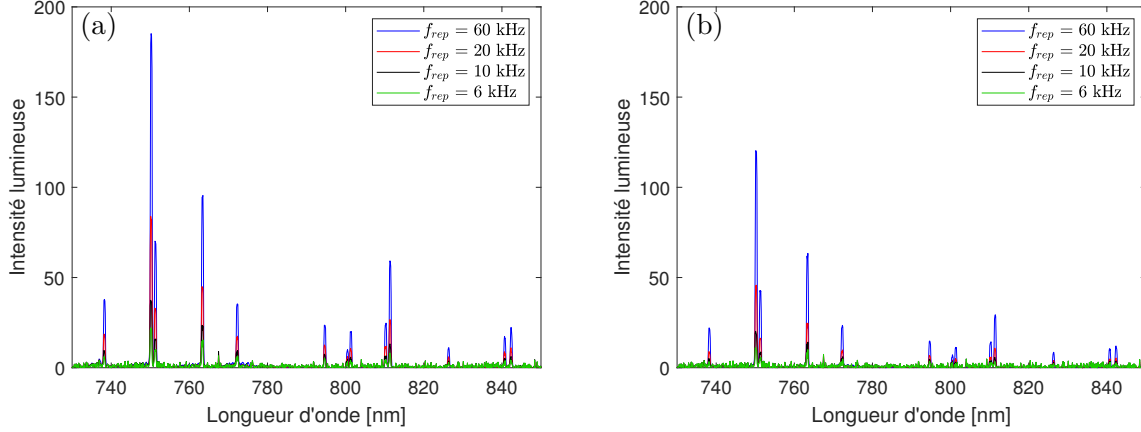


FIGURE P.2 – Spectres d’émissions des plasmas amorcés par RT autour des raies d’argon. (a) $p = 1.5$ torr. (b) $p = 3$ torr.

La Figure P.2 montre les spectres d’émissions des plasmas amorcés par RT présentés dans le chapitre 4. Le spectromètre était positionné devant le hublot faradisé permettant de voir à l’intérieur de la cavité. Sur la Figure P.2, les spectres sont montrés autour des raies d’argon pour différentes fréquences de répétition et pour deux pressions, 1.5 et 3 torr. Il est alors possible, en utilisant un modèle collisionnel-radiatif de l’argon tel que celui développé par Antoine Durocher-Jean, Edouard Desjardins et Luc Stafford dans [DJDS19], d’obtenir certaines propriétés des plasmas amorcés par RT pour différentes configurations (fréquences de répétition et pressions). En faisant cela, avec les spectres que nous avons mesurés, Luc Stafford et ses collègues ont trouvé par exemple des températures électroniques variant entre 1 et 2 eV, selon la configuration.

Cependant, ces spectres ont été obtenus en intégrant la lumière émise sur plusieurs impulsions de RT. C’est à dire que la lumière mesurée par le spectromètre est une moyenne de la lumière émise par plusieurs décharges plasma. Or, nous avons vu, à l’aide des mesures faites avec une caméra rapide, que l’intensité lumineuse varie de façon substantielle autour du pic de RT. Ainsi, ces spectres ne permettent pas une description fine de la physique des plasmas amorcés par RT. Pour ce faire il faut des mesures d’émission lumineuse résolues en temps tel que cela est fait par exemple dans [Car+14], afin de pouvoir suivre l’évolution temporelle des différentes raies autour du pic de RT. Un spectromètre permettant des mesures d’évolution temporelle doit donc être mis en place et synchronisé avec le dispositif expérimental. Cette étude fait l’objet d’un sujet de Master 2 qui débutera en 2021 au laboratoire LAPLACE (sur le même dispositif expérimental que celui présenté dans le chapitre 4 de cette thèse), co-encadré par Laurent Liard de l’université Paul Sabatier et Luc Stafford de l’université de Montréal.

II - Vers un pinceau plasma ondulatoire

Nous avons entre autres identifié dans le troisième défi le décrochage du plasma de l'initiateur. **Cette étape est incontournable et essentielle pour le développement du pinceau plasma ondulatoire.** Le point le plus compliqué réside dans l'obtention de l'information de la propagation des ondes (réponse impulsionnelle) à des endroits dans la cavité hors des initiateurs. En effet, si l'initiateur (monopole dans nos expériences) est enlevé de la cavité entre la phase initiale et la phase inverse du RT, le comportement des ondes change. En conséquence, la focalisation par RT n'a pas lieu. Il est donc nécessaire de faire appel à des méthodes plus élaborées pour atteindre cet objectif. Pour terminer ce manuscrit, quelques pistes envisageables sont proposées.

La sonde non intrusive. C'est sûrement la solution qui semble la plus évidente pour résoudre ce problème. L'idée est d'utiliser une sonde qui soit invisible aux fréquences microondes d'intérêt. Une technique basée sur les sondes non intrusives a été développée en acoustique avec un laser (non-intrusif) [STJ04]. Dans le domaine microonde, les sondes les moins intrusives sont de type électro-optique, ce qui les rend onéreuses. Elles ont été utilisées dans des applications de RT par exemple par Andréa Cozza et Florian Monsef pour le développement du RT généralisé, permettant de façonner des fronts d'onde en cavité [CM17]. Avec une sonde non intrusive, lors de la première phase du RT on pourrait mesurer l'information de la propagation des ondes entre l'antenne utilisée pour l'émission de la puissance (antenne de l'annexe dans nos expériences) et n'importe quel point de la cavité et notamment pour des points entourant un initiateur (monopole dans nos expériences). La sonde non intrusive peut alors être enlevée de la cavité entre les deux phases du RT sans modifier le comportement des ondes. Lors de la seconde phase, un plasma peut tout d'abord être amorcé sur l'initiateur, suivant la même méthode que celle suivie dans le chapitre 4 de cette thèse. De cette façon, le niveau du champ de claquage est diminué dans la cavité et notamment proche de l'initiateur. Il faudrait alors focaliser les ondes sur les points proches de l'initiateur, de sorte à tirer le plasma et à le décrocher. Ensuite, on pourrait alors imaginer contrôler le plasma partout dans la cavité. Notons pour finir que le champ électrique sera toujours fortement augmenté au niveau de l'initiateur. Des techniques plus élaborées que le RT pourraient alors être utilisées pour simultanément maintenir un champ assez faible sur l'initiateur et focaliser ailleurs dans la cavité. Une discussion sur ces techniques fait l'objet du point suivant.

Les méthodes de contrôle des ondes élaborées. Nous avons montré dans cette thèse que le RT est un bon candidat pour le développement de la nouvelle source plasma que nous proposons. Il permet en effet de focaliser à un instant et à un endroit les ondes dans des cavités. Cependant, il ne permet pas de contrôler les ondes ailleurs dans la cavité et à d'autres moments. Il serait alors intéressant d'utiliser des techniques de contrôle des ondes plus élaborées. Pour répondre à cet objectif la méthode de combinaison linéaire de la configuration de champ (en anglais "Linear Combination of Configuration Field", LCCF) développée par Jaume Benoit, Cédric Chauvière et Pierre Bonnet [BCB12] ; [BCB15] est

une très bonne candidate. Elle permet de déterminer la source temporelle à transmettre à la cavité qui donne un champ électromagnétique préalablement défini à une ou plusieurs positions sur un certain intervalle temporel. Tout d'abord, comme lors d'une expérience de RT, la réponse impulsionnelle est mesurée entre le point d'émission de la source (antenne de l'annexe dans nos expériences) et les positions auxquelles on souhaite contrôler le champ. Ensuite, la méthode LCCF trouve une combinaison linéaire de ces réponses pour calculer la source désirée. L'émission de cette source donne l'évolution temporelle du champ électromagnétique souhaitée aux positions souhaitées. Nous avons montré, en collaboration avec Ali Al Ibrahim, Cédric Chauvière et Pierre Bonnet de l'université Clermont Auvergne, que la méthode LCCF permet d'éviter des décharges parasites dans la cavité, tout en permettant l'amorçage d'un plasma à l'endroit souhaité [Maz+20a]. De la même façon, cette méthode serait aussi très efficace pour minimiser le champ au niveau de l'initiateur lorsque l'on souhaite déplacer le plasma dans la cavité. Enfin, en utilisant des techniques de contrôle des ondes élaborées telles que la LCCF, il serait aussi peut être possible d'interpoler l'information de la propagation des ondes entre deux initiateurs afin de déplacer le plasma entre ces derniers.

La simulation. Le principe serait de simuler la cavité expérimentale de sorte à obtenir la réponse impulsionnelle entre n'importe quels points choisis arbitrairement dans la cavité. On pourrait pour cela penser à utiliser un code de simulation FDTD tel que celui utilisé dans cette thèse. Cependant, dans ce cas l'influence de la dispersion numérique doit être négligeable pour obtenir une information fidèle. Il faudrait alors rendre la dimension des cellules très petites. En conséquence le temps de calcul augmenterait alors, d'autant plus dans un cas 3D. Pour contourner ce problème, une approche semiclassique est utilisée dans [Xia+16]. Ce faisant, les auteurs arrivent sur une cavité chaotique 2D et à des fréquences microondes, à obtenir le pic de focalisation de RT à des endroits arbitraires de la cavité. Notons que cette technique fonctionne tant que la géométrie de la cavité reste relativement simple à implémenter numériquement. De plus le passage en 3D demanderait sûrement des efforts conséquents de modélisation numérique, et finirait peut être par être trop coûteux en temps de calcul.

Finalement, l'histoire des plasmas amorcés par retournement temporel ne fait que commencer. Cette thèse vient apporter la première pierre à l'édifice, en proposant les premières études théoriques et les premières démonstrations expérimentales. Il reste encore de belles choses à faire sur ce sujet, comme par exemple la compréhension fine de la physique qui fait l'objet du stage de Master 2 que nous avons mentionné précédemment ou son extension à d'autres applications comme dans le domaine de la combustion, qui font actuellement l'objet de la thèse de Béatrice Fragge entre le laboratoire LAPLACE et l'ONERA, intitulée "Allumage d'une chambre de combustion par retournement temporel microonde". L'étape la plus difficile à franchir pour atteindre le fameux pinceau plasma ondulatoire reste le décrochage du plasma, afin d'obtenir dans la cavité une boule plasma pilotée spatio-temporellement par des ondes électromagnétiques.

Publications

Revues internationales

V. Mazières^{1,2}, R. Pascaud², L. Liard¹, S. Dap¹, R. Clergereaux¹, et O. Pascal¹,
“Plasma generation using time reversal of microwaves”,

Applied Physics Letters, 115.15 (oct. 2019), p. 154101, doi : 10.1063/1.5126198,

¹Laboratoire LAPLACE, Toulouse

²ISAE – SUPAERO

Disponible sur : <https://aip.scitation.org/doi/full/10.1063/1.5126198>

V. Mazières^{1,2}, A. Al Ibrahim³, C. Chauvière⁴, P. Bonnet³, R. Pascaud², R. Clergereaux¹, S. Dap¹, L. Liard¹ et O. Pascal¹,

“Transient Electric Field Shaping With the Linear Combination of Configuration Field Method for Enhanced Spatial Control of Microwave Plasmas”,

IEEE Access, 8 (2020), p. 177084-177091, doi : 10.1109/ACCESS.2020.3025366.

¹Laboratoire LAPLACE, Toulouse

²ISAE – SUPAERO, Toulouse

³CNRS, SIGMA Clermont, Institut Pascal, Université Clermont Auvergne, 63000 Clermont-Ferrand, France

⁴CNRS, UMR 6620, Laboratoire de Mathématiques, 63171 Clermont-Ferrand, France

Disponible sur : <https://ieeexplore.ieee.org/document/9201280>

V. Mazières^{1,2}, R. Pascaud², O. Pascal¹, R. Clergereaux¹, L. Stafford³, S. Dap¹ et L. Liard¹,
“Spatio-temporal dynamics of a nanosecond pulsed microwave plasma ignited by time reversal”,

Plasma Sources Sci. Technol., 29 125017, 2020, doi : 10.1088/1361-6595/abc9ff.

¹Laboratoire LAPLACE, Toulouse

²ISAE – SUPAERO, Toulouse

³Département de physique, Université de Montréal, Montréal, Québec H3C 3J7, Canada

Disponible sur : <https://iopscience.iop.org/article/10.1088/1361-6595/abc9ff>

Conférence internationale

V. Mazières^{1,*}, R. Pascaud², L. Stafford³, P. Bonnet⁴, A. Al Ibrahim⁴, C. Chauvière⁵, L. Liard¹, S. Dap¹, R. Clergereaux¹ et O. Pascal¹,

“Wave plasma brush : Synthesis and Outlooks”,

URSI GASS 2021, Rome. À présenter en août 2021 (en attente d’acceptation).

¹Laboratoire LAPLACE, Toulouse

²ISAE – SUPAERO, Toulouse

³Département de physique, Université de Montréal, Montréal, Québec H3C 3J7, Canada

⁴CNRS, SIGMA Clermont, Institut Pascal, Université Clermont Auvergne, Clermont-Ferrand, France

⁵CNRS, UMR 6620, Laboratoire de Mathématiques, Clermont-Ferrand, France

*Orateur

Conférences nationales

V. Mazières^{1,2,*}, R. Pascaud², L. Liard¹, S. Dap¹, R. Clergereaux¹, et O. Pascal¹,
 “Contrôle spatial et temporel d’un claquage micro-onde par retournement temporel”,
JCMM 2020 - 16^{èmes} Journées de Caractérisation et Microondes et Matériaux, Toulouse (nov. 2020).

Prix de la meilleure présentation orale.

V. Mazières^{1,2}, R. Pascaud², L. Liard^{1,*}, S. Dap¹, R. Clergereaux^{1,*}, et O. Pascal¹,
 “Contrôle Spatial et temporel d’un claquage micro-onde par retournement temporel”,
Rencontres scientifiques plasmas froids, Saint Dié des Vosges (sept. 2020), France.

V. Mazières^{1,2,*}, R. Pascaud², L. Liard¹, S. Dap¹, R. Clergereaux¹, et O. Pascal¹,
 “Contrôle spatial et temporel d’un claquage micro-onde en cavité métallique à basse pression”,
1^{ères} rencontres scientifiques plasmas froids et lasers, Toulouse (nov. 2019), France. (hal-02526868)

V. Mazières^{1,2,*}, R. Pascaud², L. Liard¹, S. Dap¹, R. Clergereaux¹, et O. Pascal¹,
 “Contrôle spatial et temporel d’un claquage micro-onde par retournement temporel”,
Journée Electromagnétisme et Guerre Electronique, ONERA Toulouse (nov. 2019), France.

V. Mazières^{1,2,*}, R. Pascaud², L. Liard¹, S. Dap¹, R. Clergereaux¹, et O. Pascal¹,
 “Contrôle spatial et temporel d’un claquage micro-onde par retournement temporel”,
Assemblée Générale du GdR Ondes 2019, Gif-sur-Yvette (sept. 2019), France. (hal-02526856)

¹Laboratoire LAPLACE, Toulouse

²ISAE – SUPAERO, Toulouse

*Orateur

Détails sur la modélisation FDTD

2D

A.1 Discrétisation des équations de Maxwell

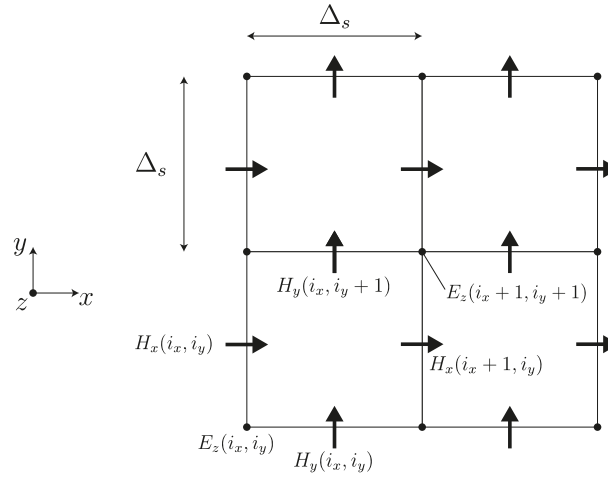


FIGURE A.1 – Illustration des composantes FDTD 2D dans le cas TE.

Les équations 3.24, 3.25 et 3.26 sont discrétisées selon un schéma classique de différences finies centrées. Pour cela, posons une grille composée de cellules rectangulaires (maillage de Yee) comme représenté sur la Figure A.1. Les dimensions de ces cellules sont prises égales selon x et y : $\Delta_s = \Delta_x = \Delta_y$. La position des champs discrétisés sera noté selon x avec les indices $i_x; i_x + 1 \dots$ et selon y avec les indices $i_y; i_y + 1 \dots$. La relation entre les composantes des champs électrique et magnétique à travers les dérivées spatiales impose que ces champs ne sont pas définis aux mêmes endroits de la grille : le champ électrique est ainsi localisé sur les nœuds des cellules et le champ magnétique entre ces nœuds. De plus, la présence des dérivées temporelles entre les composantes des champs électrique et magnétique implique que E_z , H_x et H_y ne sont pas calculées aux mêmes instants. Nous allons calculer le champ électrique E_z aux temps $t = [0; \Delta_t; 2\Delta_t \dots] = n\Delta_t$ et les composantes H_x et H_y aux temps $t = [\frac{1}{2}\Delta_t; \frac{3}{2}\Delta_t \dots] = (n + 1/2)\Delta_t$ (avec $n = [0; 1; 2 \dots]$ un entier positif). On notera ainsi E_z^n la composante du champ électrique au temps n . Les dérivées spatiales et temporelles discrétisées peuvent alors s'écrire :

- pour l'équation 3.24 (avec comme point central $E_z^{n+\frac{1}{2}}(i_x, i_y)$) :

$$\frac{\partial E_z(\mathbf{r}, t)}{\partial t} \rightarrow \frac{E_z^{n+1}(i_x, i_y) - E_z^n(i_x, i_y)}{\Delta_t} \quad (\text{A.1})$$

$$\frac{\partial H_x(\mathbf{r}, t)}{\partial y} \rightarrow \frac{H_x^{n+\frac{1}{2}}(i_x, i_y) - H_x^{n+\frac{1}{2}}(i_x, i_y - 1)}{\Delta_s} \quad (\text{A.2})$$

$$\frac{\partial H_y(\mathbf{r}, t)}{\partial x} \rightarrow \frac{H_y^{n+\frac{1}{2}}(i_x, i_y) - H_y^{n+\frac{1}{2}}(i_x - 1, i_y)}{\Delta_s} \quad (\text{A.3})$$

- pour l'équation 3.25 (avec comme point central $H_x^n(i_x, i_y)$) :

$$\frac{\partial H_x(\mathbf{r}, t)}{\partial t} \rightarrow \frac{H_x^{n+\frac{1}{2}}(i_x, i_y) - H_x^{n-\frac{1}{2}}(i_x, i_y)}{\Delta_t} \quad (\text{A.4})$$

$$\frac{\partial E_z(\mathbf{r}, t)}{\partial y} \rightarrow \frac{E_z^n(i_x, i_y + 1) - E_z^n(i_x, i_y)}{\Delta_s} \quad (\text{A.5})$$

$$(\text{A.6})$$

- pour l'équation 3.26 (avec comme point central $H_y^n(i_x, i_y)$) :

$$\frac{\partial H_y(\mathbf{r}, t)}{\partial t} \rightarrow \frac{H_y^{n+\frac{1}{2}}(i_x, i_y) - H_y^{n-\frac{1}{2}}(i_x, i_y)}{\Delta_t} \quad (\text{A.7})$$

$$\frac{\partial E_z(\mathbf{r}, t)}{\partial x} \rightarrow \frac{E_z^n(i_x + 1, i_y) - E_z^n(i_x, i_y)}{\Delta_s} \quad (\text{A.8})$$

$$(\text{A.9})$$

En utilisant ces règles de discrétisation, on trouve la version discrétisée des équations 3.24, 3.25 et 3.26 [ED15a] :

$$E_z^{n+1}(i_x, i_y) = \frac{2\varepsilon_0 - \Delta_t \sigma(i_x, i_y)}{2\varepsilon_0 + \Delta_t \sigma(i_x, i_y)} E_z^n(i_x, i_y) + \frac{2\Delta_t}{\Delta_s(2\varepsilon_0 + \Delta_t \sigma(i_x, i_y))} \left(H_y^{n+\frac{1}{2}}(i_x, i_y) - H_y^{n+\frac{1}{2}}(i_x - 1, i_y) - H_x^{n+\frac{1}{2}}(i_x, i_y) + H_x^{n+\frac{1}{2}}(i_x, i_y - 1) \right) \quad (\text{A.10})$$

$$H_x^{n+\frac{1}{2}}(i_x, i_y) = H_x^{n-\frac{1}{2}}(i_x, i_y) - \frac{\Delta_t}{\mu_0 \Delta_s} (E_z^n(i_x, i_y + 1) - E_z^n(i_x, i_y)) \quad (\text{A.11})$$

$$H_y^{n+\frac{1}{2}}(i_x, i_y) = H_y^{n-\frac{1}{2}}(i_x, i_y) + \frac{\Delta_t}{\mu_0 \Delta_s} (E_z^n(i_x + 1, i_y) - E_z^n(i_x, i_y)) \quad (\text{A.12})$$

A.2 Nombre de mode TE en cavité 2D

Pour une cavité 2D rectangulaire de dimensions L_x et L_y (et de surface $S = L_x L_y$) on trouve les fréquences de résonance f_{mn} auxquelles le champ peut exister dans la cavité :

$$f_{mn} = \frac{c_0}{2} \sqrt{\left(\frac{m}{L_x}\right)^2 + \left(\frac{n}{L_y}\right)^2} \quad (\text{A.13})$$

avec m et n des entiers positifs (et non tous nuls). Pour compter le nombre de modes, la dégénérescence de ceux-ci doit être prise en compte, c'est à dire que lorsque aucun des indices m et n n'est nul, il y a deux types de modes, transverse électrique TE_{mn} et transverse magnétique TM_{mn} , pour chaque fréquence propre [Hil98].

Le raisonnement présenté ici est l'analogie en 2D de celui utilisé par M. Dupré en 3D [Dup15]. Nous nous plaçons alors dans l'espace réciproque $(\omega, \mathbf{k})$ et nous cherchons le nombre de modes contenus dans la surface de rayon $|k| = \omega/c$. Dans cet espace, la surface de chacun des modes peut s'écrire (le facteur 1/2 comptant pour les modes TE et TM) :

$$S_{1_{\text{mode}}} = \frac{1}{2} S_{\text{carré}} = \frac{\pi^2}{2S} \quad (\text{A.14})$$

Et la surface totale des modes s'écrit :

$$S_{\text{total}} = \frac{1}{4} S_{\text{disque}} = \frac{1}{4} \pi \left(\frac{\omega}{c}\right)^2 = \frac{\pi^3}{c^2} f^2 \quad (\text{A.15})$$

Le nombre de modes en fonction de la fréquence peut ainsi s'écrire :

$$N_{2D}(f) = \frac{S_{\text{total}}}{S_{1_{\text{mode}}}} = \frac{2\pi S}{c^2} f^2 \quad (\text{A.16})$$

Et dans le cas TE, la surface d'un mode devient :

$$S_{1_{\text{mode}}}^{\text{TE}} = \frac{1}{2} S_{\text{carré}} = \frac{\pi^2}{2S} \quad (\text{A.17})$$

Et le nombre de modes TE dans une cavité 2D en fonction de la fréquence est donc :

$$N_{2D}^{TE}(f) = \frac{\pi S}{c^2} f^2 \quad (\text{A.18})$$

A.3 Équations du modèle plasma fluide

Le plasma peut être modélisé par la densité de courant $\mathbf{J}_e$ des électrons dans le vide de la façon suivante (les dépendances temporelles et spatiales sont omises pour les paramètres dépendant du champ électrique $\mathbf{E}(\mathbf{r}, t)$, puisqu'elles sont implicitement contenues dans ce dernier) :

$$\nabla \times \mathbf{E}(\mathbf{r}, t) = -\mu_0 \frac{\partial \mathbf{H}(\mathbf{r}, t)}{\partial t} \quad (\text{A.19})$$

$$\nabla \times \mathbf{H}(\mathbf{r}, t) = \varepsilon_0 \frac{\partial \mathbf{E}(\mathbf{r}, t)}{\partial t} + \mathbf{J}_c(\mathbf{r}, t) + \mathbf{J}_p(\mathbf{E}) \quad (\text{A.20})$$

$$\mathbf{J}_e(\mathbf{E}) = -en_e(\mathbf{E})\mathbf{v}_e(\mathbf{E}) \quad (\text{A.21})$$

où :

$$\frac{\partial \mathbf{v}_e(\mathbf{E})}{\partial t} = -\frac{e\mathbf{E}(\mathbf{r}, t)}{m_e} - \nu_m(\mathbf{E})\mathbf{v}_e(\mathbf{E}) \quad (\text{A.22})$$

$$\frac{\partial n_e(\mathbf{E})}{\partial t} - \nabla \cdot (D_{eff}(\mathbf{E})\nabla n_e(\mathbf{E})) = s_e(\mathbf{E}) \quad (\text{A.23})$$

avec :

$$D_{eff}(\mathbf{E}) = \frac{\xi(\mathbf{E})D_e(\mathbf{E}) + D_a(\mathbf{E})}{\xi(\mathbf{E}) + 1} \quad (\text{A.24})$$

avec $\xi(\mathbf{E}) = \nu_i(\mathbf{E})\tau_M(\mathbf{E})$ et $\tau_M(\mathbf{E})$ le temps de relaxation diélectrique défini comme $\tau_M(\mathbf{E}) = \varepsilon_0/[en_e(\mathbf{E})(\mu_e(\mathbf{E}) + \mu_i(\mathbf{E}))]$. Ce coefficient de diffusion effectif $D_{eff}(\mathbf{E})$ est une combinaison linéaire du coefficient de diffusion libre et ambipolaire, il rend compte du changement entre ces deux régimes en fonction de la densité électronique : $D_{eff} \xrightarrow{n_e \rightarrow 0} D_e$ et $D_{eff} \xrightarrow{n_e \rightarrow \infty} D_a$. Le coefficient de diffusion ambipolaire est défini comme $D_a(\mathbf{E}) = (\mu_i(\mathbf{E})/\mu_e(\mathbf{E}))D_e(\mathbf{E})$ [CB10]. La mobilité ionique autour du torr dans l'argon peut être approximée comme $\mu_i(\mathbf{E}) = \mu_i = 0.1/p$ m²/V/s avec p la pression en torr [Bas+00]. s_e est le taux de production net d'électrons, il inclut l'ionisation, l'attachement et les recombinaisons [CB10]. Dans l'argon il n'y a pas d'attachement et nous négligerons la recombinaison, dont les temps caractéristiques sont bien plus grands que les temps de simulation du RT. Le taux de production net d'électrons s'écrit donc : $s_e(\mathbf{E}) = n_e(\mathbf{E})\nu_i(\mathbf{E})$.

Comportement des antennes en émission et en réception

Sur le dispositif expérimental, la mesure et l'émission des signaux se font entre deux antennes ou plus. Plus précisément, ces opérations sont contrôlées au niveau des accès des antennes, c'est à dire au bout du câble coaxial relié à l'antenne, accessible à l'extérieur de la cavité. Un schéma d'une antenne en transmission et en réception avec les relations entre les tensions, courants et champs électriques dans le domaine fréquentiel (repris de [Bal05]) est représenté sur la Figure B.1. En transmission (Figure B.1(a)), une tension $V_t(\omega)$ est appliquée au bout du câble coaxial lié à l'antenne. Il en résulte un courant $I_t(\omega)$ au niveau du brin de l'antenne, relié à la tension par l'impédance $Z_t(\omega)$ du système : $I_t(\omega) = V_t(\omega)/Z_t(\omega)$. De ce courant résulte un champ électromagnétique rayonné en champ lointain $\mathbf{E}_t(\mathbf{r}, \omega)$ au point $\mathbf{r}$. Dans le domaine fréquentiel, ce champ est proportionnel au produit scalaire du courant $I_t(\omega)$ et du facteur $j\omega\mathbf{l}_e(\omega)$ avec $\mathbf{l}_e(\omega)$ le vecteur de longueur effective de l'antenne [Bal05]. En réception (Figure B.1(b)), la tension $V_r(\omega)$ mesurée au bout du câble coaxial est donnée par le produit du champ électrique arrivant sur l'antenne $\mathbf{E}_r(\mathbf{r}, \omega)$ et du vecteur de longueur effective $\mathbf{l}_e(\omega)$. Posons ensuite $\mathbf{H}_t(\omega) = j\omega\mathbf{l}_e(\omega)$, représentant la fonction de transfert entre le courant $I_t(\omega)$ et le champ rayonné $\mathbf{E}_t(\mathbf{r}, \omega)$, et $\mathbf{H}_r(\omega) = \mathbf{l}_e(\omega)$, représentant la fonction de transfert entre le champ arrivant sur une l'antenne $\mathbf{E}_r(\mathbf{r}, \omega)$ et la tension $V_r(\omega)$ mesurée au bout du câble coaxial. Dans le domaine temporel $\mathbf{h}_t(t)$ et $\mathbf{h}_r(t)$ peuvent être vus comme les réponses impulsionnelles de l'antenne respectivement en transmission et en réception et les produits scalaires des équations de la Figure B.1 deviennent des produits de convolution. On trouve alors que ces deux réponses impulsionnelles sont liées par une relation de dérivée temporelle : $\mathbf{h}_t(t) = \partial\mathbf{h}_r(t)/\partial t$ [SHK97] ; [Kun07]. La relation entre le champ qui est rayonné par une antenne soumise à une tension au bout du câble coaxial n'est donc pas la même que celle entre la tension mesurée au bout de ce câble coaxial et le champ incident sur cette antenne. Cependant, dans le cas de signaux bande étroite $\Delta f/f_c \ll 1$, ces effets peuvent être négligés, $\mathbf{h}_t(t)$ et $\mathbf{h}_r(t)$ étant alors simplement proportionnels et déphasés [Kun07]. On peut s'en rendre compte facilement en remarquant qu'en bande étroite, le vecteur de longueur effective est constant en fréquence, ce qui donne juste un facteur ω et un déphasage j entre $\mathbf{H}_t(\omega)$ et $\mathbf{H}_r(\omega)$.

Pour les expériences réalisées durant cette thèse $\Delta f/f_c = 250 \text{ MHz}/2.4 \text{ GHz} = 0.1$. Nous considérerons alors être dans le cas de signaux à bande étroite et négligerons donc la différence des fonctions de transfert d'émission et de réception. De ce fait, nous supposerons un rapport de proportionnalité (et un déphasage) entre la tension appliquée au bout du câble coaxial

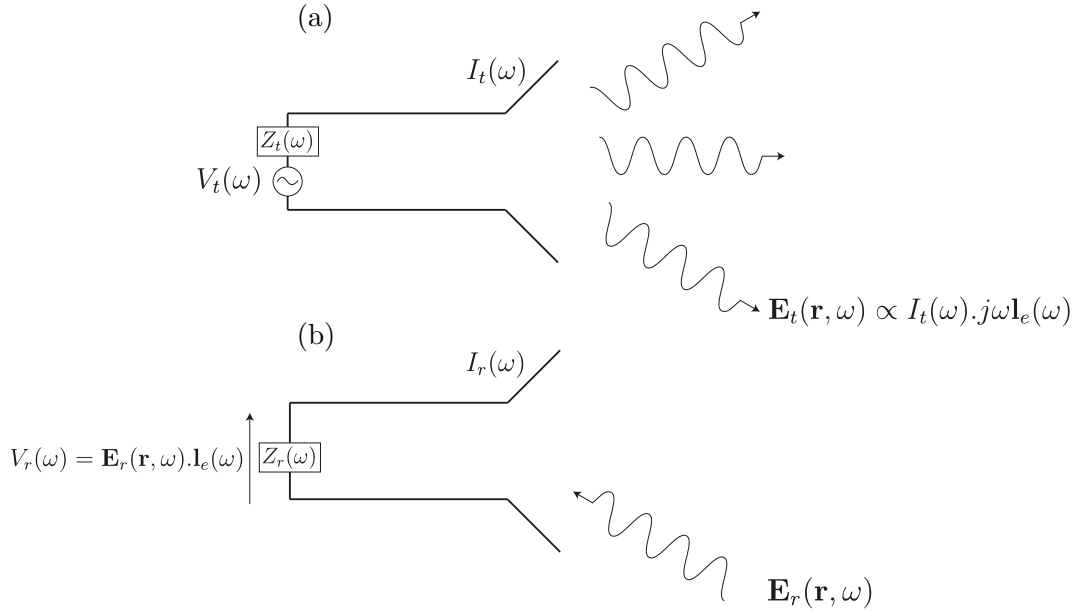


FIGURE B.1 – Schéma d’une antenne (a) en transmission et (b) en réception (relations prises de [Bal05]-p.88).

et le champ généré dans la cavité ainsi qu’entre le champ électrique arrivant au niveau de l’antenne et la tension mesurée à l’oscilloscope. De cette façon nous pouvons nous attendre, en émission, à ce que le champ électromagnétique rayonné suive bien l’évolution du signal transmis au niveau du câble coaxial et, en réception, que le signal mesuré au niveau du câble coaxial soit une image fidèle du champ électromagnétique arrivant sur l’antenne. Notons quand même que ces réactions différentes des antennes en émission et en réception (même si cet effet est négligé ici) ne fait qu’éloigner encore plus la réalité de la vision du “film passé à l’envers” lors d’un RT. Nous avons déjà vu au chapitre précédent que l’émission dans toutes les directions lors de la phase inverse et la présence d’un environnement réverbérant ont pour conséquence la superposition du film passé à l’envers avec d’autres, non présents lors de la phase initiale.

Bibliographie

- [61003] IEC 61000-4-21. : *Electromagnetic compatibility (EMC) – Part 4-21 : Testing and measurement techniques – Reverberation chamber test methods. International Electrotechnical Commission (IEC) International standard*. 2003 (cf. p. 98).
- [ADR82] Alain ASPECT, Jean DALIBARD et Gérard ROGER. « Experimental Test of Bell's Inequalities Using Time-Varying Analyzers ». In : *Phys. Rev. Lett.* 49.25 (déc. 1982). Publisher : American Physical Society, p. 1804-1807 (cf. p. 65).
- [Alb15] Stefano ALBERTI. « Physique des Plasmas I ». In : *EPFL - CRPP* (2015) (cf. p. 14).
- [Alb+19] Alejo ALBERTI et al. « Effect of chaos in a one-channel time-reversal acoustic mirror ». In : *Phys. Rev. E* 100.1 (juil. 2019), p. 012208 (cf. p. 90, 94, 112, 115, 123).
- [AS+01] A. I. AL-SHAMMA\TEXTQUOTESINGLE et al. « Design and construction of a 2.45 GHz waveguide-based microwave plasma jet at atmospheric pressure for material processing ». en. In : *J. Phys. D : Appl. Phys.* 34.18 (sept. 2001). Publisher : IOP Publishing, p. 2734-2741 (cf. p. 172).
- [Asa+20] V. S. ASADCHY et al. « Tutorial on Electromagnetic Nonreciprocity and its Origins ». In : *Proceedings of the IEEE* 108.10 (oct. 2020), p. 1684-1727 (cf. p. 66-68, 70-72, 76, 80, 82).
- [Asm+74] J. ASMUSSEN et al. « The design of a microwave plasma cavity ». In : *Proceedings of the IEEE* 62.1 (jan. 1974), p. 109-117 (cf. p. 29, 30).
- [Asm89] Jes ASMUSSEN. « Electron cyclotron resonance microwave discharges for etching and thin-film deposition ». In : *Journal of Vacuum Science & Technology A* 7.3 (mai 1989), p. 883-893 (cf. p. 30).
- [Asp19] Alain ASPECT. *Alain Aspect - Le photon onde ou particule ? L'étrangeté quantique mise en lumière*. Mai 2019 (cf. p. 65).
- [Ast05] Alexandre ASTIER. *Kaamelott*. S02E68 Stargate. Déc. 2005 (cf. p. 1).
- [Aug97] Saint AUGUSTIN. *Les Confessions*. fr. 397 (cf. p. 51).
- [Bab+00] A. I. BABARITSKII et al. « The repetitive microwave discharge as a catalyst for a chemical reaction ». en. In : *Tech. Phys.* 45.11 (nov. 2000), p. 1411-1416 (cf. p. 170, 172).
- [Bal05] Constantine A. BALANIS. *Antenna Theory : Analysis and Design*. en. John Wiley & Sons, avr. 2005 (cf. p. 207, 208).
- [Bar61] Bernard BARBER. « Resistance by Scientists to Scientific Discovery ». en. In : *Science* 134.3479 (sept. 1961), p. 596-602 (cf. p. 56).

- [Bas+00] E. BASURTO et al. « Mobility of He+, Ne+, Ar+, N+2, O+2, and CO+2 in their parent gas ». In : *Phys. Rev. E* 61.3 (mar. 2000), p. 3053-3057 (cf. p. 206).
- [BCB12] J. BENOIT, C. CHAUVIÈRE et P. BONNET. « Source identification in time domain electromagnetics ». en. In : *Journal of Computational Physics* 231.8 (avr. 2012), p. 3446-3456 (cf. p. 198).
- [BCB15] J. BENOIT, C. CHAUVIÈRE et P. BONNET. « Time-dependent current source identification for numerical simulations of Maxwell's equations ». en. In : *Journal of Computational Physics* 289 (mai 2015), p. 116-128 (cf. p. 198).
- [BCZ10] Jean-Pierre BOEUF, Bhaskar CHAUDHURY et Guo Qiang ZHU. « Theory and Modeling of Self-Organization and Propagation of Filamentary Plasma Arrays in Microwave Breakdown at Atmospheric Pressure ». In : *Phys. Rev. Lett.* 104.1 (jan. 2010), p. 015002 (cf. p. 21).
- [Boe16] Jean-Pierre BOEUF. « Plasma froids et applications ». In : *Cours ENSEEIHT 3GE TEMA* (2016) (cf. p. 2, 10, 18, 29).
- [Bol+02] Ludwig BOLTZMANN et al. *L. Boltzmann, Leçons sur la théorie des gaz, traduites par A. Gallotti,... avec une introduction et des notes de M. Brillouin*. Gauthier-Villars, 1902 (cf. p. 57).
- [Bor25] Émile BOREL. *Mécanique statistique classique*. es. Gauthier-Villars, 1925 (cf. p. 57).
- [Bou+06] A. BOUSQUET et al. « Influence of plasma pulsing on the deposition kinetics and film structure in low pressure oxygen/hexamethyldisiloxane radiofrequency plasmas ». en. In : *Thin Solid Films* 514.1 (août 2006), p. 45-51 (cf. p. 45).
- [BVY87] L. M. BIBERMAN, V. S. VOROB'EV et I. T. YAKUBOV. *Kinetics of Nonequilibrium Low-Temperature Plasmas*. en. Springer US, 1987 (cf. p. 190).
- [Cal+18] Christophe CALOZ et al. « Electromagnetic Nonreciprocity ». In : *Phys. Rev. Applied* 10.4 (oct. 2018), p. 047001 (cf. p. 62, 66-68, 70-74, 76, 82).
- [Car13] E. a. D. CARBONE. « Laser diagnostics and modelling of microwave plasmas ». English. In : (2013) (cf. p. 13).
- [Car+14] Emile CARBONE et al. « Spatio-temporal dynamics of a pulsed microwave argon plasma : ignition and afterglow ». en. In : *Plasma Sources Sci. Technol.* 24.1 (déc. 2014), p. 015015 (cf. p. 45, 175, 190, 197).
- [CB10] Bhaskar CHAUDHURY et Jean-Pierre BOEUF. « Computational Studies of Filamentary Pattern Formation in a High Power Microwave Breakdown Generated Air Plasma ». In : *IEEE Transactions on Plasma Science* 38.9 (sept. 2010), p. 2281-2288 (cf. p. 21, 127, 206).
- [CDK15] Emile CARBONE, Jan van DIJK et Gerrit KROESEN. « Experimental evidence of resonant energy collisional transfers between argon 1s and 2p states and ground state H atoms by laser collisional induced fluorescence ». en. In : *Plasma Sources Sci. Technol.* 24.2 (avr. 2015), p. 025036 (cf. p. 190).

- [CF92] D. CASSEREAU et M. FINK. « Time-reversal of ultrasonic fields. III. Theory of the closed time-reversal cavity ». In : *IEEE Transactions on Ultrasonics, Ferroelectrics, and Frequency Control* 39.5 (sept. 1992), p. 579-592 (cf. p. 60).
- [Che74] Francis F. CHEN. *Introduction to Plasma Physics*. en. Springer US, 1974 (cf. p. 18).
- [CKW03] Robert R. CALDWELL, Marc KAMIONKOWSKI et Nevin N. WEINBERG. « Phantom Energy : Dark Energy with w\textless 1 Causes a Cosmic Doomsday ». In : *Phys. Rev. Lett.* 91.7 (août 2003), p. 071301 (cf. p. 54).
- [CM17] Andrea COZZA et Florian MONSEF. « Steering Focusing Waves in a Reverberation Chamber With Generalized Time Reversal ». In : *IEEE Transactions on Antennas and Propagation* 65.3 (mar. 2017). Conference Name : IEEE Transactions on Antennas and Propagation, p. 1349-1356 (cf. p. 198).
- [Cot] Francis COTTET. *Traitement des signaux et acquisition de données. Cours et exercices corrigés 4e édition*. fr (cf. p. 94, 100, 150, 153, 156).
- [Coz11] Andrea COZZA. « The Role of Losses in the Definition of the Overmoded Condition for Reverberation Chambers and Their Statistics ». In : *IEEE Transactions on Electromagnetic Compatibility* 53.2 (mai 2011), p. 296-307 (cf. p. 100).
- [CS00] H. CONRADTS et M. SCHMIDT. « Plasma generation and plasma sources ». en. In : *Plasma Sources Sci. Technol.* 9.4 (oct. 2000), p. 441-454 (cf. p. 29).
- [Cun+12] G. CUNGE et al. « Measurement of free radical kinetics in pulsed plasmas by UV and VUV absorption spectroscopy and by modulated beam mass spectrometry ». en. In : *Plasma Sources Sci. Technol.* 21.2 (avr. 2012), p. 024006 (cf. p. 45).
- [DAF99] Carsten DRAEGER, Jean-Christian AIME et Mathias FINK. « One-channel time-reversal in chaotic cavities : Experimental results ». In : *The Journal of the Acoustical Society of America* 105.2 (jan. 1999), p. 618-625 (cf. p. 107, 120, 124).
- [Dav10] Matthieu DAVY. « Application du retournement temporel en micro-ondes à l'amplification d'impulsions et l'imagerie ». fr. Thèse de doct. Université Paris-Diderot - Paris VII, oct. 2010 (cf. p. 42).
- [Der94] Arnaud DERODE. « La coherence des ondes ultrasonores en milieu heterogene ». These de doctorat. Paris 7, jan. 1994 (cf. p. 41).
- [DF97] Carsten DRAEGER et Mathias FINK. « One-Channel Time Reversal of Elastic Waves in a Chaotic 2D-Silicon Cavity ». In : *Phys. Rev. Lett.* 79.3 (juil. 1997), p. 407-410 (cf. p. 42, 60, 85, 88, 90, 112).
- [DF99] Carsten DRAEGER et Mathias FINK. « One-channel time-reversal in chaotic cavities : Theoretical limits ». In : *The Journal of the Acoustical Society of America* 105.2 (jan. 1999), p. 611-617 (cf. p. 89, 118).

- [DJDS19] Antoine DUROCHER-JEAN, Edouard DESJARDINS et Luc STAFFORD. « Characterization of a microwave argon plasma column at atmospheric pressure by optical emission and absorption spectroscopy coupled with collisional-radiative modelling ». In : *Physics of Plasmas* 26.6 (juin 2019), p. 063516 (cf. p. 189, 197).
- [DL73] E. F. DAWSON et S. LEDERMAN. « Pulsed microwave breakdown in gases with a low degree of preionization ». In : *Journal of Applied Physics* 44.7 (juil. 1973), p. 3066-3073 (cf. p. 179).
- [DRA97] CARSTEN DRAEGER. « Ondes elastiques et reversibilite ». These de doctorat. Paris 11, jan. 1997 (cf. p. 41, 60, 85).
- [DRF95] Arnaud DERODE, Philippe ROUX et Mathias FINK. « Robust Acoustic Time Reversal with High-Order Multiple Scattering ». In : *Phys. Rev. Lett.* 75.23 (déc. 1995), p. 4206-4209 (cf. p. 41, 88).
- [Dru00a] P. DRUDE. « Zur Elektronentheorie der Metalle ». In : *Annalen der Physik* 306.3 (jan. 1900), p. 566-613 (cf. p. 27).
- [Dru00b] P. DRUDE. « Zur Elektronentheorie der Metalle; II. Teil. Galvanomagnetische und thermomagnetische Effecte ». en. In : *Annalen der Physik* 308.11 (1900), p. 369-402 (cf. p. 27).
- [DTF99] Arnaud DERODE, Arnaud TOURIN et Mathias FINK. « Ultrasonic pulse compression with one-bit time reversal through multiple scattering ». In : *Journal of Applied Physics* 85.9 (avr. 1999), p. 6343-6352 (cf. p. 42).
- [Dug59] René DUGAS. *La Théorie physique au sens de Boltzmann et ses prolongements modernes*. 1959 (cf. p. 55-57).
- [Dum54] Michael DUMMETT. « Can an Effect Precede its Cause? » In : *Aristotelian Society Proceedings Supplement* 28 (1954) (cf. p. 54).
- [Dup15] Matthieu DUPRÉ. « Contrôle des micro-ondes en milieux réverbérants ». fr. Thèse de doct. Université Paris 7, déc. 2015 (cf. p. 43, 99, 101, 205).
- [DW98] Bernhard DISCHLER et Christoph WILD. *Low-pressure Synthetic Diamond : Manufacturing and Applications*. en. Springer, jan. 1998 (cf. p. 3, 30).
- [ED15a] Atef Z. ELSHERBENI et Veysel DEMIR. *The Finite-Difference Time-Domain Method for Electromagnetics with MATLAB® Simulations*. English. 2 edition. Edison, NJ : Scitech Publishing, nov. 2015 (cf. p. 116, 117, 204).
- [ED15b] R. ERNST et J. DUAL. « Quantitative guided wave testing by applying the time reversal principle on dispersive waves in beams ». en. In : *Wave Motion* 58 (nov. 2015), p. 259-280 (cf. p. 118).
- [Ein09] Albert EINSTEIN. « On the Development of Our Views Concerning the Nature and Constitution of Radiation ». In : *Presented at the session of the Division of Physics of the 81st Meeting of German Scientists and Physicians in Salzburg on September 21, 1909* (1909) (cf. p. 65).

- [FA11] F. FOROOZAN et A. ASIF. « Time reversal based active array source localization ». In : *IEEE Trans. Signal Process.* 59.6 (juin 2011), p. 2655-2668 (cf. p. 42).
- [FEB65] F. C. FEHSENFELD, K. M. EVENSON et H. P. BROIDA. « Microwave discharge cavities operating at 2450 MHz ». In : *Rev. Sci. Instrum.* 36.3 (mar. 1965), p. 294-298 (cf. p. 30).
- [Fel66] Peter FELSENTHAL. « Nanosecond-Pulse Microwave Breakdown in Air ». In : *Journal of Applied Physics* 37.12 (nov. 1966), p. 4557-4560 (cf. p. 45, 130, 136, 170).
- [Fin09] Johannes Karl FINK. *Physical Chemistry in Depth*. en. Berlin Heidelberg : Springer-Verlag, 2009 (cf. p. 56).
- [Fin15] Mathias FINK. *Conférence : "Démon de Loschmidt, retournement temporel et holographie" avec Mathias Fink*. Nov. 2015 (cf. p. 56, 59).
- [Fin+89] M. FINK et al. « Self focusing in inhomogeneous media with time reversal acoustic mirrors ». In : *Proceedings., IEEE Ultrasonics Symposium*, oct. 1989, 681-686 vol.2 (cf. p. 40).
- [FM93] C. M. FERREIRA et M. MOISAN, éd. *Microwave Discharges : Fundamentals and Applications*. Springer US, 1993 (cf. p. 3, 30).
- [Fra+13a] Matthew FRAZIER et al. « Nonlinear Time Reversal in a Wave Chaotic System ». In : *Phys. Rev. Lett.* 110.6 (fév. 2013), p. 063902 (cf. p. 174).
- [Fra+13b] Matthew FRAZIER et al. « Nonlinear time reversal of classical waves : Experiment and model ». In : *Phys. Rev. E* 88.6 (déc. 2013), p. 062910 (cf. p. 174).
- [FWK98] M. FÜNER, C. WILD et P. KOIDL. « Novel microwave plasma reactor for diamond synthesis ». In : *Appl. Phys. Lett.* 72.10 (mar. 1998), p. 1149-1151 (cf. p. 30).
- [Gar02] Fred GARDIOL. *Traité d'électricité, volume 3 : Electromagnétisme*. Français. Lausanne : Presses Polytechniques et Universitaires Romandes, 2002 (cf. p. 65).
- [Gar67] F.E. GARDIOL. « On the thermodynamic paradox in ferrite-loaded waveguides ». In : *Proceedings of the IEEE* 55.9 (sept. 1967), p. 1616-1617 (cf. p. 63, 65).
- [GB05] D. A. GURNETT et A. BHATTACHARJEE. *Introduction to Plasma Physics : With Space and Laboratory Applications*. Cambridge : Cambridge University Press, 2005 (cf. p. 3).
- [GK08] Dan M. GOEBEL et Ira KATZ. *Fundamentals of Electric Propulsion : Ion and Hall Thrusters*. 2008 (cf. p. 3).
- [GR56] Lawrence GOULD et Louis W. ROBERTS. « Breakdown of Air at Microwave Frequencies ». In : *Journal of Applied Physics* 27.10 (oct. 1956). Publisher : American Institute of Physics, p. 1162-1170 (cf. p. 45, 128).
- [Gri94] Alfred GRILL. *Cold Plasma in Materials Fabrication : From Fundamentals to Applications*. 1994 (cf. p. 3, 30).

- [Gro14] Jean-Baptiste GROS. « Statistiques spatiales des cavités chaotiques ouvertes : applications aux cavités électromagnétiques ». These de doctorat. Nice, déc. 2014 (cf. p. 95-98, 109, 112, 113, 148, 164).
- [Hü+12] S. HÜBNER et al. « A power pulsed low-pressure argon microwave plasma investigated by Thomson scattering : evidence for molecular assisted recombination ». en. In : *J. Phys. D : Appl. Phys.* 45.5 (jan. 2012), p. 055203 (cf. p. 45, 175).
- [Has+01] K. HASSOUNI et al. « Investigation of chemical kinetics and energy transfer in a pulsed microwave H₂/CH₄ plasma ». en. In : *Plasma Sources Sci. Technol.* 10.1 (jan. 2001), p. 61-75 (cf. p. 45).
- [HHG04] G. J. M. HAGELAAR, K. HASSOUNI et A. GICQUEL. « Interaction between the electromagnetic fields and the plasma in a microwave plasma reactor ». In : *Journal of Applied Physics* 96.4 (août 2004), p. 1819-1828 (cf. p. 29).
- [Hil98] David A. HILL. « Electromagnetic Theory of Reverberation Chambers ». en. In : 1506 (déc. 1998) (cf. p. 99, 147, 148, 205).
- [HLH14] Charles HUDIN, José LOZADA et Vincent HAYWARD. « Spatial, temporal, and thermal contributions to focusing contrast by time reversal in a cavity ». en. In : *Journal of Sound and Vibration* 333.6 (mar. 2014), p. 1818-1832 (cf. p. 112).
- [HMK02] M. HEINTZE, M. MAGUREANU et M. KETTLITZ. « Mechanism of C₂ hydrocarbon formation from methane in a pulsed microwave plasma ». In : *Journal of Applied Physics* 92.12 (nov. 2002), p. 7022-7031 (cf. p. 45).
- [HP05] G. J. M. HAGELAAR et L. C. PITCHFORD. « Solving the Boltzmann equation to obtain electron transport coefficients and rate coefficients for fluid models ». en. In : *Plasma Sources Sci. Technol.* 14.4 (oct. 2005), p. 722-733 (cf. p. 26, 44, 128).
- [Ibr+16] R. IBRAHIM et al. « Experiments of time-reversed pulse waves for wireless power transmission in an indoor environment ». In : *IEEE Trans. Microw. Theory Techn.* 64.7 (juil. 2016), p. 2159-2170 (cf. p. 42).
- [Ish17] Akira ISHIMARU. *Electromagnetic Wave Propagation, Radiation, and Scattering : From Fundamentals to Applications*. English. 2 edition. Piscataway, NJ : Wiley-IEEE Press, sept. 2017 (cf. p. 65, 74).
- [Jac+07] Vincent JACQUES et al. « Experimental Realization of Wheeler's Delayed-Choice Gedanken Experiment ». en. In : *Science* 315.5814 (fév. 2007). Publisher : American Association for the Advancement of Science Section : Report, p. 966-968 (cf. p. 65).
- [Jac98] John David JACKSON. *Classical Electrodynamics Third Edition*. English. 3 edition. New York : Wiley, août 1998 (cf. p. 64).
- [Jer77] A.J. JERRI. « The Shannon sampling theorem—Its various extensions and applications : A tutorial review ». In : *Proceedings of the IEEE* 65.11 (nov. 1977). Conference Name : Proceedings of the IEEE, p. 1565-1596 (cf. p. 60).

- [Kem11] B. A. KEMP. « Resolution of the Abraham-Minkowski debate : Implications for the electromagnetic wave theory of light in matter ». In : *Journal of Applied Physics* 109.11 (juin 2011), p. 111101 (cf. p. 62).
- [Ker87] Donald E. KERR. *Propagation of Short Radio Waves, Chap 8 : The relation between absorption and dispersion*. en. IET, 1987 (cf. p. 65).
- [Kha10] A. KHALEGHI. « Measurement and analysis of ultra-wideband time reversal for indoor propagation channels ». In : *Wireless Pers. Commun.* 54 (juil. 2010), p. 307-320 (cf. p. 42).
- [Kle07] Étienne KLEIN. *Le facteur temps ne sonne jamais deux fois*. fr. 2007 (cf. p. 54-58).
- [Kle10] Etienne KLEIN. *D'où vient le temps ?* 2010 (cf. p. 52).
- [Kon00] Jin Au KONG. *Electromagnetic Wave Theory*. English. Cambridge, MA : E M W Pub, juin 2000 (cf. p. 60, 75, 79).
- [Kor+96] D. KORZEC et al. « Scaling of microwave slot antenna (SLAN) : a concept for efficient plasma generation ». en. In : *Plasma Sources Sci. Technol.* 5.2 (mai 1996), p. 216-234 (cf. p. 29).
- [Kra27] H.A. KRAMERS. « La diffusion de la lumiere par les atomes ». In : *in Congresso Internazionale de Fisici, Vol. 2, Como (Bologna : Nicolo, 1927), pp. 545-57* (1927) (cf. p. 64).
- [Kuh62] Thomas S. KUHN. *La structure des révolutions scientifiques*. fr. 1962 (cf. p. 55-57, 64, 65).
- [Kun07] Jurgen KUNISCH. « Implications of Lorentz Reciprocity for Ultra-Wideband Antennas ». In : *2007 IEEE International Conference on Ultra-Wideband*. Sept. 2007, p. 214-219 (cf. p. 148, 207).
- [LCM83] Bing-Hope LIU, David C. CHANG et Mark T. MA. *Eigenmodes and the Composite Quality Factor of a Reverberating Chamber*. en. National Bureau of Standards, 1983 (cf. p. 99).
- [Leb15] Y. A. LEBEDEV. « Microwave discharges at low pressures and peculiarities of the processes in strongly non-uniform plasma ». In : *Plasma Sources Sci. Technol.* 24.5 (août 2015), p. 053001 (cf. p. 29-31, 37).
- [Lem11] Fabrice LEMOULT. « Focalisation et contrôle des ondes en milieux complexes et localement résonants ». These de doctorat. Paris 7, jan. 2011 (cf. p. 40-43).
- [Leo09] Ulf LEONHARDT. « Perfect imaging without negative refraction ». en. In : *New J. Phys.* 11.9 (sept. 2009). Publisher : IOP Publishing, p. 093040 (cf. p. 196).
- [Ler+04] G. LEROSEY et al. « Time Reversal of Electromagnetic Waves ». In : *Phys. Rev. Lett.* 92.19 (mai 2004), p. 193904 (cf. p. 42, 79, 88, 124, 184).
- [Ler+06] G. LEROSEY et al. « Time reversal of wideband microwaves ». In : *Appl. Phys. Lett.* 88.15 (avr. 2006), p. 154101 (cf. p. 99).

- [Ler06] Geoffroy LEROSEY. « Retournement temporel d'ondes électromagnétiques et application à la télécommunication en milieux complexes ». en. Thèse de doct. ESPCI ParisTECH, déc. 2006 (cf. p. 42, 90, 94, 99, 105, 145, 148, 149, 153, 155, 166).
- [Ler+07] Geoffroy LEROSEY et al. « Focusing Beyond the Diffraction Limit with Far-Field Time Reversal ». en. In : *Science* 315.5815 (fév. 2007), p. 1120-1122 (cf. p. 196).
- [Li+14] Y. F. LI et al. « Design of a new TM021 mode cavity type MPCVD reactor for diamond film deposition ». In : *Diamond Relat. Mater.* 44 (avr. 2014), p. 88-94 (cf. p. 30).
- [Li+19] B. LI et al. « Wireless power transfer based on microwaves and time reversal for indoor environments ». In : *IEEE Access* 7 (2019), p. 114897-114908 (cf. p. 42).
- [Lid03] Andrew LIDDLE. *An Introduction to Modern Cosmology*. 2003 (cf. p. 2).
- [LL05] Michael A. LIEBERMAN et Alan J. LICHTENBERG. *Principles of Plasma Discharges and Materials Processing , 2nd Edition*. English. 2 edition. Hoboken, N.J : Wiley-Interscience, avr. 2005 (cf. p. 3, 15, 16, 18-20, 23-25, 30, 178, 189).
- [Loo+12] Yoke Leng LOO et al. « Broadband microwave Luneburg lens made of gradient index metamaterials ». EN. In : *J. Opt. Soc. Am. A, JOSAA* 29.4 (avr. 2012). Publisher : Optical Society of America, p. 426-430 (cf. p. 196).
- [LPA97] T. LAGARDE, J. PELLETIER et Y. ARNAL. « Influence of the multipolar magnetic field configuration on the density of distributed electron cyclotron resonance plasmas ». en. In : *Plasma Sources Sci. Technol.* 6.1 (fév. 1997), p. 53-60 (cf. p. 30).
- [LW00] O. I. LOBKIS et R. L. WEAVER. « Complex modal statistics in a reverberant dissipative body ». In : *The Journal of the Acoustical Society of America* 108.4 (sept. 2000), p. 1480-1485 (cf. p. 109).
- [LY19] W. LEI et L. YAO. « Performance analysis of time reversal communication systems ». In : *IEEE Commun. Lett.* 23.4 (avr. 2019), p. 680-683 (cf. p. 42).
- [Mar09] Peter M MARTIN. *Handbook of deposition technologies for films and coatings : science, applications and technology*. English. Norwich, N.Y. ; Oxford : William Andrew ; Elsevier Science [distributeur, 2009 (cf. p. 45).
- [Max65] James Clerk MAXWELL. « VIII. A dynamical theory of the electromagnetic field ». In : *Philosophical Transactions of the Royal Society of London* 155 (jan. 1865), p. 459-512 (cf. p. 64).
- [Maz+19] Valentin MAZIÈRES et al. « Plasma generation using time reversal of microwaves ». In : *Appl. Phys. Lett.* 115.15 (oct. 2019), p. 154101 (cf. p. 169, 176, 177).
- [Maz+20a] V. MAZIÈRES et al. « Transient Electric Field Shaping With the Linear Combination of Configuration Field Method for Enhanced Spatial Control of Microwave Plasmas ». In : *IEEE Access* 8 (2020), p. 177084-177091 (cf. p. 199).

- [Maz+20b] Valentin MAZIÈRES et al. « Spatio-temporal dynamics of a nanosecond pulsed microwave plasma ignited by time reversal ». en. In : *Plasma Sources Sci. Technol.* 29.12 (déc. 2020). Publisher : IOP Publishing, p. 125017 (cf. p. 175).
- [MB68] F. J. MEHR et Manfred A. BIONDI. « Electron-Temperature Dependence of Electron-Ion Recombination in Argon ». In : *Phys. Rev.* 176.1 (déc. 1968), p. 322-326 (cf. p. 190).
- [MC14] Florian MONSEF et Andrea COZZA. « Average Number of Significant Modes Excited in a Mode-Stirred Reverberation Chamber ». In : *IEEE Transactions on Electromagnetic Compatibility* 56.2 (avr. 2014), p. 259-265 (cf. p. 95, 99, 100).
- [McC+11] David J. MCCABE et al. « Spatio-temporal focusing of an ultrafast pulse through a multiply scattering medium ». en. In : *Nature Communications* 2.1 (août 2011), p. 447 (cf. p. 42).
- [Mic09] Claire MICHEL. « Chaos ondulatoire en optique guidée : amplificateur fibré double-gaine pour la génération de modes "Scar" ». These de doctorat. Nice, jan. 2009 (cf. p. 97, 120).
- [Mil67] H. Craig MILLER. « Change in Field Intensification Factor of an Electrode Projection (Whisker) at Short Gap Lengths ». In : *Journal of Applied Physics* 38.11 (oct. 1967). Publisher : American Institute of Physics, p. 4501-4504 (cf. p. 170).
- [MP06] Michel MOISAN et Jacques PELLETIER. *Physique des plasmas collisionnels : Application aux décharges haute fréquence*. Français. Les Ulis : EDP Sciences, 2006 (cf. p. 31, 32).
- [MP12] Michel MOISAN et Jacques PELLETIER. *Physics of Collisional Plasmas : Introduction to High-Frequency Discharges*. Anglais. 2012 ed. Dordrecht ; New York : Springer, 2012 (cf. p. 20, 22, 24, 178).
- [MUB98] K. MURAOKA, K. UCHINO et M. D. BOWDEN. « Diagnostics of low-density glow discharge plasmas using Thomson scattering ». en. In : *Plasma Phys. Control. Fusion* 40.7 (juil. 1998). Publisher : IOP Publishing, p. 1221-1239 (cf. p. 196).
- [MZ91] M. MOISAN et Z. ZAKRZEWSKI. « Plasma sources based on the propagation of electromagnetic surface waves ». en. In : *J. Phys. D : Appl. Phys.* 24.7 (juil. 1991), p. 1025-1048 (cf. p. 29).
- [MZP79] M. MOISAN, Z. ZAKRZEWSKI et R. PANTEL. « The theory and characteristics of an efficient surface wave launcher (surfatron) producing long plasma columns ». en. In : *J. Phys. D : Appl. Phys.* 12.2 (fév. 1979), p. 219-237 (cf. p. 29).
- [NBEZ09] I. H. NAQVI, P. BESNIER et G. EL ZEIN. « Effects of Time Variant Channel on a Time Reversal UWB System ». In : *GLOBECOM 2009 - 2009 IEEE Global Telecommunications Conference*. Nov. 2009, p. 1-6 (cf. p. 174).

- [NPC13] Heeso NOH, Sébastien M. POPOFF et Hui CAO. « Broadband subwavelength focusing of light using a passive sink ». EN. In : *Opt. Express*, OE 21.15 (juil. 2013). Publisher : Optical Society of America, p. 17435-17446 (cf. p. 196).
- [Pas04] Blaise PASCAL. *Les Provinciales. Pensées et opuscles divers*. fr. 2004 (cf. p. 51).
- [PKS09] Hyun Woo PARK, Seung Bum KIM et Hoon SOHN. « Understanding a time reversal process in Lamb wave propagation ». en. In : *Wave Motion* 46.7 (nov. 2009), p. 451-467 (cf. p. 118).
- [Pod+06] S. PODENOK et al. « Electric field enhancement factors around a metallic, end-capped cylinder ». In : *NANO* 01.01 (juil. 2006). Publisher : World Scientific Publishing Co., p. 87-93 (cf. p. 170).
- [Pra91] Claire PRADA. « Retournement temporel des ondes ultrasonores. Application à la focalisation ». These de doctorat. Paris 7, jan. 1991 (cf. p. 40).
- [Pri12] Huw PRICE. « Does Time-Symmetry Imply Retrocausality ? How the Quantum World Says "Maybe" ? » In : *Studies in History and Philosophy of Science Part B : Studies in History and Philosophy of Modern Physics* 43.2 (2012), p. 75-83 (cf. p. 54).
- [QTMTH12] O. QUEVEDO-TERUEL, R. C. MITCHELL-THOMAS et Y. HAO. « Frequency dependence and passive drains in fish-eye lenses ». In : *Phys. Rev. A* 86.5 (nov. 2012). Publisher : American Physical Society, p. 053817 (cf. p. 196).
- [Qui04] Nicolas QUIEFFIN. « Etude du rayonnement acoustique de structures solides : vers un système d'imagerie haute résolution ». These de doctorat. Paris 6, jan. 2004 (cf. p. 103, 104, 107, 112, 124).
- [Rai13] Jean-Luc RAIMBAULT. « Modélisations fluides des plasmas ». In : *Cours de master "Plasmas : de l'espace au laboratoire"* (2013) (cf. p. 24).
- [Rai57] S. RAIMES. « The theory of plasma oscillations in metals ». en. In : *Rep. Prog. Phys.* 20.1 (jan. 1957), p. 1-37 (cf. p. 27).
- [Rai91] Yuri P. RAIZER. *Gas Discharge Physics*. en. Sous la dir. de John E. ALLEN. Berlin Heidelberg : Springer-Verlag, 1991 (cf. p. 19, 22, 188).
- [Rax05] Jean-Marcel RAX. *Physique des plasmas - Cours et applications : Cours et applications*. Français. Paris : Dunod, 2005 (cf. p. 9, 15, 17-19).
- [RD02] Bertrand RUSSELL et Philippe DEVAUX. *La Méthode scientifique en philosophie*. Français. Paris : Payot, 2002 (cf. p. 54).
- [RF01] J. de ROSNY et M. FINK. « An experiment below the wave diffraction limits with an acoustic sink : the ideal time reversal experiment ». In : *The Journal of the Acoustical Society of America* 109.5 (mai 2001), p. 2437-2437 (cf. p. 79, 196).
- [Ros00] Julien de ROSNY. « Milieux réverbérants et réversibilité ». These de doctorat. Paris 6, jan. 2000 (cf. p. 42, 89, 90, 94, 103, 104, 107, 109, 112, 120, 123).

- [Rou+94] A. ROUSSEAU et al. « Pulsed microwave discharge : a very efficient H atom source ». en. In : *J. Phys. D : Appl. Phys.* 27.11 (nov. 1994), p. 2439-2441 (cf. p. 45).
- [Rum54] V. H. RUMSEY. « Reaction Concept in Electromagnetic Theory ». In : *Phys. Rev.* 94.6 (juin 1954), p. 1483-1491 (cf. p. 73).
- [Rus+01] V. D. RUSANOV et al. « Properties of a catalytically active pulsed microwave discharge at atmospheric pressure ». en. In : *Dokl. Phys.* 46.4 (avr. 2001), p. 242-246 (cf. p. 170, 172).
- [Sam+12] Seiji SAMUKAWA et al. « The 2012 Plasma Roadmap ». en. In : *J. Phys. D : Appl. Phys.* 45.25 (juin 2012), p. 253001 (cf. p. 3, 29, 38).
- [Sch20] Delmore SCHWARTZ. *Calmly We Walk through This April's Day*. en. Mai 2020 (cf. p. 51).
- [Sch87] Manfred R. SCHROEDER. « Statistical Parameters of the Frequency Response Curves of Large Rooms ». English. In : *JAES* 35.5 (mai 1987), p. 299-306 (cf. p. 100, 129, 164).
- [Sel14] Kamardine SELEMANI. « Analyse et optimisation des chambres réverbérantes à l'aide du concept de cavité chaotique ouverte ». These de doctorat. Paris Est, fév. 2014 (cf. p. 98, 148, 164).
- [SHK97] A. SHLIVINSKI, E. HEYMAN et R. KASTNER. « Antenna characterization in the time domain ». In : *IEEE Transactions on Antennas and Propagation* 45.7 (juil. 1997), p. 1140-1149 (cf. p. 148, 207).
- [Sil+09] F. SILVA et al. « Microwave engineering of plasma-assisted CVD reactors for diamond deposition ». In : *J. Phys. : Condens. Matter* 21.36 (août 2009), p. 364202 (cf. p. 30, 36).
- [Sil19] François SILVA. *Quelle Alimentation pour quel Plasma : Générateurs Microondes pour la production de plasma*. GREMI (Orléans), 2019 (cf. p. 32, 33, 37).
- [Sim18] Antoine SIMON. « Étude de dispositifs de limitation de puissance microonde en technologie circuit imprimé exploitant des plasmas de décharge ». These de doctorat. Toulouse, ISAE, déc. 2018 (cf. p. 138).
- [SS90] H.-J. STÖCKMANN et J. STEIN. « "Quantum" chaos in billiards studied by microwave absorption ». In : *Phys. Rev. Lett.* 64.19 (mai 1990). Publisher : American Physical Society, p. 2215-2218 (cf. p. 109).
- [St9] Hans-Jürgen STÖCKMANN. *Quantum Chaos : An Introduction*. Cambridge : Cambridge University Press, 1999 (cf. p. 148, 164).
- [Ste+14] Ilija STEFANOVIĆ et al. « Argon metastable dynamics and lifetimes in a direct current microdischarge ». In : *Journal of Applied Physics* 116.11 (sept. 2014), p. 113302 (cf. p. 190).
- [STJ04] Alexander M. SUTIN, James A. TENCATE et Paul A. JOHNSON. « Single-channel time reversal in elastic solids ». In : *The Journal of the Acoustical Society of America* 116.5 (nov. 2004), p. 2779-2784 (cf. p. 198).

- [Stu+14] Guido S. J. STURM et al. « Microwaves and microreactors : Design challenges and remedies ». en. In : *Chemical Engineering Journal* 243 (mai 2014), p. 147-158 (cf. p. 38).
- [Su+14] J. J. SU et al. « A novel microwave plasma reactor with a unique structure for chemical vapor deposition of diamond films ». en. In : *Diamond and Related Materials* 42 (fév. 2014), p. 28-32 (cf. p. 3, 30, 34, 35).
- [TH05] Allen TAFLOVE et Susan C. HAGNESS. *Computational Electrodynamics : The Finite-Difference Time-Domain Method, Third Edition*. English. 3 edition. Boston : Artech House, mai 2005 (cf. p. 5, 89, 118).
- [TK20] Sajjad TARAVATI et Ahmed A. KISHK. « Space-Time Modulation : Principles and Applications ». In : *IEEE Microwave Magazine* 21.4 (avr. 2020), p. 30-56 (cf. p. 172).
- [Val16] Henri VALLON. « Focusing high-power electromagnetic waves using time-reversal ». These de doctorat. Université Paris-Saclay (ComUE), mar. 2016 (cf. p. 147, 148).
- [VIS85] A. L. VIKHAREV, O. A. IVANOV et A. N. STEPANOV. « Multiple-pulse UHF breakdown in intersecting wave beams ». In : *Radiophysics and Quantum Electronics* 28 (jan. 1985), p. 26-31 (cf. p. 179).
- [Vya+16] Sandeep Kumar VYAS et al. « Review of Magnetron Developments ». en. In : *Frequenz* 70.9-10 (sept. 2016). Publisher : De Gruyter Section : Frequenz, p. 455-462 (cf. p. 38).
- [WBW15] J. WINTER, R. BRANDENBURG et K.-D. WELTMANN. « Atmospheric pressure plasma jets : an overview of devices and new directions ». en. In : *Plasma Sources Sci. Technol.* 24.6 (oct. 2015). Publisher : IOP Publishing, p. 064001 (cf. p. 172).
- [Whe89] John Archibald WHEELER. « Information, Physics, Quantum : The Search for Links ». In : *Proceedings III International Symposium on Foundations of Quantum Mechanics*. 1989, p. 354-358 (cf. p. 52).
- [Wik20] WIKIPÉDIA. *The Persistence of Memory*. en. Avr. 2020 (cf. p. 53).
- [Xia+16] Bo XIAO et al. « Focusing waves at arbitrary locations in a ray-chaotic enclosure using time-reversed synthetic sonas ». In : *Phys. Rev. E* 93.5 (mai 2016). Publisher : American Physical Society, p. 052205 (cf. p. 199).
- [ZP08] Xi-Ming ZHU et Yi-Kang PU. « Using OES to determine electron temperature and density in low-pressure nitrogen and argon plasmas ». en. In : *Plasma Sources Sci. Technol.* 17.2 (mai 2008). Publisher : IOP Publishing, p. 024002 (cf. p. 196).

Résumé — Un des défis auxquels sont confrontées les technologies plasmas est le développement de concepts de sources de plasma qui peuvent être transposés à des dimensions plus importantes, avec pour objectif par exemple le traitement de surface d’objets de grandes dimensions (panneaux solaires, écrans plats, grand wafer...). Cependant, plus la surface à traiter est grande, plus il devient difficile de la traiter. En fait, lorsqu’on augmente la surface, ces technologies se heurtent à des limites physiques et technologiques. Pour les sources plasmas microondes, l’incapacité à contrôler les plasmas vient du caractère multi-mode des grandes cavités, qui rend le contrôle de la distribution spatiale du champ électrique (et donc du plasma) en leur sein très difficile avec les technologies conventionnelles.

L’objectif de la thèse est de développer une source plasma micro-onde répondant à ce besoin de contrôle des plasmas dans les grandes cavités (multimodes). Nous avons alors introduit un concept innovant de source de plasma micro-ondes : le “pinceau plasma ondulatoire”. Le principe de cette source plasma consiste à contrôler dynamiquement la position du plasma dans une cavité en jouant sur la forme d’onde du signal transmis à la cavité. L’idée est alors d’utiliser le “Retournement Temporel”, qui permet de focaliser spatio-temporellement l’énergie électromagnétique dans les cavités de grandes dimensions.

Cette thèse propose les premières études théoriques et les premières démonstrations expérimentales du concept de “pinceau plasma ondulatoire”.

Mots clés : Source plasma microonde, Retournement temporel, cavité réverbérante.

Abstract — One of the challenges that face plasma technologies is the development of plasma source concepts that can be scaled to larger dimensions, for example for the surface processing of large objects (solar panels, flat screens, large wafers, etc.). However, the more the area to be processed is large, the more it is difficult to process it. In fact, when increasing the area these technologies face physical and technological limitations. For microwave plasma sources, the inability to control plasmas comes from the multi-mode nature of large cavities, which makes the control of the spatial distribution of the electric field (and thus of the plasma) inside very difficult with conventional technologies.

The objective of this thesis is to develop a microwave plasma source addressing this need for plasma control in large cavities (multimode). We then introduced an innovative concept of microwave plasma source : the “wave plasma brush”. The principle of this plasma source consists in dynamically controlling the position of the plasma in a cavity by playing on the waveform of the transmitted signal to the cavity. The idea is then to use “Time Reversal”, which allows a spatio-temporal focusing of the electromagnetic energy in large cavities.

This thesis proposes the first theoretical studies and the first experimental demonstrations of the concept of “wave plasma brush”.

Keywords : Microwave plasma source, Time reversal, Reverberant cavity.
