



Retournement temporel d'ondes électromagnétiques et application à la télécommunication en milieux complexes

Geoffroy Lerosey

► To cite this version:

Geoffroy Lerosey. Retournement temporel d'ondes électromagnétiques et application à la télécommunication en milieux complexes. Physics [physics]. ESPCI ParisTECH, 2006. English. NNT : . pastel-00003585

HAL Id: pastel-00003585

<https://pastel.hal.science/pastel-00003585>

Submitted on 9 Jun 2008

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

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

Thèse de Doctorat de l'Université Paris 7 – Denis Diderot

UFR de Physique

Spécialité :

Physique Macroscopique

Présentée par :

Geoffroy Lerosey

Pour obtenir le grade de Docteur de l'Université Paris 7

Retournement temporel d'ondes
électromagnétiques et application à la
télécommunication en milieux complexes

sous la direction de Mathias Fink

Soutenance le 13 Décembre à 15h

Devant le jury composé de :

Claude Boccara

Rémi Carminati Rapporteur

Ghaïs El Zein Rapporteur

Mathias Fink

Roger Maynard

Anne Sentenac

Julien de Rosny

Table des matières

I	Le Retournement temporel : de l'acoustique à l'électromagnétisme	6
I.1	Ondes acoustiques et retournement temporel	8
I.1.1	Réversibilité et réciprocity des équations	8
I.1.2	La cavité à retournement temporel	10
I.1.3	Miroirs à retournement temporel et milieux multidiffuseurs	15
I.1.4	Cas particulier d'un seul transducteur en cavité réverbérante	21
I.2	De l'acoustique à l'électromagnétisme	27
I.2.1	Réversibilité et réciprocity des ondes électromagnétiques	27
I.2.2	La cavité à retournement temporel électromagnétique	29
I.2.3	Application aux antennes	33
I.2.4	Influence du réseau d'antennes réceptrices lors de la mesure du champ retourné temporellement	38
I.3	Un premier miroir à retournement temporel dans les domaine des micro-ondes .	44
I.3.1	Retournement temporel d'un signal sur porteuse	44
I.3.2	Réponses impulsionnelles en bande de base	48
I.3.3	Dispositif expérimental et premiers essais	51
I.3.4	Retournement temporel électromagnétique en cavité réverbérante : les premiers résultats probants	55
I.4	Passage à un miroir multivoies large bande passante	60
I.4.1	Dispositif expérimental	60
I.4.2	Retournement temporel large bande	63
I.4.3	Influence de la bande passante et du nombre d'antennes du miroir sur le retournement temporel	68
I.4.4	Mise en évidence de la focalisation spatiale	72
I.5	Et les ondes évanescences ? Une première expérience de focalisation sub-longueur d'onde	78
I.5.1	Des mesures étonnantes	78
I.5.2	Champ proche, champ lointain et information spatiale	81
I.5.3	L'imagerie en champ proche	84
I.5.4	Résultats expérimentaux et premières justifications	87
I.6	Conclusion	96
II	Télécommunications et retournement temporel	98
II.1	Les télécommunications modernes : problèmes et techniques actuelles	100
II.1.1	Les systèmes Multiple-Input Multiple-Output (MIMO)	100
II.1.2	Les communications Ultra Large Bande (UWB)	105
II.1.3	L'apport du retournement temporel aux communications sans fils	110

II.1.4	Première expérience en ultrasons	112
II.2	Télécommunications MISO Ultra large bande en environnement indoor à petite échelle	118
II.2.1	Principe des expériences et dispositif expérimental	118
II.2.2	Résultats expérimentaux	123
II.2.3	Modélisation et discussion	127
II.2.4	Vers les communications MIMO-Mu Ultra Large Bande	131
II.3	Filtrage inverse ou filtrage adapté ?	135
II.3.1	Dispositif expérimental et principe du retournement temporel itératif . .	135
II.3.2	Comparaison des méthodes dans le contexte des télécommunications . . .	140
II.3.3	Comparaison des méthodes en fonction du bruit	144
II.4	Communications dans le gigahertz en cavité réverbérante	150
II.4.1	Principe des expériences et dispositif expérimental	150
II.4.2	Influence du nombre d'antennes dans le miroir et du nombre de récepteurs.	153
II.4.3	Télécommunications dans un canal bruité	156
II.4.4	Focalisation sub-longueur d'onde et télécommunications	161
II.5	Conclusion	165

Introduction

Les communications sans fil sont devenues ces dernières années un véritable phénomène de société. Les acteurs du marché, en proposant des produits toujours plus perfectionnés ainsi que les services qui leurs sont associés, ont transformé les téléphones mobiles en véritable plateforme de loisir. Ainsi, les utilisateurs peuvent désormais regarder des émissions de télévision sur leurs mobiles, téléphoner en visioconférence, ou encore télécharger le dernier titre de Madonna en exclusivité. De même, en ce qui concerne Internet, de plus en plus de consommateurs se tournent à l'heure actuelle vers les réseaux locaux sans fils, qui permettent une connection pratique à la toile quel que soit l'endroit où ils se trouvent. Ce besoin de connectivité sans fil concerne tous les composants électroniques de notre quotidien, de la télévision au téléphone portable, en passant par l'appareil photo et la chaîne haute fidélité. Les volumes de données transmises connaissent une croissance exponentielle, augmentant la complexité des terminaux d'émission et de réception et les bandes passantes utilisées.

Cependant, ces besoins technologiques sont bloqués par des réalités physiques et normatives. En particulier, la bande passante utilisable pour chaque type de produit est réglementée par des institutions nationales ou internationales, ce qui oblige les opérateurs à limiter la partie du spectre allouée à chacun des utilisateurs. Il est alors possible d'inventer des modulations plus complexes qui permettent de transmettre plus d'information sur une bande passante égale, mais là encore une limitation existe. En effet, les énergies d'émission sont elles aussi restreintes et le bruit présent dans l'environnement interdit des modulations trop complexes.

Parmi les autres problèmes rencontrés lorsque la quantité d'information transmise augmente, l'effet de l'environnement sur les signaux est très important. En effet, les ondes subissent des réflexions sur les divers obstacles rencontrés pendant leur propagation. Ce phénomène se traduit, lorsque la bande passante est faible, par des déphasages sur les symboles émis, et donc par une augmentation des erreurs commises à la réception. Il est possible, moyennant une connaissance du canal de propagation que l'on peut obtenir par apprentissage, de corriger en temps réel les

effets néfastes du milieu. Cette méthode est d'ailleurs utilisée dans la plupart des systèmes de communication modernes. Cependant, à mesure que l'on augmente le débit d'information et donc la bande passante, la complexité des méthodes de correction croît de façon exponentielle, notamment à cause de l'allongement temporel des symboles émis.

Peu avant le début de cette thèse, deux techniques ont été proposées afin d'améliorer l'efficacité spectrale des systèmes de communication, c'est-à-dire d'augmenter, à bande passante constante, la quantité d'information transmise. La première, publiée par Andrews et ses collaborateurs dans la revue *Nature* [1], proposait d'utiliser les différentes polarisations des ondes électromagnétiques afin de tripler la quantité d'information des systèmes de communication sans fils. Cette méthode, bien que très astucieuse, n'a pour l'instant pas trouvé d'applications du fait de sa difficulté de mise en oeuvre. Dans le même temps, Moustakas et ses collaborateurs ont montré dans la revue *Science* [2], qu'il est possible de profiter de la complexité d'un environnement de propagation afin de transmettre plus d'information, ceci en profitant de canaux de communication indépendants.

Dans ce contexte, le Laboratoire Ondes et Acoustique, sous l'impulsion de Mathias Fink, s'est lancé dans le retournement temporel d'ondes électromagnétiques et dans son application aux télécommunications. Le retournement temporel est une technique qui permet de focaliser spatialement et temporellement une onde à partir de la connaissance de la réponse impulsionnelle entre deux points d'un milieu complexe [3]. Dans un premier chapitre, nous traiterons de l'application du retournement temporel aux ondes électromagnétiques théoriquement et expérimentalement, après avoir rappelé certains résultats obtenus en acoustique. Le deuxième chapitre de ce manuscrit sera consacré à l'application du retournement temporel aux télécommunications. Nous verrons, après avoir introduit la technique dans le contexte des communications modernes, quels pourraient être les intérêts apportés par le retournement temporel, au travers d'expériences réalisées avec des ondes ultrasonores puis électromagnétiques.

Chapitre I

Le Retournement temporel : de l'acoustique à l'électromagnétisme

Le problème de la réversibilité du temps dans un processus peut être illustré par l'expérience suivante : un bloc de matière explose en de nombreux fragments et l'on veut créer la scène inverse afin de reconstituer le bloc. Conceptuellement, il est possible d'envisager qu'après avoir mesuré la vitesse et la position de chacun de ces fragments sur une sphère, on les renvoie dans la direction exacte d'où ils viennent et avec la même vitesse. Les fragments convergent alors vers le point d'explosion initial, comme si l'on avait filmé le phénomène et passé la bande en sens inverse. En effet, les équations qui gouvernent le mouvement de chacune des particules sont invariantes par renversement du temps.

Cette expérience de pensée, bien que physiquement acceptable, est en fait irréalisable. En premier lieu, le nombre de particules mises en cause est beaucoup trop grand pour avoir toutes les informations les concernant, et ainsi recréer la scène à l'envers. De plus, ce système chaotique est très sensible aux conditions initiales : sur ce type de mouvements divergents, une erreur commise sur un vecteur vitesse initial se propage exponentiellement lors de la réémission du fait des multiples collisions entre les particules. Pour cette raison, le retournement du temps en mécanique classique est impossible et c'est donc vers un type de physique moins sensible qu'il faut se tourner.

En physique ondulatoire, au contraire, une quantité d'information finie permet de décrire parfaitement un champ d'ondes. En effet le plus petit détail utile à l'expérimentateur pour définir le champ est de l'ordre de la plus petite longueur d'onde contenue dans le système. De plus, conséquence de la linéarité des systèmes ondulatoires, les erreurs commises à l'émission de cer-

taines ondes ne vont pas se répercuter sur le reste de l'information comme c'était le cas en physique corpusculaire, ce qui garantit une sensibilité bien moindre aux conditions initiales. Pour ces raisons, le retournement temporel d'ondes est possible même dans des systèmes complexes.

Dans un premier temps nous verrons dans ce chapitre comment le retournement temporel a pu être appliqué d'abord en acoustique, et plus particulièrement dans le domaine des ultrasons. Nous rappellerons brièvement pourquoi, et dans quelles conditions, le retournement temporel (RT) est applicable avec ce type d'ondes, et introduirons la notion de miroir à retournement temporel (MRT) dont nous dégagerons les propriétés en milieu hétérogène, ainsi qu'en cavité réverbérante.

Dans une deuxième partie, nous montrerons comment ces concepts développés en acoustique ultrasonore peuvent être étendus au domaine des ondes électromagnétiques grâce à la similarité des équations de propagation qui régissent ces champs ondulatoires. Le caractère vectoriel des champs électromagnétiques sera pris en compte. Nous démontrerons la théorie du renversement du temps en nous appuyant sur le théorème de réciprocité de Lorentz.

Puis nous décrirons deux réalisations expérimentales du retournement temporel d'ondes électromagnétiques dans le gigahertz. Dans notre premier montage, nous introduirons une technique de modulation/démodulation afin de pouvoir utiliser du matériel d'acquisition et de génération numérique basse fréquence. Ce prototype, dont la bande passante est limitée à 5 MHz, nous a permis d'obtenir nos premiers résultats.

Nous décrirons ensuite la conception du deuxième dispositif expérimental multi-voies et plus large bande avec lequel il nous a été possible d'étudier plus quantitativement qu'auparavant l'efficacité du procédé. Nous exposerons alors des résultats expérimentaux et nous les comparerons à ceux qu'avaient obtenus les acousticiens quelques années auparavant. Nous nous attacherons également à mettre en relief les différences que l'on a rencontrées notamment en ce qui concerne les réseaux de capteurs.

Enfin, dans une dernière partie, nous expliquerons comment certains résultats surprenants de prime abord nous ont conduits à étudier l'influence d'objets placés dans le champs proche de la source initiale. La possibilité d'une participation des ondes évanescences à la focalisation sera invoquée au travers d'une analogie avec l'imagerie en champ proche, et les premières expériences de focalisation d'énergie sur une taille petite devant la longueur d'onde seront présentées et discutées.

I.1 Ondes acoustiques et retournement temporel

I.1.1 Réversibilité et réciprocity des équations

En acoustique la propagation d'une onde en milieu fluide hétérogène et non dissipatif est gouvernée par l'équation suivante :

$$\rho_0(\mathbf{r}) \nabla \cdot \left(\frac{1}{\rho_0(\mathbf{r})} \nabla \Phi(\mathbf{r}, t) \right) = \frac{1}{c_0^2(\mathbf{r})} \frac{\partial^2 \Phi(\mathbf{r}, t)}{\partial t^2} \quad (\text{I.1})$$

où $\Phi(\mathbf{r}, t)$ est le potentiel acoustique de l'onde au point $\mathbf{r}$, $c_0(\mathbf{r})$ correspond à la distribution spatiale de célérité du son du milieu et $\rho_0(\mathbf{r})$ à la distribution de densité du milieu. Cette équation présente la propriété d'être invariante par renversement du temps. En effet, si un potentiel $\Phi(\mathbf{r}, t)$ en est solution, alors $\Phi(\mathbf{r}, -t)$ l'est également car nous sommes uniquement en présence de dérivés d'ordre pair en temps. Cette propriété implique que pour toute onde divergente $\Phi(\mathbf{r}, t)$, il existe une onde $\Phi(\mathbf{r}, -t)$ qui converge vers sa source acoustique.

L'équation précédente est très pratique car elle permet de déduire à la fois la vitesse de l'onde et la pression à partir d'un potentiel unique qui n'a pas de signification physique mais permet un traitement mathématique efficace. Cependant, on préférera dans la suite manipuler la pression et la vitesse des ondes plutôt que le potentiel comme cela est fait habituellement, afin de faciliter le parallèle établi avec les ondes électromagnétiques. Dans un premier temps, nous allons démontrer le principe de réciprocity de Helmholtz-Kirchhoff à partir du système d'équations linéaires et couplées qui régit l'évolution du champ acoustique (équation de conservation de la masse, équation d'Euler, équation d'Etat). Dans le cas de l'acoustique linéaire, ce système peut s'écrire, sous l'hypothèse d'adiabaticité, sans terme source et sans écoulement de la façon suivante :

$$\begin{cases} \frac{\partial}{\partial t} \rho(\mathbf{r}, t) &= -\nabla \cdot (\rho_0(\mathbf{r}) \mathbf{v}(\mathbf{r}, t)) \\ \rho_0(\mathbf{r}) \cdot \frac{\partial}{\partial t} \mathbf{v}(\mathbf{r}, t) &= -\nabla p(\mathbf{r}, t) \\ p(\mathbf{r}, t) &= c_0^2(\mathbf{r}) \cdot \rho(\mathbf{r}, t) \end{cases} \quad (\text{I.2})$$

où l'on fait intervenir $\rho(\mathbf{r}, t)$ la variation de densité, $\mathbf{v}(\mathbf{r}, t)$, la vitesse particulière et $p(\mathbf{r}, t)$ la variation de pression du milieu au point $\mathbf{r}$ et au temps t .

Ici si l'on change la variable t en $-t$, la vitesse particulière $\mathbf{v}_a(\mathbf{r}, t)$ qui correspond à l'onde divergente deviendra $-\mathbf{v}_a(\mathbf{r}, -t)$ dans le cas de l'onde convergente (ceci peut s'expliquer physiquement par le fait que le vecteur vitesse est une dérivé temporelle du vecteur position). Ce

système présentant deux dérivées temporelles d'ordre 1, il est possible de réaliser un retournement temporel instantané en enregistrant à un instant donné dans tout l'espace à la fois la pression et la vitesse et en renvoyant la pression inchangée et la vitesse en sens opposé. Bien sûr, cette approche en trois dimensions n'est pas techniquement réalisable. En fait, nous allons voir que la connaissance de la pression et de la vitesse normale sur une surface fermée (2D) suffit à décrire entièrement le champ dans tout le volume englobé par cette surface.

Ce résultat découle du théorème de Helmholtz-Kirchhoff qui peut être établi à partir des équations de l'acoustique linéaire de la façon suivante. Supposons deux champs de pression, vitesse et densité $\{p_1, \mathbf{v}_1, \rho_1\}$ et $\{p_2, \mathbf{v}_2, \rho_2\}$ qui résultent de deux sources s_1 et s_2 . Ces champs sont solutions du système d'équation I.2, dans lequel on inclut les sources. La première équation du système s'écrit alors en régime harmonique pour chacune des sources :

$$\begin{cases} i\omega \frac{p_1(\mathbf{r}, \omega)}{c_0^2(\mathbf{r})} = -\nabla \cdot (\rho_0(\mathbf{r}) \mathbf{v}_1(\mathbf{r}, \omega)) + s_1(\mathbf{r}) \\ i\omega \frac{p_2(\mathbf{r}, \omega)}{c_0^2(\mathbf{r})} = -\nabla \cdot (\rho_0(\mathbf{r}) \mathbf{v}_2(\mathbf{r}, \omega)) + s_2(\mathbf{r}) \end{cases} \quad (\text{I.3})$$

La première équation est multipliée par la pression p_2 et la deuxième par p_1 puis les deux équations obtenues sont soustraites, on obtient alors :

$$-p_2(\mathbf{r}, \omega) \nabla \cdot (\rho_0(\mathbf{r}) \mathbf{v}_1(\mathbf{r}, \omega)) + p_2(\mathbf{r}, \omega) s_1(\mathbf{r}) = -p_1(\mathbf{r}, \omega) \nabla \cdot (\rho_0(\mathbf{r}) \mathbf{v}_2(\mathbf{r}, \omega)) + p_1(\mathbf{r}, \omega) s_2(\mathbf{r}) \quad (\text{I.4})$$

soit encore :

$$p_1(\mathbf{r}, \omega) \nabla \cdot (\rho_0(\mathbf{r}) \mathbf{v}_2(\mathbf{r}, \omega)) - p_2(\mathbf{r}, \omega) \nabla \cdot (\rho_0(\mathbf{r}) \mathbf{v}_1(\mathbf{r}, \omega)) = p_1(\mathbf{r}, \omega) s_2(\mathbf{r}) - p_2(\mathbf{r}, \omega) s_1(\mathbf{r}) \quad (\text{I.5})$$

Afin d'obtenir le théorème voulu, il reste maintenant à simplifier le membre de droite de cette équation en tenant compte du fait que $\mathbf{v}_1(\mathbf{r}, \omega) \cdot \nabla p_2(\mathbf{r}, \omega) = \mathbf{v}_2(\mathbf{r}, \omega) \cdot \nabla p_1(\mathbf{r}, \omega)$ (cette relation est une conséquence de l'équation de conservation de la masse de (I.2)). Puis on intègre sur un volume V en utilisant le théorème de Gauss relatif aux flux du gradient d'un champ vectoriel qui traverse la surface S englobant le volume V :

$$\oint_S \rho_0(\mathbf{r}) \{p_2(\mathbf{r}, \omega) \mathbf{v}_1(\mathbf{r}, \omega) - p_1(\mathbf{r}, \omega) \mathbf{v}_2(\mathbf{r}, \omega)\} \cdot \mathbf{dS} = \iiint_V (p_1(\mathbf{r}, \omega) s_2 - p_2(\mathbf{r}, \omega) s_1) dV \quad (\text{I.6})$$

qui constitue le théorème de réciprocité de Helmholtz-Kirchhoff reliant deux champs de pressions et vitesses aux sources qui les ont créés.

Ce théorème va nous permettre dans la suite de montrer comment le retournement temporel est envisageable de façon formelle. D'autre part il nous permet de rappeler un résultat essentiel concernant les expériences de retournement temporel qui ont été faites en acoustique.

Prenons en effet comme sources s_1 et s_2 deux Diracs spatiaux $\delta(\mathbf{r} - \mathbf{r}_1)$ et $\delta(\mathbf{r} - \mathbf{r}_2)$. Les champs de pressions solutions, $p_1(\mathbf{r}, \omega)$ et $p_2(\mathbf{r}, \omega)$, sont alors appelés fonctions de Green et sont notés $G(\mathbf{r}_1, \mathbf{r}, \omega)$ et $G(\mathbf{r}_2, \mathbf{r}, \omega)$. Il n'est en général pas facile de calculer analytiquement ces fonctions de Green mais on peut montrer qu'elles décroissent, quand r est grand, en $\frac{1}{r}$. Il est également possible de prouver en prenant comme volume d'intégration une sphère de rayon infiniment grand, que le membre de gauche de cette équation est nul. On obtient alors :

$$G(\mathbf{r}_1, \mathbf{r}_2, \omega) = G(\mathbf{r}_2, \mathbf{r}_1, \omega) \quad (\text{I.7})$$

Cette dernière équation démontre, en régime harmonique, la réciprocité d'un milieu hétérogène : c'est-à-dire que le champ produit en $\mathbf{r}_2$ par une source ponctuelle en $\mathbf{r}_1$ est égal à celui qui est créé en $\mathbf{r}_1$ par une source ponctuelle en $\mathbf{r}_2$, le champ réciproque étant ici la pression. En pratique, dans une expérience de retournement temporel, les émetteurs et récepteurs pourront être interchangés car leurs rôles sont réciproques.

I.1.2 La cavité à retournement temporel

Le concept de cavité à retournement temporel a été développé par D. Cassereau et M. Fink [4]. Il repose sur les théorèmes énoncés précédemment. En effet, les auteurs ont démontré qu'en exploitant le principe de Helmholtz-Kirchoff, l'opération de Retournement Temporel d'un champ ne consiste plus à inverser le champ acoustique en tout point du volume considéré à un instant donné, mais seulement le champ et sa dérivée normale sur la surface délimitant ce volume. Il est alors possible, du moins en pensée, d'imaginer une surface de transducteurs acoustiques englobant complètement un milieu hétérogène donné. Si, au sein de ce milieu, une source émet une onde acoustique divergente, on sait alors, puisque l'on mesure le champ sur toute la surface, que nous avons une connaissance totale du champ dans le milieu. De plus, ce champ étant solution de l'équation de propagation des ondes acoustiques, nous pouvons fabriquer le champ "Retourné Temporel" qui lui est associé. D'après ce que nous avons vu dans la première partie, ce champ correspond alors à l'onde convergente associée à l'onde qui avait été initialement émise. Cette cavité permet donc littéralement de renvoyer vers sa source une onde initialement divergente qui a été émise dans le milieu. Dès lors, une expérience de

retournement temporel peut se décrire en deux phases :

- Une phase, dite d'enregistrement, durant laquelle la source émet une impulsion de durée donnée¹, la cavité formée de transducteurs enregistrant alors le champ et la dérivée normale du champ.
- Une phase, dite de réémission, durant laquelle chaque transducteur de la cavité se comporte comme une source et émet le champ ainsi que sa dérivée normale en sens opposé, ces deux grandeurs étant émises dans une chronologie inversée par rapport à la réception.

Le principe d'une telle expérience peut se représenter sous la forme de deux schémas qui permettent une approche visuelle du phénomène. Ces schémas sont représentés dans les figures I.1 et I.2.

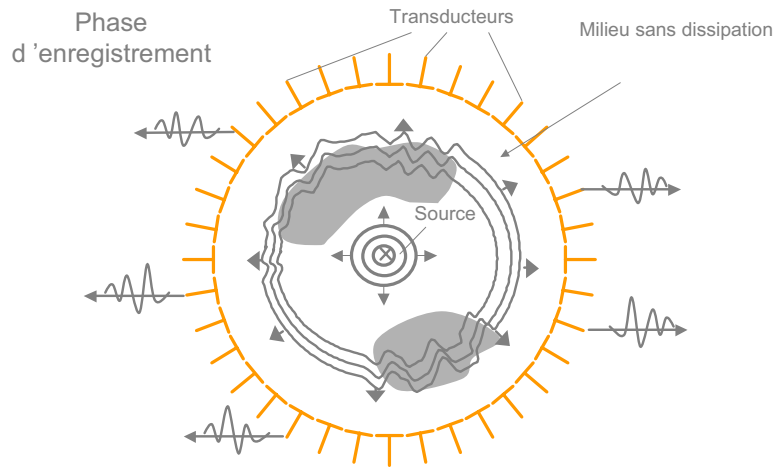


FIG. I.1 – Phase d'enregistrement du champ

Une source située à l'intérieur de la cavité émet une impulsion acoustique dans un milieu qui peut être hétérogène. L'onde sphérique générée se réfléchit, diffuse et diffracte de manière complexe lors de son passage dans le milieu hétérogène. Lorsque l'onde atteint la surface S , le champ de pression $p(\mathbf{r}, t)$ et la vitesse normale $\mathbf{v}(\mathbf{r}, t)$ sont enregistrés par des transducteurs, puis mis en mémoire. La phase d'acquisition s'arrête lorsqu'il n'y a plus d'énergie à l'intérieur de la cavité.

¹Les théorèmes sont tous démontrés en régime harmonique pour des raisons pratiques mais sont généralisables aux signaux quelconques, ainsi la référence [5] propose une étude temporelle de ce formalisme.

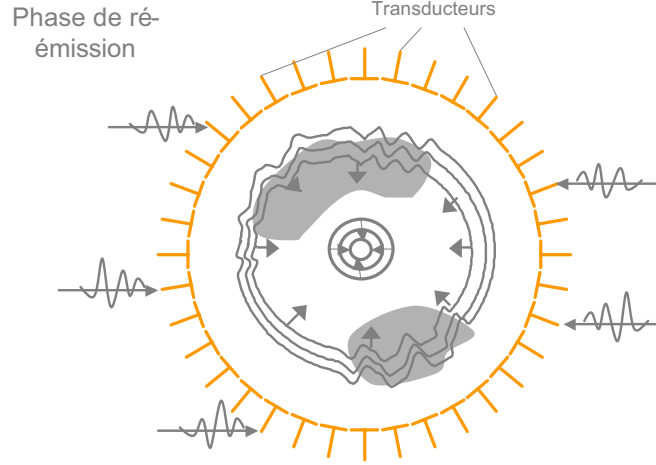


FIG. I.2 – Phase de réémission des signaux retournés temporellement

Durant cette deuxième phase, les versions retournées temporellement du champ de pression $p(\mathbf{r}, -t)$ et de la vitesse normale $-\mathbf{v}(\mathbf{r}, -t)$ sont réémises par chaque transducteur de la surface dans le milieu. Didier Cassereau et Mathias Fink ont montré que le champ créé dans tout le volume de la cavité est alors égal à celui de l'onde émise initialement mais retournée temporellement [4]. Cette onde converge donc vers son point source initial. Les auteurs se sont également posé la question de la nature de l'onde au point d'émission initial. En effet, dans les raisonnements précédents, tout donne l'impression que l'on recrée au sein du milieu la réplique exacte du champ qui avait été émis lors de la phase d'enregistrement, et donc la source initiale. On pourrait donc s'attendre à obtenir une finesse de focalisation qui dépend uniquement de la taille du point source. En réalité la taille de la tache focale, qui est définie comme la largeur à mi-hauteur de la zone sur laquelle l'énergie se concentre, est limitée par la diffraction. Sa dimension caractéristique est donc de l'ordre de la demi longueur d'onde associée aux signaux utilisés ($\lambda/2$).

Ce dernier résultat peut également être démontré simplement en utilisant le théorème de Helmholtz-Kirchhoff démontré auparavant. Pour ce faire, nous allons considérer que deux points placés en $\mathbf{r}_1$ et $\mathbf{r}_2$ au sein de la cavité ont émis chacun une onde divergente monochromatique créant ainsi les champs $\{p_1, \mathbf{v}_1\}$ et $\{p_2, \mathbf{v}_2\}$. Nous pouvons alors relier le conjugué du premier champ $\{p_1^*, -\mathbf{v}_1^*\}$ au deuxième champ par le théorème de Helmholtz-Kirchoff, ceci s'écrit :

$$\begin{aligned} \oint_S \rho_0(\mathbf{r}) \{ -p_2(\mathbf{r}, \omega) \mathbf{v}_1^*(\mathbf{r}, \omega) - p_1^*(\mathbf{r}, \omega) \mathbf{v}_2(\mathbf{r}, \omega) \} d\mathbf{S} \\ = \iiint_V (p_1^*(\mathbf{r}, \omega) \delta(\mathbf{r} - \mathbf{r}_2) - p_2(\mathbf{r}, \omega) \delta^*(\mathbf{r} - \mathbf{r}_1)) dV \quad (\text{I.8}) \end{aligned}$$

où $\mathbf{r}$ correspond à un point des différents domaines d'intégration.

On peut alors exprimer les pressions et vitesses grâce aux fonctions de Green qui s'écrivent pour le champ créé par le point source $\mathbf{r}_1$ (resp. pour le point source $\mathbf{r}_2$) : $p_1(\mathbf{r}, \omega) = G(\mathbf{r}_1, \mathbf{r}, \omega)$ et $\mathbf{v}_1(\mathbf{r}, \omega) = -\frac{1}{i\rho_0(\mathbf{r})\omega} \nabla G(\mathbf{r}_1, \mathbf{r}, \omega)$.

$$\begin{aligned} \oint_S \{ G(\mathbf{r}_2, \mathbf{r}, \omega) \nabla G^*(\mathbf{r}_1, \mathbf{r}, \omega) - G^*(\mathbf{r}_1, \mathbf{r}, \omega) \nabla G(\mathbf{r}_2, \mathbf{r}, \omega) \} d\mathbf{S} \\ = i\omega \iiint_V G^*(\mathbf{r}_1, \mathbf{r}, \omega) \delta(\mathbf{r} - \mathbf{r}_2) - G(\mathbf{r}_2, \mathbf{r}_1, \omega) \delta(\mathbf{r} - \mathbf{r}_1) dV \quad (\text{I.9}) \end{aligned}$$

A présent nous considérons que les sources sont réelles, soit que $\delta^*(\mathbf{r} - \mathbf{r}_1) = \delta(\mathbf{r} - \mathbf{r}_1)$, et on utilise également la réciprocité du milieu qui peut s'écrire $G(\mathbf{r}_1, \mathbf{r}_2, \omega) = G(\mathbf{r}_2, \mathbf{r}_1, \omega)$. En effectuant les derniers calculs nous obtenons :

$$\oint_S \{ G(\mathbf{r}, \mathbf{r}_2, \omega) \nabla G^*(\mathbf{r}_1, \mathbf{r}, \omega) - G^*(\mathbf{r}_1, \mathbf{r}, \omega) \nabla G(\mathbf{r}, \mathbf{r}_2, \omega) \} d\mathbf{S} = -2i\omega \Im \{ G(\mathbf{r}_1, \mathbf{r}_2, \omega) \} \quad (\text{I.10})$$

où l'opérateur $\Im$ représente la partie imaginaire.

A présent nous pouvons expliciter le membre de gauche de cette expression. Nous voyons apparaître des fonctions de Green qui vont de $\mathbf{r}_1$ à $\mathbf{r}$ et qui sont conjuguées : ce sont les signaux créés au point $\mathbf{r}_1$ que notre cavité à retournement temporel a enregistré en tout point $\mathbf{r}$ de la surface fermée et qu'il renvoie dans une chronologie inversée. Les fonctions de Green qui vont de $\mathbf{r}$ à $\mathbf{r}_2$ relient, elles, les signaux émis par la cavité à retournement temporel au champ créé en un point $\mathbf{r}_2$ du volume à l'intérieur de la cavité. Cette équation est donc la traduction mathématique du champ créé par retournement temporel d'une source placée initialement en $\mathbf{r}_1$, lorsque celui-ci est observé en $\mathbf{r}_2$. La conclusion que l'on peut en tirer est la suivante : lorsque l'on réalise le retournement temporel d'une onde émise par un point initial, cela a pour résultat de générer dans le milieu la partie imaginaire de la fonction de Green qui lie le point source de l'émission à tout autre point du milieu. Cette remarque est particulièrement claire si l'on

suppose le milieu homogène, ce qui implique qu'une fonction de Green entre deux point $\mathbf{r}_1$ et $\mathbf{r}_2$ peut s'écrire $G(\mathbf{r}_1, \mathbf{r}_2, \omega) = \frac{\exp(i\mathbf{k} \cdot (\mathbf{r}_2 - \mathbf{r}_1))}{\mathbf{k} \cdot (\mathbf{r}_2 - \mathbf{r}_1)}$. Le champ obtenu par retournement temporel est alors un sinus cardinal de largeur à mi-hauteur typique $\lambda/2$, comme l'avaient montré Didier Cassereau et Mathias Fink :

$$\Phi_{RT}(\mathbf{r}_2, \omega) = -2i\omega \frac{\sin(\mathbf{k} \cdot (\mathbf{r}_2 - \mathbf{r}_1))}{\mathbf{k} \cdot (\mathbf{r}_2 - \mathbf{r}_1)} \quad (\text{I.11})$$

Ce développement en régime monochromatique de la cavité à retournement temporel se révèle pratique pour justifier le résultat obtenu pour un retournement temporel parfait, mais cache une autre propriété fondamentale du retournement temporel qui nous sera très utile par la suite et que nous allons présenter ici : le retournement temporel est un filtre adapté au sens du traitement du signal. Afin d'introduire cette propriété, il est nécessaire de considérer que les sources utilisées ne sont plus monochromatiques mais impulsionnelles : nous allons donc prendre en compte la variable temporelle. De même, pour se rapprocher des expériences réelles qui font intervenir des transducteurs monopolaires, seule la pression sera considérée lors de la phase de réémission des signaux.

Du point de vue du traitement du signal, dans le cas des systèmes linéaires et invariants par translation dans le temps, la réponse entre un point situé en $\mathbf{r}_1$ et un point situé en $\mathbf{r}$ est définie par sa réponse impulsionnelle $h(\mathbf{r}_1, \mathbf{r}, t)$. Quel que soit alors le signal $f(t)$ émis au point $\mathbf{r}_1$, on peut toujours écrire le signal reçu en $\mathbf{r}$ comme $h(\mathbf{r}_1, \mathbf{r}, t) \otimes f(t)$ et la réponse impulsionnelle prend ainsi en compte tous les éléments de la chaîne électro-acousto-électrique². Si l'on fait de plus l'hypothèse que les réponses des instruments de mesure et de génération de signaux sont parfaites, c'est-à-dire égales à un Dirac temporel, la réponse impulsionnelle entre un couple émetteur/récepteur monopolaire est alors directement proportionnelle à la fonction de Green :

$$h(\mathbf{r}_1, \mathbf{r}, t) \propto G(\mathbf{r}_1, \mathbf{r}, t) \quad (\text{I.12})$$

Avec cette notation, et compte tenu des hypothèses faites, on peut alors écrire le résultat du retournement temporel d'une source qui a émis en $\mathbf{r}_1$ un signal $f(t)$. En effet, nous pouvons considérer que chacune des sources élémentaires de la cavité émet alors le signal qu'elle a reçu et ceci dans une chronologie inversée : $h(\mathbf{r}_1, \mathbf{r}, -t) \otimes f(-t)$. Chaque point $\mathbf{r}_2$ reçoit alors de la part de toutes les sources tapissant la cavité le signal précédent qui est à nouveau convolué par la réponse impulsionnelle de la source au récepteur $h(\mathbf{r}, \mathbf{r}_2, t)$. Ainsi pour un point quelconque

²Ici et dans la suite le symbole $\otimes$ représente le produit de convolution.

situé en $\mathbf{r}_2$ le signal reçu peut s'écrire comme l'intégrale sur une surface de la résultante de toutes ces sources :

$$S_{RT}(\mathbf{r}_2, t) \propto \oint_S h(\mathbf{r}, \mathbf{r}_2, t) \otimes h(\mathbf{r}_1, \mathbf{r}, -t) \otimes f(-t) dS \quad (\text{I.13})$$

où $\mathbf{r}$ est à nouveau un point de la surface d'intégration.

Si on cherche à considérer à présent le champ engendré au point source initial, il suffit de choisir $\mathbf{r}_2 = \mathbf{r}_1$ pour obtenir le résultat du champ en temporel, au point de focalisation :

$$S_{RT}(\mathbf{r}_1, t) \propto \oint_S h(\mathbf{r}, \mathbf{r}_1, t) \otimes h(\mathbf{r}_1, \mathbf{r}, -t) \otimes f(-t) dS \quad (\text{I.14})$$

On retrouve alors la définition même du filtre adapté comme utilisé en traitement du signal. Ce résultat s'explique simplement "avec les mains". Lors de la phase d'enregistrement du retournement temporel, le signal émis initialement est filtré par le milieu de propagation qui peut être aberrateur et dispersif. Ainsi, chaque composante spectrale des réponses impulsionnelles présente une amplitude et une phase particulières. Le retournement temporel ne fait alors rien de plus que de renvoyer dans le milieu chaque réponse impulsionnelle sans en modifier les amplitudes relatives des diverses composantes fréquentielles, afin qu'elles soient exactement adaptées au milieu, mais tout en compensant leurs phases afin qu'elles se somment toutes de façon cohérente en un point et à un temps donnés.

I.1.3 Miroirs à retournement temporel et milieux multidiffuseurs

Si le concept de Cavité à Retournement Temporel permet théoriquement aux ondes acoustiques de parcourir leurs trajectoires initiales à l'inverse, en pratique cette cavité s'avère difficilement réalisable. D'après le critère de Shannon, un champ ultrasonore et sa dérivée normale sont correctement échantillonnés spatialement si la distance qui sépare deux transducteurs est inférieure à la demi-longueur d'onde. Supposons que l'on désire créer une petite cavité sphérique, immergée dans l'eau, dont la surface est recouverte de transducteurs fonctionnant à la fréquence centrale de 1 MHz. La longueur d'onde moyenne correspondant à ces fréquences est de 1,5 mm. Ainsi les transducteurs ne devront pas être espacés de plus de 0,75 mm, ce qui implique, pour une cavité d'une dizaine de centimètres de diamètre, la présence de plus de 17000 transducteurs reliés chacun à une électronique propre ! Pour réaliser concrètement des expériences de Retournement Temporel, la surface sur laquelle est enregistrée l'onde acoustique doit donc être réduite. Au

laboratoire Ondes et Acoustique, les expériences de Retournement Temporel sont réalisées au moyen de sondes acoustiques classiques, composées de plusieurs transducteurs, dont l'ouverture totale est limitée. Une telle sonde associée à une électronique spécifique, c'est-à-dire capable de piloter chaque transducteur indépendamment des autres, forme alors un Miroir à Retournement Temporel (MRT).

Si la diminution de l'ouverture angulaire permet la réalisation pratique de tels miroirs, elle en est également la principale limitation. Tandis que la cavité perçoit l'onde acoustique divergente sur 4π stéradians pendant la première phase du retournement temporel, le miroir, lui, enregistre l'onde sur l'ouverture finie du réseau de transducteurs ce qui induit inévitablement une perte d'information. Lors de la phase de réémission, seule une partie de l'onde peut être retournée temporellement et la qualité de focalisation se dégrade.

Nous disposons au laboratoire de différents types de miroirs à retournement temporel qui peuvent être plans ou préfocalisés, à une ou deux dimensions. Leurs principales limitations sont les suivantes :

- Dans l'approximation de Fresnel, la résolution spatiale optimale du MRT est donnée en milieu homogène, par la largeur à -6 dB de la tache focale soit :

$$\delta \approx \frac{\lambda F}{D}$$

où D est la taille du miroir et F la distance focale.

- Dans le domaine temporel, la bande passante limitée des transducteurs introduit un étalement temporel des signaux. Le retournement temporel d'un pic de Dirac n'est en effet pas un pic de Dirac mais l'autocorrélation de la réponse acousto-électrique des transducteurs. En d'autres termes, le signal initial est filtré deux fois par les transducteurs et c'est ce qui définit la compression optimale du MRT.
- L'échantillonnage spatial du MRT est à l'origine de "lobes de réseaux" autour de la tache focale. Ces lobes sont minimisés lorsque le pas est égal à la taille des transducteurs.
- L'échantillonnage temporel du signal sur le miroir est responsable de lobes secondaires dans la focalisation spatiale. Un échantillonnage temporel de $\frac{T}{8}$, où T est la période centrale de l'onde acoustique, suffit à limiter ces lobes à -30 dB [6].

L'utilisation de miroirs à retournement temporel conduit donc à d'importantes limitations dans les milieux parfaitement homogènes. Comme nous allons le voir, leur utilisation est bien plus efficace dans les milieux complexes. En effet, si les miroirs à retournement temporel ont une

ouverture angulaire trop petite pour bénéficier d’une bonne focalisation, il est toutefois possible d’augmenter virtuellement leur taille dans certaines conditions.

Les premières expériences de retournement temporel dans des milieux dits désordonnés ont été réalisées par Arnaud Derode, Philippe Roux et Mathias Fink en 1995 [7]. Le principe de l’expérience est le suivant. Le milieu désordonné consiste en une collection de tiges identiques parallèles immergées dans l’eau. L’objectif est d’étudier le retournement temporel d’une onde émise par une source située d’un côté de la forêt de tiges à l’aide d’un miroir à retournement temporel situé de l’autre côté.

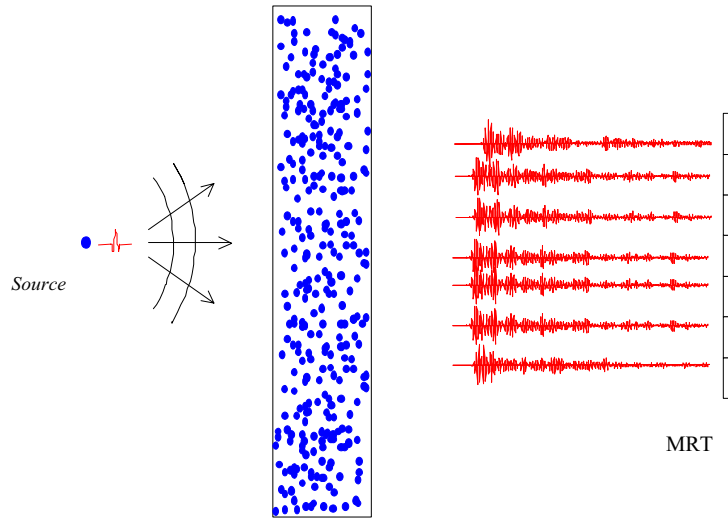


FIG. I.3 – Phase d’enregistrement du champ

Dans la phase d’enregistrement du miroir à retournement temporel (Fig. I.3), la source, un transducteur mono-élément, émet une impulsion brève ($\approx 1 \mu s$) en direction des tiges. Chacun des éléments du MRT enregistre le champ qui ressort du milieu. On peut par exemple observer sur la première voie du miroir qu’après l’arrivée de la première impulsion il existe toute une “coda” qui s’étend sur plus de $160 \mu s$ (Fig. I.5.a)). Cette coda provient des diffusions successives de l’onde initiale sur les tiges et est différente sur chacune des voies de réception du miroir. Plus le signal est éloigné du premier front d’onde, plus il aura subi de diffusions au sein de la forêt de tiges. Toutes les codas des différentes voies sont stockées dans des mémoires numériques qui peuvent être lues à l’endroit et à l’envers.

Lorsque le MRT ne reçoit plus d’énergie de la part du milieu, il passe de la phase d’enregistrement à la phase de réémission (Fig. I.4). Le signal reçu sur chaque voie est inversé chronologiquement puis émis. L’onde se propage à nouveau à travers la forêt de tiges et comme à l’aller,

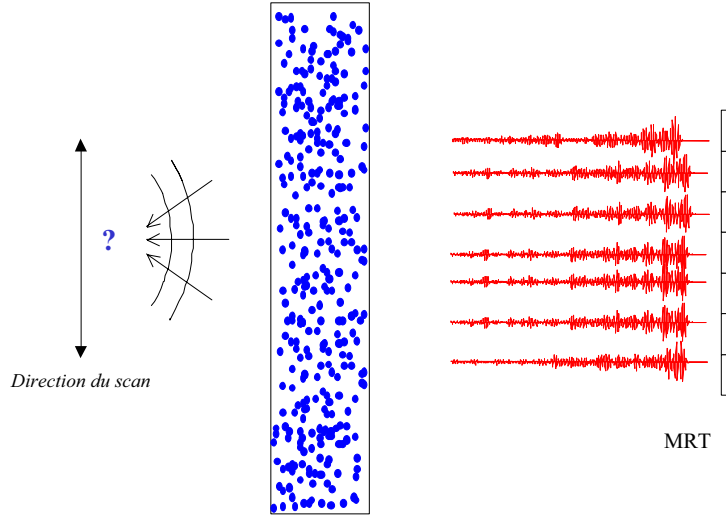


FIG. I.4 – phase de réémission du champ retourné temporellement

elle y est multidiffusée.

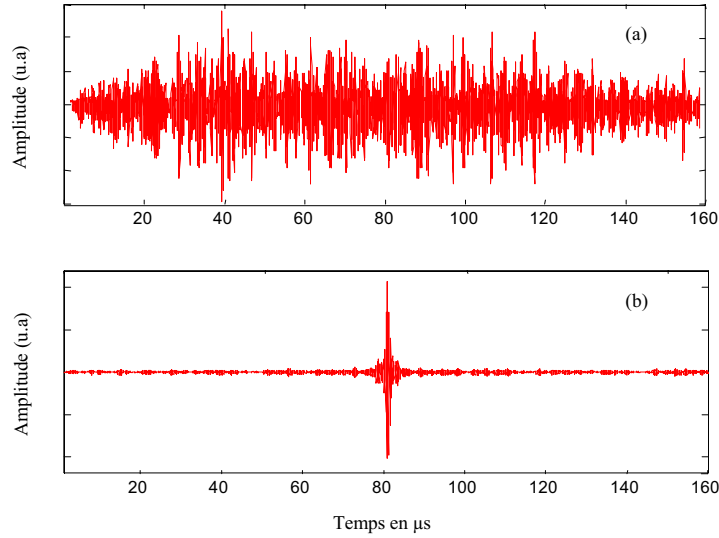


FIG. I.5 – (a) Signal reçu sur la voie 1. (b) Compression temporelle au point focal.

On observe alors qu'un front d'onde sort du milieu pour focaliser en temps (Fig. I.5.(b)) et en espace au point où se trouvait initialement la source (Fig. I.6). On utilise alors le transducteur qui a servi à l'émission pour aller mesurer le champ retourné temporellement autour de la position où l'on avait émis l'onde initiale et parallèlement à la forêt de tiges. La tache focale, qui peut être définie en retenant le maximum de la valeur absolue des signaux temporels enregistrés sur chaque position, caractérise la qualité spatiale de la focalisation. Elle est comparée à la tache focale d'une expérience de retournement temporel similaire mais réalisée sans la forêt de tiges, c'est-à-dire en milieu homogène (Fig. I.6).

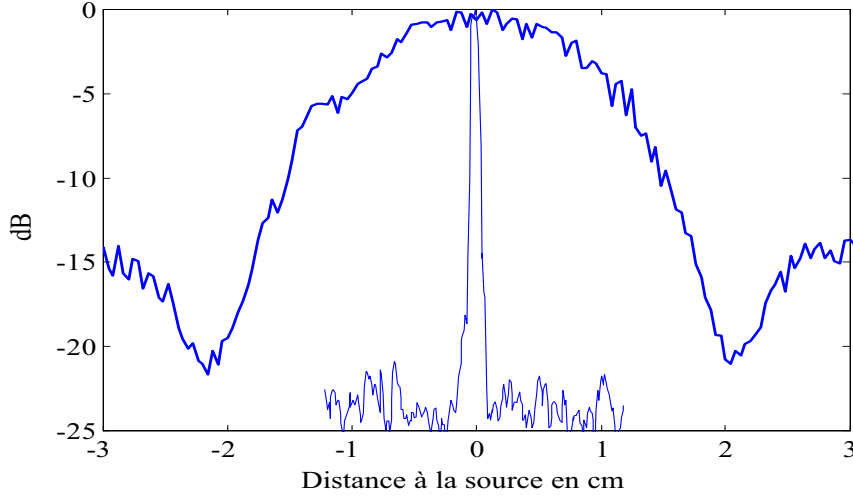


FIG. I.6 – Focalisation spatiale par RT dans l’eau (trait épais) et à travers la forêt de tiges (trait fin) (d’après [7])

La tache focale obtenue à travers le milieu multidiffuseur est dans cet exemple plus fine d’un facteur 20 par rapport à celle obtenue dans l’eau sans la forêt de tiges. Ce résultat semble contre-intuitif : on pourrait même s’attendre à ce qu’elle soit moins bonne, parce que l’onde reçue sur le MRT a perdu la mémoire de l’endroit où elle a été émise, et que le miroir n’a recueilli qu’une petite partie du champ engendré par la source. Ces deux remarques sont en fait fausses. On peut en effet user d’arguments géométriques pour justifier la finesse de la tache focale obtenue en présence de la forêt de tiges. Lorsque le miroir est placé dans l’eau, la dimension de la tache focale obtenue est limitée par l’ouverture angulaire du réseau, ce qui conduit à une dimension caractéristique $\delta = \frac{\lambda F}{D}$. Si on prend une fréquence de 1 MHz, une distance focale de 60 cm et sachant que la largeur du réseau est de 5 cm, on obtient une largeur à -6 dB qui vaut $\delta \approx 2$ cm, valeur en accord avec la courbe expérimentale. En revanche, dans le cas où la forêt de tiges est présente, le miroir profite de toutes les réflexions qui ont eu lieu. La largeur de la tache focale est alors limitée par l’ouverture angulaire interceptée par la forêt de tiges lors de la phase d’enregistrement. Ainsi, c’est la largeur angulaire de la forêt de tiges qui va définir la largeur de la tache focale, contrairement au cas des milieux homogènes où celle-ci est limitée par la taille du MRT. Ici l’échantillon mesure 10 cm de côté et est placé à 3 cm de la position de la source : avec une telle ouverture angulaire on obtient une tache focale dont la largeur est de l’ordre de la longueur d’onde, soit 1.5 mm à environ. Tout se passe donc comme si le milieu multidiffuseur avait pour effet de rendre le MRT plus grand, en rabattant sur lui des ondes qu’il n’aurait normalement pas pu enregistrer et qui n’auraient donc pas pu participer à la focalisation. Nous verrons par la suite que ce principe est généralisable à d’autres configurations et qu’il va se

montrer très utile.

Comme nous l'avons vu, dans un milieu fortement désordonné, la dimension de la tache focale que l'on peut obtenir par retournement temporel n'est plus gouvernée par la dimension du MRT. Dès lors, on peut se poser la question de l'utilité d'un miroir composé de plusieurs transducteurs. En réalité, plus le miroir est large, plus il comporte de transducteurs et cela va considérablement changer certaines de ses performances. La figure I.7 compare la focalisation spatiale obtenue dans le même milieu de tiges que précédemment avec 1 transducteur, soit un miroir d'étendue nulle, et avec 128 transducteurs.

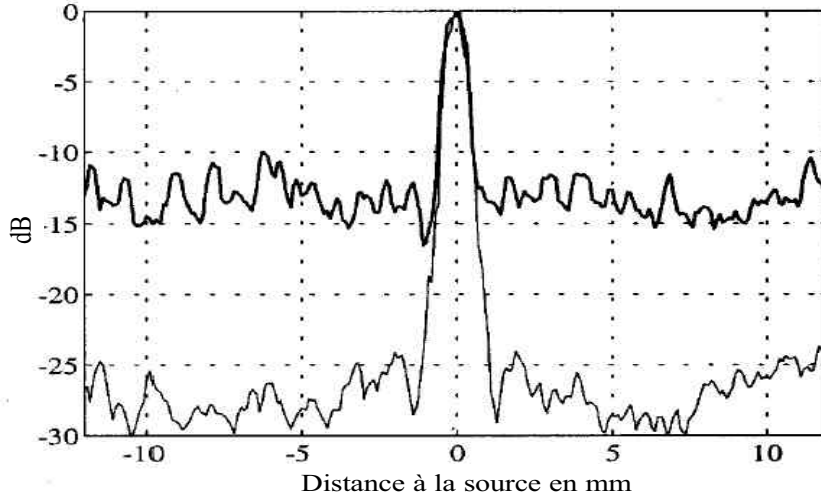


FIG. I.7 – Focalisation spatiale par RT à travers une forêt de tiges pour 1 transducteur (trait épais) et pour 128 transducteurs (trait fin). Cette figure est reproduite de [8].

Comme on peut le constater, que le miroir soit composé de 128 transducteurs ou d'un unique transducteur, la forme de la tache focale est la même, et ceci est en accord avec le fait que cette dernière est essentiellement définie par l'ouverture angulaire du milieu diffuseur. Par contre il est intéressant de noter que le gain de focalisation, lequel est défini comme le rapport entre l'amplitude maximale au point focal et l'amplitude maximale autour, est amélioré de 20 dB lorsque l'on utilise 128 transducteurs au lieu d'un seul. Arnaud Derode et ses collaborateurs ont montré dans [8] que l'amélioration est du même ordre de grandeur si l'on s'intéresse au rapport entre l'amplitude de la compression temporelle et celle des lobes secondaires. Pour expliquer ceci, il est utile d'invoquer l'équation I.13 qui a été écrite pour le cas de la cavité mais reste valable dans le cas d'un miroir à N transducteurs et donne le champ en tout point $\mathbf{r}$ du milieu :

$$S_{RT}(\mathbf{r}, t) \propto \sum_{i=1}^N h(\mathbf{r}_i, \mathbf{r}, t) \otimes h(\mathbf{r}_a, \mathbf{r}_i, -t) \otimes f(-t) \quad (\text{I.15})$$

où $\mathbf{r}_a$ est la position initiale de la source, $f(t)$ le signal qu'elle a émis et $\mathbf{r}_{i,i=\{1..N\}}$ les positions des transducteurs du miroir.

Si le milieu est fortement désordonné, les réponses impulsionnelles correspondant aux différents éléments du miroir vont être décorrélées. Ainsi, si l'on s'intéresse à la formule précédente en $\mathbf{r} = \mathbf{r}_a$ et pour l'instant $t = 0$, les N réponses se somment de façon cohérente pour donner un gain en amplitude proportionnel à N . Par contre, si l'on regarde à un autre instant, ou si l'on se place ailleurs dans le milieu, les N réponses se somment de façon incohérente et ainsi le gain n'est que de $\sqrt{N}$. Dès lors, le gain sur le rapport de l'amplitude au point focal à l'amplitude moyenne ailleurs est environ de $\sqrt{N}$, ce qui avec $N = 128$ donne un gain de 20 dB, valeur que l'on retrouve expérimentalement pour la compression temporelle. Par ailleurs dans [8], les auteurs ont également montré que pour un nombre de transducteurs supérieur à 130, ce gain sature car le miroir semble ne plus enregistrer d'information supplémentaire même si on l'agrandit. Ceci est dû au fait qu'il existe des corrélations dites "longue portée" qui ont pour effet de limiter le nombre de signaux décorrélés que l'on peut enregistrer dans un milieu multi-diffuseur.

Si l'on explique ici le rôle du nombre de transducteurs, en comparant les résultats de la compression avec un capteur et avec 128, on n'explique cependant pas pourquoi le retournement temporel fonctionne même avec un seul capteur. La partie suivante est ainsi consacrée au retournement temporel monovoie dans une cavité réverbérante chaotique.

I.1.4 Cas particulier d'un seul transducteur en cavité réverbérante

La propagation du son dans les milieux clos réverbérants a d'abord été étudiée en acoustique architecturale. Un son émis à l'intérieur d'une salle se propage, est réfléchi et diffusé sur les murs et les obstacles. Le champ en un point est la superposition des échos diffusés par les divers éléments. Sabine [9] caractérise cette décroissance du champ acoustique par le temps de réverbération, qui joue un rôle important dans l'étude de l'acoustique des salles. Suite aux multiples réflexions et autres diffusions, le champ acoustique réverbéré prend l'apparence d'un signal dont l'amplitude suit une distribution aléatoire. Néanmoins cette coda est parfaitement reproductible. Nous allons voir que les informations contenues dans cette coda peuvent être exploitées lors d'une expérience de Retournement Temporel. Carsten Draeger et Mathias Fink [10] ont observé pour la première fois des codas dont la durée devient "infiniment" longue devant la durée de l'impulsion dans des plaques de silicium en forme de disque tronqué.

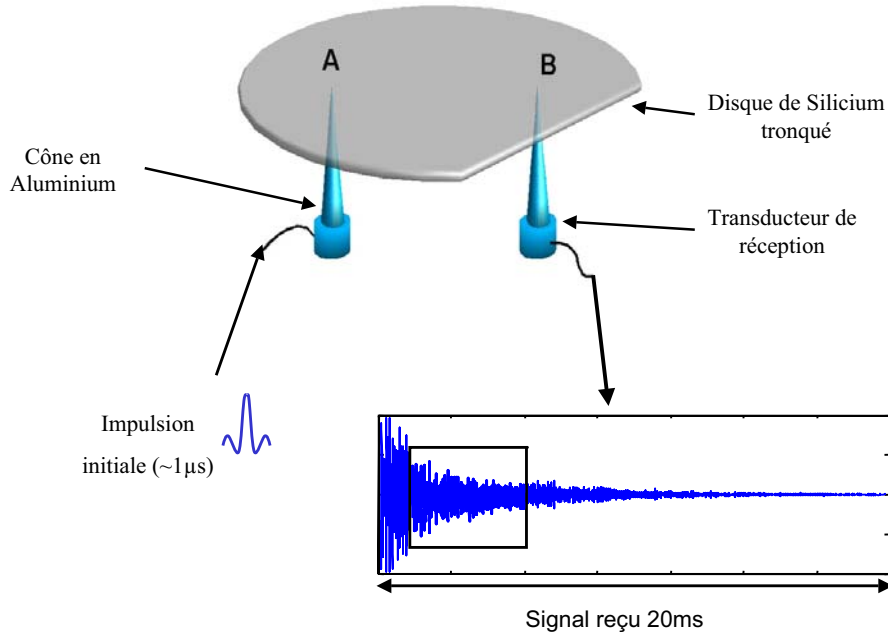


FIG. I.8 – Dispositif expérimental du retournement temporel dans une cavité chaotique

Comme le montre la figure I.8, les ondes élastiques sont générées et détectées par des transducteurs transverses couplés à des pointes en aluminium. Un des transducteurs émet une courte impulsion ultrasonore (1 période d'une sinusoïde à 1 MHz) au point **A** qui se propage et se réfléchit sur les parois de la cavité. L'énergie acoustique est piégée à l'intérieur de cette cavité car la rupture d'impédance silice/air empêche presque tout rayonnement de cette énergie dans le milieu extérieur. L'autre transducteur collecte les échos au point **B** et fournit ainsi la réponse impulsionnelle qui relie le point d'émission **A** et le point de réception **B** : $h(\mathbf{A}, \mathbf{B}, t)$. Grâce aux multiples réflexions de l'onde sur les bords de la cavité, la coda s'étend sur plus de 20 ms alors que l'impulsion émise initialement ne dure pas plus d'une microseconde.

Puis une partie de cette réponse impulsionnelle est retournée temporellement (Fig I.9) et réémise au point **B**. Le champ résultant de cette émission est mesuré grâce à un interféromètre optique hétérodyne (développé par Daniel Royer [11]) sur un carré de 10 mm de côté autour du point d'émission de l'impulsion initiale (point **A**). La figure I.10 montre l'évolution spatio-temporelle de ce champ acoustique. Une focalisation des ondes ultrasonores est observée au point d'émission de l'impulsion initiale (point **A**) au temps $t = 1200 \mu\text{s}$, ce qui correspond à la longueur de la fenêtre retournée. La focalisation atteint la limite de diffraction : le diamètre à mi-hauteur du pic de focalisation tend vers la demi-longueur d'onde $\frac{\lambda}{2}$. Ce pic de focalisation est d'ailleurs de symétrie circulaire, ce qui indique que le front d'onde convergent qui donne naissance à ce pic, provient de tout l'espace entourant le point focal comme c'était le cas lors de la focalisation

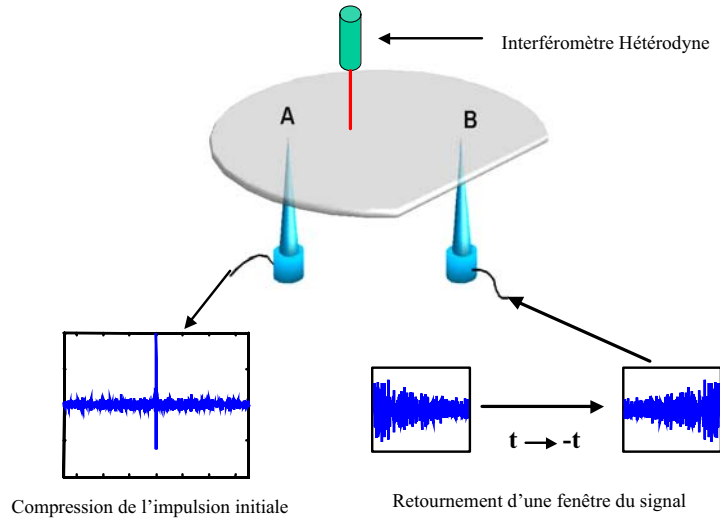


FIG. I.9 – Retournement temporel et cartographie du champ autour du point **A**

“idéale” dans la Cavit      Retournement Temporel. Bien qu’une perte d’information spatiale soit in  vitablement induite par l’utilisation d’un seul transducteur, elle est compens  e par un accroissement de l’information temporelle, d     ux multiples r  flexions de l’onde sur les parois de la cavit  . Ainsi le concept th  orique de Cavit      Retournement Temporel s’est r  v  l   envisageable gr  ce    l’utilisation de cavit  s r  verb  rantes.

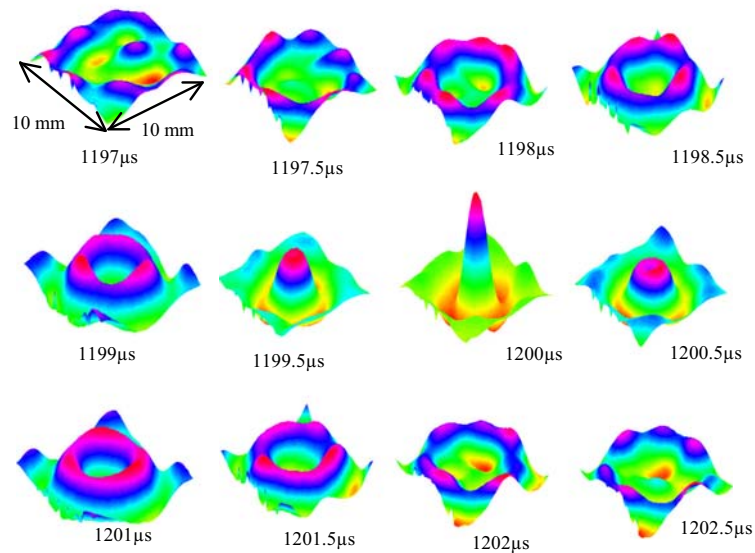


FIG. I.10 – Scan spatial du champ    diff  rents temps autour du point **A**

L’onde acoustique a subi plusieurs milliers de r  flexions sur les parois de la cavit  . La r  ponse impulsionnelle peut alors se d  composer en la superposition d’impulsions provenant de plusieurs milliers de transducteurs virtuels, images de l’unique transducteur r  el, qui agissent comme une v  ritable cavit      retournement temporel : c’est l’effet “kal  idoscope ultrasonore”.

Cependant on ne parle ici que d'impulsions. Que se passe-t-il quand un signal monochromatique est envoyé ? On ne peut alors plus s'attendre à une quelconque focalisation de l'onde retournée temporellement. En effet, si on émet une onde monochromatique, celle-ci va se propager sur toute la surface créant noeuds et ventres en fonction des dimensions de la cavité mais elle ne convergera en aucun cas : c'est la limite d'une telle approche.

Au contraire, comme on l'a vu, une cavité à retournement temporel peut focaliser une onde monochromatique sur une tache focale de diamètre $\frac{\lambda}{2}$. De même, en régime impulsionnel, si on considère un miroir à retournement temporel derrière une forêt de tiges, celui-ci peut focaliser sur une tache focale fine un champ même s'il n'est constitué que d'un unique capteur. Nous voyons ainsi apparaître une donnée essentielle qui n'a pas été prise en compte précédemment : la bande passante influence grandement le retournement temporel. Afin d'expliquer ceci qualitativement nous allons confondre ici les lobes secondaires créés par le retournement temporel, qu'ils soient spatiaux ou temporels. En effet, ceux-ci sont strictement équivalents comme cela a été démontré dans [12].

L'augmentation du rapport entre l'amplitude du pic et l'amplitude moyenne du bruit en fonction de la bande passante utilisée a d'abord été observée à travers les milieux multidiffuseurs [12]. Nous avons volontairement laissé de côté cet aspect là dans la partie précédente pour le présenter sur le cas plus explicite qu'est le miroir monovoie en cavité réverbérante. Arnaud Derode et Mathias Fink ont expliqué ce phénomène en termes de "grains d'information". Dans ce modèle, la réponse impulsionnelle du milieu peut être assimilée à une succession de grains d'information décorrelés de durée δt , où δt correspond à la période de l'impulsion initiale. Arnaud Derode a montré que l'amplitude du pic après retournement temporel est proportionnelle à la somme cohérente des N grains d'information contenus dans la fenêtre de retournement temporel de longueur T , soit $N = T/\delta t$. Le bruit, quant à lui, provient de la somme incohérente de ces mêmes N grains. Ainsi le rapport entre le bruit lors d'une expérience de retournement temporel, que celui-ci soit temporel ou spatial, et l'amplitude du pic de focalisation est proportionnel à $\sqrt{N}$. Cependant ce modèle n'est valable que dans le cas du régime diffusif, c'est-à-dire quand tous les grains d'information sont décorrelés. Il est possible d'avoir une approche plus générale en considérant l'aspect fréquentiel. On peut alors définir le nombre de grains d'information comme le rapport entre la bande passante utilisée $\delta\nu$ et la fréquence de corrélation du milieu $\delta\omega$. Le rapport entre le pic et le bruit est alors donné par $\sqrt{\frac{\delta\nu}{\delta\omega}}$. La fréquence de corrélation dépend uniquement des caractéristiques du milieu de propagation et des fréquences de l'onde utilisée. Ainsi, dans une cavité, si l'absorption est faible et donc le temps de réverbération est très grand,

la fréquence de corrélation est donnée par l'inverse de la distance entre les modes propres de la cavité. Par contre, si le temps de réverbération est plus faible que l'inverse de la distance entre deux modes de la cavité, ceux-ci ne sont pas résolus : c'est alors ce temps τ qui limite la compression par retournement temporel, car on peut écrire la fréquence de corrélation du milieu comme $\delta\omega = 1/\tau$. Le rapport pic/bruit prend ainsi la forme $\sqrt{\delta\nu\tau}^3$. Dans la plupart des travaux de cette thèse nous travaillerons avec cette dernière formule. En effet, les expériences réalisées en électromagnétisme l'ont toutes été dans des cavités réverbérantes dans lesquelles le temps d'Heisenberg, soit l'inverse de la distance entre deux modes, était toujours très supérieur au temps de réverbération.

³Ces résultats sont valables si tout le signal enregistré est retourné temporellement. Pour une étude exhaustive du rapport signal-sur-bruit qui prend en compte tous ces paramètres le lecteur pourra se référer à la discussion sur le “facteur de contraste” en retournement temporel dans [13]

Dans cette première partie, la théorie du retournement temporel appliquée au domaine de l'acoustique a été rappelée. Nous avons ainsi vu qu'une cavité à retournement temporel est capable de renvoyer vers sa source une onde en la focalisant sur une tache dont l'étendue spatiale est donnée par la partie imaginaire de la fonction de Green du milieu, soit la demi-longueur d'onde en milieu homogène. Il a également été montré qu'en utilisant des miroirs à retournement temporel dans des milieux complexes nous sommes en mesure de compresser temporellement une impulsion fortement réverbérée jusqu'à lui redonner sa forme initiale. Cette impulsion est de plus concentrée sur une tache focale dont la dimension est gouvernée par l'ouverture angulaire du milieu diffuseur plutôt que par celle du miroir, jusqu'à atteindre la limite de diffraction dans le cas du retournement temporel en cavité réverbérante. La qualité de la compression temporelle et spatiale a été quantifiée d'abord en terme de nombre de capteurs. Nous avons souligné que le rapport entre l'amplitude du pic de retournement temporel et l'amplitude moyenne du "bruit" créé est proportionnel à $\sqrt{N}$, où N est le nombre de transducteurs du miroir à retournement temporel. Puis, le rôle de la bande passante dans la compression temporelle a été discuté et nous avons abouti à la conclusion que dans le cas d'un milieu réverbérant avec de l'absorption, le rapport signal-sur-bruit en amplitude du retournement temporel est proportionnel à $\sqrt{\delta\nu\tau}$, où τ est le temps de réverbération du milieu et $\delta\nu$ la bande passante du signal.

Ces résultats vont être à la base des études concernant les télécommunications haut débit par retournement temporel qui ont été réalisées pendant cette thèse. Cependant, avant d'étudier ces concepts, il a été nécessaire de prouver la faisabilité du retournement temporel avec des ondes électromagnétiques. La partie suivante traite donc de la réversibilité des ondes électromagnétiques ainsi que de la possibilité de réaliser des expériences de retournement temporel avec de telles ondes. Le principe de réciprocité de Lorentz sera appliqué afin de développer le concept de cavité à retournement temporel électromagnétique ; un lien avec les antennes, outils de mesure du champ électrique, sera présenté.

I.2 De l'acoustique à l'électromagnétisme

I.2.1 Réversibilité et réciprocity des ondes électromagnétiques

A l'instar des ondes acoustiques, l'équation de Helmholtz est très pratique pour traiter des problèmes d'électromagnétisme car elle permet une étude mathématique aisée, comme dans le cas de la propagation guidée des ondes. Cependant, dans les situations que nous allons étudier par la suite, il va être indispensable de distinguer champs électrique et magnétique. Nous allons donc préférer partir des équations de Maxwell en milieu hétérogène et sans sources (équation de Maxwell-Gauss, conservation du flux, Maxwell-Faraday, Maxwell-Ampère) :

$$\left\{ \begin{array}{l} \nabla \cdot \mathbf{D}(\mathbf{r}, t) = 0 \\ \nabla \cdot \mathbf{B}(\mathbf{r}, t) = 0 \\ \nabla \times \mathbf{E}(\mathbf{r}, t) = -\frac{\partial}{\partial t} \mathbf{B}(\mathbf{r}, t) \\ \nabla \times \mathbf{H}(\mathbf{r}, t) = \frac{\partial}{\partial t} \mathbf{D}(\mathbf{r}, t) \end{array} \right. \quad (\text{I.16})$$

où l'on a introduit les vecteurs déplacement électrique $\mathbf{D}$ et excitation magnétique $\mathbf{H}$ tels que :

$$\left\{ \begin{array}{l} \mathbf{D}(\mathbf{r}, t) = \epsilon_0 \overleftrightarrow{\epsilon}(\mathbf{r}, \omega) \mathbf{E}(\mathbf{r}, t) \\ \mathbf{B}(\mathbf{r}, t) = \mu_0 \overleftrightarrow{\mu}(\mathbf{r}, \omega) \mathbf{H}(\mathbf{r}, t) \end{array} \right. \quad (\text{I.17})$$

avec $\overleftrightarrow{\epsilon}(\mathbf{r}, \omega)$ et $\overleftrightarrow{\mu}(\mathbf{r}, \omega)$ les tenseurs de permittivité diélectrique et de perméabilité magnétique du milieu, qui sont supposés symétriques.

Dans le système d'équations précédent, si l'on change la variable temporelle t en $-t$, le vecteur champ électrique $\mathbf{E}(\mathbf{r}, t)$ va se changer en $\mathbf{E}(\mathbf{r}, -t)$, car c'est un "vrai" vecteur. Le champ magnétique $\mathbf{B}(\mathbf{r}, t)$ va lui devenir $-\mathbf{B}(\mathbf{r}, -t)$ car c'est un "pseudo vecteur". En effet celui-ci dérive directement du courant au travers de la loi de Biot et Savart. Ces deux remarques montrent que le système d'équation précédent est réversible. Comme dans le cas des ondes acoustiques, on peut imaginer la possibilité de réaliser un retournement temporel d'une onde électromagnétique en générant dans tout un volume le champ électrique, l'opposé du champ magnétique, le tout dans une chronologie inversée et pour tout instant t . De nouveau cette technique est en pratique irréalisable mais nous allons voir qu'elle peut être remplacée par le retournement temporel du champ électromagnétique sur une surface fermée entourant la source initiale.

Afin de démontrer cette propriété nous allons utiliser le théorème de réciprocity de Lorentz en milieu hétérogène comme cela est fait dans [14]. Pour redémontrer ce théorème nous écrivons les

équations de Maxwell en régime monochromatique et dans lesquelles nous incluons des sources électriques ⁴. Pour des raisons de clarté, les dépendances en $\mathbf{r}$ et ω seront implicites. Considérons à présent deux sources qui présentent des densités de courant $\mathbf{J}_1$ et $\mathbf{J}_2$ et qui donnent naissance aux champs $\{\mathbf{E}_1, \mathbf{D}_1, \mathbf{B}_1, \mathbf{H}_1\}$ et $\{\mathbf{E}_2, \mathbf{D}_2, \mathbf{B}_2, \mathbf{H}_2\}$. Ces deux champs vérifient les équations de Maxwell-Faraday et Maxwell-Ampère et on peut écrire le système d'équations suivant :

$$\left\{ \begin{array}{ll} \nabla \times \mathbf{E}_1 = -i\omega \mathbf{B}_1 & \text{et} \quad \nabla \times \mathbf{E}_2 = -i\omega \mathbf{B}_2 \\ \nabla \times \mathbf{H}_1 = \mathbf{J}_1 + i\omega \mathbf{D}_1 & \text{et} \quad \nabla \times \mathbf{H}_2 = \mathbf{J}_2 + i\omega \mathbf{D}_2 \end{array} \right. \quad (\text{I.18})$$

En suivant la même logique que lorsque nous avons démontré le théorème d'Helmholtz-Kirchhoff en acoustique, on projète les équations relatives au champ électrique sur le champ magnétique et réciproquement pour les équations du champ magnétique que nous projetons sur le champ électrique. Nous obtenons alors le système suivant :

$$\left\{ \begin{array}{ll} \mathbf{H}_2 \cdot (\nabla \times \mathbf{E}_1) = -i\omega \mathbf{H}_2 \cdot \mathbf{B}_1 & \text{et} \quad \mathbf{H}_1 \cdot (\nabla \times \mathbf{E}_2) = -i\omega \mathbf{H}_1 \cdot \mathbf{B}_2 \\ \mathbf{E}_2 \cdot (\nabla \times \mathbf{H}_1) = \mathbf{E}_2 \cdot (\mathbf{J}_1 + i\omega \mathbf{D}_1) & \text{et} \quad \mathbf{E}_1 \cdot (\nabla \times \mathbf{H}_2) = \mathbf{E}_1 \cdot (\mathbf{J}_2 + i\omega \mathbf{D}_2) \end{array} \right. \quad (\text{I.19})$$

La suite des opérations est similaire à la démonstration effectuée en acoustique : nous faisons les sommes et différences de ces 4 équations afin d'obtenir l'expression suivante :

$$\begin{aligned} & \{\mathbf{H}_2 \cdot \nabla \times \mathbf{E}_1 - \mathbf{E}_1 \cdot \nabla \times \mathbf{H}_2\} + \{\mathbf{E}_2 \cdot \nabla \times \mathbf{H}_1 - \mathbf{H}_1 \cdot \nabla \times \mathbf{E}_2\} \\ &= i\omega \{\mathbf{H}_2 \cdot \mathbf{B}_1 - \mathbf{H}_1 \cdot \mathbf{B}_2\} - i\omega \{\mathbf{D}_1 \cdot \mathbf{E}_2 - \mathbf{E}_1 \cdot \mathbf{D}_2\} + \mathbf{J}_1 \cdot \mathbf{E}_2 - \mathbf{J}_2 \cdot \mathbf{E}_1 \end{aligned} \quad (\text{I.20})$$

Le terme de gauche se réécrit alors sous la forme $\nabla \cdot \{\mathbf{E}_1 \times \mathbf{H}_2 - \mathbf{E}_2 \times \mathbf{H}_1\}$. De plus en utilisant les expressions des vecteurs $\mathbf{D}$ et $\mathbf{H}$ et en tenant compte du fait que les tenseurs de permittivité et de perméabilité sont symétriques, les deux premiers termes du membre de droite s'annulent, ce qui donne :

$$\nabla \cdot \{\mathbf{E}_1 \times \mathbf{H}_2 - \mathbf{E}_2 \times \mathbf{H}_1\} = \mathbf{J}_2 \cdot \mathbf{E}_1 - \mathbf{J}_1 \cdot \mathbf{E}_2 \quad (\text{I.21})$$

Enfin il reste à intégrer l'équation précédente sur un volume et à utiliser le théorème de Green-Ostrogradsky afin d'obtenir :

⁴Le cas des sources magnétiques ne nous intéressera pas et sera donc laissé de coté. Cependant il existe un principe de réciprocité incluant les sources de type magnétique.

$$\oint_S \{\mathbf{E}_1 \times \mathbf{H}_2 - \mathbf{E}_2 \times \mathbf{H}_1\} \cdot d\mathbf{S} = \iiint_V \{\mathbf{J}_2 \cdot \mathbf{E}_1 - \mathbf{J}_1 \cdot \mathbf{E}_2\} dV \quad (\text{I.22})$$

Cette dernière relation, qui constitue le théorème de réciprocité de Lorentz généralisé avec sources électriques, lie les champs électriques et magnétiques aux sources de courant qui les ont créés.

Elle permet de démontrer deux principes qui vont nous être très utiles pour nos expériences de retournement temporel et de télécommunications. Tout d'abord, grâce à ce théorème, nous serons en mesure par la suite de transposer au cas des ondes électromagnétiques le concept de cavité à retournement temporel. Mais auparavant, il permet d'établir la réciprocité des ondes électromagnétiques en milieu hétérogène. En effet, si l'on choisit comme volume d'intégration une sphère de rayon infiniment grand, les expressions asymptotiques des champs électriques et magnétiques sont proportionnelles à $\frac{1}{r}$, où r est le rayon de la sphère. Dans ce cas on peut montrer que le membre de gauche de l'équation précédente tend vers zéro, ce qui conduit au principe de réciprocité avec sources de courant $\mathbf{J}_1$ et $\mathbf{J}_2$:

$$\iiint_V \mathbf{J}_1(\mathbf{r}) \cdot \mathbf{E}_2(\mathbf{r}, \omega) dV = \iiint_V \mathbf{J}_2(\mathbf{r}) \cdot \mathbf{E}_1(\mathbf{r}, \omega) dV \quad (\text{I.23})$$

Si on considère que les sources sont des dipôles élémentaires, c'est-à-dire que $\mathbf{J}_1(\mathbf{r}) = -i\omega \mathbf{p}_1 \delta(\mathbf{r} - \mathbf{r}_1)$ et $\mathbf{J}_2(\mathbf{r}) = -i\omega \mathbf{p}_2 \delta(\mathbf{r} - \mathbf{r}_2)$, où $\mathbf{p}_{1,2}$ est le vecteur polarisation des dipôles, la formule précédente s'écrit :

$$\mathbf{p}_1 \cdot \mathbf{E}_2(\mathbf{r}_1, \omega) = \mathbf{p}_2 \cdot \mathbf{E}_1(\mathbf{r}_2, \omega) \quad (\text{I.24})$$

Cette dernière expression est la version la plus explicite du théorème de réciprocité avec sources. En effet, il apparaît clairement ici que le champ électrique dans la direction de polarisation de la source est inchangé quand les positions de l'émetteur et du récepteur sont interchangées. C'est également cette formulation qui est utilisée en théorie des antennes.

I.2.2 La cavité à retournement temporel électromagnétique

Comme dans les cas des ondes acoustiques, nous allons partir du principe de réciprocité généralisé, afin d'introduire la notion de cavité à retournement temporel électromagnétique. A nouveau nous allons traiter ceci en régime harmonique, afin de faciliter les calculs. Pour ce faire, on considère deux vecteurs courants $\mathbf{J}_1(\mathbf{r}, \omega)$ et $\mathbf{J}_2(\mathbf{r}, \omega)$ dont les champs s'écrivent

$\{\mathbf{E}_1(\mathbf{r}, \omega), \mathbf{H}_1(\mathbf{r}, \omega)\}$ et $\{\mathbf{E}_2(\mathbf{r}, \omega), \mathbf{H}_2(\mathbf{r}, \omega)\}$. Nous allons alors appliquer le principe de r  ciprocit   de Lorentz qui lie le complexe conjugu   du premier champ $\{\mathbf{E}_1^*(\mathbf{r}, \omega), -\mathbf{H}_1^*(\mathbf{r}, \omega)\}$, au champ cr     par la source de courant $\mathbf{J}_2$. Cette relation s'  crit :

$$\oint_S \{\mathbf{E}_1^*(\mathbf{r}, \omega) \times \mathbf{H}_2(\mathbf{r}, \omega) + \mathbf{E}_2(\mathbf{r}, \omega) \times \mathbf{H}_1^*(\mathbf{r}, \omega)\} d\mathbf{S} = - \iiint_V \{\mathbf{J}_2(\mathbf{r}, \omega) \cdot \mathbf{E}_1^*(\mathbf{r}, \omega) - \mathbf{J}_1^*(\mathbf{r}, \omega) \cdot \mathbf{E}_2(\mathbf{r}, \omega)\} dV \quad (\text{I.25})$$

Puis, il est n  cessaire de faire quelques approximations afin de pouvoir simplifier cette derni  re   quation. Tout d'abord, nous allons consid  rer que la surface est tr  s   loign  e de la source et faire une approximation paraxiale, afin de pouvoir   crire que $\mathbf{H}(\mathbf{r}, \omega) = \frac{\mathbf{n} \times \mathbf{E}(\mathbf{r}, \omega)}{\eta}$, o   $\eta = \sqrt{\frac{\mu}{\epsilon}}$ est l'imp  dance du milieu et $\mathbf{n}$ la normale    la surface de la cavit  . La relation de r  ciprocit   de Lorentz prend alors la forme suivante :

$$\frac{2}{\eta} \oint_S \mathbf{E}_1^*(\mathbf{r}, \omega) \times \mathbf{E}_2(\mathbf{r}, \omega) d\mathbf{S} = - \iiint_V \{\mathbf{J}_2(\mathbf{r}, \omega) \cdot \mathbf{E}_1^*(\mathbf{r}, \omega) - \mathbf{J}_1^*(\mathbf{r}, \omega) \cdot \mathbf{E}_2(\mathbf{r}, \omega)\} dV \quad (\text{I.26})$$

On consid  re ensuite que les sources sont des dip  les   l  mentaires de sorte que l'on peut les   crire comme $\mathbf{J}_1(\mathbf{r}) = i\omega p_1 \delta(\mathbf{r} - \mathbf{r}_1) \mathbf{n}_1$ et $\mathbf{J}_2(\mathbf{r}) = i\omega p_2 \delta(\mathbf{r} - \mathbf{r}_2) \mathbf{n}_2$, avec $\mathbf{n}_1$ et $\mathbf{n}_2$ des vecteurs norm  s. Le membre de droite de l'  quation pr  c  dente prend alors la forme simple :

$$- \iiint_V \{\mathbf{J}_2(\mathbf{r}, \omega) \cdot \mathbf{E}_1^*(\mathbf{r}, \omega) - \mathbf{J}_1^*(\mathbf{r}, \omega) \cdot \mathbf{E}_2(\mathbf{r}, \omega)\} dV = -(i\omega p_2 \mathbf{n}_2 \cdot \mathbf{E}_1^*(\mathbf{r}_2, \omega) + i\omega p_1 \mathbf{n}_1 \cdot \mathbf{E}_2(\mathbf{r}_1, \omega)) \quad (\text{I.27})$$

Ce dernier terme peut encore   tre modifi   en utilisant le th  or  me de r  ciprocit   I.24 explicit   pr  c  demment :

$$- \iiint_V \{\mathbf{J}_2(\mathbf{r}, \omega) \cdot \mathbf{E}_1^*(\mathbf{r}, \omega) - \mathbf{J}_1^*(\mathbf{r}, \omega) \cdot \mathbf{E}_2(\mathbf{r}, \omega)\} dV = -i\omega p_2 \mathbf{n}_2 \cdot (\mathbf{E}_1^*(\mathbf{r}_2, \omega) + \mathbf{E}_1(\mathbf{r}_2, \omega)) \quad (\text{I.28})$$

Finalement, l'  quation I.26 peut donc   tre reformul  e de la fa  on suivante :

$$\frac{1}{\eta} \oint_S \mathbf{E}_1^*(\mathbf{r}, \omega) \times \mathbf{E}_2(\mathbf{r}, \omega) d\mathbf{S} = -i\omega p_2 \mathbf{n}_2 \cdot \Re \{\mathbf{E}_1(\mathbf{r}_2, \omega)\} \quad (\text{I.29})$$

Pour faciliter l'interpr  tation de ces r  sultats, comme en acoustique, il est utile d'introduire la notion de fonction de Green dyadique qui est l'analogue de la fonction de Green en acoustique

(i.e., des tenseurs d'ordre 2 solutions d'un Dirac spatial). Par définition [15], le champ électrique s'exprime comme :

$$\mathbf{E}(\mathbf{r}, \omega) = i\omega\mu \iiint_V \overleftrightarrow{G}(\mathbf{R}, \mathbf{r}, \omega) \cdot \mathbf{J}(\mathbf{R}, \omega) dV \quad (\text{I.30})$$

où $\mathbf{R}$ est un point du volume d'intégration.

Compte tenu des sources utilisées, les champs $\mathbf{E}_1$ et $\mathbf{E}_2$ prennent alors la forme suivante :

$$\mathbf{E}_1(\mathbf{r}, \omega) = -\omega^2 p_1 \mu \overleftrightarrow{G}(\mathbf{r}_1, \mathbf{r}, \omega) \mathbf{n}_1 \quad \mathbf{E}_2(\mathbf{r}, \omega) = -\omega^2 p_2 \mu \overleftrightarrow{G}(\mathbf{r}_2, \mathbf{r}, \omega) \mathbf{n}_2 \quad (\text{I.31})$$

En tenant compte du fait que $k = \frac{\omega\mu}{\eta}$, et en remplaçant les champs par les expressions précédentes, l'équation I.29 s'écrit :

$$k \oint_S \overleftrightarrow{G}^*(\mathbf{r}_1, \mathbf{r}, \omega) \mathbf{n}_1 \overleftrightarrow{G}(\mathbf{r}_2, \mathbf{r}, \omega) \mathbf{n}_2 d\mathbf{S} = \mathbf{n}_2 \Im \left\{ \overleftrightarrow{G}(\mathbf{r}_1, \mathbf{r}_2, \omega) \right\} \mathbf{n}_1 \quad (\text{I.32})$$

Grâce à la réciprocité des fonctions dyadiques, et en remarquant que cette dernière expression est valable quels que soient les vecteurs directeurs $(\mathbf{n}_1, \mathbf{n}_2)$ choisis, on obtient finalement une expression proche de celle que l'on avait obtenu dans la cas des ondes acoustiques :

$$k \oint_S \overleftrightarrow{G}(\mathbf{r}, \mathbf{r}_2, \omega) \overleftrightarrow{G}^*(\mathbf{r}_1, \mathbf{r}, \omega) d\mathbf{S} = \Im \left\{ \overleftrightarrow{G}(\mathbf{r}_1, \mathbf{r}_2, \omega) \right\} \quad (\text{I.33})$$

Cette équation s'interprète de la façon suivante, une source émet du point $\mathbf{r}_1$ avec une certaine polarisation et le champ vectoriel est enregistré sur une surface fermée. Cette surface fermée est tapissée de petits dipôles électriques qui vont réémettre le complexe conjugué du champ avec la même polarisation. Le champ ainsi créé en tout point $\mathbf{r}_2$ est alors la partie imaginaire de la fonction de Green dyadique qui relie les points $\mathbf{r}_1$ et $\mathbf{r}_2$. Ce résultat est très similaire à celui que nous avons obtenu en acoustique. A nouveau, si l'on se place en milieu homogène, la formule précédente nous donne une focalisation du champ retourné temporellement sur une tache en sinus cardinal.

Dès lors, le principe de la cavité à retournement temporel appliqué au cas des ondes électromagnétiques sur une surface fermée peut être modélisé, comme cela avait été fait en acoustique, par la séquence suivante :

Une source de courant dont la direction de polarisation est selon $\mathbf{n}_1$, et qui est située à l'intérieur de la cavité, émet une impulsion électromagnétique brève. Le milieu de propagation peut être éventuellement hétérogène. L'onde sphérique générée se réfléchit et diffracte de manière complexe lors de son passage dans le milieu hétérogène. Lorsque l'onde atteint la surface S , le

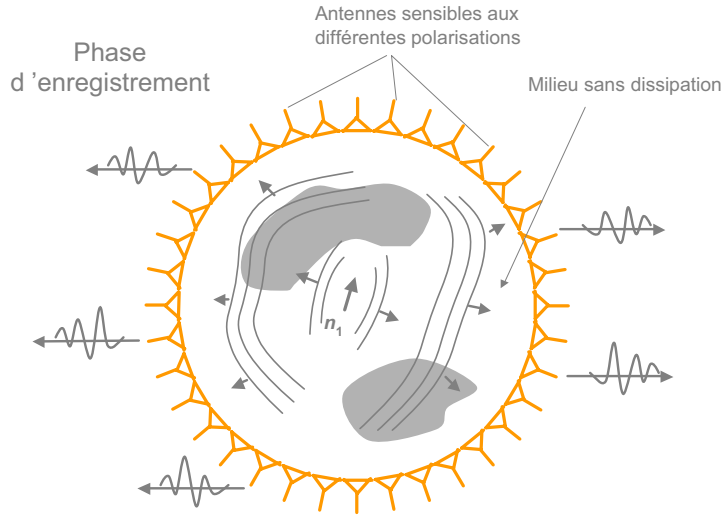


FIG. I.11 – Phase d'enregistrement du champ électromagnétique

champ électrique $\mathbf{E}(\mathbf{r}, t)$ et le champ magnétique $\mathbf{B}(\mathbf{r}, t)$ sont enregistrés par des antennes. Ces signaux sont ensuite numérisés puis stockés dans des mémoires. La phase d'acquisition s'arrête lorsqu'il n'y a plus d'énergie à l'intérieur de la cavité.

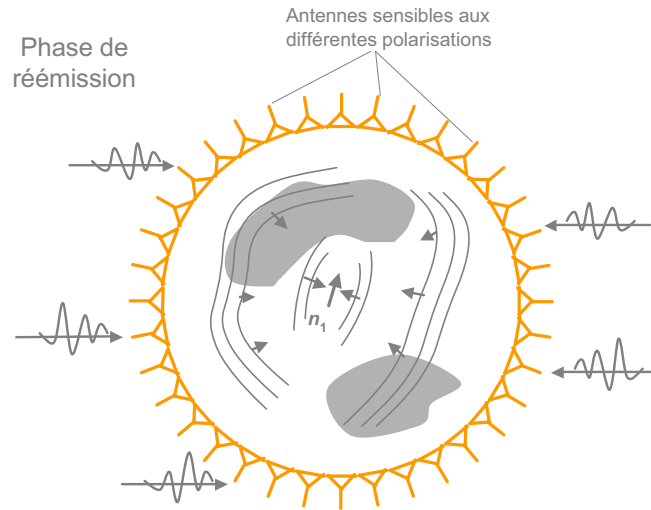


FIG. I.12 – Phase de réémission des signaux retournés temporellement

Durant cette deuxième phase les versions retournées temporellement du champ électrique $\mathbf{E}(\mathbf{r}, -t)$ et du champ magnétique $-\mathbf{B}(\mathbf{r}, -t)$ sont réémises par chaque antenne de la surface dans le milieu. Les ondes ainsi créées étant solutions des équations de Maxwell, le champ engendré dans tout le volume de la cavité est égal à celui de l'onde émise initialement mais

retournée temporellement, cette onde converge donc vers son point source initial.

I.2.3 Application aux antennes

Dans la partie précédente, la présentation du fonctionnement de la cavité à retournement temporel a été faite en parlant uniquement du rôle des champs électromagnétiques. Nous allons présenter, dans ce paragraphe, une autre approche plus adaptée au domaine des ondes décimétriques. Dans cette gamme de fréquences, l'outil de mesure est l'antenne qui, selon qu'elle est électrique ou magnétique, mesure une différence de potentiel ou un courant, lesquels sont tous deux provoqués par un champ électromagnétique. Nous nous limiterons dans notre approche au cas d'antennes électriques sensibles uniquement à une composante du champ, dont l'exemple le plus connu est le dipôle [16]. Afin de traiter des problèmes d'antennes, il est courant de faire appel à la théorie des systèmes électriques linéaires. Dans ce formalisme, toute transmission entre deux antennes est complètement déterminée par la notion d'impédance qui relie tensions et courants dans un circuit.

On peut montrer, dans le cas où deux antennes 1 et A sont assez éloignées, qu'une tension V_0 aux bornes de l'antenne 1 va engendrer sur l'antenne A une tension :

$$V_A \propto Z_{1A} V_0 \quad (\text{I.34})$$

où Z_{1A} est l'impédance mutuelle entre l'antenne 1 et l'antenne A .

Lorsque nous allons réaliser une opération de retournement temporel, nous allons d'abord enregistrer une telle tension à l'aide d'un oscilloscope, puis la retourner temporellement et la réémettre. La tension reçue sur l'antenne 1 après retournement temporel, et en régime harmonique, se met ainsi sous la forme :

$$V_1^{RT} \propto Z_{A1} Z_{1A}^* V_0^* \quad (\text{I.35})$$

où Z_{A1} est l'impédance mutuelle entre l'antenne A et l'antenne 1.

De la même façon, la tension reçue aux bornes d'une antenne différente de celle qui avait émis initialement, et que l'on va nommer 2, peut s'écrire :

$$V_2^{RT} \propto Z_{A2} Z_{1A}^* V_0^* \quad (\text{I.36})$$

où cette fois apparaît l'impédance mutuelle entre l'antenne A et l'antenne 2.

A présent si l'on veut connaître le résultat d'un retournement temporel parfait, il suffit d'écrire que le retournement temporel a été effectué non plus avec une unique antenne A mais avec un ensemble d'antennes sensibles aux différentes polarisations du champ et dont la densité surfacique est σ . La tension résultante sur une antenne quelconque 2 à l'intérieur de la cavité prend alors la forme :

$$V_2^{RT} \propto \sigma V_0^* \iint_S Z_{A2} Z_{1A}^* dS \quad (\text{I.37})$$

Dans cette équation, il est entendu que l'intégration porte sur les antennes de la cavité notées A . Ainsi, si une antenne 1 a été soumise à une tension V_0 , que l'on procède à un retournement temporel parfait, et que le champ est mesuré par une antenne 2 quelconque dans la cavité, le résultat dépendra uniquement de la grandeur Ξ définie comme :

$$\Xi = \sigma \iint_S Z_{A2} Z_{1A}^* dS \quad (\text{I.38})$$

Afin de rendre ce résultat plus maniable, nous allons expliciter les divers termes de cette formule en utilisant la définition de l'impédance mutuelle telle qu'elle est énoncée dans [16]. Celle-ci s'écrit :

$$Z_{1A} = \frac{V(1 \rightarrow A)}{I_1} \Big|_{I_A \rightarrow 0} \quad (\text{I.39})$$

Cette impédance est ainsi définie comme la tension créée par une antenne 1 sur une antenne A , notée $V(1 \rightarrow A)$, divisée par le courant aux bornes de l'antenne 1, lorsque l'antenne A est en circuit "quasi-ouvert". Nous pouvons alors calculer la tension $V(1 \rightarrow A)$ en utilisant un calcul de force électromotrice induite. Cette tension est liée à la distribution de courant $\mathbf{J}_A$ de l'antenne A , et au champ électrique $\mathbf{E}_1$ provenant de l'antenne 1 par la relation :

$$V(1 \rightarrow A) = - \frac{\int \mathbf{E}_1(\mathbf{r}) \cdot \mathbf{J}_A(\mathbf{r}) d^3\mathbf{r}}{I_A} \quad (\text{I.40})$$

où intervient également le courant en entrée de l'antenne A (que l'on suppose très faible).

On peut donc exprimer à nouveau l'impédance mutuelle Z_{1A} en utilisant cette dernière expression pour obtenir :

$$Z_{1A} = \frac{-1}{I_1 \cdot I_A} \int_V \mathbf{E}_1(\mathbf{r}) \cdot \mathbf{J}_A(\mathbf{r}) d^3\mathbf{r} \quad (\text{I.41})$$

De manière à simplifier ce calcul en utilisant les résultats de la partie précédente, nous allons exprimer le champ électrique en faisant intervenir les fonctions de Green dyadiques définies précédemment, ce qui donne :

$$Z_{1A} = \frac{-i\mu\omega}{I_1 \cdot I_A} \int_V \int_V \mathbf{J}_A(\mathbf{R}) \overleftrightarrow{G}(\mathbf{R}, \mathbf{r}) \cdot \mathbf{J}_1(\mathbf{r}) d^3\mathbf{R} d^3\mathbf{r} \quad (\text{I.42})$$

Afin d'exprimer la quantité Ξ , nous écrivons de la même manière Z_{A2} :

$$Z_{A2} = \frac{-i\mu\omega}{I_A \cdot I_2} \int_V \int_V \mathbf{J}_2(\mathbf{R}) \overleftrightarrow{G}(\mathbf{R}, \mathbf{r}) \cdot \mathbf{J}_A(\mathbf{r}) d^3\mathbf{R} d^3\mathbf{r} \quad (\text{I.43})$$

Afin d'utiliser la relation I.33 obtenue précédemment, nous allons à présent considérer une cavité formée de petits dipôles de sorte que leurs distributions de courant prennent la forme : $\mathbf{J}_A(\mathbf{r}) = l_A I_A \delta(\mathbf{r} - \mathbf{r}_A) \cdot \mathbf{n}_A$, avec l_A la longueur des antennes, $\mathbf{r}_A$ leurs positions et $\mathbf{n}_A$ leurs orientations. En tenant compte de cette hypothèse, nous pouvons reformuler l'équation I.38 :

$$\Xi = -\sigma \oint_S \frac{(i\mu\omega l_A)^2}{I_1^* I_2} \int_V \int_V \mathbf{J}_2(\mathbf{r}) \overleftrightarrow{G}^*(\mathbf{r}_A, \mathbf{r}) \mathbf{n}_A \mathbf{n}_A \overleftrightarrow{G}(\mathbf{R}, \mathbf{r}_A) \mathbf{J}_1^*(\mathbf{R}) d^3\mathbf{R} d^3\mathbf{r} d^2\mathbf{r}_A \quad (\text{I.44})$$

Cette expression peut être modifiée compte tenu du fait que l'intégrale de surface s'applique uniquement aux fonctions de Green, ce qui donne :

$$\Xi = -\sigma \frac{(i\mu\omega l_A)^2}{I_1^* I_2} \int_V \int_V \mathbf{J}_2(\mathbf{r}) \oint_S \left\{ \overleftrightarrow{G}^*(\mathbf{r}_A, \mathbf{r}) \mathbf{n}_A \mathbf{n}_A \overleftrightarrow{G}(\mathbf{R}, \mathbf{r}_A) \right\} d^2\mathbf{r}_A \mathbf{J}_1^*(\mathbf{R}) d^3\mathbf{R} d^3\mathbf{r} \quad (\text{I.45})$$

On reconnaît le terme de gauche de l'équation I.33, que l'on remplace par le terme de droite afin d'obtenir :

$$\Xi = -\sigma \frac{(i\mu\omega l_A)^2}{k I_1^* I_2} \int_V \int_V \mathbf{J}_2(\mathbf{r}) \Im \left\{ \overleftrightarrow{G}(\mathbf{R}, \mathbf{r}) \right\} \mathbf{J}_1^*(\mathbf{R}) d^3\mathbf{R} d^3\mathbf{r} \quad (\text{I.46})$$

Ici nous allons faire à nouveau une simplification qui sera toujours vérifiée dans le cas de nos expériences. Nous n'allons considérer que des antennes à ondes stationnaires, ce qui est le cas des antennes filaires. Dans ce cas, on peut démontrer que la densité de courant est en phase avec le courant dans l'antenne, ce qui a comme conséquence que $\frac{\mathbf{J}_1^*}{I_1^*}$ est réel ainsi que $\frac{\mathbf{J}_2}{I_2}$. On peut alors s'affranchir des conjugués pour les courants et faire sortir la partie imaginaire de la double somme. Celle-ci est remplacée par une partie réelle car le facteur est imaginaire pur.

$$\Xi = \sigma l_A^2 \frac{i\mu\omega}{k} \Re \left\{ \int_V \int_V \frac{-i\mu\omega}{I_1 I_2} \mathbf{J}_2(\mathbf{r}) \overleftrightarrow{G}(\mathbf{R}, \mathbf{r}) \mathbf{J}_1(\mathbf{R}) d^3\mathbf{R} d^3\mathbf{r} \right\} \quad (\text{I.47})$$

On reconnaît alors l'expression de l'impédance mutuelle entre l'antenne 1 et l'antenne 2. Si on remplace le facteur $\frac{\omega\mu}{k}$ qui est égal à l'impédance du vide η , cette dernière équation se simplifie alors en :

$$\Xi = \sigma l_A^2 \eta \Re \{ Z_{12} \} \quad (\text{I.48})$$

On peut alors exprimer la grandeur Ξ à l'aide des formules I.38 et I.48 :

$$\Xi = \sigma l_A^2 \eta \Re \{ Z_{12} \} = \sigma \oint_S Z_{A2} Z_{1A}^* dS \quad (\text{I.49})$$

La conclusion que l'on peut en tirer est la suivante : si une antenne 1 est initialement soumise à une tension, que l'on enregistre les tensions résultantes sur une cavité fermée en prenant en compte toutes les composantes du champ électrique, que l'on retourne alors temporellement les tensions et qu'on les génère sur toute cette surface, la tension que mesurée sur une antenne 2 quelconque dans la cavité est proportionnelle à la partie réelle de l'impédance mutuelle Z_{12} entre ces 2 antennes. Cette dernière formulation est beaucoup plus adaptée à notre problème que la version formelle de la cavité à retournement temporel qui fait intervenir les fonctions de Green dyadiques.

Afin de valider cette formule, nous avons réalisé des simulations sur le logiciel de simulation d'antennes "4Nec2" couramment utilisé dans la communauté et qui a été développé par Arie Voors [17]. Nous avons donc simulé une cavité à retournement temporel qui entoure deux antennes 1 et 2 placées dans le volume et dont nous avons fait varier l'espacement. Pour chaque distance, nous avons comparé les deux expressions de la grandeur Ξ à l'aide des formules I.38 et I.48. Le résultat correspondant a ensuite été normalisé par le coefficient $\sigma \eta l_A^2$. Les antennes placées en 1 et 2 sont des dipôles de longueur 0.54λ et de diamètre 1 mm. La taille des petits dipôles qui tapissent la cavité est de $\lambda/40$, et nous en avons placé 126. Le rayon de la cavité mesure 6 longueurs d'ondes. La figure I.13 montre le résultat de cette comparaison pour des espacements entre antennes allant de 0 à 2λ , puis de $-\lambda$ à λ .

Il y a un très bon accord entre les deux courbes, bien que le nombre d'antennes qui recouvrent la cavité soit assez faible ici (126). Pour un espacement nul on retrouve la partie réelle de

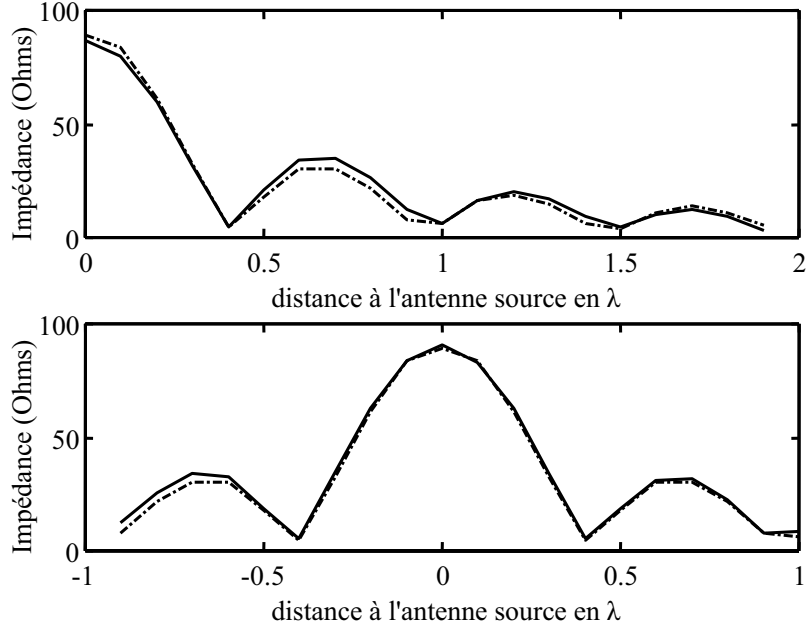


FIG. I.13 – Comparaisons entre les deux expressions de Ξ données par la relation I.38 (trait plein) et la relation I.48 (pointillés), lorsque la distance entre les antennes 1 et 2 varie.

l'impédance propre des antennes 1 et 2 qui vaut à peu près 90 Ohms pour la taille d'antenne que nous avons choisie. Afin de vérifier que le facteur de proportionnalité dans l'équation I.48 est correct, nous avons également réalisé une série de simulation imposant une distance inter-antennes constante tout en faisant varier la longueur des dipôles qui forment la cavité à retournement temporel. A nouveau nous comparons les deux expressions de Ξ , en normalisant le résultat de l'équation I.48 par $\sigma\eta$. Le résultats de cette simulation pour laquelle l_2 varie de 0 à 0.35λ est représentée sur la figure I.14.

Comme on peut le voir les 2 courbes se superposent à nouveau de façon claire. Il est intéressant de noter que bien que nous ayons fait une approximation, dans la partie précédente, en considérant une cavité tapissée de dipôles infiniment petits, le résultat reste valable pour des tailles non négligeables.

Ce formalisme va nous être utile pour caractériser le retournement temporel en ne faisant intervenir que les impédances mutuelles de nos capteurs. Cependant, l'impédance mutuelle entre deux antennes dépend également de l'environnement immédiat de celles-ci. Or, nous allons voir par la suite que pour mesurer une tache focale créée par retournement temporel, il est obligatoire de mesurer celle-ci avec un réseau d'antennes plutôt qu'en déplaçant un unique capteur. Il va donc être nécessaire de compléter ce formalisme de la cavité à retournement tem-

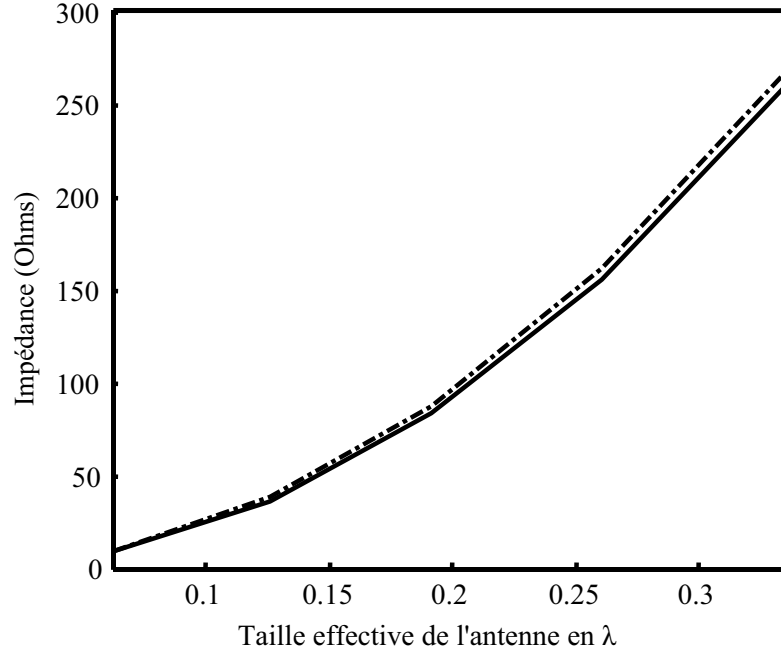


FIG. I.14 – Comparaisons entre les deux expressions de Ξ données par la relation I.38 (trait plein) et la relation I.48 (pointillés), lorsque la taille des petits dipôles qui forment la cavité varie.

porel électromagnétique en prenant en compte un réseau d'antennes de réception. Nous allons également prendre en compte un paramètre important : l'impédance de charge des antennes utilisées.

I.2.4 Influence du réseau d'antennes réceptrices lors de la mesure du champ retourné temporellement

Afin de prendre en compte l'effet du réseau de réception sur la mesure du champ retourné temporellement, nous allons à nouveau employer une modélisation des antennes fondée sur l'utilisation des impédances mutuelles. Ainsi un réseau de N antennes est complètement déterminé par une matrice d'impédances qui s'écrit :

$$\mathbf{Z} = \begin{bmatrix} Z_{11} & Z_{12} & \dots & Z_{1N} \\ Z_{21} & Z_{22} & \dots & \dots \\ \dots & \dots & \dots & \dots \\ Z_{N1} & \dots & \dots & Z_{NN} \end{bmatrix} \quad (\text{I.50})$$

où les impédances mutuelles sont toujours définies par la relation :

$$Z_{ij} = \frac{V(j)}{I_i \quad I_j \rightarrow 0, \forall i \neq j} \quad (\text{I.51})$$

Les tensions mesurées aux bornes de chaque antenne du réseau peuvent alors s'écrire en fonction des courants qui y circulent par une relation matricielle du type :

$$\mathbf{V} = \mathbf{Z}\mathbf{I} \quad (\text{I.52})$$

où nous avons introduit les vecteurs courant, $\mathbf{I} = \{I_1, I_2, \dots, I_N\}$ et tension, $\mathbf{V} = \{V_1, V_2, \dots, V_N\}$. Si on suppose que toutes les antennes du réseau ont la même impédance de charge Z_c (qui vaut 50Ω en pratique pour un amplificateur d'émission ou de réception), on peut alors modéliser l'émission d'un courant i_0 par une antenne 0 loin du réseau par le schéma qui est représenté sur la figure I.15 :

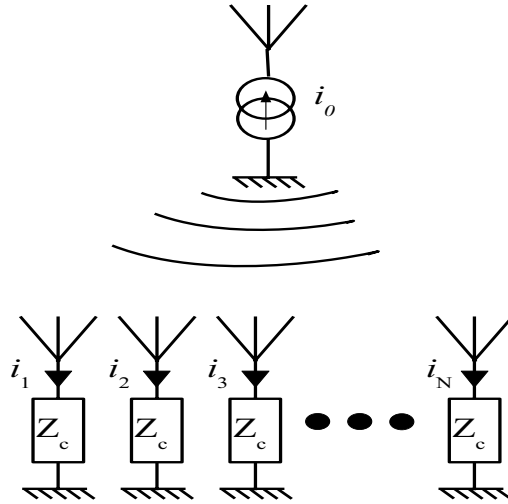


FIG. I.15 – Modélisation d'un réseau d'antennes soumis à une antenne parcouru par un courant i_0 .

Ainsi la tension créée aux bornes du dispositif d'acquisition, peut à la fois s'écrire grâce à la loi d'Ohm aux bornes de l'impédance de charge et en utilisant la matrice des impédances mutuelles :

$$U_k = -Z_c i_k = \sum_{l=1}^N Z_{kl} i_l + Z_{k0} i_0 \quad (\text{I.53})$$

où Z_{k0} est l'impédance mutuelle entre l'antenne k du réseau et l'antenne émettrice 0 (lointaine) et ce pour $1 \leq k \leq N$.

Afin de passer à une notation matricielle, nous introduisons le symbole de Kronecker $\delta_{i,j}$ qui vaut 0 si $i \neq j$ et 1 dans le cas contraire. On reformule alors l'équation précédente :

$$-Z_{k0}i_0 = \sum_{l=1}^N (Z_{kl} - \delta_{k,l}Z_c) i_l \quad (\text{I.54})$$

Le terme $(Z_{kl} - \delta_{k,l}Z_c i_l)$ peut s'écrire sous la forme matricielle suivante :

$$(Z_{kl} - \delta_{k,l}Z_c) = (\mathbf{Z} - Z_c \mathbf{1})_{k,l} \quad (\text{I.55})$$

où $\mathbf{1}$ est la matrice identité de dimensions N .

On peut alors définir son inverse et l'utiliser afin d'inverser l'équation I.54 pour obtenir :

$$i_l = - \sum_{k=1}^N (\mathbf{Z} - Z_c \mathbf{1})_{k,l}^{-1} Z_{k0} i_0 \quad (\text{I.56})$$

A présent nous allons utiliser la réciprocité afin de faciliter nos calculs. En effet, afin de modéliser une expérience de retournement temporel, nous aurions du générer initialement un courant i_0 sur une antenne du réseau, enregistrer les courants correspondants sur toutes les antennes de la cavité et appliquer à ces antennes les courants conjugués. Au lieu de cela, nous allons générer initialement un courant sur une unique antenne de la cavité, puis utiliser la réciprocité pour calculer les courants créés par retournement temporel grâce à cette unique antenne, sur notre réseau de réception.

Supposons donc que cette antenne émette un courant i_0 et que l'on enregistre alors les courants i_l créés sur chaque antenne du réseau. Nous avons vu que physiquement mesurer la résultat du retournement temporel sur une antenne l' lorsque l'on a appris à focaliser sur une antenne l revient à un coefficient près à mesurer $V_{l'} V_l^*$ si nous n'avons qu'une seule antenne dans notre miroir. Cette grandeur n'est autre que $i_{l'} i_l^*$, au produit des impédances de charge près. On peut donc écrire que la résultante du retournement temporel par une unique antenne de la cavité, mesurée sur notre réseau de réception, est proportionnelle à :

$$i_{l'} i_l^* = \sum_{k=1}^N \sum_{k'=1}^N (\mathbf{Z} - Z_c \mathbf{1})_{k,l'}^{-1} Z_{k0} i_0 i_0^* Z_{k'0}^* [(\mathbf{Z} - Z_c \mathbf{1})_{k',l}^{-1}]^\dagger \quad (\text{I.57})$$

A présent, on considère que nous n'avons plus seulement une antenne qui émet mais toute une cavité dont la surface est tapissée d'antennes comme dans la partie précédente. On peut alors considérer que, par linéarité, on doit sommer le résultat du retournement temporel pour chaque antenne 0 appartenant à la surface. Les courant i_0 restent les mêmes car le signal d'émission ne dépend pas de l'antenne considérée dans la phase d'enregistrement. On peut alors, si l'on

suppose que les antennes sont densément réparties mais sans couplage⁵, remplacer la somme par une intégrale de surface :

$$i_{l'} i_l^* = i_0 i_0^* \sigma \oint_S \sum_{k=1}^N \sum_{k'=1}^N (\mathbf{Z} - Z_c \mathbf{1})_{k,l'}^{-1} Z_{k0} Z_{k'0}^* [(\mathbf{Z} - Z_c \mathbf{1})_{k',l}^{-1}]^\dagger d^2 \mathbf{r}_0 \quad (\text{I.58})$$

où σ est à nouveau la densité surfacique d'antennes sensibles aux différentes polarisations du champ électrique qui couvrent toute la cavité et $\mathbf{r}_0$ représente leurs positions.

Ceci peut alors se simplifier en tenant compte du fait que seul le produit des impédances mutuelles $Z_{k0} Z_{k'0}^*$ dépend de la position $\mathbf{r}_0$ des antennes sur la cavité :

$$i_{l'} i_l^* = i_0 i_0^* \sigma \sum_{k=1}^N \sum_{k'=1}^N (\mathbf{Z} - Z_c \mathbf{1})_{k,l'}^{-1} \oint_S Z_{k0} Z_{k'0}^* d^2 \mathbf{r}_0 [(\mathbf{Z} - Z_c \mathbf{1})_{k',l}^{-1}]^\dagger \quad (\text{I.59})$$

On utilise alors le résultat obtenu dans la partie précédente pour un unique couple d'antenne et cette dernière formule s'exprime comme :

$$i_{l'} i_l^* = i_0 i_0^* \sigma l_0^2 \eta \sum_{k=1}^N \sum_{k'=1}^N (\mathbf{Z} - Z_c \mathbf{1})_{k,l'}^{-1} \Re \{Z_{kk'}\} [(\mathbf{Z} - Z_c \mathbf{1})_{k',l}^{-1}]^\dagger \quad (\text{I.60})$$

où on a introduit la longueur l_0 supposée petite des antennes qui tapissent la cavité et η l'impédance du vide.

Enfin, ceci étant à nouveau valable pour tout $l, l' \in [1..N]$, on peut reformuler cette dernière équation sous la forme d'une notation matricielle et on trouve alors que les vecteurs courant sur chaque antenne vérifient :

$$\mathbf{II}^* = i_0 i_0^* \sigma l_0^2 \eta (\mathbf{Z} - Z_c \mathbf{1})^{-1} \Re \{\mathbf{Z}\} [(\mathbf{Z} - Z_c \mathbf{1})^{-1}]^\dagger \quad (\text{I.61})$$

A l'aide de cette dernière formule, il est possible, lorsqu'une antenne d'un réseau a émis initialement un champ électromagnétique, de calculer les signaux reçus après retournement temporel sur chacune des antennes du réseau. Ce résultat ne s'exprime qu'en terme d'impédances mutuelle, qui sont des grandeurs facilement calculables numériquement dans le cas d'antennes simples. Nous allons donc pouvoir prédire, moyennant un calcul numérique, ce que produirait le retournement temporel en fonction du réseau de réception utilisé, à condition que ses antennes soient assez simples. De plus, cette formule tient compte du milieu proche qui entoure

⁵Cette hypothèse n'est pas contradictoire car la cavité est placée dans le champ lointain du réseau de réception. Une répartition dense d'antennes sous-entend donc que le champ est assez bien échantillonné, par exemple toutes les longueurs d'ondes, ce qui compte tenu que les antennes de la cavité sont petites, n'implique pas un couplage élevé

les antennes au travers des impédances mutuelles. Dans le cas où celui-ci est complexe, on ne pourra pas calculer ces impédances mutuelles numériquement, mais une mesure réalisée avec un analyseur de réseaux permettra d'y avoir accès. Ceci permettra également d'en déduire les résultats du retournement temporel. Ces derniers résultats généralisent l'approche traitée dans la partie précédente pour 2 antennes uniquement. Par la suite nous montrons comment nous avons confirmé cette formule par l'expérience.

Dans cette partie nous avons transposé les principes qui régissent le retournement temporel au cas des ondes électromagnétiques. En effet les équations de Maxwell sont réversibles temporellement si l'on néglige la dissipation et la propagation dans un milieu même hétérogène est réciproque. Ainsi on a pu mettre en évidence que l'utilisation d'une cavité à retournement temporel "tapissée" d'antennes sensibles à toutes les composantes d'un champ électromagnétique permet de renvoyer vers sa source une onde, pour donner un champ spatial gouverné par la partie imaginaire de la fonction de Green dyadique du milieu. Afin d'appliquer ce principe formel à nos expériences, nous avons formalisé le retournement temporel appliqué aux antennes en terme d'impédances mutuelles et montré que lorsque l'on retourne temporellement un champ qui a été émis par une antenne, la tension mesuré après retournement temporel sur une autre antenne est proportionnelle à la partie réelle de l'impédance mutuelle entre ces deux antennes. Enfin nous avons généralisé ce résultat au cas d'un réseau d'antennes et montré que l'on peut écrire une équation matricielle qui donne les courants induits après retournement temporel sur toutes les antennes de ce réseau, et ce uniquement en terme d'impédances mutuelles. Cette partie nous a donc permis de mieux comprendre le retournement temporel appliqué aux antennes. Nous avons de plus développé des outils qui permettent d'en prédire le comportement. Il est dès lors légitime de se demander pourquoi ce principe n'a jamais été appliqué aux ondes électromagnétiques bien qu'il soit réalisable et malgré tous les résultats intéressants obtenus dans le domaine ultrasonore. La réponse à cette question est purement technologique : les ondes électromagnétiques dans la gamme du gigahertz posent des problèmes d'instrumentation que l'on ne rencontre pas dans le domaine de l'acoustique ultrasonore usuel, où les fréquences sont de l'ordre du mégahertz. Dans la partie suivante nous allons les exposer et voir comment certaines solutions ont pu être mises en oeuvre afin de les contourner. Nous détaillerons par la suite le dispositif expérimental du premier miroir à retournement temporel monovoie à micro-ondes et nous en montrerons les résultats.

I.3 Un premier miroir à retournement temporel dans le domaine des micro-ondes

I.3.1 Retournement temporel d'un signal sur porteuse

L'une des difficultés rencontrée avec les ondes électromagnétiques, relativement aux ondes acoustiques, tient à leur célérité, très grande dans les milieux usuels, et en particulier dans l'air. S'il est "facile" de générer et de manipuler des ondes de basse fréquence, de l'ordre du MHz, celles-ci ont alors des longueurs d'ondes très élevées et leur utilisation est malaisée en laboratoire. En outre, l'un des intérêts du retournement temporel résidant dans la possibilité de focaliser des ondes sur des tailles aussi fines que la longueur d'onde, on comprend que les basses fréquences présentent assez peu d'intérêt pour nous. A titre d'exemple, considérons un signal acoustique à transmettre $s(t)$, dont la bande passante est de 20 kHz. Si l'on voulait réaliser une antenne électromagnétique pour transmettre ce signal, celle-ci devrait mesurer autour de la demi-longueur d'onde pour être efficace, soit $\lambda/2 = c/f \approx 7$ km de haut, avec la vitesse de la lumière $c = 3.10^8$ m.s⁻¹. Afin de manier des longueurs d'ondes d'un ordre de grandeur plus raisonnable, il est préférable de travailler dans des bandes de fréquence bien plus hautes. Cependant, lorsque l'on monte en fréquence, les composants électroniques deviennent très onéreux, si toutefois ils existent.

Pour pallier ces problèmes, on utilise des techniques de modulation, procédés qui sont utilisés dans de nombreux domaines, comme les télécommunications ou l'imagerie, et dont le principe est basé sur le décalage en fréquence d'un signal. Employer une modulation consiste à appliquer une transformation de la forme :

$$s(t) \rightarrow S(t) \cdot \exp(j2\pi f_0 t) \quad (\text{I.62})$$

où f_0 est la fréquence porteuse du signal, qui peut être bien plus élevée que la fréquence maximale du signal initial $s(t)$ et $S(t)$ est l'enveloppe complexe du signal qui porte l'information contenue dans $s(t)$. L'intérêt réside dans le fait que le signal est ainsi déplacé vers les fréquences élevées d'une valeur f_0 ; en effet, si l'on considère cette transformation dans le domaine de Fourier on a :

$$\tilde{s}(f) \rightarrow \tilde{S}(f) \otimes \delta(f - f_0) \quad (\text{I.63})$$

en ne considérant que les fréquences positives et en omettant les constantes. Il est clair, ici, que le spectre $\tilde{S}(f)$ de $S(t)$ a été translaté de la fréquence porteuse f_0 . Ainsi, si l'on veut

maintenant envoyer un signal acoustique, il suffit de le moduler sur une porteuse à 100 MHz par exemple, ce qui donne des tailles d'antennes de l'ordre du mètre : c'est le principe de la radio FM. L'opération est réalisée par le biais d'un modulateur. Du côté de la réception, un démodulateur permet de translater en sens inverse le spectre du signal reçu afin de manipuler à nouveau des composantes basses fréquences appelées "signaux en bande de base". Ainsi, à l'émission comme à la réception, l'appareillage électronique est assez simple, la complexité étant alors dans le modulateur et le démodulateur.

C'est de cette observation qu'est née l'idée de la première démonstration du retournement temporelle avec des ondes électromagnétiques. En effet, pour rester dans des domaines manipulables en laboratoire, les ondes utilisées devaient être de l'ordre du GHz. De plus, pour des raisons de coût, il était intéressant d'utiliser des composants électroniques courants donc peu onéreux. Le choix s'est porté sur la bande de fréquence utilisée par les systèmes WIFI et Bluetooth, en pleine expansion à ce moment là, dont la fréquence porteuse est de 2.45 GHz et donc la longueur d'onde associée vaut 12.25 cm. Dès lors s'est posée la question de la technique à employer pour réaliser un retournement temporel à ces fréquences. Nous avons vu que lors de la phase d'enregistrement du champ, il est nécessaire de numériser les signaux mesurés et de les stocker en mémoire. Afin de respecter le critère de Shannon, des convertisseurs analogique-numérique fonctionnant au moins à 5 Giga échantillons par seconde sont nécessaires pour assurer un bon échantillonnage du signal. Lors de la phase de réémission, il est nécessaire d'utiliser des convertisseurs numérique-analogique fonctionnant à ces fréquences : ces appareils n'existent toujours pas à ce jour.

En réalité, un signal modulé peut toujours s'écrire de la façon suivante :

$$e(t) = A(t)\exp(i\Phi(t)) \quad (\text{I.64})$$

Où $A(t)$ est l'amplitude du signal, et $\Phi(t)$ la phase du signal.

Afin de retourner temporellement ce signal, il faut alors retourner temporellement à la fois la phase de ce signal et l'amplitude associée. Dans notre cas la modulation employée sera de type IQ⁶. Un signal modulé de type IQ peut s'écrire de la façon suivante :

$$e(t) = \{E_I(t) \cos(\omega_0 t) + E_Q(t) \sin(\omega_0 t)\} \quad (\text{I.65})$$

où $E_I(t)$ est le signal en phase car il est multiplié dans le modulateur par un cosinus à la

⁶In-phase and Quadrature

fréquence porteuse, $E_Q(t)$ est le signal en quadrature car il est multiplié par un sinus, et ω_0 est la pulsation associée à la fréquence porteuse. Si on appelle B la bande passante des signaux en bande de base, c'est-à-dire la fréquence maximale de ces signaux, celle-ci est toujours inférieure à la fréquence porteuse.

Si l'on veut transmettre dans un milieu donné une onde électromagnétique modulée en phase et en quadrature, le signal à émettre s'écrit comme l'équation I.65. On peut alors toujours exprimer en n'importe quel point du milieu le signal réel reçu $r(t)$ comme :

$$r(t) = \{R_I(t) \cos(\omega_0 t) + R_Q(t) \sin(\omega_0 t)\} \quad (\text{I.66})$$

où $R_I(t)$ et $R_Q(t)$ sont les signaux reçus en bande de base.

Il est important de noter ici que la seule connaissance des signaux reçus en bande de base contient toute l'information de la propagation de l'onde dans le milieu. Afin de retourner temporellement cette onde, il faut maintenant générer $r(-t)$ qui peut s'écrire dans le formalisme IQ comme :

$$r(-t) = \{R_I(-t) \cos(\omega_0 t) - R_Q(-t) \sin(\omega_0 t)\} \quad (\text{I.67})$$

Cette dernière équation est la définition même du retournement temporel d'un signal modulé en phase et en quadrature. La seule connaissance des composantes en bande de base du signal permet de réaliser l'opération, qui consiste à :

- Retourner temporellement les composantes en phase et en quadrature du signal en bande de base.
- Conjuguer la phase de la fréquence porteuse du signal modulé c'est-à-dire multiplier la composante en quadrature par un signe“-”.

Ainsi apparaît l'avantage de travailler sur des signaux modulés pour faire du retournement temporel : avec des composants électroniques d'acquisition et de génération arbitraire capables de traiter des signaux basse fréquence, on va pouvoir focaliser en temps et en espace des ondes de très haute fréquence. Un exemple pratique sera donné dans la partie suivante qui est consacrée à la description de notre premier miroir à retournement temporel monovoie dans le domaine des micro-ondes. Mais auparavant, il est nécessaire de d'étudier plus précisément et de formaliser un peu cette technique. En effet, nous avons vu en acoustique que la compression temporelle obtenue par retournement temporel dépend de la bande passante utilisée et que dans certaines conditions la focalisation spatiale peut atteindre une demi-longueur d'onde : qu'en est il pour

le retournement temporel sur fréquence porteuse ?

Pour avoir une idée de la focalisation spatiale, on écrit d’abord le résultat d’une opération de retournement temporel sur modulation, en gardant comme signal d’émission initial le signal $e(t)$ et en utilisant une bande passante négligeable. On écrit alors le champ recréé grâce à la formule I.33 démontrée pour la cavité à retournement temporel électromagnétique. Celui-ci sera proportionnel à la partie imaginaire de la fonction de Green dyadique du milieu à la pulsation porteuse ω_0 . En milieu homogène, on en conclut que la focalisation sera limitée par la longueur d’onde associée à la fréquence porteuse. Nous avons là une conséquence importante du retournement temporel sur porteuse : bien que l’on ne manipule que des fréquences bien inférieures à la fréquence porteuse du signal, la largeur de la tache de focalisation est conditionnée par la fréquence de la modulation. On peut ainsi obtenir la finesse des figures de diffraction observées en ultrasons : il suffit d’augmenter la fréquence de la modulation sans modifier le reste du dispositif.

Concernant la compression temporelle, le problème est un peu plus compliqué. Nous avons vu que le fait de moduler le signal sur une fréquence porteuse provoque une translation du spectre des signaux en bande de base vers la fréquence porteuse. Parallèlement les études d’acoustique ont montré que la qualité de la compression temporelle, c’est-à-dire le rapport de l’amplitude du pic de retournement temporel à l’amplitude moyenne des lobes secondaires, est gouvernée le nombre de “grains d’informations” qui sont présents dans la bande passante du signal, soit par le rapport entre la bande passante des signaux et la fréquence de corrélation du milieu de propagation. La fréquence de corrélation d’un milieu dont l’absorption est nulle est d’autant plus faible que l’on se place sur une plage de fréquences élevée⁷. Cependant, dans le cas d’un milieu absorbant, la compression est limitée par le temps de réverbération, qui comme nous l’avons dit précédemment, empêche de résoudre les modes du milieu. Cette limitation ne dépend pas de la fréquence porteuse mais uniquement de la bande passante utilisée et des caractéristiques du milieu. En conclusion, il peut être utile d’augmenter la fréquence porteuse afin de bénéficier de plus de grains d’informations et donc d’une meilleure compression temporelle lors d’une expérience de retournement temporel. En revanche, celle-ci sera tout de même limitée par le produit $\Delta\nu\tau$, avec τ le temps réverbération du milieu et $\Delta\nu$ la bande passante des signaux utilisés.

⁷A titre d’exemple, dans une cavité, le nombre de mode est directement proportionnel à la fréquence porteuse des ondes

I.3.2 Réponses impulsionnelles en bande de base

Afin de formuler plus précisément le retournement temporel en bande de base, il est utile de généraliser le concept de réponse impulsionnelle en bande de base. Ces réponses relient les entrées IQ du modulateur aux sorties du démodulateur et prennent en compte la propagation à la fréquence porteuse. Afin de définir ces réponses, nous allons écrire le signal reçu après propagation, en fonction de la réponse impulsionnelle du milieu. Pour simplifier on assimilera les réponses des dispositifs électroniques à des Diracs afin de ne pas les faire intervenir dans le calcul : seuls les effets de la propagation seront pris en compte. Supposons donc que l'on émette un signal en bande de base $\{E_I(t), E_Q(t)\}$, modulé à la pulsation ω_0 . En notant $h_{RF}(t)$ la réponse impulsionnelle réelle du milieu entre l'émetteur et le récepteur, on écrit alors le signal $s(t)$ reçu comme :

$$s(t) = \{E_I(t) \cos(\omega_0 t) + E_Q(t) \sin(\omega_0 t)\} \otimes h_{RF}(t) \quad (\text{I.68})$$

On définit alors les réponses impulsionnelles en bande de base $H_{I \rightarrow I}(t)$, $H_{I \rightarrow Q}(t)$, $H_{Q \rightarrow I}(t)$ et $H_{Q \rightarrow Q}(t)$, afin de réécrire ce même signal $s(t)$:

$$\begin{aligned} s(t) = & \{E_I(t) \otimes H_{I \rightarrow I}(t) + E_Q(t) \otimes H_{Q \rightarrow I}(t)\} \cos(\omega_0 t) \\ & + \{E_Q(t) \otimes H_{Q \rightarrow Q}(t) + E_I(t) \otimes H_{I \rightarrow Q}(t)\} \sin(\omega_0 t) \end{aligned} \quad (\text{I.69})$$

Le but est maintenant d'exprimer la première équation de manière à faire apparaître les dépendances en cosinus et en sinus en facteur. Pour ce faire, nous allons développer les produits de convolution dans l'équation I.68 :

$$\begin{aligned} s(t) = & \int_{-\infty}^{\infty} E_I(t - \tau) \cos(\omega_0(t - \tau)) h_{RF}(\tau) d\tau \\ & + \int_{-\infty}^{\infty} E_Q(t - \tau) \sin(\omega_0(t - \tau)) h_{RF}(\tau) d\tau \end{aligned} \quad (\text{I.70})$$

En développant les termes en sinus et cosinus dans les deux intégrales précédentes on obtient alors :

$$\begin{aligned} s(t) = & \cos(\omega_0 t) \left\{ \int_{-\infty}^{\infty} E_I(t - \tau) \cos(\omega_0 \tau) h_{RF}(\tau) d\tau - \int_{-\infty}^{\infty} E_Q(t - \tau) \sin(\omega_0 \tau) h_{RF}(\tau) d\tau \right\} \\ & + \sin(\omega_0 t) \left\{ \int_{-\infty}^{\infty} E_Q(t - \tau) \cos(\omega_0 \tau) h_{RF}(\tau) d\tau + \int_{-\infty}^{\infty} E_I(t - \tau) \sin(\omega_0 \tau) h_{RF}(\tau) d\tau \right\} \end{aligned} \quad (\text{I.71})$$

Cette dernière expression peut alors être reformulée en utilisant les produits de convolution :

$$\begin{aligned}
s(t) = & \cos(\omega_0 t) \{E_I(t) \otimes \cos(\omega_0 t)h_{RF}(t) - E_Q(t) \otimes \sin(\omega_0 t)h_{RF}(t)\} \\
& \sin(\omega_0 t) \{E_Q(t) \otimes \cos(\omega_0 t)h_{RF}(t) + E_I(t) \otimes \sin(\omega_0 t)h_{RF}(t)\}
\end{aligned} \tag{I.72}$$

Il reste alors à identifier les différents termes des équation I.69 et I.72, pour obtenir les 4 réponses impulsionnelles en bande de base. On observe qu'en réalité seules 2 de ces 4 réponses sont indépendantes :

$$\begin{aligned}
H_I(t) = H_{I \rightarrow I}(t) = H_{Q \rightarrow Q}(t) &= h_{RF}(t) \cdot \cos(\omega_0 t) \\
H_Q(t) = -H_{I \rightarrow Q}(t) = H_{Q \rightarrow I}(t) &= -h_{RF}(t) \cdot \sin(\omega_0 t)
\end{aligned} \tag{I.73}$$

Les signaux en réception pourront être déduit des signaux en émission par le produit de convolution matriciel suivant :

$$s(t) = \begin{bmatrix} E_I^{out}(t) \\ E_Q^{out}(t) \end{bmatrix} = \begin{bmatrix} H_I(t) & H_Q(t) \\ -H_Q(t) & H_I(t) \end{bmatrix} \otimes \begin{bmatrix} E_I^{in}(t) \\ E_Q^{in}(t) \end{bmatrix} \tag{I.74}$$

Plusieurs conclusions peuvent être tirées des équations I.73 et I.74. Tout d'abord, il est suffisant d'envoyer un signal bref sur la voie I de notre modulateur pour obtenir les 4 réponses impulsionnelles du système, car seules 2 sont indépendantes. C'est-à-dire qu'il n'est pas nécessaire d'avoir recours à deux phases d'enregistrement pour réaliser une expérience de retournement temporel.

D'autre part, si l'on considère l'émission d'un signal E_I sur la voie I du modulateur de la carte d'émission, la formule I.74 nous donne les signaux en sortie du démodulateur en réception :

$$s^I(t) = \begin{bmatrix} H_I(t) & H_Q(t) \\ -H_Q(t) & H_I(t) \end{bmatrix} \otimes \begin{bmatrix} E_I(t) \\ 0 \end{bmatrix} = \begin{bmatrix} H_I(t) \otimes E_I(t) \\ -H_Q(t) \otimes E_I(t) \end{bmatrix} \tag{I.75}$$

Si on veut alors faire du retournement temporel, il suffit de renvoyer les signaux en bande de base retournés temporellement tout en conjugant la phase de la porteuse. On mesure alors la tension aux bornes de l'antenne qui a émis initialement, celle-ci est proportionnelle à :

$$s_{RT}^I(t) = \begin{bmatrix} H_I(t) & H_Q(t) \\ -H_Q(t) & H_I(t) \end{bmatrix} \otimes \begin{bmatrix} H_I(-t) \otimes E_I(-t) \\ H_Q(-t) \otimes E_I(-t) \end{bmatrix} \tag{I.76}$$

On peut alors calculer le produit matriciel pour obtenir les signaux résultant du retournement temporel d'un signal émis sur la voie I :

$$s_{RT}^I(t) = \begin{bmatrix} (H_I(t) \otimes H_I(-t) + H_Q(t) \otimes H_Q(-t)) \otimes E_I(-t) \\ (-H_Q(t) \otimes H_I(-t) + H_I(t) \otimes H_Q(-t)) \otimes E_I(-t) \end{bmatrix} \quad (I.77)$$

De la même façon, si l'on considère que l'on a émis initialement une impulsion sur la voie Q, le résultat du retournement temporel donne :

$$s_{RT}^Q(t) = \begin{bmatrix} (-H_I(t) \otimes H_Q(-t) + H_Q(t) \otimes H_I(-t)) \otimes E_Q(-t) \\ (-H_I(t) \otimes H_I(-t) - H_Q(t) \otimes H_Q(-t)) \otimes E_Q(-t) \end{bmatrix} \quad (I.78)$$

Que l'on regarde les signaux s_{RT}^I ou s_{RT}^Q , il est clair que la voie de modulation excitée initialement comporte la somme des autocorrélations des réponses impulsionnelles en bande de base tandis que l'autre voie comporte la différence des intercorrélations de ces mêmes réponses. Par conséquent le signal reçu sur la voie d'émission initiale va bien comporter un maximum à l'instant $t = 0$, car les deux autocorrélations vont être maximales à cet instant précis. De plus, on peut remarquer que le signal présent sur la deuxième voie, qui est une différence d'intercorrélations, ne va pas être toujours nul : cela va produire des lobes secondaires. Par contre nous sommes certains qu'à l'instant $t = 0$ défini comme le temps du maximum du pic de retournement temporel, ce signal va être nul car les intercorrélations sont égales ($H_I(0) \otimes H_Q(0) = H_Q(0) \otimes H_I(0)$). A ce temps précis on a donc un maximum sur la voie de modulation initiale et un zéro sur l'autre. Par contre, en dehors de ce temps, des lobes secondaire d'une amplitude équivalente vont être présents sur les deux voies du démodulateur. On peut également remarquer que le retournement temporel agit de façon asymétrique sur les voies en phase et en quadrature. Ainsi un signal retourné temporellement sur la voie I arrive avec un signe positif, mais il va être naturellement négatif⁸ si on apprend à focaliser par retournement temporel sur la voie Q. Enfin, on peut souligner qu'il est possible d'envoyer des signaux différents en même temps sur les deux voies I et Q, en utilisant la propriété de linéarité des ondes électromagnétiques. Le résultat est alors une combinaison linéaire des 2 signaux s_{RT}^I et s_{RT}^Q .

Toutes ces observations nous permettent d'aborder les expériences de retournement temporel en bande de base dans le paragraphe suivant. Nous y présentons le premier dispositif

⁸Il suffit alors d'envoyer le signal à focaliser avec un signe opposé pour recevoir un résultat positif.

expérimental qui nous a permis de réaliser un retournement temporel d'ondes électromagnétiques sur fréquence porteuse.

I.3.3 Dispositif expérimental et premiers essais

Le premier dispositif expérimental de miroir à retournement temporel a été réalisé au tout début de cette thèse, avec un matériel simple et bon marché. Le choix s'est porté naturellement sur des émetteurs/récepteurs utilisés dans les réseaux locaux, fonctionnant dans la gamme de fréquence des systèmes WIFI. L'intérêt de ces cartes est que la fréquence porteuse est assez élevée, 2.45 GHz, sans que le reste des éléments mis en jeu soient très sophistiqués. En effet, la largeur de la bande passante théorique de ces cartes est de 10 MHz et la fréquence porteuse est générée directement en interne. La figure I.16 montre le dispositif qui a été construit et dont nous allons expliciter les différents éléments.

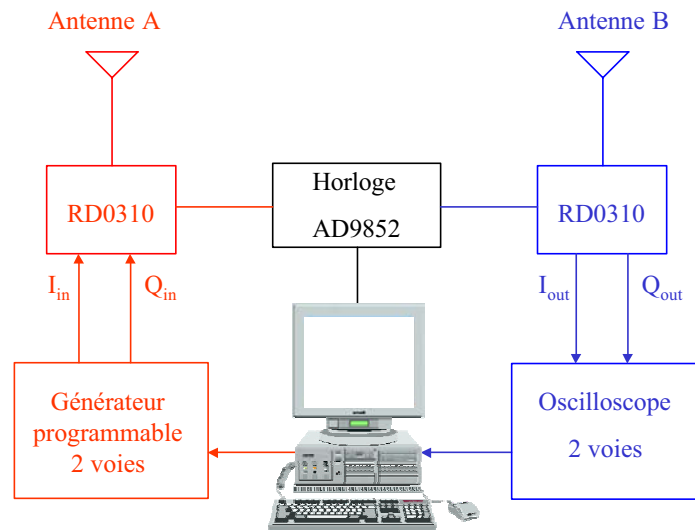


FIG. I.16 – Dispositif expérimental du premier MRT monovoie.

Les cartes émettrices/réceptrices choisies sont des RD0310 du fabricant RF micro device. Ces cartes, dont la fréquence porteuse est générée par une boucle à verrouillage de phase, ne nécessitent qu'une horloge d'assez basse fréquence : 44 MHz. Le grand intérêt de ces modèles est d'être vendus sans protocole de communication, c'est-à-dire que l'utilisateur a un accès direct aux signaux analogiques en phase et en quadrature, en entrée ou en sortie. Un autre avantage est que l'amplification est incluse à l'origine, avec une puissance de 16 dBm en émission et un gain de 10 dB en réception, ce qui est suffisant pour des expériences en laboratoire. Afin de pouvoir

les utiliser, il a simplement fallu créer un circuit d'alimentation et de contrôle pour basculer du mode émission au mode réception, ainsi que pour régler les gains des différents amplificateurs. Puis les émetteurs/récepteurs, l'un fonctionnant en mode miroir à retournement temporel (en rouge sur la figure) et l'autre fonctionnant en mode récepteur ont été placés dans des boîtiers blindés afin d'éviter les radiation parasites venant de l'extérieur.

Un générateur de créniaux à double sortie a été utilisé afin que les cartes d'émission et de réception partagent la même horloge. Nous avons employé un synthétiseur de fréquence de marque Analog Device, la carte AD9852 de la figure précédente. Cette carte ayant besoin d'une entrée à 10 MHz, nous avons utilisé pour l'alimenter un générateur de signaux traditionnel, que l'on a omis volontairement sur le schéma pour des raisons de clarté. Une fois le circuit d'alimentation réalisé et la carte mise dans un boîtier blindé, celle-ci dispose de deux sorties synchronisées très stable en fréquence et en phase. Cette dernière caractéristique est indispensable, car une erreur sur l'horloge à 44 MHz va être multipliée par un facteur 50 sur la porteuse à 2.45 GHz. Par ailleurs tous les paramètres de ce synthétiseur sont programmés par un ordinateur via un port série.

La méthodologie utilisée pour le retournement temporel est la suivante :

- un signal bref en bande de base est envoyé sur les canaux IQ de l'émetteur par un générateur arbitraire de signaux à deux voies synchronisées. Cet appareil a une fréquence d'échantillonnage de 40 MHz, soit une bande passante de 15 MHz. Ceci est suffisant pour les 10 MHz des carte RD0310.
- Ce signal est modulé sur la fréquence porteuse par le modulateur de la carte émettrice, puis amplifié et enfin transmis à l'antenne d'émission A qui est de type omnidirectionnelle.
- Les ondes se propagent dans le milieu et une partie du champ est captée par l'antenne de réception B qui est, elle aussi, omnidirectionnelle.
- Le signal ainsi créé est alors démodulé par la carte réceptrice donnant ainsi naissance aux signaux IQ de sortie.
- On réalise l'acquisition de ces signaux à l'aide d'un oscilloscope à 2 voies. Les signaux sont alors numérisés sur 8 bits. Des moyennes peuvent éventuellement être effectuées si besoin est.
- Par le biais d'un port GPIB, on transfère les signaux vers un ordinateur. Ceux-ci sont alors filtrés et retournés temporellement, comme décrit précédemment, sous Matlab.
- Une autre connexion GPIB permet ensuite de charger ces signaux dans les mémoires numériques du générateur arbitraire, afin de passer à la phase de réémission du retournement temporel.

- Après modulation, amplification et émission des signaux retournés temporellement, les signaux I et Q reçus sur l'antenne A sont mesurés par l'oscilloscope à l'aide duquel on peut en faire l'acquisition.

Les premiers essais ont été réalisés alors que les deux antennes étaient séparées de quelques mètres au sein du laboratoire. Comme on le voit sur la figure I.17, le signal envoyé est composé de 2 arches de sinusoïde à 1 MHz sur la voie I et est nul sur la voie Q.

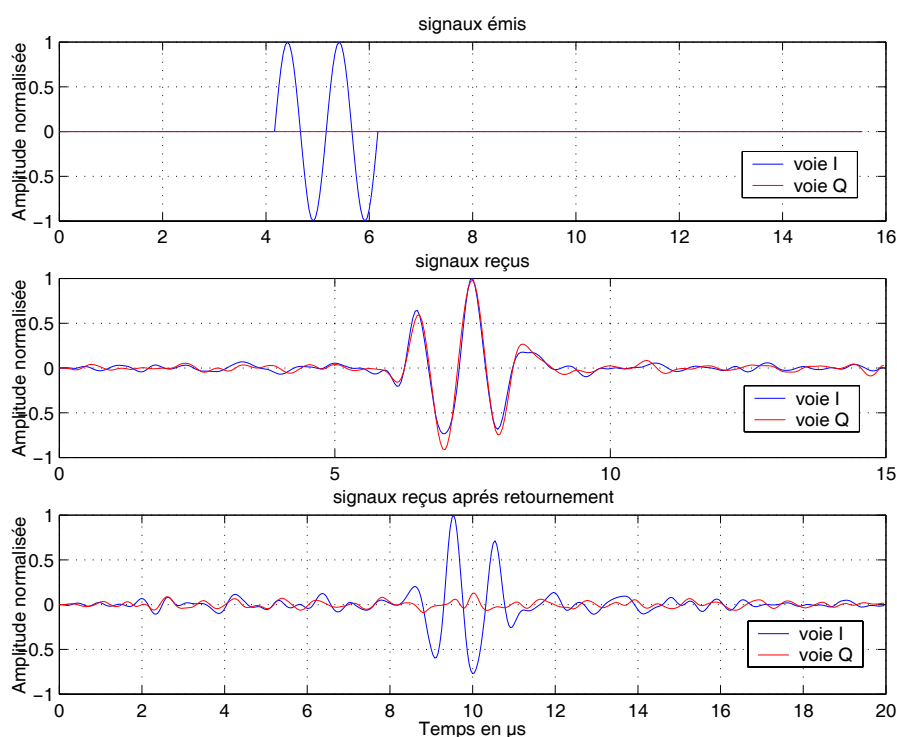


FIG. I.17 – Retournement temporel d'un signal envoyé sur la voie I.

En observant les réponses en phase et en quadrature, on peut s'apercevoir de l'absence d'allongement du signal : on observe uniquement le signal initial déphasé. La raison en est assez simple : l'impulsion dure $2 \mu\text{s}$. Pour voir un étalement significatif du signal, soit au moins une arche de sinusoïde supplémentaire, un temps de réverbération de $1 \mu\text{s}$ est nécessaire. Il faudrait ainsi que l'onde ait parcouru grâce aux diverses réflexions au moins 300 m dans le laboratoire ! Il peut éventuellement exister de telles ondes, mais le coefficient de réflexions des murs étant trop faible, celles-ci ont une amplitude négligeable. Il est clair que ce n'est pas le meilleur environnement pour observer une quelconque compression temporelle et donc pour prouver les avantages du retournement temporel. En revanche, on remarque que le signal reçu est conforme au formalisme impulsionnel en bande de base : l'impulsions est reçue dans une chronologie

inversée sur le voie I, et le signal sur la voie Q est nul. Nous avons réalisé la même opération en utilisant cette fois à l'émission une impulsion sur la voie Q et un signal nul sur la voie I (Fig I.18).

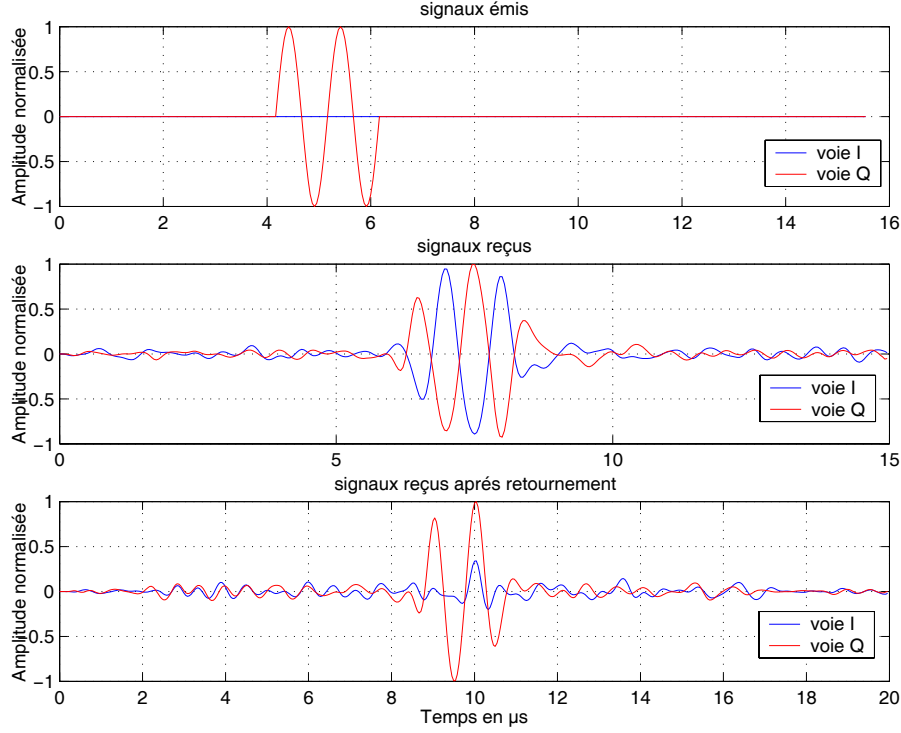


FIG. I.18 – Retournement temporel d'un signal envoyé sur la voie Q.

Bien évidemment, nous n'observons toujours aucun étalement de l'impulsion initiale. Par contre le déphasage induit par la propagation est différent, ceci étant du à l'asymétrie de la matrice des réponses impulsionnelles. Sur les courbes correspondant aux signaux reçus après retournement temporel, on voit que l'impulsion n'est présente que sur la voie Q. De plus, on peut remarquer que les 2 arches de sinusoïde sont dans le même sens que lors de l'émission : ceci est une nouvelle fois en accord avec les formules du retournement temporel en bande de base. En effet, l'impulsion a été retournée temporellement mais elle arrive avec un signe négatif et on retrouve ainsi exactement l'impulsion initiale, car notre signal initial est antisymétrique. Enfin, que l'on ait focalisé en I ou en Q, les deux figures précédentes montrent que sur la voie non choisie le signal est exactement nul au milieu de l'impulsion qui est le $t = 0$ du retournement temporel. Les résultats précédents, s'ils permettent de vérifier certains aspects du retournement temporel sur fréquence porteuse, ne suffisent pas à démontrer le bien fondé de la technique, en terme de compression temporelle ou spatiale. En réalité, la bande passante de notre système étant

bien trop faible⁹ pour observer un quelconque étalement temporel, l'opération de retournement temporel effectuée s'apparente à une opération de conjugaison de phase.

I.3.4 Retournement temporel électromagnétique en cavité réverbérante : les premiers résultats probants

Compte tenu de la faible bande passante de notre premier système, nous n'avons observé qu'un très faible étalement des ondes émises dans des environnements normaux tels que notre laboratoire. Nous sommes très loin des expériences de Carsten Draeger faites dans des cavités acoustiques très réverbérantes [10]. En effet, dans ces expériences, il obtenait des temps de réverbération de l'ordre de 20 ms pour des impulsions initiale de $1 \mu s$. Notons tout de même une différence entre les deux dispositifs : Carsten Draeger disposait d'une bande passante mesurée de l'ordre de 60% alors que la nôtre est uniquement de l'ordre du millième. Nous n'espérons donc pas obtenir des étalement temporels aussi élevés.

Notre montage a cependant été transportée à l'école Supélec, dans le laboratoire de Signaux et Systèmes, dans lequel Alain Azoulay et Vikass Monheburhun nous ont gracieusement donné accès à une chambre réverbérante. Les dimensions de cette cavité, dont tous les murs sont métalliques, sont données sur le schéma de la Fig I.19. En outre, on peut voir sur cette figure que la cavité dispose d'un tableau de connecteurs externes, ce qui nous permet de changer les branchements sans modifier l'intérieur de la cavité et donc sans modifier les conditions aux limites, de manière à pouvoir tester la réciprocité ou encore vérifier la focalisation spatiale obtenue par retournement temporel.

Ainsi, tout le matériel utilisé, émetteurs/récepteurs, ordinateur, générateur de signaux et oscilloscope a été installé en dehors de la cavité, la connection avec les antennes dans la cavité étant réalisée par l'intermédiaire du tableau de connecteurs externes. Le protocole expérimental qui avait été utilisé au laboratoire a été conservé. Dans une première phase, nous émettons donc une impulsion uniquement sur la voie I du modulateur de la carte d'émission et le signal sur la voie Q reste égal à zéro. L'impulsion utilisée ici, comme on peut le voir sur la figure I.20.(a) est de type gaussien et dure à peu près $1 \mu s$.

Le signal, modulé et amplifié, est transmis par le biais du connecteur externe à une antenne de type omnidirectionnelle, A, qui se trouve dans la cavité. Une autre antenne B, de même type, est placée à une distance d'à peu près 2 m de la première. Cette antenne reçoit une

⁹Nous avons en réalité estimé une bande passante à -6 dB de l'ordre de 4 MHz

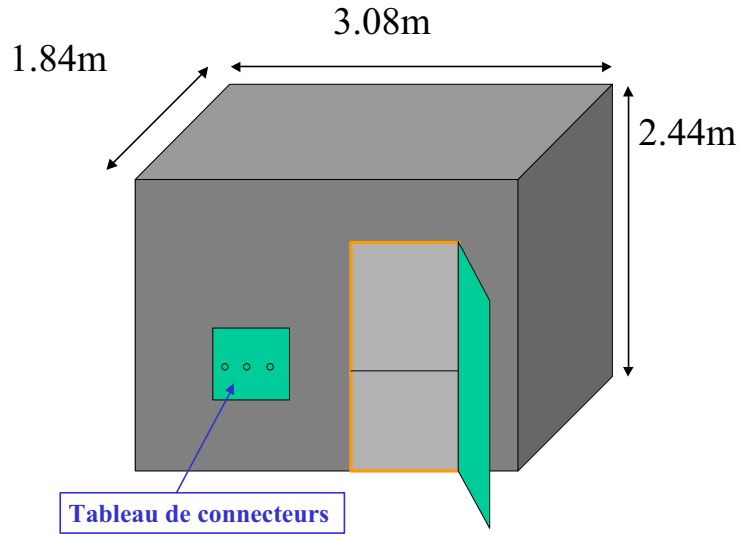


FIG. I.19 – Schéma de la cavité réverbérante utilisée à Supélec.

partie du champ, générant à ses bornes un signal qui est démodulé. On procède alors grâce à l'oscilloscope à l'acquisition des réponses en bande de base que l'on peut voir sur la figure I.20.(b). Nous observons un étalement de l'impulsion initiale sur près de $8 \mu s$, ce qui représente un étalement d'un facteur 8 par rapport à sa durée initiale, ou encore à peu près 20000 périodes de la fréquence porteuse. Les ondes émises par l'antenne A ont donc parcouru plus de 2 km au sein de la chambre réverbérante de $14 m^3$: elles ont été réfléchies 6000 fois en moyenne. Il faut noter qu'ici l'atténuation de l'onde est presque uniquement due à l'effet de peau sur les parois métalliques de la cavité.

Dans un premier temps, à l'aide du tableau de connecteurs externes, nous avons pu vérifier la réciprocité des réponses impulsionnelles sans changer la configuration à l'intérieur de la chambre. Pour ce faire, il a suffi de brancher l'émetteur sur le connecteur correspondant à l'antenne B et réciproquement. Les réponses impulsionnelles en bande de base ainsi mesurées, que l'on note, $\{H_I(A \rightarrow B), H_Q(A \rightarrow B)\}$ et $\{H_I(B \rightarrow A), H_Q(B \rightarrow A)\}$ ont pu être comparées, et après filtrage du bruit, nous avons mesuré un coefficient de corrélation de l'ordre de 0.98.

Dans la phase de réémission, les réponses en bande de base sont retournées temporellement et la porteuse est conjuguée en multipliant la partie en quadrature par un signe $-$. On émet ces signaux par l'antenne B et on acquiert les signaux en bande de base provenant du champ mesuré sur l'antenne A. Ces signaux sont représentés sur la figure I.21.(a). Il faut noter ici que l'origine des temps a été choisie arbitrairement afin de centrer les courbes. Comme on peut le voir, bien que nous ne disposions que d'un miroir à retournement temporel à une seule voie,

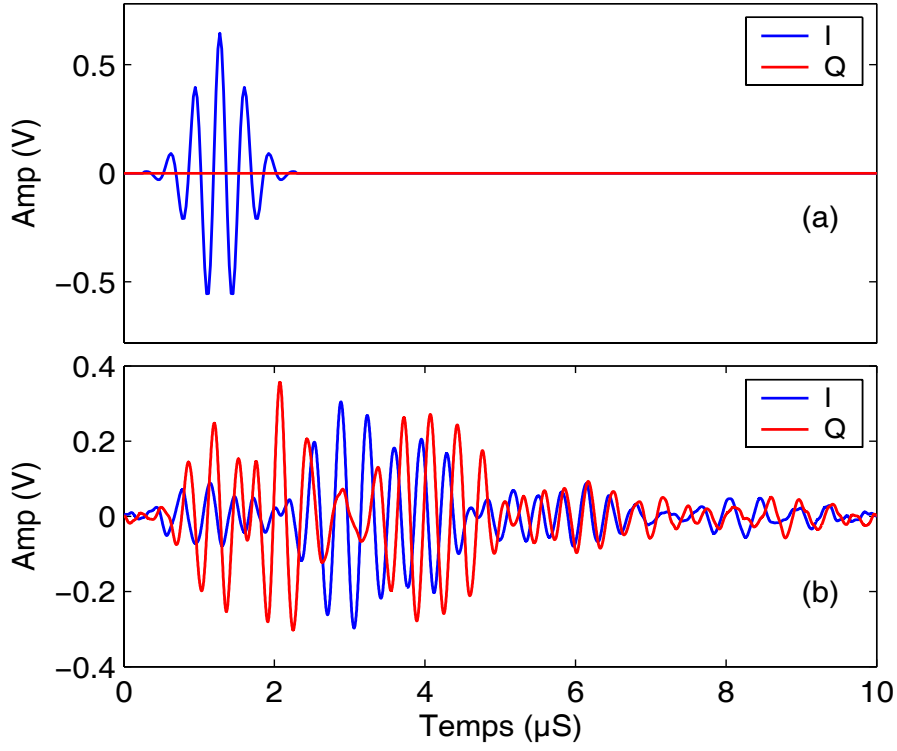


FIG. I.20 – (a) Signaux IQ lors de la phase d’émission en A. (b) signaux IQ reçus par l’antenne B.

on observe une compression temporelle des signaux envoyés. On retrouve ainsi l’impulsion que l’on avait émise initialement sur la voie I.

Comme nous l’avions prédit, le signal est maximal sur la voie I à $t = 0$ (le milieu de la fenêtre ici) et il est exactement nul sur la voie Q. De plus, l’impulsion reçue sur la voie I n’est pas parfaite, on peut observer des lobes secondaires. De même, du bruit est présent sur la voie Q : ceci est en accord avec les résultats discutés précédemment en ultrasons. En effet, le miroir n’étant constitué que d’une seule voie, le gain d’antenne¹⁰ est nul dans cette situation. La compression temporelle résulte donc uniquement du caractère large bande du signal émis initialement. Nous avons vu que dans un tel cas le rapport entre l’amplitude du pic et l’amplitude moyenne du bruit est donné par $\sqrt{\frac{\delta\nu}{\delta\omega}}$, $\delta\nu$ étant la bande passante utilisée et $\delta\omega$ la fréquence de corrélation du milieu. Or, ici le temps d’Heisenberg de la cavité, c’est-à-dire l’inverse de la distance entre les modes de la cavité, est de $80 \mu\text{s}$ ce qui est bien supérieur au temps d’atténuation. Le rapport signal-sur-bruit est donc limité par le temps de réverbération de la chambre réverbérante τ , et s’écrit donc $\sqrt{\delta\nu\tau}$. Celui-ci est donné par le constructeur, il vaut à peu près $3.6 \mu\text{s}$, et la bande

¹⁰Dans tout ce manuscrit le terme “gain d’antenne” représente un gain d’énergie dû à la focalisation d’ondes par un réseau de capteurs. Il n’a donc rien à voir avec le gain d’une antenne au sens de la théorie des antennes

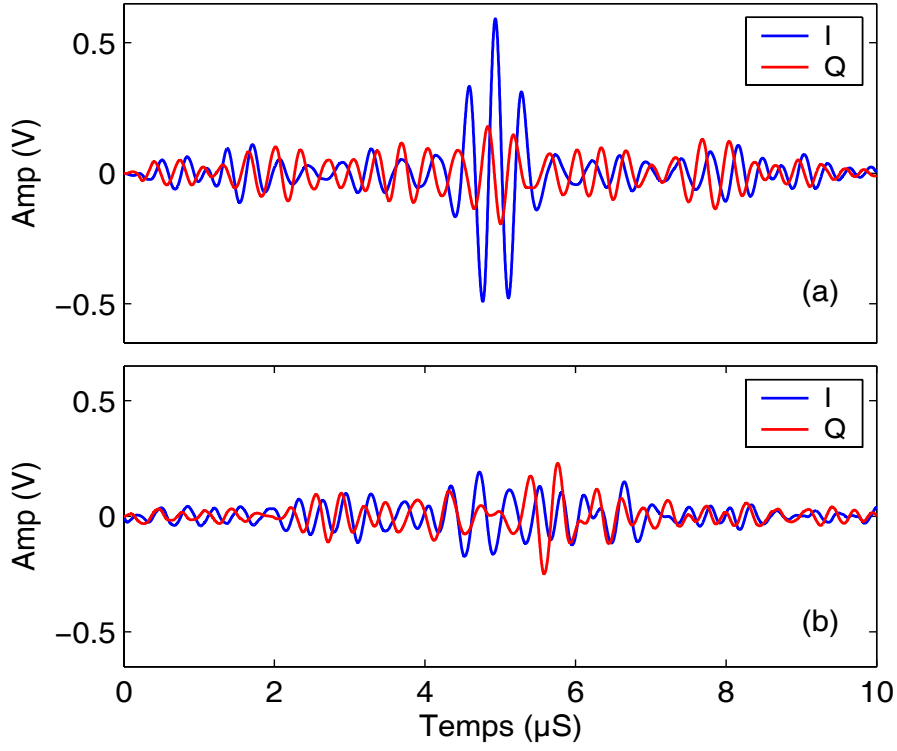


FIG. I.21 – (a) Signaux IQ reçus après retournement temporel en A. (b) signaux IQ après retournement temporel en C.

passante utilisée pour l'impulsion initiale est de 2 MHz. Ceci donne un rapport signal-sur-bruit de $\sqrt{\delta\nu\tau} \approx 3$, résultat qualitativement en accord avec l'amplitude du pic de retournement temporel et le niveau moyen de bruit obtenus expérimentalement.

Enfin, bien que nous ne disposions pas à l'époque de réseau d'antennes nous permettant de mesurer de façon quantitative la focalisation spatiale obtenue par retournement temporel, nous avons tout de même été en mesure de l'estimer. Pour cela, une troisième antenne, C, est également présente dans la cavité. Elle est située à peu près à 2 m (i.e., 15 longueurs d'onde) de l'antenne d'émission A et de l'antenne de réception B. A l'aide du tableau externe, il est alors possible de brancher la carte de réception sur l'antenne C, sans modifier les conditions à l'intérieur de la cavité. Nous avons donc pu mesurer le champ au point C, alors que l'antenne B émettait le champ retourné temporellement correspondant à l'antenne A. Les signaux en bande de base reçus en C sont représentés sur la figure I.21.(b). Nous constatons alors uniquement la présence de lobes secondaires, dont l'amplitude est sensiblement égale à celle des lobes secondaires que l'on avait en dehors de l'impulsion sur l'antenne A. Ceci est une première mise en évidence de la focalisation spatiale que l'on peut réaliser par retournement temporel dans un milieu réverbérant, avec un miroir monovoie.

Dans cette partie, nous avons montré que le retournement temporel en bande de base permet de contourner nombre de problèmes liés aux fréquences élevées des ondes décimétriques électromagnétiques. Nous avons également développé un formalisme impulsionnel en bande de base qui permet d'interpréter les résultats obtenus. En exploitant ce concept, nous avons réalisé le premier prototype de miroir à retournement temporel monovoie à 2.45 GHz, en utilisant des cartes de type WIFI, et un appareillage de mesure et de génération assez rudimentaire. Nous avons ainsi présenté pour la première fois des résultats de compression temporelle en cavité réverbérante. Notre formalisme impulsionnel en bande de base a été vérifié à l'aide des résultats expérimentaux. Puis, nous avons été en mesure de présenter une première preuve de la focalisation spatiale que procure le retournement temporel en électromagnétisme, sans pouvoir toutefois la quantifier. Ces résultats, les premiers sur le sujet à l'époque, ont fait l'objet d'une publication dans la revue *Physical Review Letters* [18] ainsi que d'un dépôt de brevet international [19]. Cependant aucune étude quantitative n'a pu être réalisée à cause de la bande passante très limitée, du caractère monovoie du MRT et de la nécessité de réaliser ces expériences dans une cavité réverbérante. Nous avons donc décidé de réaliser un prototype multivoie, et plus large bande, de manière à pouvoir le manipuler dans des environnements moins réverbérants que nous serions à même de réaliser au laboratoire. Ce montage nous a permis d'effectuer une étude quantitative du retournement temporel d'ondes électromagnétiques, que nous présentons dans la partie suivante.

I.4 Passage à un miroir multivoies large bande passante

I.4.1 Dispositif expérimental

Le dispositif expérimental présenté dans la partie précédente n'a pas permis d'étudier de façon quantitative les propriétés de compression temporelle et de focalisation spatiale associées au retournement temporel. Or, c'est sur ces aspects du retournement temporel que sont fondées nombre d'applications, comme on peut le voir dans le domaine de l'acoustique, que ce soit en télécommunications sous-marines [20], en imagerie médicale [21] ou encore en thérapie médicale [22]. De plus, le fait d'être obligé de se placer dans une chambre réverbérante à facteur de qualité élevé n'était ni pratique, le laboratoire n'en disposant pas, ni réaliste en termes d'applications potentielles. Nous nous sommes donc lancés dans la construction d'un nouveau prototype. Les principales caractéristiques que nous voulions sont répertoriées ci-après.

- Tout d'abord, pour des raisons pratiques, la fréquence porteuse devait rester à 2.45 GHz, afin de pouvoir utiliser le même appareillage hyperfréquences de base, en particulier les antennes.
- Afin de pouvoir réaliser les expériences au laboratoire tout en étudiant l'influence de la bande passante, nous avons décidé de passer à une bande passante supérieure à 100 MHz, celle de notre ancien prototype étant en réalité limitée à 4 MHz.
- Parallèlement, il nous fallait disposer d'un miroir multivoies ainsi que d'un réseau d'antennes en réception afin de pouvoir étudier l'influence du nombre de capteurs dans le MRT, ainsi que la qualité de la focalisation spatiale.
- Enfin, la conception d'un tel prototype par une société spécialisée était trop onéreuse. Nous avons décidé d'acheter des pièces séparées et de réaliser le MRT nous-mêmes.

Le mode impulsionnel n'étant pas très utilisé dans le domaine des télécommunications, nous avons toutefois été limités dans nos choix par le matériel disponible. Le modulateur le plus large bande (250 MHz) que nous avons pu acheter à 2.45 GHz est le RF2480 de la marque RF micro device dont nous détaillerons le fonctionnement par la suite. De plus, aucun démodulateur ne possédait une bande passante aussi large que celle du RF2480, le démodulateur le plus large bande offrant 80 MHz. Nous avons donc du adapter notre protocole expérimental à cette nouvelle spécification : nous avons opté pour une démodulation numérique des signaux radiofréquences. Ceci n'a pas été très gênant dans la mesure où depuis la construction du premier MRT, les oscilloscopes numériques avaient beaucoup évolué, offrant des bandes passantes allant jusqu'à 6 GHz, soit plus que ce dont nous avons besoin. En revanche les générateurs de signaux étaient, et sont toujours, beaucoup plus limités, mais il existait à ce moment un générateur ar-

bitraire de signaux à 2 voies échantillonnées à 1 GHz, soit avec 300 MHz de bande passante. Finalement, le dispositif expérimental que nous avons développé est schématisé sur la figure I.22.

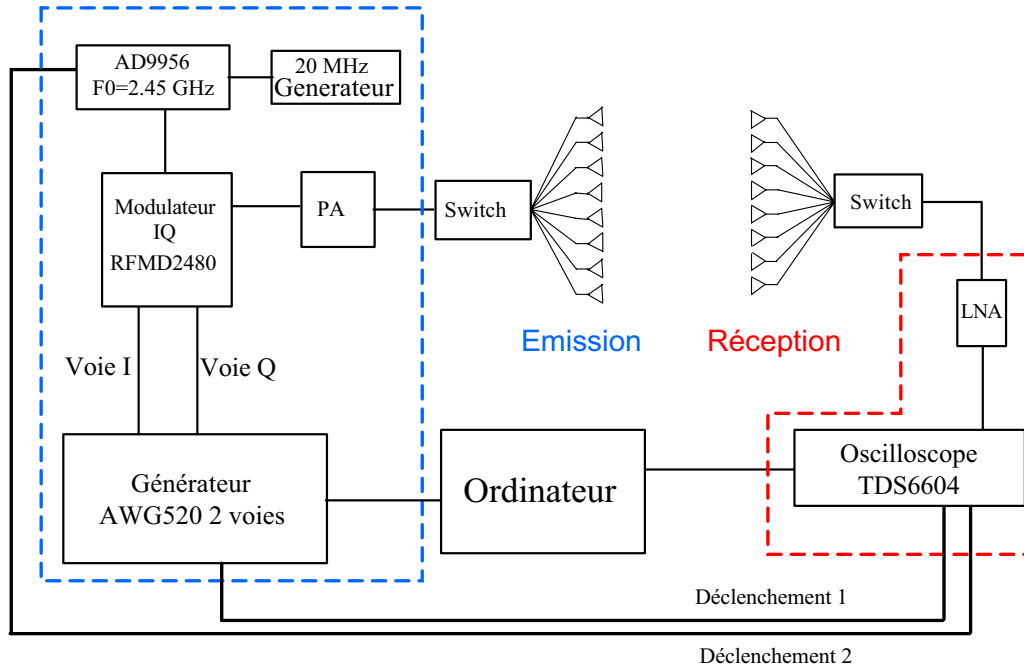


FIG. I.22 – Schéma du dispositif expérimental.

Partie émission :

- Le générateur de signaux arbitraires AWG520 (Tektronix) possède 2 voies synchronisées avec 300 MHz de bande passante, ce sont elles qui vont générer la modulation IQ. Ce générateur est relié à l'ordinateur par le réseau du laboratoire, ce qui rend les transferts d'information très rapides.
- Le modulateur RF2480 a une bande passante IQ de 250 MHz. La modulation se fait sur une porteuse à 2.45 GHz qui est externe. Les signaux IQ doivent rentrer avec une composante continue de 3V (l'alimentation du modulateur). Nous avons donc réalisé 2 circuits basés sur des transistors qui permettent d'ajouter aux signaux IQ qui sortent du générateur une composante à 3V. Le tout est placé dans un boîtier blindé avec un circuit d'alimentation. En sortie, on dispose d'un signal RF modulé avec un taux de rejection de porteuse de l'ordre de 5%.
- La fréquence porteuse est générée par une carte Analog Device (AD9956) à partir d'une référence à 20 MHz qui est fournie par un générateur sinusoïdal externe. Cette carte est munie d'un quartz à 2.45 GHz et d'une boucle à verrouillage de phase ainsi que d'un diviseur

de fréquence. L'horloge produite par le quartz est divisée puis comparée avec la fréquence référence à 20 MHz. Une boucle de rétroaction corrige alors en temps réel le décalage en fréquence afin de produire une horloge très stable qui sert de porteuse au modulateur après filtrage.

- Un amplificateur assure ensuite une puissance de sortie de l'ordre de 20 dBm.
- Un multiplexeur (noté switch sur le schéma) permet de basculer entre les 8 antennes d'un réseau d'émission. En effet, la réalisation d'un vrai miroir à 8 voies n'étant envisageable pour des raisons d'encombrement et de coût, on utilisera la propriété de linéarité des ondes. Chaque antenne du MRT sera ainsi utilisée tour à tour grâce au multiplexeur et le résultat du retournement sauvegardé dans l'ordinateur. Afin de simuler un MRT à 8 voies, il suffit alors de sommer les signaux obtenus avec chacune des antennes. Enfin le multiplexeur est relié à l'ordinateur via une interface USB/digital National Instruments.

Partie réception :

- Nous disposons là encore d'un multiplexeur 8 voies, ce qui va nous permettre de recevoir le champ sur un réseau de 8 antennes, en basculant d'une antenne à l'autre. Lui aussi est relié à l'ordinateur et est pilotable par Matlab.
- Les signaux en réception passent ensuite par un amplificateur faible bruit qui nous assure un gain de l'ordre de 10 dB.
- Comme nous ne disposons pas de démodulateur, l'acquisition des signaux est réalisée avec un oscilloscope numérique TDS6604 (Tektronix) qui échantillonne à 20 GS/s, soit à peu près 8 points par période de porteuse.

Synchronisation du prototype :

Le retournement temporel étant très sensible à la phase, une grande stabilité est nécessaire sur la porteuse du signal. La synchronisation est donc double : l'oscilloscope est d'abord synchronisé par une sortie de la carte AD9956, qui est à la fréquence porteuse divisée par 8, soit à peu près 308 MHz. Ensuite, pour que le générateur émette les signaux IQ également en phase, ce dernier est déclenché par une sortie de synchronisation de l'oscilloscope. De la sorte, le déphasage IQ est toujours constant vis-à-vis de la porteuse, car l'oscilloscope garde lui-même un déphasage constant vis-à-vis de la porteuse lors de ses acquisitions.

Amplification des signaux :

L'amplification à l'émission est réalisée à l'aide d'un amplificateur de puissance RF2126 de marque RF micro device, qui offre une puissance de sortie de 20 dBm, et qui a été placé également en boîtier blindé. La tension de sortie maximale est de l'ordre de 8 Vcc. La figure I.23 quantifie la linéarité de la chaîne constituée du modulateur et de l'amplificateur d'émission. On a mesuré la sortie de l'amplificateur en fonction de la sortie du générateur pour une impulsion de type gaussien et d'une durée de 10 ns. Le résultat n'est pas parfait, ceci étant dû au fait que le matériel utilisé n'est pas destiné à un usage impulsif. Les signaux de faible amplitude seront donc moins amplifiés par le système que les signaux de grande amplitude. Cependant ce problème n'est pas très gênant car le retournement temporel est très sensible à la phase mais beaucoup moins à la dynamique des signaux comme l'a montré Arnaud Derode [23] en réalisant un retournement temporel avec des signaux quantifiés sur 1 bit.

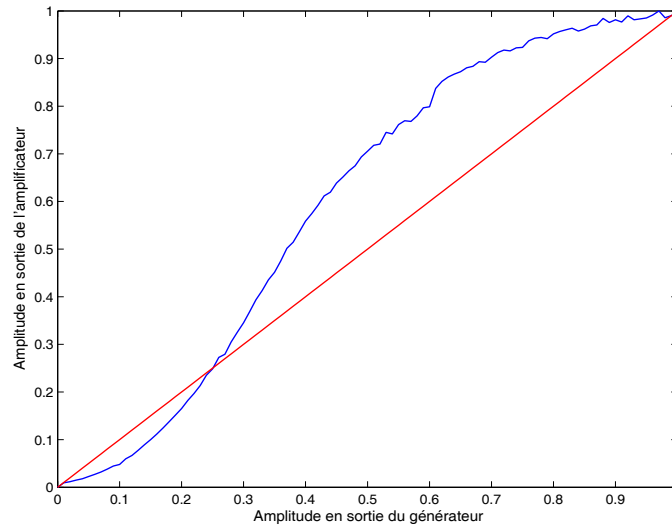


FIG. I.23 – Etude de la linéarité du Miroir à retournement temporel.

Concernant la partie réception, nous disposons d'un amplificateur faible bruit linéaire, dont le gain est de 10 dB, que nous utiliserons ou pas selon les conditions expérimentales. En effet nous allons voir que dans la plupart des cas les expériences ont été réalisées dans une cavité de fabrication "maison" dans laquelle il n'était pas nécessaire d'amplifier les signaux reçus.

I.4.2 Retournement temporel large bande

Afin d'étudier le retournement temporel d'ondes électromagnétiques dans des conditions se rapprochant le plus possible de celles qu'avaient connu les acousticiens, nous avons conçu un environnement réverbérant en partant d'un matériel très courant au laboratoire : une cuve.

L'intérieur de cette cuve d'à peu près 1 m^3 a été recouvert de papier d'aluminium. Elle a été inclinée sur le côté de manière à ce que la partie ouverte soit verticale. Afin de pouvoir fermer cette cavité artisanale, une cloison amovible, elle aussi recouverte d'aluminium, a été fabriquée. La figure I.24 montre l'intérieur de cette cavité dans laquelle on distingue les 2 réseaux de 8 antennes, ainsi que les multiplexeurs auxquels elles sont reliées.

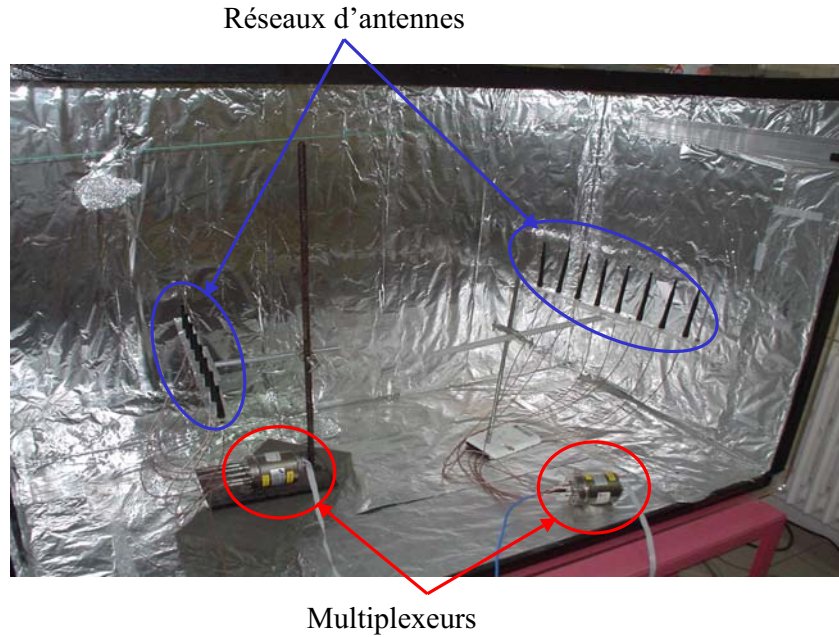


FIG. I.24 – Photo de la cavité fabriquée au laboratoire.

Le protocole opératoire utilisé ici est à peu près le même que celui que l'on respectait avec l'ancien dispositif. Une impulsion est envoyée sur la partie en phase (voie I) du modulateur. L'impulsion d'une durée de 10 ns est de type gaussien, elle est générée numériquement à une fréquence d'échantillonnage de 1 GS/s avec Matlab, puis elle est filtrée numériquement entre 10 MHz et 250 MHz. Puis ce signal est chargé dans la mémoire du générateur de signaux arbitraires et envoyé en voie I. La voie Q, quant à elle, émet un signal de même durée mais constant et égal à zéro. Après modulation et amplification du signal, celui-ci est transmis par une antenne au sein de la cavité. Le champ ondulatoire ainsi généré est réfléchi et diffracté au contact des parois. Une partie du champ est reçue par une autre antenne. Nous disposons de 8 antennes sur notre miroir à retournement temporel et de 8 antennes sur notre réseau de réception, ce qui donne 64 couples (i, j) différents, donc la possibilité d'enregistrer 64 réponses impulsionnelles $h_{i,j}(t)$ différentes. A la réception, le signal est directement acquis à l'oscilloscope numérique qui l'échantillonne à une fréquence de 20 GS/s. La figure I.25 montre un exemple de réponse

impulsionnelle entre deux antennes.

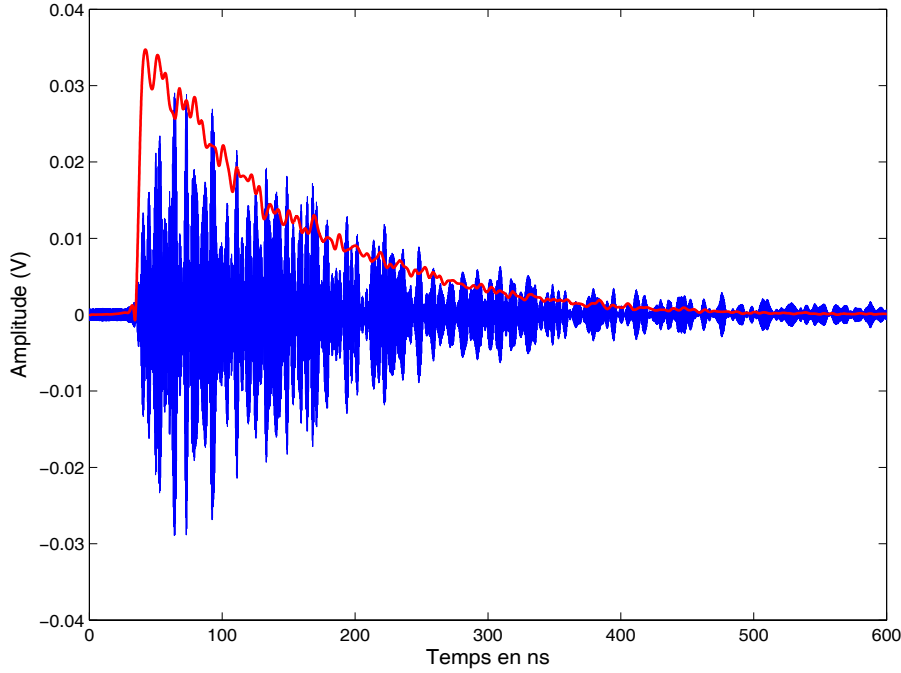


FIG. I.25 – Exemple de réponse impulsionnelle entre 2 antennes (bleu). Enveloppe moyenne des 64 réponses (Rouge).

Nous avons également tracé sur cette figure la moyenne des enveloppes des 64 réponses enregistrées, afin avoir une idée du temps de réverbération à l'intérieur de la cavité. A l'oeil, il est possible d'estimer que le temps pour que l'amplitude décroisse de 2/3 est de l'ordre de 150 ns. Afin d'être plus précis, on peut définir un temps de réverbération τ , au sens de la valeur efficace du signal, de la façon suivante :

$$\tau = \sqrt{\frac{\langle \int t^2 \cdot h_{i,j}^2(t) dt \rangle}{\langle \int h_{i,j}^2(t) dt \rangle}} \quad (I.79)$$

où le signe $\langle \rangle$ représente la moyenne sur tous les couples d'antennes. Le résultat de ce calcul donne $\tau = 160$ ns.

Bien sûr, ce temps de réverbération est très inférieur à celui que l'on obtenait dans la chambre réverbérante du LSS à Supélec qui était de $3.6 \mu s$. Nous avons donc perdu un rapport de l'ordre de 20. Cependant, la bande passante du nouveau MRT étant à peu près 100 fois supérieure à celle de l'ancien MRT, nous avons tout de même un facteur d'étalement proche de 20. Cet étalement ne rivalise pas encore avec les expériences ultrasonores mais s'en rapproche tout de même beaucoup.

Après avoir fait l'acquisition de ce signal, et sauvegardé le fichier sur l'ordinateur de l'oscilloscope, le fichier est copié dans l'ordinateur central qui commande tous les instruments. Une

démodulation numérique est alors réalisée sous Matlab, en utilisant une projection sur la base cosinus/sinus à la fréquence porteuse, puis un filtrage passe bas numérique. La figure I.26 montre le résultat de la démodulation, c'est-à-dire les réponses impulsionnelles en bande de base ainsi obtenues.

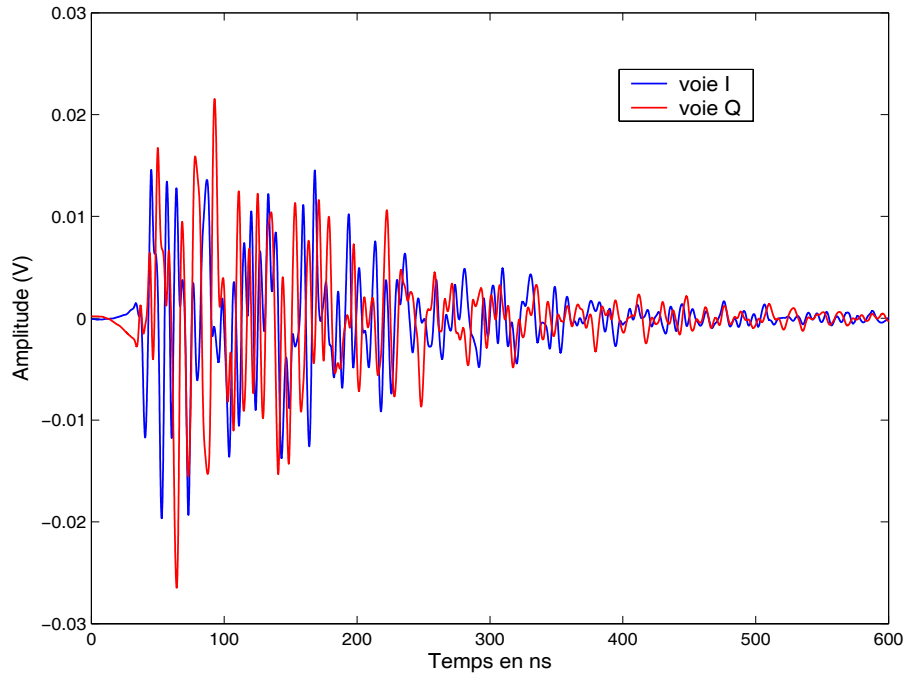


FIG. I.26 – Exemple de réponses impulsionnelles en bande de base.

Afin de pouvoir renvoyer ces réponses retournées temporellement par le biais du générateur arbitraire, il est nécessaire ensuite de les sous-échantillonner pour passer de 20 Giga-points par nanoseconde à 1 Giga-points par nanoseconde. Puis, on suit la procédure habituelle de retournement temporel sur fréquence porteuse, les 2 signaux sont inversés chronologiquement et on conjugue la porteuse. Une fois ces opérations réalisées les signaux sont d'abord normalisés par le maximum en valeur absolu des signaux I et Q, puis ils sont chargés dans les mémoires du générateur arbitraire avant d'être émis, modulés et transmis par l'antenne émettrice de l'impulsion initiale. L'onde retournée temporellement se propage dans la cavité. L'antenne de réception transmet alors la partie du champ qu'elle reçoit à l'oscilloscope. Le signal ainsi acquis est tracé sur la figure I.27 sur laquelle on peut voir que l'impulsion initiale a été compressée. Une nouvelle fois l'origine des temps a été choisie de manière arbitraire de façon à centrer le pic sur la fenêtre.

Nous pouvons tirer plusieurs conclusions de l'analyse de cette courbe. Tout d'abord comme dans

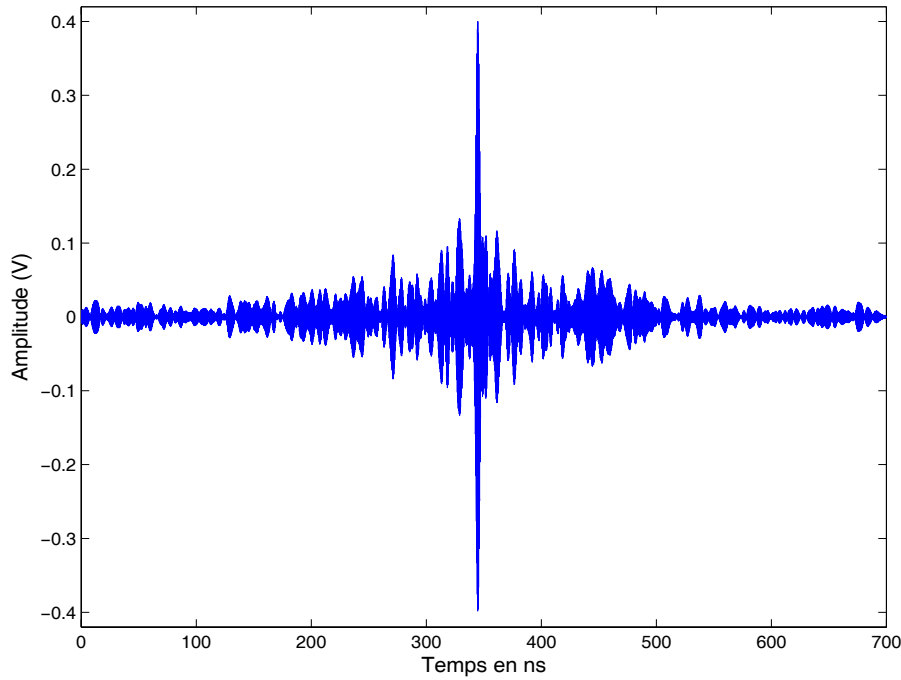


FIG. I.27 – Signal reçu après retournement temporel avec 1 antenne.

le cas des expériences conduites dans la chambre réverbérante de Supélec, nous sommes limités par le temps de réverbération du signal et non pas par le nombre de modes. En effet, le volume de la cavité étant proche de 1 m^3 , le temps de Heisenberg est de $5.6 \text{ } \mu\text{s}$. Le temps d'absorption étant autour de 160 ns , c'est lui qui va limiter la qualité de la compression temporelle. Le facteur de compression théorique est donc de l'ordre de $\sqrt{\delta\nu\tau} \approx \sqrt{200 \cdot 10^6 * 160 \cdot 10^{-9}} \approx 5.5$, or ici on peut l'estimer expérimentalement à environ 5 ce qui est du même ordre de grandeur. Il est également intéressant de remarquer le gain en amplitude que procure la compression temporelle. Lorsque nous avons enregistré la réponse impulsionnelle sur la figure I.25 nous avons mesuré une amplitude crête-crête moyenne de l'ordre de 30 mV . Ici l'amplitude absolue de l'impulsion est de 800 mVcc , l'opération de retournement temporel a engendré un gain de l'ordre de 26 en amplitude. Ceci est dû au fait que la réponse impulsionnelle a été normalisée avant la réémission par le MRT : l'amplitude du pic est donc proportionnelle à $\delta\nu\tau \approx 30$. La dernière remarque que nous pouvons faire concerne la forme de l'impulsion : celle-ci n'est pas symétrique par rapport à $t = 0$ comme on s'y attend. Ici c'est la non linéarité de l'électronique d'amplification employée, représentée sur la figure I.23, qui provoque cette asymétrie.

On peut aussi repasser en bande de base afin d'observer l'impulsion que l'on avait envoyée initialement sur la voie I du modulateur d'émission. C'est ce que l'on représente sur la figure I.28.

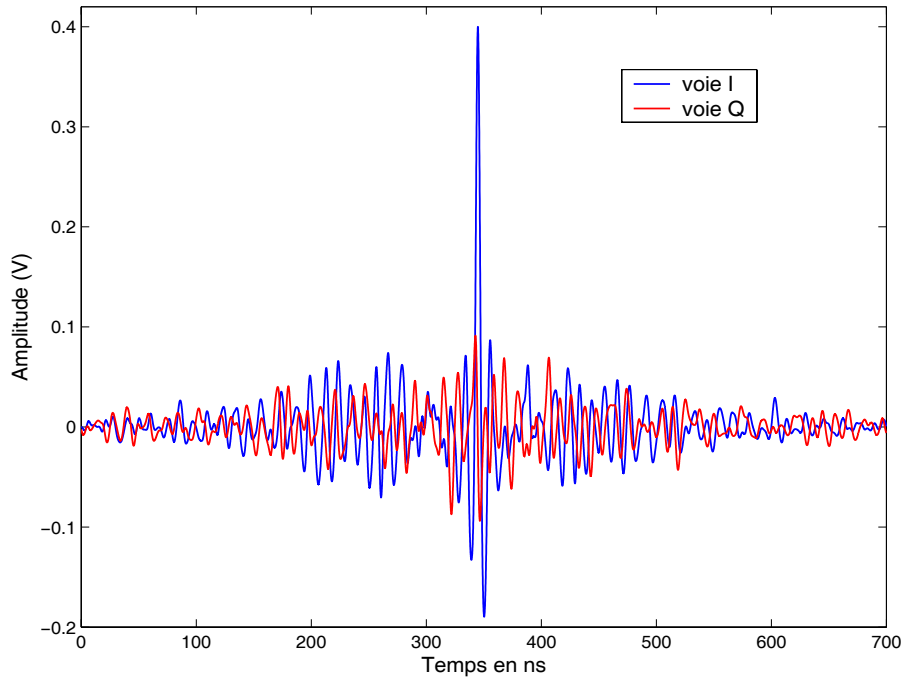


FIG. I.28 – Signal reçu après retournement temporel avec 1 antenne, en bande de base.

On retrouve bien une impulsion d’une durée approximative de 10 ns sur la voie I, et des lobes secondaires autour de celle-ci ainsi que sur la voie Q. Ces lobes secondaires ont eux aussi une amplitude qui peut être reliée à celle du pic par le facteur de compression que l’on a calculé pour le signal sur fréquence porteuse. De plus, on peut à nouveau vérifier que lorsque le pic sur la voie en phase atteint son maximum (ici autour de 350 ns), le signal sur la voie en quadrature est exactement égal à zéro comme nous l’avons démontré précédemment.

Ici encore nous n’avons utilisé qu’un miroir à retournement temporel monovoie, en profitant uniquement de la largeur de bande passante donc nous disposons. C’est la raison pour laquelle la qualité de la compression temporelle peut paraître un peu médiocre. Nous allons voir dans la partie suivante que comme dans le cas de l’acoustique, le RT va être d’autant meilleur que le nombre d’éléments du MRT est élevé et la bande passante utilisée large.

I.4.3 Influence de la bande passante et du nombre d’antennes du miroir sur le retournement temporel

Le retournement temporel réalise donc une compression temporelle du signal réverbéré et permet d’obtenir une impulsion d’une durée équivalente à celle qui avait été émise initialement. Cependant, autour de ce pic de nombreux lobes secondaires sont présents et nous allons voir que leur énergie moyenne est d’autant plus faible relativement à celle du pic que la bande passante est large. De manière à ce que cette étude ne prenne pas un temps trop considérable, nous

avons décidé de réaliser ceci numériquement, à partir de signaux expérimentaux. Afin de pouvoir réaliser des moyennes, nous avons utilisé 8 antennes de réception séparées de $\lambda/2$ ainsi qu'un miroir à retournement temporel composé de 8 antennes. Les 64 couples de réponses impulsionnelles $h_{ij}(t)$ ont été enregistrés et nous avons procédé au retournement temporel de chacun de ces signaux. Après acquisition, une matrice de 64 signaux temporels $S_{i,j}^{RT}(t) = h_{ij}(-t) \otimes h_{ij}(t)$ a donc été acquise correspondant aux compressions temporelles sur chacune des antennes de réception j , réalisées avec les antennes d'émission i . Tous ces signaux expérimentaux ont ensuite été filtrés à l'aide de filtres passe-bande numériques, de façon à limiter la bande passante de 40 MHz à 200 MHz, et ce par pas de 2 MHz. Pour chaque compression temporelle filtrée, on calcule une estimation du rapport signal-sur-bruit¹¹, soit :

$$RSB(\Delta\Omega) = \frac{\langle \text{Max}(BP_{\Delta\Omega}(S_{i,j}^{RT}(t)))^2 \rangle}{\langle \frac{1}{\tau} \int_T^{T+\tau} (BP_{\Delta\Omega}(S_{i,j}^{RT}(t)))^2 dt \rangle} \quad (\text{I.80})$$

où $BP_{\Delta\Omega}$ est un filtre passe-bande de largeur $\Delta\Omega$, $\langle \rangle$ représente une moyenne sur les 64 compressions temporelles, τ est le temps sur lequel on calcule l'énergie moyenne des lobes secondaires et T une origine arbitraire.

Cette grandeur représente le rapport entre le carré de l'amplitude maximale du pic de retournement temporel et l'énergie moyenne des lobes secondaires créés pour une bande passante donnée. Nous soulignons ici que celle-ci quantifie à la fois le niveau des lobes secondaires obtenus en dehors du pic de retournement temporel et le niveau de bruit créé par le retournement temporel en tout point de la cavité. Son évolution en fonction de la bande passante est tracée sur la figure I.29, sur laquelle nous avons également superposé une régression linéaire.

Comme on peut le voir, l'évolution de ce rapport signal-sur-bruit en intensité est linéaire par rapport à la bande passante utilisée, ce qui est bien en accord avec les prédictions d'une évolution de l'amplitude en $\sqrt{\delta\nu\tau}$, où τ est le temps de réverbération du milieu. Par ailleurs, à l'aide d'un ajustement linéaire il est possible d'estimer le temps de réverbération τ qui doit être égal au coefficient directeur de la droite. Ici on a un coefficient directeur $\alpha = 0.19 \mu\text{s}$ ce qui est du même ordre de grandeur que le τ estimé dans la partie précédente (160 ns). La dernière information que nous déduisons de cette courbe concerne le dispositif expérimental lui-même. Nous avons dit précédemment que la bande passante du prototype réalisé est limitée par celle du modulateur à 250 MHz mais ceci est démenti ici : comme on le voit sur la courbe, au delà de 150 MHz la courbe du rapport signal-sur-bruit commence à saturer. La bande passante de notre système

¹¹Les lobes secondaires sont ici qualifiés de bruit, car ils sont indésirables.

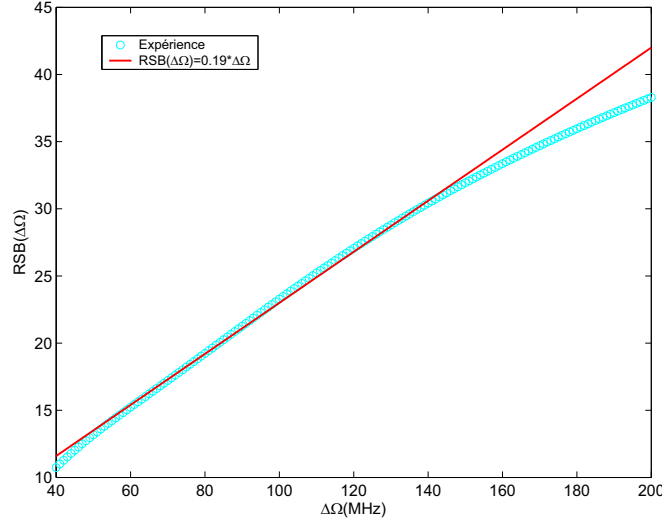


FIG. I.29 – Influence de la bande passante sur le rapport signal-sur-bruit.

doit donc être revue à la baisse. Cela provient du fait que le signal est filtré par tous les éléments du montage dont certains comme les amplificateurs ne sont pas très large bande.

La bande passante d'un MRT a une influence sur le retournement temporel. Mais il est également possible de focaliser spatialement une onde monochromatique avec un miroir à retournement temporel. Le retournement temporel se réduit alors à de la conjugaison de phase et c'est le nombre d'antennes qui définit la qualité de la focalisation. En effet, nous avons souligné lors de notre discussion sur les résultats en acoustique que les lobes secondaires créés par le retournement temporel, qu'ils soient temporels ou spatiaux, croissent en $\sqrt{N}$ quand l'amplitude du pic croît elle en N , ce qui donne un gain en rapport signal-sur-bruit en amplitude de $\sqrt{N}$. Afin de vérifier qu'il en est de même dans nos expériences de retournement temporel d'ondes électromagnétiques, nous avons réalisé 3 séries d'expériences. Dans ces 3 expériences, le MRT est constitué d'antennes commerciales $\lambda/2$ typiques et seule la distance entre celles-ci est modifiée, passant de $\lambda/8$ à $\lambda/2$ puis λ . Comme précédemment, les 64 réponses impulsionnelles correspondant aux couples d'antennes émettrices/réceptrices sont enregistrées tour à tour. Puis on réalise le retournement temporel de ces signaux ainsi que l'acquisition des champs retournés temporellement à l'aide du réseau de réception. Une nouvelle fois, on enregistre les 64 compressions temporelles $S_{i,j}^{RT}(t)$ ainsi créées, et ce pour chacun des MRTs étudiés. Puis nous calculons l'énergie moyenne des lobes secondaires temporels générés par l'expérience ainsi que le carré de l'amplitude maximale du pic de retournement temporel, et ce en fonction du nombre d'antennes dont est composé le MRT, pour chaque espacement inter-antennes choisi. De même, nous allons mesurer l'amplitude maximale obtenue avec chaque MRT pour un nombre d'éléments dans le

miroir variant de 1 à N afin d'estimer le rapport :

$$RSB(N) = \frac{\langle \text{Max}(\sum_{i=1}^N (S_{i,j}^{RT}(t)))^2 \rangle}{\langle \frac{1}{\tau} \int_T^{T+\tau} (\sum_{i=1}^N (S_{i,j}^{RT}(t)))^2 dt \rangle} \quad (\text{I.81})$$

où cette fois $\langle \rangle$ est une moyenne sur les antennes de réception, ce qui permet une meilleure estimation du rapport signal-sur-bruit.

Le rapport signal-sur-bruit obtenu pour un nombre d'antenne N est alors normalisé par le rapport signal-sur-bruit du miroir monovoie, afin d'éliminer le gain qui est du à la bande passante, pour se concentrer sur le gain d'antennes. La figure I.30 représente ce dernier ratio pour les 3 MRT différents que l'on a étudiés.

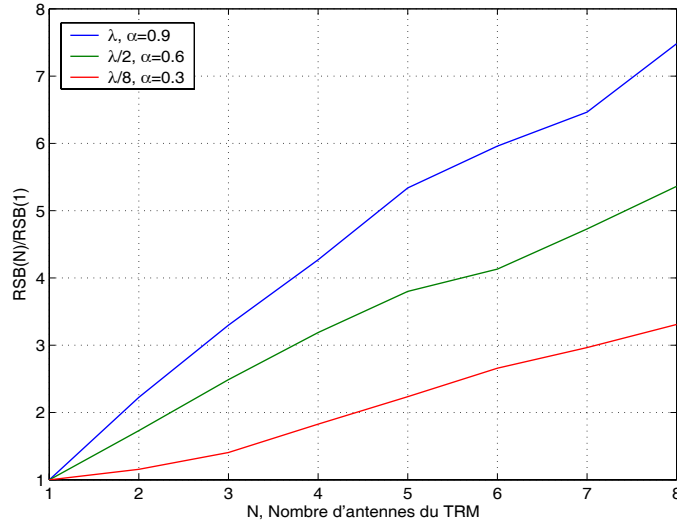


FIG. I.30 – Influence du nombre d'antennes sur le rapport signal-sur-bruit.

Comme on s'y attendait, la dépendance du rapport signal-sur-bruit en fonction du nombre d'antennes est à nouveau linéaire pour les 3 MRT étudiés. Sur la légende de la courbe sont également montrés les coefficients directeurs des différentes droites obtenues, selon la distance inter-antennes choisie. Les valeurs de ces coefficients directeurs croissent en fonction de la distance entre les antennes : de fait, plus les antennes sont proches, plus les champs qu'elles enregistrent sont spatialement corrélés. En effet, si les réponses impulsionnelles étaient parfaitement indépendantes, la pente mesurée devrait être unitaire. Remarquons enfin que pour un espacement de $\lambda/2$, valeur typiquement utilisée en acoustique ultrasonore, nous observons une pente inférieure à 1. Ceci montre que la corrélation spatiale des champs à une telle distance n'est pas nulle : nous allons vérifier ceci dans la partie suivante en estimant la focalisation spa-

tiale obtenue par retournement temporel, qui est aussi un estimateur de la corrélation spatiale d'un champ ondulatoire au sein d'un milieu [12].

I.4.4 Mise en évidence de la focalisation spatiale

Les expériences ultrasonores en cavité ont permis de montrer une focalisation d'un champ acoustique sur une tache dont la taille est de l'ordre de la demi-longueur d'onde. De telles mesures en acoustique ont été réalisées soit :

- en déplaçant derrière une forêt de tiges un capteur mono-élément autour du point où il avait émis l'impulsion initiale.
- en scannant à l'aide d'un interféromètre hétérodyne la surface d'une cavité de silicium plane, autour du point sur lequel on avait émis l'impulsion initiale.
- en focalisant derrière un milieu multidiffuseur par retournement temporel sur une voie centrale d'une sonde échographique, afin d'observer la focalisation spatiale sur les voies adjacentes.

Chacune de ces techniques est bien adaptée à la mesure d'une tache focale : dans tous les cas la mesure est non-perturbative en ce sens qu'elle ne modifie pas le champ acoustique créé. En effet le capteur mono-élément est très fin, le laser totalement passif et la sonde échographique statique. En hyperfréquences, le problème est un peu plus complexe car une antenne est reliée à un câble qui transfère le signal mesuré au système d'acquisition. Dans la cavité réverbérante, la mesure ne peut pas être réalisée en déplaçant cette antenne car le mouvement de celle-ci et du câble va modifier les conditions aux limites dans la cavité, et donc le champ électromagnétique au sein de la cavité. La mesure d'une tache focale ne va donc pouvoir se faire qu'au moyen d'un réseau d'antennes fixes, analogue à une sonde échographique. Le fait que le blindage des câbles diffracte les ondes électromagnétiques pose un autre problème : cela modifie le champ créé par retournement temporel. Enfin, chaque antenne d'un réseau étant un diffuseur pour les antennes adjacentes, la mesure va fortement dépendre des antennes utilisées. Ce phénomène sera considéré plus tard au travers du couplage ; nous nous servirons alors des résultats obtenus grâce au formalisme des impédances mutuelle entre antennes. Afin d'estimer au mieux la focalisation spatiale obtenue par retournement temporel, nous avons voulu fabriquer un réseau d'antennes dont l'effet sur le champ soit minimal et contrôlable en suivant les recommandations suivantes :

- Tout d'abord, il est nécessaire de réaliser des antennes sur un plan de masse afin de limiter l'influence des câbles. Les antennes seront donc fixées à leurs bases sur une plaque de cuivre,

les câbles étant de l'autre côté de la plaque qui agit ainsi comme un écran.

- Afin de minimiser ensuite le couplage entre antennes, nous avons décidé de travailler avec des antennes physiquement petites.
- Enfin, de façon à pouvoir aisément tenir compte du couplage entre antennes, nous avons utilisé des antennes filaires, dont les matrices d'impédances mutuelles se calculent assez facilement numériquement [16].

Tenant compte de ces obligations, nous avons donc opté pour des connecteurs filaires qui possèdent une base à la masse et dont l'âme est protégée par une couche de plastique. Cet élément rayonnant a une longueur de 2.1 cm et un rayon de 1 mm. La couche de plastique sert à éviter que celui-ci soit en contact avec le plan de masse. Une plaque de cuivre a ensuite été trouée de manière à pouvoir glisser les antennes pour souder leur base sur le côté métallique de celle-ci et faire dépasser leur âme du côté qui est recouvert par un isolant. Enfin, nous avons choisit de fixer l'espacement entre antennes à $\lambda/8$ afin de mesurer la tache focale sur une longueur d'onde. La figure I.31 montre une photo du réseau d'antennes ainsi constitué.

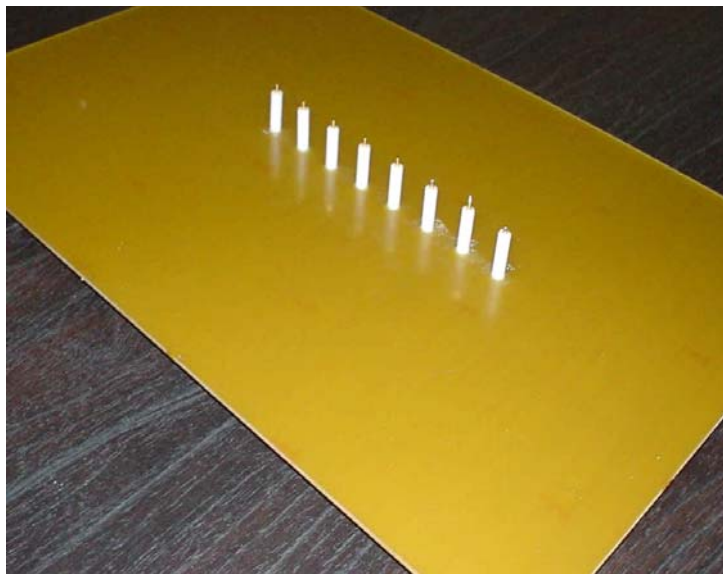


FIG. I.31 – Photo du réseau d'antennes utilisé lors des mesures de tache focale.

Les antennes de ce réseau sont donc des monopôles sur un plan de masse. Celui-ci étant d'une dimension très supérieure à la hauteur des antennes et plus grande que la longueur d'onde, le comportement des antennes peut être assimilé à celui de dipôles, l'image du monopôle étant créée par la plaque de cuivre.

Afin de mesurer une tache focale engendrée par retournement temporel nous suivons la procédure suivante. On enregistre d'abord les réponses impulsionnelles entre une antenne du bord du

réseau, à qui l'on attribue le numéro 1, et les 8 éléments du MRT. Puis on effectue la phase de réémission du retournement temporel avec ces 8 réponses, en se servant des multiplexeurs. On enregistre alors sur les 8 antennes de réception le signal mesuré pour chacune des antennes d'émission ce qui nous donne 64 signaux qui peuvent s'écrire :

$$S_{i,j}^1(t) = h_{i,j}(t) \otimes h_{i,1}(-t) \otimes E(-t) \quad (\text{I.82})$$

où i est un élément du réseau d'émission, j un élément du réseau de réception et $E(t)$ l'impulsion émise initialement. Afin d'avoir accès à une mesure plus précise, nous moyennons ensuite sur les éléments du MRT de manière à n'avoir plus qu'un signal moyen sur chaque antenne de réception, c'est-à-dire que nous simulons un MRT à 8 voies. Le maximum de ces signaux va nous donner accès à la tache focale qui est alors définie par la relation suivante :

$$F(j) = \text{Max} \left\{ \left| \sum_{i=1}^8 S_{i,j}^1(t) \right| \right\}_t \quad (\text{I.83})$$

Nous avons vérifié que la courbe obtenue est symétrique selon que l'on veuille focaliser sur l'antenne du bord gauche ou sur l'antenne du bord droit. Puis nous avons symétrisé la courbe de manière à obtenir une tache focale sur l'intervalle $[-\frac{7}{8}\lambda, \frac{7}{8}\lambda]$. Celle-ci est représentée en bleu sur la figure I.32.

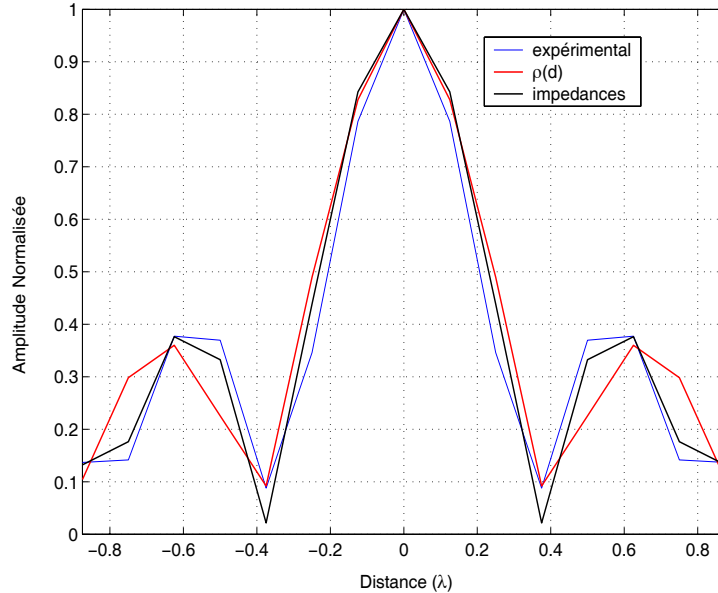


FIG. I.32 – Tache focale obtenue expérimentalement et ajustements théoriques.

Comme l'ont démontré Arnaud Derode et ses collaborateurs, le retournement temporel est un estimateur de la corrélation spatiale du champ ondulatoire dans un milieu donné [12]. Nos

antennes étant polarisées verticalement, notre mesure est directement reliée à la composante E_z du champ électromagnétique. Nous avons donc tracé la courbe théorique donnant la corrélation spatiale $\rho(d)$ de E_z grace à une formule analytique qui a été calculée dans le cas des cavités réverbérantes [24] :

$$\rho(d) = \frac{\langle E_z(0).E_z(d) \rangle}{\langle E_z^2(0) \rangle} = \frac{3}{2} \frac{\sin(kd)}{kd} \left(1 - \frac{1}{(kd)^2} \right) + \frac{3}{2} \frac{\cos(kd)}{(kd)^2} \quad (\text{I.84})$$

avec k le nombre d'onde et d la distance entre deux points.

Nous constatons que la mesure expérimentale est bien en accord avec cette formule, en particulier en ce qui concerne la position du premier minimum du signal. Cependant nous remarquons que les lobes secondaires sont déformés par rapport à la courbe théorique : ceci est dû au couplage entre les antennes. Nous avons également tracé sur cette figure la courbe obtenue en utilisant la formule I.61 qui permet de prédire les courants $\mathbf{\Pi}^*$ sur les antennes de réception lors d'une opération de retournement temporel. Pour ce faire, nous avons mesuré les dimensions exactes de nos antennes et calculé les impédances mutuelles de notre réseau afin de calculer $\mathbf{\Pi}^*$. Une nouvelle fois, la courbe se superpose très bien et de plus les lobes secondaires se superposent mieux aux lobes mesurés expérimentalement. Ceci tend à valider notre approche théorique concernant le retournement temporel électromagnétique appliqué aux réseaux d'antennes.

Concernant la forme de la courbe elle-même il faut faire plusieurs remarques. Tout d'abord, on observe que la distance du premier zéro n'est pas $\lambda/2$ comme dans le cas des expériences ultrasonores mais 0.4λ . De plus, il faut noter que cette valeur est dépendante du type d'antennes utilisées, à cause du couplage inter-éléments. Ceci est intéressant : afin de fabriquer des miroirs à retournement temporel efficaces il va falloir adapter la distance entre les antennes aux antennes elles-mêmes. Dans le cas d'antennes simples, le formalisme développé à partir des impédances mutuelles pourra donner une valeur d'espacement entre antennes à choisir. Pour des antennes plus compliquées on pourra mesurer une tache focale sur le réseau que l'on veut caractériser à l'aide d'un MRT calibré et la courbe expérimentale nous donnera une valeur de distance inter-antennes afin d'optimiser le réseau à calibrer.

Enfin, cette dernière mesure permet d'expliquer la courbe I.30 qui montrait l'évolution du rapport signal-sur-bruit du retournement temporel en énergie en fonction du nombre d'éléments du MRT. En effet nous y avons vu que pour une séparation des antennes d'émission de $\lambda/2$, qui est la distance usuelle en ultrasons, le coefficient directeur n'est pas égal à 1. Ceci s'explique ici car le minimum de la corrélation entre antennes n'est pas obtenu pour une telle séparation

mais pour 0.4λ . En fait, les coefficient directeurs des droites obtenues sont directement reliés a la tache focale : plus les signaux entre 2 antennes sont décorrélés, plus ce coefficient va être élevé (sans pour autant dépasser 1). Ce degré de corrélation entre les signaux va être donné par la valeur de la tache focale à une distance donnée. Ainsi, à $\lambda/8$, on a une corrélation spatiale des signaux de l'ordre de 0.7, ce qui explique le faible coefficient directeur de la droite obtenue sur la figure I.30 qui est de 0.3. De même, à $\lambda/2$, on a $\alpha = 0.6$ quand la tache focale vaut 0.4 et à λ , $\alpha = 0.9$ pour une valeur de corrélation spatiale de l'ordre de 0.1.

Dans cette partie, nous avons montré que le retournement temporel d'ondes électromagnétiques se comporte comme son analogue acoustique. A cette fin, nous avons construit un prototype de miroir à retournement temporel à large bande passante et multivoies. L'étude de l'influence de la bande passante nous a permis de vérifier que le rapport entre l'amplitude du pic de retournement temporel et l'amplitude moyenne des lobes secondaires suit bien une évolution en $\sqrt{\delta\nu\tau}$, avec $\delta\nu$ la bande passante du système et τ le temps de réverbération. Le caractère multivoies du prototype nous a également permis d'étudier l'influence du nombre d'éléments présents dans le MRT. Nous avons ainsi vérifié que l'ajout d'antennes permet d'obtenir un gain de rapport pic/lobes en intensité qui est linéaire en fonction du nombre d'antennes du miroir à retournement temporel. Le coefficient directeur de cette droite, quant à lui, dépend de la distance entre les antennes du MRT. Enfin, en utilisant un réseau d'antennes en réception, nous avons procédé à la mesure d'une tache de focalisation obtenue par retournement temporel. Nous avons vérifié que sa dimension est de l'ordre de la demi-longueur d'onde est qu'elle est dépendante de la distance inter-éléments du réseau de réception et des antennes utilisées au travers du couplage inter-antennes. Ces résultats ont fait l'objet d'une publication dans la revue Applied Physics Letters [25]. Il n'a cependant pas été facile d'effectuer ces expériences, notamment parce que les capteurs influent sur la mesure en électromagnétisme, les câbles et les antennes utilisés étant eux même des diffuseurs. Nous avons donc dû nous affranchir de ces problèmes en fabriquant nous même un réseau d'antennes idéal. Nous allons voir dans la partie suivante comment certaines expériences de retournement temporel réalisées nous ont conduit à considérer le problème de la diffusion des champs électromagnétiques dans le champ proche d'une antenne. Nous montrerons qu'il est possible d'obtenir des résultats surprenants, en terme de focalisation spatiale, grâce à l'utilisation d'antennes très particulières.

I.5 Et les ondes évanescentes ? Une première expérience de focalisation sub-longueur d'onde

I.5.1 Des mesures étonnantes

Comme nous l'avons noté dans la partie précédente, la mesure d'une tache focale "propre" n'a pas été évidente en électromagnétisme. En effet, scanner le champ retourné temporellement en déplaçant l'antenne qui a servi lors de la phase d'enregistrement ainsi que son câble modifie les conditions aux limites et donc le champ à mesurer. Initialement, nous avons conçu un réseau composé de 8 antennes achetées dans le commerce afin de procéder à une mesure spatiale du champ généré par retournement temporel, de manière à ne pas avoir à changer les conditions aux limites après l'acquisition des réponses impulsionnelles. Ces antennes, des monopôles de fort diamètre mais de hauteur faible, étaient vendues pour être utilisées directement comme élément rayonnant lors de transmissions de type WIFI. Nous avons ensuite utilisé une plaque de plexiglass que nous avons percée 8 fois en espaçant les trous de $\lambda/8$. Les antennes y ont alors été placées parallèlement les unes aux autres. Le miroir à retournement temporel quant à lui est constitué d'antennes commerciales de type dipolaire de hauteur $\lambda/2$ qui sont séparées d'une distance de $\lambda/2$. Comme précédemment, la mesure d'une tache focale se déroule de la façon suivante :

- une impulsion d'une durée de 10 ns est générée par le générateur de signal sur la voie "en phase" du modulateur d'émission et aucun signal n'est envoyé sur la voie "en quadrature".
- Le signal modulé et amplifié est émis successivement par les 8 antennes du miroir à retournement temporel.
- Les réponses impulsionnelles entre ces antennes et l'antenne sur laquelle on cherche à focaliser, notée k sont alors enregistrées. Ces signaux sont ensuite démodulés numériquement afin de les retourner temporellement et de les émettre tour à tour par chacune des antennes (i) du MRT.
- Pour chacun des champs retournés temporellement que l'on crée, on mesure alors sur chaque antenne du réseau de réception (j) à $\lambda/8$ les 8 signaux temporels :

$$S_{i,j}^k(t) = h_{i,j}(t) \otimes h_{i,k}(-t) \otimes E(-t) \quad (\text{I.85})$$

- On mesure alors la tache focale obtenue comme lors des mesures avec le réseau idéal :

$$F(j) = \text{Max} \left\{ \left| \sum_{i=1}^8 S_{i,j}^k(t) \right| \right\}_t \quad (\text{I.86})$$

Le résultat de cette mesure lorsque l'on veut focaliser sur la première antenne du réseau est représenté sur la courbe bleue de la figure I.33. Nous avons symétrisé cette courbe après avoir vérifié que les champs scannés étaient identiques, que l'on cherche à focaliser sur l'antenne 1 ou sur l'antenne 8.

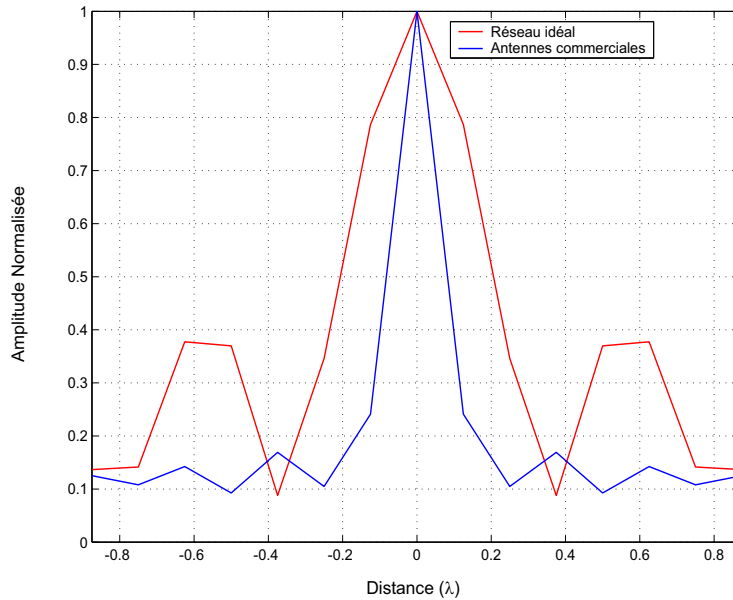


FIG. I.33 – Tache focale obtenue expérimentalement avec le réseau idéal et les antennes commerciales.

A titre comparatif, nous représentons également sur cette figure la tache de focalisation que nous avons obtenue avec notre réseau idéal, en rouge. Il est clair que dans le cas du réseau constitué d'antennes commerciales la largeur à mi-hauteur du pic est bien inférieure : on peut l'estimer à 0.16λ . Cependant, on ne peut pas savoir si cette valeur est plus faible en réalité car l'espacement utilisé pour les antennes est du même ordre de grandeur que cette largeur. Nous avons alors réalisé plusieurs séries d'expériences en faisant varier la géométrie du réseau, le nombre d'antennes de ce réseau et l'espacement entre les antennes. A chaque expérience réalisée avec ces antennes commerciales, le résultat fut le même : le signal semble se focaliser sur une seule antenne et ce quelle que soit la configuration choisie.

Afin de comprendre les phénomènes physiques qui gouvernent la focalisation par retournement temporel sur ce réseau d'antennes, celui-ci étant impossible à simuler simplement à cause de la complexité des antennes commerciales, nous avons dû envisager une hypothèse : les antennes ne sont pas faites pour être utilisées sans plan de masse. En effet, étant donné que ce sont des monopôles, il paraît logique qu'elles aient besoin d'un plan de masse afin que celui-ci crée leur "image" pour donner un dipôle. La logique aurait alors voulu que sans plan de masse

ces antennes soient inefficaces. Nous en avons conclu que sans plan de masse ces antennes “cherchent” un objet métallique qui peut alors générer l’image du monopôle afin de respecter la neutralité de la charge. Si cet objet métallique est constitué par le blindage du câble qui relie les antennes au multiplexeur, il est possible que la mesure des réponses impulsionnelles soit biaisée. En effet, les antennes ainsi formées sont bien plus grandes que la longueur d’onde utilisée : les câbles ont une longueur de 50 cm, soit 4 longueurs d’ondes. Comme la mesure résulte d’une intégration du champ électromagnétique sur l’extension spatiale de l’antenne, cela peut avoir pour effet de “décorrélér” les signaux. Nous avons vérifié cette hypothèse en réalisant la même expérience, mais cette fois en plaçant le réseau de 8 antennes sur un plan de masse de manière à créer l’image des monopôles. Sur la figure I.34 nous représentons le résultat de cette mesure qui est superposé à la courbe obtenue sans plan de masse.

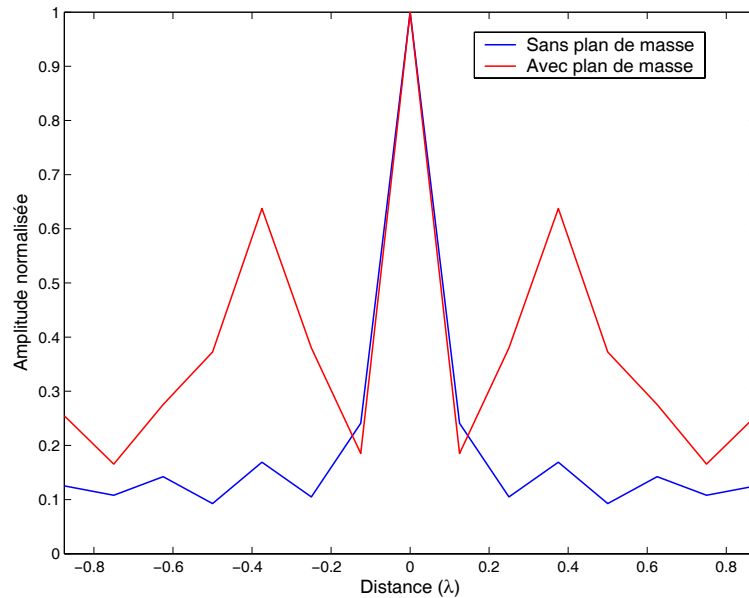


FIG. I.34 – Tache focale obtenue expérimentalement pour les antennes commerciales avec ou sans plan de masse.

Sur cette mesure, le résultat obtenu est beaucoup moins probant que celui de la figure précédente. En effet, 3 antennes sur les 8 reçoivent un signal dont l’amplitude est supérieure à 0.5. Par contre, cette fois-ci, le résultat n’est pas le même quelle que soit l’antenne sur laquelle on focalise sur le réseau de réception. On observe à chaque fois un maximum sur l’antenne choisie mais les amplitudes sur les autres antennes varient assez brutalement selon la configuration étudiée. On ne peut donc pas tirer de conclusion définitive concernant l’effet des câbles dans l’expérience réalisée sans plan de masse. Nous pouvons cependant faire 2 remarques ; la première est que le fait d’ajouter un plan de masse à pour effet de coupler les antennes, ce qui diminue l’amplitude

re ue sur les antennes du centre du r seau au profit de celles du bord (“effet de r seau”). La deuxi me observation est que m me si la tache de diffraction n’est pas aussi fine que lors de la mesure sans plan de masse, elle pr sente des variations dont la taille reste inf rieure   la longueur d’onde des signaux utilis s. Nous allons voir dans la partie suivante que ceci a une importance sur ce que nous pouvons conclure de nos exp riences.

I.5.2 Champ proche, champ lointain et information spatiale

Afin de comprendre le lien entre l’information spatiale port e par une onde et la distance   laquelle on l’observe   partir de sa source, nous allons adopter un formalisme assez usuel dans le domaine de l’optique [26] qui est fond  sur la d composition en ondes planes d’un champ  lectromagn tique. Une partie de cette d monstration provient de la r f rence [27]. Nous allons d buter notre analyse sur un probl me scalaire   2 dimensions (Ox, Oy) . Les r sultats seront ensuite g n ralis s au cas des champs 3D vectoriels. Nous supposons donc un champ  lectromagn tique E_z confin  lat ralement qui se propage selon l’axe Oy et dont la valeur en $y = 0$ est connue pour tout x . Nous consid rons un champ monochromatique d’amplitude complexe $E_z(x, y)$ et de pulsation ω dont nous omettrons la d pendance $\exp(i\omega t)$ pour des raisons de clart  lors de la d monstration. Ce champ v rifie alors l’ quation de propagation des ondes  lectromagn tiques en r gime monochromatique, ou  quation de Helmholtz :

$$(\partial_x^2 + \partial_y^2)E_z(x, y) + k_z^2 E_z(x, y) = 0 \quad (\text{I.87})$$

o  $k_z = 2\pi/\lambda = \omega/c$, λ  tant la longueur d’onde du champ dans le milieu et c la c l rit  des ondes.

Ce champ respecte les conditions initiales suivantes : il est connu sur la droite $y = 0$ o  il vaut $E_z(x, y = 0)$ et il se propage selon la direction des y croissants. Nous allons ensuite  crire la solution de cette  quation sous la forme d’un d veloppement en ondes planes ou d veloppement en spectre angulaire. Pour ce faire, on introduit tout d’abord la transform e de Fourier du champ consid r  par rapport   la variable x , ce qui s’ crit :

$$E_z(x, y) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \tilde{E}_z(\alpha, y) \exp(i\alpha x) d\alpha \quad (\text{I.88})$$

L’ quation de Helmholtz devient avec cette nouvelle expression du champ  lectrique :

$$\frac{\partial^2 \tilde{E}_z(\alpha, y)}{\partial y^2} + \gamma^2 \tilde{E}_z(\alpha, y) = 0 \quad (\text{I.89})$$

où nous avons introduit la variable γ telle que $\gamma^2 = k^2 - \alpha^2$.

La solution de ce champ se met sous la forme suivante :

$$\tilde{E}_z(\alpha, y) = A(\alpha) \exp(i\gamma y) + B(\alpha) \exp(-i\gamma y) \quad (\text{I.90})$$

pour laquelle la variable γ dépend de α comme :

$$\begin{aligned} \gamma &= \sqrt{k^2 - \alpha^2} \text{ si } |\alpha| \leq k \\ \gamma &= i\sqrt{\alpha^2 - k^2} \text{ si } |\alpha| > k \end{aligned} \quad (\text{I.91})$$

Par la suite, on peut utiliser les conditions aux limites de notre problème pour en déduire que $B(\alpha)$ est nul, car la propagation se fait dans le sens $y > 0$, et que $A(\alpha) = \tilde{E}_z(\alpha, y = 0)$. Cette dernière égalité permet d'obtenir le champ $E_z(x, y > 0)$ en fonction de la condition en $y = 0$, ce qui s'écrit :

$$E_z(x, y > 0) = \frac{1}{2\pi} \int_{-\infty}^{\infty} \tilde{E}_z(\alpha, y = 0) \exp(i\alpha x + i\gamma(\alpha)y) d\alpha \quad (\text{I.92})$$

Il apparaît ici que la composante selon y du vecteur d'onde que nous avons noté γ dépend de α , notamment au travers de la différence $k^2 - \alpha^2$.

Ce résultat démontré avec une onde scalaire à 2 dimensions se généralise au cas des champs vectoriels à 3 dimensions. Un tel champ est solution de l'équation de Helmholtz vectorielle :

$$\Delta \mathbf{E} + k^2 \mathbf{E} = \mathbf{0} \quad (\text{I.93})$$

et pour tout $y > 0$ on peut écrire que le champ vérifie :

$$\mathbf{E}(\mathbf{r}) = \frac{1}{4\pi^2} \int_{-\infty}^{\infty} \tilde{\mathbf{E}}(\mathbf{K}, y = 0) \exp(i\mathbf{K} \cdot \mathbf{R} + i\gamma(\mathbf{K})y) d^2\mathbf{K} \quad (\text{I.94})$$

avec $\mathbf{r}$ un point de l'espace, $\mathbf{R}$ un point du plan $y = 0$ et toujours les conditions sur γ qui sont données par :

$$\begin{aligned} \gamma(\mathbf{K}) &= \sqrt{k^2 - |\mathbf{K}|^2} \text{ si } |\mathbf{K}| \leq k \\ \gamma(\mathbf{K}) &= i\sqrt{|\mathbf{K}|^2 - k^2} \text{ si } |\mathbf{K}| > k \end{aligned} \quad (\text{I.95})$$

On peut alors expliciter l'équation I.94 de la façon suivante : le champ est exprimé comme une superposition linéaire d'ondes planes de vecteurs d'ondes $\mathbf{k} = \{\mathbf{K}, \gamma\}$ dont l'amplitude

représentée par $\tilde{\mathbf{E}}(\mathbf{K}, y = 0)$ est la transformée de Fourier du champ dans le plan $y = 0$. La variable $\mathbf{K}$ apparaît alors comme la pulsation spatiale associée aux variations du champ considéré sur le plan $y = 0$. On peut alors distinguer deux types d'ondes dans cette décomposition :

- lorsque la fréquence spatiale vérifie $|\mathbf{K}| \leq k$, γ est réel. L'argument de l'exponentielle dans l'équation I.94 est alors imaginaire pur et l'onde associée est dite propagative.
- Si la fréquence spatiale est plus haute (pour $|\mathbf{K}| > k$), la variable γ devient imaginaire pur et l'onde décroît de façon exponentielle en fonction de la distance parcourue sur l'axe Oy . Ce type d'onde est dit évanescent car son amplitude est négligeable après une distance de propagation de l'ordre de la longueur d'onde.

Supposons maintenant que le champ en $y = 0$ résulte d'une source ponctuelle. Cette source présente des variations spatiales couvrant tout l'espace des $\mathbf{K}$, mais seulement une partie de ces fréquences spatiales vont “survivre” à la propagation dans le milieu homogène, celles qui sont dites propagatives. La propagation se comporte comme un filtre passe-bas en fréquences spatiales : les plus petits détails du champ ondulatoire, qui sont représentés par les grandes fréquences spatiales, vont donc être perdus. Cette propriété est à l'origine de la limite de diffraction en imagerie conventionnelle.

Dans la référence [26], Joseph Goodman va plus loin au travers d'une analogie avec le traitement du signal en écrivant de 2 façons le champ $\mathbf{E}(\mathbf{r})$ en fonction de sa transformée de Fourier :

$$\begin{cases} \mathbf{E}(\mathbf{r}) &= \frac{1}{4\pi^2} \int_{-\infty}^{\infty} \tilde{\mathbf{E}}(\mathbf{K}, y = 0) \exp(i\mathbf{K} \cdot \mathbf{R} + i\gamma(\mathbf{K})y) d^2\mathbf{K} \\ \mathbf{E}(\mathbf{r}) &= \frac{1}{4\pi^2} \int_{-\infty}^{\infty} \tilde{\mathbf{E}}(\mathbf{K}, y) \exp(i\mathbf{K} \cdot \mathbf{R}) d^2\mathbf{K} \end{cases} \quad (\text{I.96})$$

Il exprime alors la transformée de Fourier du champ en un point quelconque en fonction de la transformée de Fourier en $y = 0$ par l'intermédiaire de la fonction de transfert :

$$\tilde{\mathbf{E}}(\mathbf{K}, y) = H(\mathbf{K}, y) \cdot \tilde{\mathbf{E}}(\mathbf{K}, y = 0) \quad (\text{I.97})$$

avec lorsque la propagation se fait sur une distance supérieure à la longueur d'onde :

$$\begin{aligned} H(\mathbf{K}, y) &= \exp(i\sqrt{k^2 - \mathbf{K}^2}y) & \text{si } |\mathbf{K}| \leq k \\ H(\mathbf{K}, y) &= 0 & \text{si } |\mathbf{K}| > k \end{aligned} \quad (\text{I.98})$$

Grâce à ce formalisme, le phénomène de propagation peut alors être vu comme un filtre spatial linéaire et dispersif dont la bande passante est finie. La sortie de ce filtre est nulle si le vecteur d'onde $\mathbf{K}$ est supérieur à l'inverse de la longueur d'onde donc pour des variations spatiales de dimensions inférieures à la longueur d'onde.

A l'aide de cette approche on peut expliquer le retournement temporel en milieu homogène sans introduire de fonctions de Green comme nous l'avons fait précédemment. Lors d'une expérience de retournement temporel, les ondes sont mesurées en champ lointain, c'est-à-dire à bien plus d'une longueur d'onde de la source. Le signaux enregistrés ont donc été "filtrés" par la propagation en espace libre et les détails les plus fins de la source ont été perdus. Ainsi, seules les ondes dont la variation spatiale est plus grande que la longueur d'onde principale du champ vont participer à la focalisation pour donner une tache focale dont la taille en milieu homogène est de l'ordre de cette longueur d'onde. Julien de Rosny a d'ailleurs montré que pour obtenir une tache de focalisation plus fine que la limite de diffraction, il est nécessaire de recréer les ondes évanescentes : c'est l'expérience du puits acoustique [28]. Sans cela, on ne peut pas s'attendre à priori à pouvoir créer une figure de focalisation dont la variation spatiale est inférieure à la longueur d'onde utilisée lors de la phase d'enregistrement. Dès lors comment interpréter les résultats présentés précédemment ? Nous allons voir dans la partie suivante que dans certaines conditions ces détails peuvent survivre à la propagation en espace libre : c'est la base de l'imagerie en champ proche.

I.5.3 L'imagerie en champ proche

Dans la partie précédente, nous avons mis en évidence que les détails plus fins que la longueur d'onde ne sont pas transmis en milieu homogène. Si on se sert alors des relations d'incertitudes de Heisenberg, que ce dernier avait lui même énoncées à l'origine pour expliquer le pouvoir de résolution d'un microscope, on peut écrire concernant la dimension x :

$$\Delta x \Delta k_x > 2\pi \quad (\text{I.99})$$

Cette relation indique que le plus petit détail observable en x est relié à la projection du vecteur d'onde selon la même dimension. Ainsi, après propagation, comme on a montré que Δk_x appartient à l'intervalle $[-2\pi/\lambda, 2\pi/\lambda]$, le plus petit objet que l'on pourra distinguer ne pourra pas dépasser :

$$\Delta x > \frac{2\pi}{\Delta k_x} > \frac{\lambda}{2} \quad (\text{I.100})$$

On retrouve à nouveau la limite de l'optique conventionnelle, mais cette relation de Heisenberg nous apprend autre chose : si on arrive à capter des ondes dont le vecteur de propagation est grand, alors les détails observables pourront être d'autant plus petits. L'idée de l'optique en champ proche vient de cette observation et elle a été comprise assez tôt [29] bien

que les expériences n'aient été réalisées qu'assez récemment. En effet nous avons vu dans le développement en ondes planes exprimé dans la partie précédente que les vecteurs d'ondes trop grand pour se propager sont atténués par un facteur :

$$A(\delta) = \exp(-\sqrt{\mathbf{K}^2 - k^2}\delta) \quad (\text{I.101})$$

où δ est la distance parcourue depuis le plan sur lequel nous avons fait le développement, $\mathbf{K}$ est le vecteur d'onde transverse et $k = 2\pi/\lambda$ le nombre d'onde dans le vide.

Il est clair que plus le vecteur $\mathbf{K}$ sera grand, plus l'atténuation sera grande. De même, plus la distance d'observation δ à laquelle on se place par rapport au plan $y = 0$ sera grande, plus le champ aura décru. Cependant on voit également que si l'on se place très près d'un échantillon, qui est ici en $y = 0$, on aura accès, si elles existent, à des ondes évanescentes et donc à des détails plus fins que la demi-longueur d'onde. Ce champ observable uniquement à la surface de notre échantillon est appelé le champ proche optique. Dans le domaine de l'optique, chaque surface est couverte d'un grand nombre d'ondes évanescentes. En effet une surface, aussi parfaite qu'elle puisse être, présente toujours des aspérités qui, lorsqu'elles sont illuminées, diffractent les ondes électromagnétiques. Si les dimensions de ces défauts sont petites devant la longueur d'onde, elles donneront naissance à des ondes dont les vecteurs transverses sont grands et par conséquent à des ondes évanescentes. Dans la référence [27], le champ proche est alors défini comme "la zone dans laquelle les ondes évanescentes contribuent significativement au champ électromagnétique".

On comprend dès lors que pour avoir accès aux détails les plus fins d'un échantillon, et par là même briser la limite de l'optique conventionnelle, il va falloir approcher le détecteur le plus près possible de la surface de celui-ci [30]. Une question se pose : comment est-il possible de mesurer une onde évanescente sachant que celle-ci ne se propage pas ? La réponse à cette question tient dans le principe de retour inverse de la lumière : si un objet sub-longueur d'onde peut transformer par diffraction une onde propagative en onde évanescente, il peut, de la même façon, transformer par diffraction des ondes évanescentes en ondes propagatives. Cette dernière observation est à la base de l'optique en champ proche : en plongeant un objet sub-longueur d'onde dans le champ proche de l'échantillon observé, les ondes évanescentes présentes sur sa surface vont être partiellement transformées en ondes propagatives et pouvoir se propager, avec les informations qu'elles contiennent, jusqu'au détecteur. Cette dernière remarque montre que la résolution d'un tel microscope, qui nous l'avons vu dépend de la distance objet-échantillon, dépend aussi de la dimension caractéristique de l'objet diffuseur que l'on place dans le champ

proche de l'échantillon : plus celle-ci sera petite, plus son pouvoir de "conversion" d'ondes évanescentes de très grands vecteurs d'onde sera grand et plus la résolution sera fine.

Le but de cette partie n'est pas de réaliser une étude approfondie de la microscopie optique en champ proche, et c'est pourquoi nous ne développerons pas davantage la théorie qui lui est associée. Notons cependant que plusieurs techniques permettent de détecter en champ lointain de l'information sur le champ proche optique et donc de faire de l'imagerie avec une résolution bien meilleure que celle de la microscopie classique. Parmi elles, on citera trois techniques qui peuvent être mises en oeuvre grâce à une illumination en réflexion ou en transmission :

- Le microscope à effet tunnel optique : dans cette technique, une fibre optique monomode taillée en pointe joue le rôle de diffuseur des ondes évanescente. Puis les ondes propagatives créées se propagent dans la fibre jusqu'au détecteur.
- Une autre technique consiste à utiliser des nanoparticules, généralement en or, qui vont diffuser les ondes évanescentes. Puis, une fibre optique couplée à un détecteur réalise la mesure de ce champ propagatif.
- Enfin, la technique la plus couramment utilisée est la technique de microscopie optique de champ proche à balayage. Ici, c'est une pointe métallique dont le bout est d'une dimension nanométrique, qui est chargée de diffuser le champ évanescent. Il existe alors plusieurs montages pour détecter le champ créé. Un bon exemple est donné dans [27].

Dans ces trois techniques, la finesse des détails que l'on peut observer est uniquement limitée par la distance du diffuseur à la surface à imager et par la taille du diffuseur. A titre d'exemple, la figure I.35 montre une mesure réalisée par microscopie en champ proche de la localisation d'un champ magnétique à la surface d'un film d'or semi-métallique. On peut y apercevoir que les endroits où l'énergie est localisée ont des dimensions latérales de l'ordre de la dizaine de nm, ce qui compte tenu de la longueur d'onde utilisée, est bien inférieur à la limite de diffraction. Une mesure en microscopie classique n'aurait pas permis d'observer ces zones de localisation. Si l'imagerie en champ proche est très connue dans le domaine de l'optique, elle n'en reste pas moins valable dans d'autres gammes de fréquences et de nombreuses techniques sont apparues dans le domaine des micro-ondes ces dernières années. Ainsi, dans [31], les auteurs rapportent un procédé d'imagerie en champ proche de propriétés électriques dans le GHz. La longueur d'onde utilisée est autour de 2 cm pour une résolution de 100 μm , soit à peu près $\lambda/200$. De même, dans [32], un procédé de chauffage en champ proche est expérimenté sur une surface de $0.3 * 0.5 \text{ mm}^2$, et ce en utilisant une longueur d'onde de 3 cm.

Les nombreux travaux actuels sur l'imagerie en champ proche nous ont permis d'émettre des

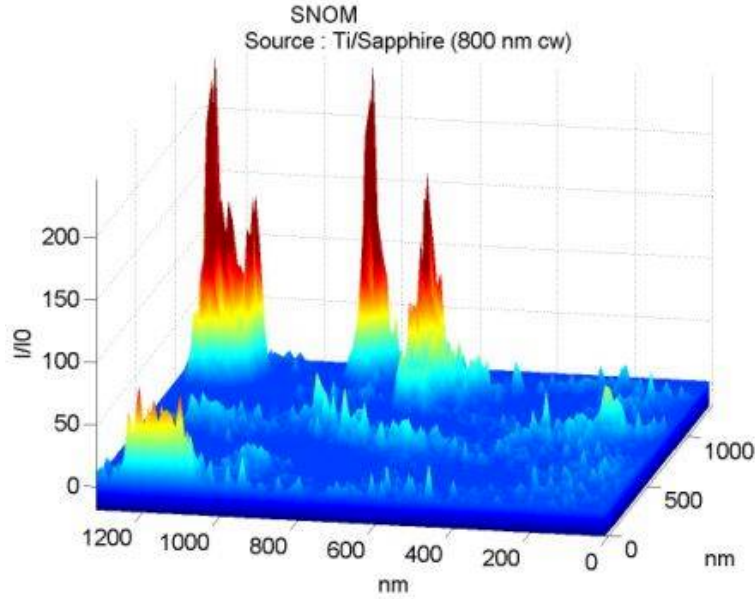


FIG. I.35 – Localisation du champ électromagnétique sur une couche d’or semi continue (source : Laboratoire d’Optique Physique, ESPCI).

hypothèses afin d’expliquer les résultats surprenants que nous avons obtenus, et nous ont guidé vers une nouvelle approche du retournement temporel d’ondes électromagnétiques.

I.5.4 Résultats expérimentaux et premières justifications

Les résultats qui vont être présentés dans cette partie ne sont pas à ce jour complètement compris et font toujours l’objet de recherches au laboratoire. En effet, suite aux résultats obtenus avec les antennes commerciales, présentés en première partie de ce chapitre, nous avons entrepris une démarche purement expérimentale visant à étudier la possibilité de focaliser des champs par retournement temporel sur des dimensions inférieures à la longueur d’onde.

L’hypothèse que nous avons choisie dès le départ est la suivante : les antennes commerciales, prévues pour être montées sur un plan de masse, sont très mal adaptées en impédance lorsqu’elles sont utilisées telles quelles. Une mesure réalisée à l’analyseur de réseau est venue conforter notre intuition puisque nous avons mesuré une impédance de $5 + 10i \, \Omega$, ce qui est bien loin des $50 \, \Omega$ des cables et circuits utilisés. De telles antennes doivent théoriquement être très inefficaces, car la puissance rayonnée dépend de la partie réelle de l’impédance, la partie imaginaire créant quant à elle un champ appelé “réactif” qui reste confiné au voisinage de l’antenne [16]. Dès lors, comment une antenne presque purement réactive peut-elle capter et émettre avec un gain du même ordre de grandeur qu’une antenne adaptée en impédance ? Nous avons répondu

à cette question en mettant en cause les éléments métalliques présents dans l'environnement proche de l'antenne, qui sont de fait des diffuseurs. Nous avons vu, en effet, dans la partie précédente qu'une onde évanescente peut être convertie en onde propagative par un diffuseur. Le champ réactif créé par une antenne dont l'impédance possède une partie imaginaire élevée étant constitué d'ondes évanescentes, nous en avons conclu que nos résultats pouvaient s'expliquer par une conversion d'ondes évanescentes en ondes propagatives. D'une part, cela explique pourquoi des antennes presque purement réactives donnent tout de même naissance à un champ propagatif en présence de diffuseurs, et donc pourquoi nos antennes sont efficaces. Mais plus intéressant encore, cela explique la finesse de la tache focale obtenue par retournement temporel lorsque nous utilisons ces antennes. En effet, d'après notre raisonnement, l'information captée en champ lointain par le miroir à retournement temporel contient des détails spatiaux très fins, grâce à la conversion évanescent/propagatif : c'est la base de l'optique en champ proche comme nous l'avons rappelé. Lors de la phase de réémission du retournement temporel, la conversion se fait dans le sens inverse : des ondes propagatives deviennent évanescentes ; la réciprocity de la cette conversion a en effet été démontré dans [14]. La conclusion est donc que les ondes évanescentes participent à la focalisation par retournement temporel, ce qui permet une concentration de l'énergie sur une tache plus fine que la limite de diffraction.

Une autre explication du phénomène, qui est strictement équivalente, tient dans la formule I.33, qui donne l'expression du champ après retournement temporel par une cavité parfaite en milieu hétérogène. En effet, nous avons indiqué que le champ ainsi créé est concentré sur une tache qui varie comme la partie imaginaire de la fonction de Green Dyadique entre deux points du milieu. Cela nous a permis d'en déduire la limite de diffraction en milieu homogène, car la partie imaginaire de la fonction de Green Dyadique est alors un sinus cardinal. Cependant, rien n'indique que dans un milieu hétérogène, cette fonction soit aussi régulière. On peut alors imaginer créer des milieux dans lesquels la partie imaginaire de la fonction de Green Dyadique entre deux points varie à une échelle très inférieure à la longueur d'onde considérée. Si l'on place alors des antennes très proches dans ce milieu, il est possible de focaliser sur chacune d'entre elles par retournement temporel : ceci constitue une autre interprétation du phénomène.

Nous avons donc décidé de fabriquer nos propres antennes afin de valider nos intuitions. La première chose à faire était de fabriquer des antennes dont l'impédance présenterait une partie réelle nulle, afin de générer très peu d'ondes propagatives en milieu homogène, et qui posséderait par contre une partie imaginaire élevée, de manière à créer un fort champ réactif. Une telle source, que l'on peut qualifier de non radiative, a fait l'objet de nombreuses recherches dans

le domaine de l'optique depuis quelques années, en particulier en ce qui concerne l'imagerie en champ proche. Dans le domaine des microondes, au contraire, elle est connue depuis longtemps : un élément rayonnant dont la dimension est petite devant la longueur d'onde en est un exemple parfait. Nous avons donc opté pour un câble coaxial, dont l'âme dépasse de quelques millimètres. Afin de rendre l'interprétation de nos résultats plus facile, nous avons de plus décidé de fabriquer un réseau d'antennes sur un plan de masse, afin d'écranter à nouveau les câbles comme nous l'avions fait pour notre réseau d'antennes "idéales". De manière à obtenir des résultats probants, nous avons réalisé un réseau dont les antennes sont espacées d'une distance très inférieure à la longueur d'onde utilisée (12.25 cm) ; le pas a été choisi à 4 mm. La figure I.36 montre une photo du réseau ainsi constitué, sur laquelle on peut voir les bouts de câble coaxial d'une longueur de 2 mm qui dépassent du plan de masse.



FIG. I.36 – Photo du réseau fait de 8 antennes non radiatives, espacées de 4 mm.

Nous avons ensuite vérifié que ces antennes placées dans l'air sont très peu efficaces, c'est-à-dire qu'elles ne produisent pratiquement pas de rayonnement électromagnétique en champ lointain. Pour ce faire, nous avons mesuré à l'aide d'une antenne $\lambda/2$ commerciale, et dans notre cavité, la moyenne de la valeur absolue de l'amplitude reçue lorsque les 8 antennes du réseau émettent tour à tour une impulsion. Comme nous nous y attendions, la tension mesurée était à peu près 50 fois inférieure à celle mesurée avec des antennes commerciales, ce qui prouve que nos antennes sont bien principalement réactives en milieu homogène. Restait alors à fabriquer des diffuseurs, que nous placerions autour de chacune des antennes, de façon à ce que le champ réactif soit converti en champ propagatif. Le diffuseur le plus élémentaire en électromagnétisme

étant le fil de cuivre, nous avons fabriqué des couronnes de fil de cuivre, d'une hauteur de 4 ou 5 cm. Ces fils de cuivre sont assemblés à leur base avec de la pate à modeler. Celle-ci permet en outre d'isoler les diffuseurs du plan de masse sur lequel ils vont être disposés. En effet, si ceux-ci sont en contact avec le plan de masse, tout courant induit dans un diffuseur par le champ réactif sera intégralement transmis au plan de masse, le rendant ainsi totalement inefficace. Enfin, nous avons entouré les bouts de câble coaxial dépassant du plan de masse des couronnes de fils de cuivre. Le réseau d'antennes ainsi constitué est présenté sur la photo de la figure I.37.



FIG. I.37 – Photo du réseau d'antennes non radiatives avec couronnes de diffuseurs, que nous qualifierons par la suite d'antennes microstructurées.

La première chose que nous avons constaté est que ces antennes étaient maintenant en mesure d'émettre ou de recevoir un champ propagatif. Nous avons donc réalisé la même opération que précédemment : à l'aide d'une antenne commerciale placée dans la cavité, nous avons mesuré la moyenne de la valeur absolue de l'amplitude reçue lorsque les 8 antennes microstructurées du réseau émettent tour à tour une impulsion. Cette fois la moyenne mesurée est de 50 mV pic à pic, soit un gain de 50 par rapport aux antennes sans les couronnes, ou soit encore la valeur typique mesurée en utilisant un réseau d'antennes commerciales. Ensuite nous avons mesuré, comme nous l'avons fait avec notre réseau d'antennes "idéales", la tache focale obtenue par retournement temporel. Le résultat qui est présenté sur la figure I.38 a été obtenu après focalisation sur la cinquième antenne du réseau.

Nous avons également représenté sur cette figure, à titre de comparaison, la tache focale obtenue

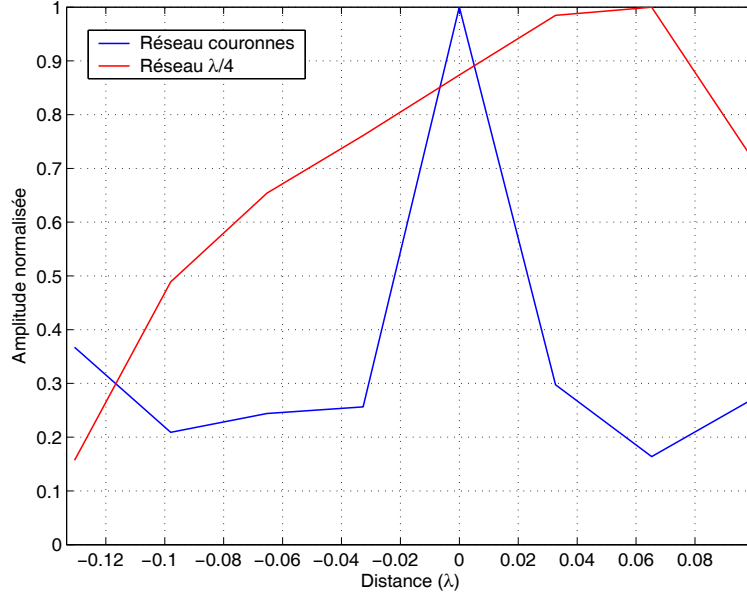


FIG. I.38 – Taches focales obtenues par retournement temporel lorsque l’on focalise sur l’antenne 5 du réseau à $\lambda/30$.

pour un réseau d’antennes de hauteur usuelle $\lambda/4$, d’un diamètre de 1 mm, et qui sont séparées de la même distance. Comme on peut le voir, focaliser sur l’antenne 5 d’un tel réseau n’est même pas possible à cause du couplage entre les antennes. On obtient un maximum d’amplitude sur les antennes 6 et 7. D’autre part, le signal a une amplitude supérieure à 0.5 sur 7 antennes des 8 antennes du réseau, la focalisation a donc échoué. Par contre, en utilisant nos antennes “couronnes”, on peut s’apercevoir que le signal est concentré sur l’antenne 5, avec une amplitude au moins 3 fois supérieure à celle qui est reçue sur les autres antennes du réseau. Le rapport entre l’amplitude obtenue sur l’antenne 5 et celle obtenue sur les autres antennes ne semble en revanche pas décroître, que le miroir à retournement temporel soit constitué de 1 antenne ou de 8 antennes. Parallèlement nous avons mesuré la phase sur les antennes du réseau, au temps $t = 0$ du retournement temporel, donc au niveau du maximum obtenu sur la voie 5. Nous avons pu observer que celle-ci évolue beaucoup, avec une variation qui est de l’ordre du pas du réseau. Enfin, il est important de noter que la position temporelle des maxima reçus sur les antennes du réseau varie elle aussi, de -15 ns à 15 ns selon l’antenne considérée. Ces observations ne sont pas toutes comprises et font toujours l’objet de recherches.

Afin de vérifier l’utilité de telles antennes, notamment en vue de réaliser de la télécommunication, nous représentons également sur la figure I.39 le résultat de la focalisation par retournement temporel lorsque l’on cherche à focaliser sur la quatrième antenne du réseau à $\lambda/30$, ainsi que sur la cinquième.

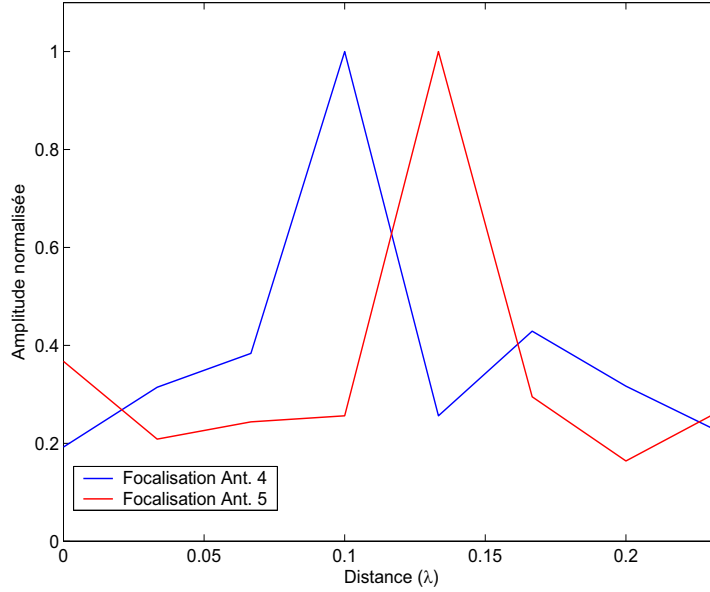


FIG. I.39 – Taches focales obtenues par retournement temporel lorsque l’on focalise sur l’antenne 4 et sur l’antenne 5 du réseau à $\lambda/30$.

Cette dernière figure montre qu’il est possible grâce à ces antennes et au retournement temporel d’adresser simultanément deux antennes séparées uniquement d’un trentième de longueur d’ondes.

Afin de valider notre théorie, nous avons également réalisé la même expérience que précédemment en utilisant un réseau d’antennes séparées cette fois de $\lambda/12$. Ces antennes sont également des monopôles sur un plan de masse et leur hauteur est de 2 cm. Nous avons mesuré à l’analyseur de réseau, pour chacune des antennes, une impédance de $10 - 110i \, \Omega$, ce qui assure une nouvelle fois que celles-ci créent un champ réactif intense. Nous avons pu observer qu’une nouvelle fois, les entourer d’une couronne de diffuseur augmente leur efficacité : en moyenne, l’amplitude croît d’un rapport 1.5 pour les antennes avec diffuseurs. D’autre part, nous avons mesuré une tache de focalisation par retournement temporel, en utilisant les antennes d’abord avec diffuseurs, puis sans diffuseurs. Ces taches focales, qui ont été normalisées, sont tracées sur la figure I.40 pour une focalisation sur l’antenne 4 du réseau.

Ici encore, on voit que lorsque l’on place des diffuseurs dans le champ réactif des antennes, la tache de focalisation est bien plus fine, ce qui est en accord avec les explications que nous avançons. Enfin, afin de vérifier que le champ réactif est bien à la source de cette focalisation, nous avons réitéré l’expérience, cette fois-ci en utilisant des antennes commerciales, adaptées en impédance, qui sont donc censées ne générer que des ondes propagatives. Nous avons mesuré l’amplitude moyenne sur une autre antenne de la cavité, quand les 8 antennes émettent tour à

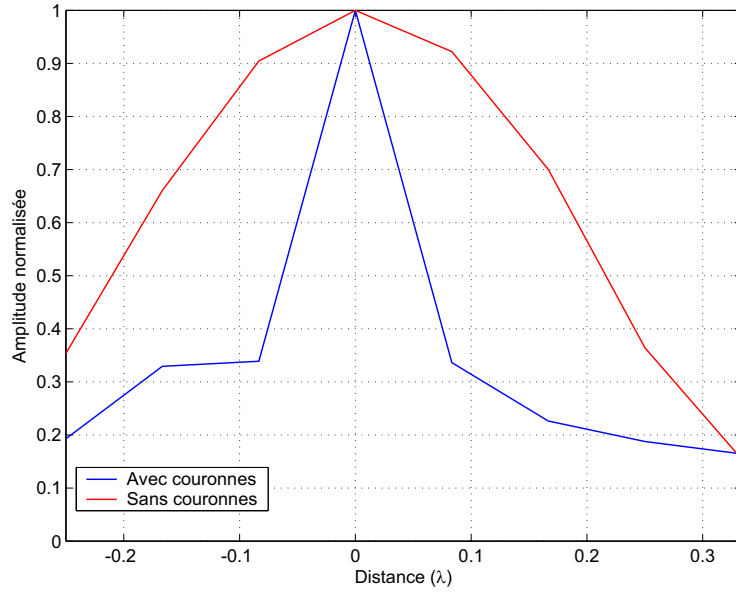


FIG. I.40 – Taches focales obtenues par retournement temporel lorsque l’on focalise sur l’antenne 4 du réseau à $\lambda/12$.

tour. Cette mesure, réalisée pour les antennes avec et sans diffuseurs, donne des résultats tout à fait équivalents.

A première vue, les diffuseurs ne jouent donc aucun rôle lorsque l’on utilise des antennes adaptées en impédance. Sur la figure I.41 sont représentées les taches focales obtenues avec et sans diffuseurs.

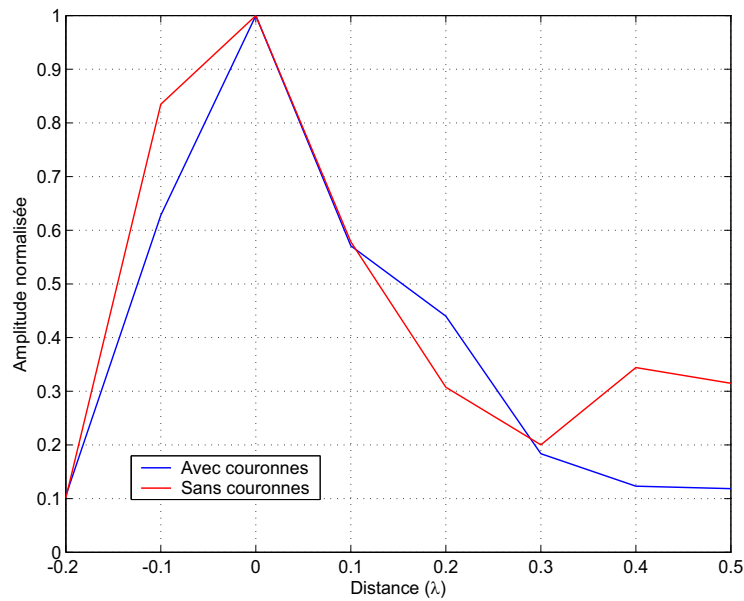


FIG. I.41 – Taches focales obtenues par retournement temporel lorsque l’on focalise sur l’antenne 3 du réseau d’antennes adaptées en impédance.

Cette fois, les taches focales sont très similaires et la largeur à mi-hauteur des 2 pics est exactement la même. La présence des diffuseurs provoque quelques variations par rapport au réseau initial, mais celles-ci sont imputables au fait que les impédances des antennes sont un peu modifiées par la présence des couronnes. En revanche, le fait que la taille de la tache de focalisation par retournement temporel soit une nouvelle fois de l'ordre de la longueur d'onde, montre que les ondes évanescentes ne jouent pas ici un rôle prépondérant. Cette observation tend à valider notre interprétation des résultats obtenus.

Lors de ces expériences, les couronnes de diffuseurs ont été réalisées à la main, et il nous a donc été difficile d'obtenir des résultats très quantitatifs et parfaitement reproductibles. La fabrication de telles antennes, microstructurées, avec des procédés systématiques permettrait de lever les dernières ambiguïtés qui peuvent subsister sur la participation des ondes évanescentes à la focalisation par retournement temporel dans les configurations étudiées.

Dans cette partie nous avons vu que selon les antennes qui sont utilisées, des résultats surprenants peuvent être obtenus lors d'une expérience de retournement temporel d'ondes électromagnétiques. En particulier en cherchant à focaliser une onde sur une antenne d'un réseau, selon les dispositifs expérimentaux étudiés, nous avons obtenu des taches focales dont la dimension caractéristique était plus faible que la longueur d'onde des signaux utilisés. Or, nous savons que pour obtenir une tache focale plus petite que la longueur d'onde, il est nécessaire de prendre en compte les ondes évanescentes dans le processus de retournement temporel. Ces ondes évanescentes sont normalement perdues lors de la propagation et la refocalisation par retournement temporel ne peut donner que des taches focales dont la dimension caractéristique est de l'ordre de la longueur d'onde. Pour expliquer les phénomènes observés nous nous sommes servis des résultats obtenus dans le domaine de l'optique de champ proche. Nous avons alors étudié l'influence de diffuseurs placés dans le champ proche de nos antennes d'émission. En utilisant le fait que de tels diffuseurs sont capables de convertir des ondes évanescentes en ondes propagatives nous avons montré, grâce à la propriété de réciprocité d'une telle conversion, que nous étions capables de focaliser, par retournement temporel en champ lointain, des ondes sur une tache focale bien plus fine que la longueur d'onde utilisée ($\lambda/30$). Ces travaux, qui sont toujours en cours d'investigation, ont probablement un bon potentiel d'application et une demande de brevet a ainsi été déposée auprès du CNRS. De même, une publication devrait suivre cette demande de brevet.

I.6 Conclusion

Nous avons consacré ce premier chapitre au retournement temporel des ondes et à ses propriétés. Pour ce faire, nous avons commencé par dresser un état de l'art des expériences réalisées en acoustique ultrasonore dans la mesure où elles nous semblaient pertinentes pour cet exposé. Nous avons ainsi rappelé que l'utilisation de miroirs à retournement temporel permet, en milieux complexes, de focaliser une onde sur une tache dont la dimension peut être aussi fine que la demi longueur d'onde. La compression temporelle que procure une opération de retournement temporel en milieu réverbérant a également été soulignée, et nous avons rappelé que le rapport entre l'amplitude d'un pic de retournement temporel et les lobes secondaires créés croît en racine du nombre de capteurs utilisés et de la bande passante du signal. Puis, nous avons montré comment ces résultats sont également applicables aux ondes électromagnétiques car celles-ci sont régies par des équations de propagation similaires à celles de l'acoustique. Nous avons toutefois évoqué les différences que l'on peut rencontrer en passant d'un domaine à l'autre, notamment en ce qui concerne les antennes utilisées en radiofréquences, et nous avons développé un formalisme qui permet de les prendre en compte. Ceci nous a permis de réaliser et d'interpréter les premières expériences de retournement temporel d'ondes électromagnétiques avec un miroir monovoie, à une fréquence de 2.45 GHz et en utilisant une bande passante de 4 MHz. Ce dispositif ne permettant pas de réaliser des mesures quantitatives, nous avons conçu un deuxième prototype de miroir à retournement temporel, centré autour de la même fréquence, mais cette fois avec 200 MHz de bande passante. L'utilisation de multiplexeurs nous a par ailleurs permis de travailler avec 8 voies en émission et 8 voies en réception. Avec ce dispositif, nous avons vérifié que les principes qui ont été établis en acoustique restent valables dans le domaine des microondes. La compression temporelle et la focalisation spatiale obtenue par retournement temporel ont été étudiés en fonction du nombre d'antennes et de la bande passante du MRT. Enfin, la participation des ondes évanescentes au processus de focalisation d'ondes électromagnétiques par retournement temporel a été étudiée au travers d'une analogie avec l'imagerie optique de champ proche. Nous avons montré qu'en utilisant des antennes à priori totalement inefficaces, couplées à des diffuseurs dans le champ proche, il est possible d'obtenir une focalisation des ondes sur une tache dont la dimension est bien plus faible que la longueur d'onde utilisée.

Le retournement temporel, étudié depuis quelques années en acoustique, est une technique assez nouvelle dans le domaine de l'électromagnétisme, mais nous allons voir qu'elle est assez

prometteuse. Pour ce faire, dans la partie suivante, nous présenterons les problèmes et enjeux des telecommunications à haut débit et montrerons comment le retournement temporel d'ondes électromagnétiques en milieu réverbérant permet de résoudre un certain nombre de ces problèmes.

Chapitre II

Télécommunications et retournement temporel

Les premières expériences de télécommunications par retournement temporel ont d'abord été réalisées avec des ondes acoustiques dans des chenaux sous-marins [20, 33, 34]. Ce type d'environnements peu profonds, dans lesquels les navires communiquent par ondes sonores, se comportent de façon analogue à des guides d'ondes. Une impulsion émise initialement brève se transforme en une succession d'impulsions à cause des multiples réflexions des ondes sur les interfaces eau/sol et air/eau. Ceci a pour conséquence de dégrader les signaux d'autant plus que les bateaux qui veulent communiquer sont éloignés. L'équipe de Kupermann au SCRIPPS de San Diego a alors eu l'idée d'appliquer le retournement temporel dans ces milieux, comme Philippe Roux et ses collaborateurs l'avaient déjà fait avec des ondes ultrasonores dans des cuves de laboratoire [35]. Cette technique permet d'un part de compenser l'allongement temporel des signaux qui est du aux réflexions, ce qui facilite la réception des signaux. D'autre part, le retournement temporel offre un codage spatial de l'information, en ce sens que cette dernière est spatialement focalisée. Ces avantages ont fait du retournement temporel un sujet d'étude particulièrement prometteur pour les communications sous-marines, notamment en vue d'application militaires.

Après avoir démontré la faisabilité du retournement temporel avec des ondes électromagnétiques, il nous reste donc à prouver l'utilité du procédé et à en quantifier les apports pour les télécommunications. Pour ce faire, nous commencerons par évoquer brièvement les techniques de communication sans fils actuelles ainsi que leurs limites. Nous introduirons alors les méthodes qui sont à l'étude actuellement pour les futures générations d'appareils de télécommunication.

Nous verrons que les deux techniques principales envisagées, l’une fondée sur les systèmes multi-antennes et l’autre sur l’utilisation de très larges bandes passantes, ont un lien très direct avec le retournement temporel. Nous décrirons alors la première expérience de communication par retournement temporel réalisée au laboratoire.

Puis, les premiers résultats ayant été obtenus dans des milieux diffusifs “idéaux”, nous nous rapprocherons de la réalité en décrivant une expérience modélisant à petite échelle un étage d’immeuble. Ces expériences, réalisées avec des ondes ultrasonores, nous donnerons l’occasion de quantifier la qualité des communications par retournement temporel en fonction des paramètres physiques du système considéré. Elles nous permettront alors de souligner les différences qui existent entre le retournement temporel et un filtrage adapté à la réception. Nous développerons également un modèle simple afin de justifier et de comprendre les résultats obtenus expérimentalement.

Nous continuerons cette étude en envisageant les avantages et inconvénients de deux méthodes : le retournement temporel, qui est un filtrage adapté physique, et le filtrage inverse. Une méthode itérative, fondée sur le retournement temporel, sera présentée. Cette technique converge vers le filtre inverse. Après avoir évalué la qualité de la communication lorsque chacune de ces méthodes est utilisée, nous décrirons un formalisme simple qui permet, selon le milieu de communication étudié, de choisir entre les deux techniques. Nous verrons que dans la plupart des cas, afin d’optimiser la robustesse de la transmission d’information, il sera possible d’utiliser une méthode hybride qui se place entre le filtrage adapté et le filtrage inverse.

Enfin, afin de vérifier la validité des résultats obtenus avec des ultrasons, nous montrerons les premières expériences de télécommunications haut débit par retournement temporel à 2.45 GHz. A cette fin, nous utiliserons le miroir à retournement temporel développé dans la partie précédente ainsi que la cavité réverbérante que nous avons fabriquée. Nous étudierons l’influence de la bande passante, du débit d’information, du nombre d’antennes utilisées dans le miroir à retournement temporel et du nombre d’utilisateurs adressés sur la qualité de la communication. Nous verrons également que lorsque le bruit présent dans le milieu est pris en compte, le retournement temporel a pour effet d’augmenter le rapport signal-sur-bruit sur chaque récepteur. Enfin nous présenterons des premières expériences de communications sub-longueur d’onde par retournement temporel.

II.1 Les télécommunications modernes : problèmes et techniques actuelles

II.1.1 Les systèmes Multiple-Input Multiple-Output (MIMO)

Lors de la propagation d'ondes électromagnétiques dans des milieux complexes, comme peuvent l'être les environnements urbains, celles-ci sont réfléchies et diffractées de nombreuses fois sur les obstacles qu'elles rencontrent. Les multiples réflexions donnent naissance à des ondes qui participent par la suite au signal total reçu en un point donné de l'espace. Dans le cas où ces ondes se somment de façon destructive, l'utilisateur placé en ce point ne pourra pas recevoir le signal qui lui était destiné : c'est l'effet de "Fading" [36]. Si la fréquence des signaux à transmettre est faible, les déphasages relatifs provoqués par les réflexions successives sont négligeables, et en un point l'onde reçue est sensiblement la même que l'onde émise. Par contre, si l'on veut transmettre une grande quantité d'information, on doit s'assurer qu'un maximum d'énergie sera reçu grâce à l'onde qui n'a pas rencontré d'obstacles. La télévision hertzienne en est un bon exemple : l'antenne d'émission est placée assez haut, de manière à ce que chaque antenne de réception, placée par exemple sur un toit de maison, voit uniquement l'onde directe. Cependant, avec l'arrivée de la téléphonie mobile et l'augmentation des débits d'information, ces phénomènes de fading sont devenus de plus en plus gênants, au point de nécessiter des terminaux de réception de plus en plus sophistiqués et donc de plus en plus onéreux.

Puis, récemment, Foschini et ses collaborateurs ont montré que loin d'être un facteur limitant, la présence d'obstacles dans un milieu de propagation peut permettre d'augmenter la quantité d'information que l'on peut transmettre d'un point à un autre [37]. Le principe est fondé sur l'utilisation d'antennes multiples en réception, en émission ou en émission/réception. La figure II.1 illustre assez simplement le concept à l'aide d'un système constitué de 2 antennes en réception et 2 antennes en émission, que l'on notera donc MIMO 2×2 , pour Multiple-Input Multiple-Output. Dans cet exemple, chaque antenne d'émission émet un signal de même fréquence mais déphasé de $\pi/2$, ces signaux sont représentés en bleu et rouge. Le réseau d'émission étant placé dans le champ lointain du réseau de réception, les 2 antennes de réception voient ces signaux arriver avec la même différence de phase lors de l'émission en espace libre (fig. II.1.(a)). En effet, l'onde est plane lorsqu'elle arrive sur le réseau de réception. Le signal reçu sur chaque antenne du réseau de réception est donc le même.

Au contraire, si le signal a subi des réflexions, celles-ci vont induire des déphasages différents

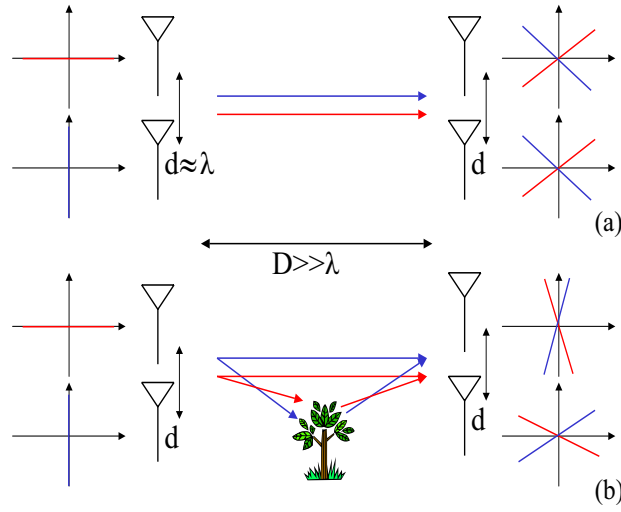


FIG. II.1 – Effet de la présence d'un obstacle sur un système MIMO 2*2.

pour les signaux venant des différentes antennes d'émission, et allant vers les 2 antennes du réseau de réception (fig. II.1.(b)). Le déphasage est d'autant plus significatif que le nombre de diffuseurs est élevé et que ceux-ci sont distribués dans l'espace. Dans ce cas, il est clair que l'information reçue est plus riche. Cet exemple simple montre l'intérêt des systèmes MIMO dans des environnements complexes. D'un point de vue pratique, la présence de diffuseurs dans le milieux de propagation va créer des canaux de communication indépendants entre les antennes d'un réseau d'émission et celles d'un réseau de réception, bien que ceux-ci soient placés à une distance bien supérieure à la longueur d'onde l'un de l'autre, et soient d'une taille comparable à la longueur d'onde. Ainsi, partant de ce principe, il va être possible d'utiliser des systèmes MIMO qui vont permettre de transmettre une quantité d'information plus importante qu'un système SISO entre 2 points de l'espace.

Afin de donner une idée du gain que peut apporter un système MIMO en terme de quantité d'information transférée, il est utile à ce stade d'introduire la notion de capacité d'un canal Gaussien. Cette grandeur, définie par Shannon en 1948 dans [38], définit la quantité d'information en bits par seconde, qu'il est possible de transmettre entre un point et un autre sans erreur, et ce pour une bande passante de 1 Hertz. Elle dépend de l'énergie de chaque symbole d'information (E_s) et de l'énergie du bruit qui est présent dans le canal de propagation (N_0), c'est à dire du bruit qui va être ajouté à un symbole lors de sa transmission entre les deux points :

$$C = \log_2\left(1 + \frac{E_s}{N_0}\right) \text{ en Bits/Hz/s} \quad (\text{II.1})$$

Cette formule a été adaptée au cas des réseaux multi-antennes par Foschini et Gans en 1998 dans [37]. Une autre démonstration a été donnée par Telatar en 1999 dans [39]. Dans ces deux publications, la capacité est définie, pour un système à M antennes d'émission et N antennes de réception qui sont reliées entre elles par la matrice de transfert $\mathbf{H}$ de dimension $M * N$, comme :

$$C = \max_{Tr(\mathbf{R}_{ss})=M} \log_2 \left[\det\left(\mathbf{I}_N + \frac{E_s}{M.N_0} \mathbf{H}^\dagger \mathbf{R}_{ss} \mathbf{H}\right) \right] \text{ en Bits/Hz/s} \quad (\text{II.2})$$

où $\mathbf{R}_{ss}$ est la matrice de covariance des signaux émis, $\mathbf{I}_N$ est la matrice identité de dimension $N * N$, et où la condition $Tr(\mathbf{R}_{ss}) = M$ assure que l'énergie totale émise est indépendante du nombre d'antennes d'émission.

Cette relation est assez peu intuitive telle qu'elle est exprimée ici. Cependant, il est possible de faire quelques approximations, afin de la rendre plus lisible. Si l'on considère que l'on ne va favoriser, lors de l'émission, aucune voie de sorte que l'énergie est équirépartie en direction des N antennes, on obtient $\mathbf{R}_{ss} = \mathbf{I}_M$. On peut alors montrer, en réalisant une décomposition en valeur singulière de la matrice $\mathbf{H}$, que la capacité prend la forme :

$$C = \sum_{i=1}^r \log_2\left(1 + \frac{E_s}{M.N_0} \lambda_i\right) \text{ en Bits/Hz/s} \quad (\text{II.3})$$

où les λ_i sont les valeurs singulières de la matrice $\mathbf{H}$, et r est son rang.

Cette dernière expression appelle la remarque suivante : plus le rang de la matrice $\mathbf{H}$ est grand, plus la capacité est élevée, c'est-à-dire plus on peut communiquer d'information. De même, plus le milieu de propagation est complexe, et plus la matrice de transfert entre les réseaux émetteur et récepteur est de rang élevé. D'un point de vue mathématique, le nombre de canaux de communications créés par l'environnement complexe est égal au rang de la matrice de transfert. Ainsi, contrairement aux systèmes SISO, dont le gain en capacité ne peut être que logarithmique, les systèmes MIMO peuvent offrir un gain linéaire en fonction du nombre d'antennes de réception.

Les systèmes MIMO, selon la façon dont ils sont utilisés, peuvent donc offrir un gain en capacité soit linéaire, soit logarithmique. Nous allons ici, sans entrer dans les détails, passer en revue les divers gains que peuvent apporter les systèmes MIMO. Les deux premiers permettent d'augmenter la capacité de façon logarithmique, les deux seconds procurent des gains linéaires

en capacité. Pour un ouvrage complet sur le sujet, on se référera à [40]

Le gain d'antennes :

Cette technique consiste à augmenter le rapport signal-sur-bruit reçu par un utilisateur. On obtient en effet une augmentation de ce rapport signal-sur-bruit en combinant de façon cohérente les signaux émis par différentes antennes d'émission et/ou les signaux reçus par les antennes de réception. On peut par exemple considérer un système SIMO (Single Input Multiple Output), dont chaque antenne du réseau de réception reçoit un signal différent en amplitude et en phase. Le récepteur peut alors sommer les signaux de façon cohérente de manière à maximiser le rapport signal-sur-bruit. Le gain en RSB est alors linéaire par rapport au nombre d'antennes du récepteur, si les signaux reçus sont parfaitement décorrélés. Cette technique nécessite de connaître le canal de communication à la réception. De manière équivalente, si l'on connaît le canal à l'émission dans un système MISO ou MIMO, il est possible de réaliser de la formation de voies, de manière à maximiser l'énergie reçue.

Le gain de diversité :

Utiliser la diversité permet de lutter contre les effets de fading que nous avons décrit précédemment. Cette technique est utilisée depuis de nombreuses années déjà pour les systèmes SIMO [41], dans lesquels chaque antenne de réception reçoit un signal qui présente des minima d'amplitudes, mais à des instants différents. Le récepteur peut alors combiner ces signaux, de manière à ce que le signal résultant présente beaucoup moins de zones de faible amplitude que chaque antenne lorsqu'elle est utilisée de façon indépendante. Pour des signaux à large bande passante, il est possible de réaliser la même opération dans le domaine fréquentiel. En ce qui concerne les systèmes MISO et MIMO, il est possible de tirer profit de la diversité en ayant ou non connaissance du canal à l'émission. Dans ce cas, c'est en fabriquant des signaux d'émission adéquats que l'on tire profit de cette dernière à la réception. Un bon exemple utilisant cette technique dans le cas de systèmes MIMO 2*2 a été publié par Alamouti dans la référence [42].

Le gain de réduction d'interférences :

Lorsque de l'information est envoyée simultanément à plusieurs antennes ou par plusieurs antennes, sur la même gamme de fréquence, nous sommes en présence d'interférences inter-

utilisateurs. Ceci se traduit par le fait que transmettre une information par une antenne va générer un signal sur une antenne qui n'est pas visée. Ces interférences sont autant de "bruits" qui sont indésirables lors de processus de communication. La réduction d'interférence peut être réalisée en fabriquant des signaux à émettre de façon à ce que les interférences créées soient minimisées. La technique dite du "Waterpourring", développée par Cover et Thomas dans [43], utilise la réduction d'interférence à l'émission.

Le gain de multiplexage spatial :

Le multiplexage spatial est uniquement possible dans le cas de systèmes MIMO. Cette technique consiste à diviser le flux de données à transmettre en plusieurs sous-flux qui vont être émis simultanément par le réseau d'émission, et captés par le réseau de réception. Si les différents couples émetteur/récepteur sont décorrélés, le récepteur, qui a connaissance du canal, peut alors différencier les signaux venant des différentes antennes. Après démodulation il est en mesure de recomposer le flux de données initial. Cette technique permet d'envoyer, dans le meilleur des cas, une quantité d'information proportionnelle au nombre de paires d'antennes émettrices/réceptrices N . Un algorithme basé sur ce principe, nommé V-Blast, a été développé par les Bell Labs dans la référence [44]. Le multiplexage spatial peut aussi être utilisé dans des systèmes MIMO-Mu, c'est-à-dire des systèmes pour lesquels une base composée de M antennes s'adresse à N utilisateurs simultanément. Si les signaux de la base vers chaque utilisateur sont décorrélés, il est possible, moyennant un précodage à l'émission, d'envoyer un signal différent à chaque récepteur. De même, si chaque récepteur émet un signal en même temps, la base, qui a connaissance du canal, peut décrypter le message envoyé par chaque utilisateur. Nous verrons que cette forme de multiplexage spatial nous intéressera particulièrement par la suite.

Notons enfin que certaines techniques fondées sur les systèmes MIMO, et pour lesquelles aucune information sur le canal n'est nécessaire, sont actuellement en cours d'étude, elles sont généralement référés au terme "blind MIMO", soit MIMO aveugle. Cependant, dans l'état actuel des recherches, il est nécessaire de connaître le canal à l'émission ou à la réception, de façon à profiter des divers gains que peut apporter un système MIMO. Cette dernière remarque mérite une précision : avoir connaissance du canal à l'émission ou à la réception n'est pas équivalent. En effet, dans le premier cas, il est possible de maximiser l'énergie que l'on veut transmettre à chaque antenne du réseau de réception, par chaque antenne du réseau d'émission. En ce sens, avoir une connaissance du canal à l'émission sera toujours plus efficace. Cette dernière remarque

est d'ailleurs visible dans la formule II.2 : à l'émission, il est possible de modifier l'énergie des signaux émis et donc la matrice $\mathbf{R}_{ss}$, afin de maximiser la capacité. Ceci est impossible par un traitement à la réception. Enfin, notons que l'utilisation d'émetteurs/récepteurs MIMO a déjà fait son apparition dans les systèmes de communication de type WIFI de dernière génération.

II.1.2 Les communications Ultra Large Bande (UWB)

Nous avons vu qu'utiliser un système MIMO permet d'augmenter la quantité d'information que l'on souhaite transmettre d'un point à un autre en environnement complexe. Cependant, la formule de Shannon (II.1) nous apprend également qu'augmenter la bande passante permet de réaliser la même opération : en effet, la capacité est définie en Bits/Hertz/s. Dès lors, pourquoi ne pas utiliser des systèmes de communication à très large étendue spectrale qui résoudraient le problème de l'augmentation des débits d'information ? Il y a plusieurs raisons à cela. La première est triviale : si chaque utilisateur se sert d'une bande passante démesurée, les autres ne peuvent plus communiquer. De plus, utiliser une large bande passante complique beaucoup la conception des circuits électroniques. Cependant, ce qui a le plus freiné la recherche sur les systèmes ultra large bande passante est le fait que le spectre des ondes électromagnétiques est réglementé dans chaque pays. Tous les systèmes de communications fonctionnent selon des normes : les opérateurs de téléphonie portable payent une redevance en fonction de la bande passante qu'ils utilisent. Quant aux systèmes locaux, tels que le WIFI, certaines plages du spectre leur ont été allouées, sous la forme d'une dérogation. Il a donc fallu attendre que les instances dirigeantes, et au premier chef la FCC (United States Federal Communication Commission) allouent un spectre permettant de réaliser des communications ultra large bande passante (UWB). La figure II.2 représente le masque autorisé par la FCC [45], c'est-à-dire la limite de puissance autorisée en dBm/MHz en fonction de la fréquence. Sur cette courbe sont également représentés les spectres autorisés pour le GSM, l'UMTS et les bandes ISM (WIFI et Bluetooth). La bande principale autorisée pour la communication se situe entre 3.1 et 10 GHz et, comme on peut le voir, la puissance y est limitée à -41 dBm/MHz. Cette valeur est 70 dBm/MHz en dessous des puissances autorisées pour la téléphonie portable, et se situe au niveau du bruit ambiant. Cependant, il est nécessaire de noter que celle-ci est définie en fonction de la bande passante utilisée, et à titre comparatif, un GSM occupe 30 KHz de bande passante, quant un système UWB est prévu pour fonctionner sur plusieurs GHz.

Cette limite de puissance va avoir plusieurs conséquences : en premier lieu, les distances de

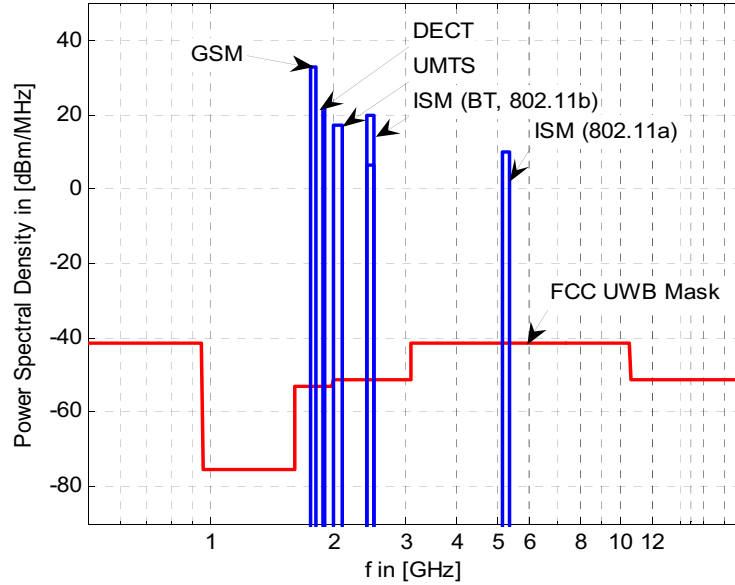


FIG. II.2 – Spectre attribué à l’UWB comparativement aux autres normes.

transmissions vont être très limitées. A l’heure actuelle, la limite est fixée à 10 mètres, ce qui est faible même pour un environnement “indoor”. De plus, cette faible puissance autorisée va favoriser des solutions qui maximisent l’énergie reçue par les utilisateurs.

En ce qui concerne la norme, la FCC définit un signal UWB comme un signal dont la bande passante à -10 dB est supérieure à 500 MHz, ou comme un signal dont la bande passante à -10 dB est supérieure à 20% de la fréquence porteuse. La fréquence minimale autorisée pour les communications étant de 3.1 GHz, en pratique une communication UWB ne pourra pas se faire avec une bande passante inférieure à 500 MHz. Enfin, concernant les détails techniques, voici les principales caractéristiques que doit respecter un système UWB [46, 47] :

- Le but étant de réaliser de la transmission haut débit, il est nécessaire que les communications assurent un débit supérieur à 100 MB/s afin de pouvoir transmettre de la vidéo haute définition sans fil.
- Il est nécessaire de pouvoir adresser plusieurs utilisateurs en même temps, sans que le débit par utilisateur en souffre.
- Le compromis entre le débit et la distance de communication doit être le plus efficace possible, tout en respectant les normes sur les puissances.
- Les système devront être robustes face aux interférences provenant des différents systèmes qui occupent une partie du spectre UWB (i.e. WIFI).
- L’une des principales recommandations concerne les terminaux de réception qui devront être les moins complexes possibles de manière à offrir des prix très compétitifs. Cette remarque

se traduit par le fait que le maximum des traitements devra se faire à l'émission des signaux, assurant ainsi une réception simplifiée.

- Enfin les terminaux de réception devront assurer une consommation faible, ce qui limite l'amplification en réception ainsi que les traitements numériques.

Parmi les problèmes que pose la technique UWB, on peut tout d'abord noter qu'un récepteur va être fortement perturbé par un signal provenant d'une source faible bande comme un réseau local de type WIFI. La réception sera d'autant plus difficile que la puissance tolérée est faible. Mais la principale préoccupation de la communauté est sans doute l'effet de la propagation des signaux UWB. En effet, utiliser des impulsions courtes permet d'augmenter le débit d'information, mais a pour conséquence de compliquer considérablement la détection du signal reçu. Ainsi, si un signal d'une nanoseconde est envoyé à travers une pièce, d'une antenne à une autre, chaque réflexion sur un objet ou un mur va créer une impulsion distincte de la première : 30 cm de propagation suffisent à décaler l'onde réfléchie de 1 ns.

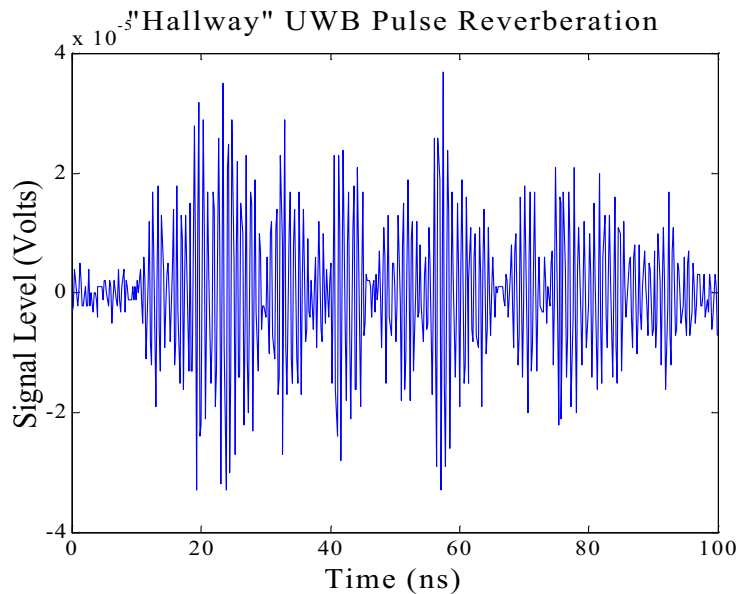


FIG. II.3 – Exemple d'une réponse impulsionnelle à 10 m en environnement indoor (Source : Multispectral Solutions, Inc.).

A titre d'exemple, la figure II.3 montre la réponse, dans un couloir, à une impulsion de 1 ns. La distance entre l'émetteur et le récepteur est de 10 m, sans vis-à-vis entre les deux antennes. L'impulsion s'est étalée sur près de 100 ns, ce qui représente une centaine d'impulsions. En terme de communication, cette dispersion temporelle va être très problématique, car pour éviter que les bits émis ne se superposent, il va falloir les séparer de la durée typique de la

réponse impulsionnelle. Ceci aura alors pour conséquence de diminuer le débit d'un facteur 100. Une autre solution consiste alors à utiliser un récepteur de type "RAKE"¹, qui avec une connaissance du canal, peut déconvoluer la réponse impulsionnelle de façon à décoder l'information. Cependant, ce type de récepteur est très coûteux en calcul, et en consommation. D'autre part, la complexité d'une telle déconvolution croît exponentiellement avec la durée de la réponse impulsionnelle [48].

Ces problèmes restent assez ouverts, comme le montrent les principales techniques à l'étude pour les communications UWB :

La méthode pulsée multi-bandes

La bande passante est divisée en sous-bandes étroites, et une technique impulsionnelle est utilisée dans chaque sous-bande. Cela permet à la fois de travailler avec des composants électroniques plus faibles bandes, qui consomment moins, et donc de fabriquer des produits moins onéreux. Par ailleurs, cette division en sous-bandes permet également une flexibilité de la bande passante, en ce sens qu'une partie celle-ci peut être inutilisée si il ya trop d'interactions avec des systèmes environnant. Par contre, le désavantage de la méthode est sa faible capacité à profiter des trajets multiples du fait du traitement incohérent de la bande passante. De plus, les récepteurs et émetteurs doivent être capable de basculer en fréquence de façon très rapide. Cette caractéristique est très problématique pour le récepteur, qui doit être aussi simple que possible [46].

Le multiplexage par division de fréquences orthogonales multi-bandes

Ici, la bande passante est à nouveau divisée en sous-bandes. Celles-ci restent assez larges : 500 MHz. Au sein de chacune de ces sous-bandes, une méthode de multiplexage par division en fréquences orthogonales (OFDM) est appliquée [46]. pour chaque sous-bande, l'émetteur et le récepteur ont en commun une séquence pseudo-aléatoire. Le signal est transmis par des changements de fréquences (les 500 MHz sont coupées en tranches de 4 MHz), qui respectent cette séquence pseudo-aléatoire. Cette méthode est très résistante au bruit et aux interférences externes. D'autre part, elle tire parti de tous les trajets multiples, assurant ainsi un bon rapport

¹Le terme RAKE signifie râteau en anglais. Ce type de récepteur est composé de plusieurs sous récepteurs (les "fingers") qui, avec une connaissance du canal, décodent chacun une composante temporelle du signal correspondant à un trajet multiple.

signal-sur-bruit. En revanche, il est nécessaire d'avoir en émission et en réception des puces capables de réaliser des transformées de Fourier rapides (FFT), ce qui va à l'encontre du cahier des charges.

L'accès multiple par répartition en codes en séquences directes

En DS-CDMA [49], une séquence pseudo-aléatoire de 0 (ou -1) et de 1 est générée pour chaque utilisateur. Chaque symbole d'information est ensuite convolué par cette séquence et transmis : chaque symbole a donc une durée égale à la durée d'une impulsion multipliée par la longueur de la séquence. A la réception, on utilise à nouveau un RAKE, afin de profiter des trajets multiples subis par les ondes. Afin de décoder l'information, le signal reçu est corrélé par la séquence pseudo-aléatoire qui est connue du récepteur. Cette technique est elle aussi très robuste aux interférences faible bande passante provenant des autres systèmes de communication. La technologie est plus simple que celle employée en OFDM, mais nécessite tout de même un corrélateur et un récepteur de type RAKE.

L'accès multiple par saut temporel

Ici, une séquence pseudo-aléatoire est à nouveau générée [50], elle est également connue de l'émetteur et du récepteur. Les symboles sont répartis sur un nombre donné d'impulsions, mais cette fois c'est l'espacement entre ces impulsions qui est défini par le code. A la réception, on corrèle à nouveau le signal reçu par la séquence d'impulsions. Cette technique est très robuste aux interférences qui proviennent des autres utilisateurs. Elle permet également de bien profiter de la bande passante, car le fait de briser la périodicité des impulsions émises évite les résonances de spectre. De plus elle présente l'avantage d'un traitement du signal assez simplifié. En revanche il est nécessaire de placer un corrélateur dans le récepteur. Enfin, un récepteur de type RAKE est inévitable pour profiter des trajet multiples.

Cependant la technique UWB est pour l'instant en perpétuel évolution et aucune norme n'a pour l'instant été adoptée définitivement. De plus, chacune des techniques étudiées prévoit une architecture de récepteur qui va contre les attentes des industriels. Nous allons voir dans la partie suivante que nombre des problèmes rencontrés avec de tels systèmes peuvent être contournés grâce au retournement temporel.

II.1.3 L'apport du retournement temporel aux communications sans fils

Nous venons de voir que les deux techniques envisagées afin d'augmenter les débits d'informations dans les communications sans fils sont basées sur les systèmes multi-antennes MIMO et les systèmes ultra-large bande passante. Nous allons ici, avant de présenter des expériences de communications par retournement temporel, justifier l'utilisation d'un tel procédé. Pour ce faire, nous allons commencer par faire le lien entre les systèmes MIMO et le retournement temporel.

Tout d'abord, le retournement temporel est une méthode qui exploite naturellement la diversité spatiale d'un milieu. Ceci est particulièrement visible si l'on écrit le résultat du retournement temporel de façon matricielle. En effet, notons $h_{ij}(t)_{(i,j) \in \{1..N\}^2}$, les réponses impulsionnelles qui relient un réseau d'antennes émettrices, à un réseau d'antennes réceptrices, que l'on va supposer de même taille N . Pour chaque fréquence, on peut définir la matrice $\mathbf{H}(f)$, qui relie chaque signal émis sur les N antennes, noté $\mathbf{S}(f)$, au signal reçu noté $\mathbf{Y}(f)$:

$$\mathbf{Y}(f) = \mathbf{H}(f)\mathbf{S}(f) \quad (\text{II.4})$$

L'opération de retournement temporel consiste alors à conjuguer le vecteur $\mathbf{Y}(f)$, et à le renvoyer. Le signal reçu sur le réseau initial s'écrit alors :

$$\mathbf{Y}_{RT}(f) = {}^t \mathbf{H}(f) \mathbf{H}^*(f) \mathbf{S}^*(f) = {}^t (\mathbf{H}(f) \mathbf{H}^\dagger(f)) \mathbf{S}^*(f) \quad (\text{II.5})$$

Ici on reconnaît la grandeur $\mathbf{H}(f) \mathbf{H}^\dagger(f)$ qui est présente dans la définition de la capacité des systèmes MIMO, que nous appelons au LOA l'opérateur de retournement temporel. Le rang de cet opérateur définit le nombre de canaux indépendants que l'on peut créer dans un milieu donné. Il définit également le nombre de tache focales que l'on peut créer par retournement temporel, c'est à dire le nombre d'antennes d'un réseau de réception qu'il est possible d'adresser indépendamment par retournement temporel [51]. La formule de la capacité ne donne pas de technique pour tirer profit de la diversité spatiale d'un milieu, elle indique simplement une limite supérieure. Il est clair ici que le retournement temporel va tirer profit naturellement de cette diversité spatiale. Nous allons donc pouvoir communiquer en se servant du gain de multiplexage spatial que nous avons évoqué lors de la discussion sur les systèmes MIMO. Une base de N antennes pourra ainsi s'adresser à M utilisateurs, et nous serons alors dans la configuration MIMO-Mu décrite précédemment.

Le retournement temporel va également tirer profit des autres gains que nous avons définis pour les systèmes MIMO :

- Le fait d'utiliser plusieurs antennes en émission permet de bénéficier du gain en diversité fréquentielle, qui est une conséquence directe de la diversité spatiale. En effet, chacune des réponses impulsionnelles entre une antenne du réseau d'émission et l'un des utilisateurs va voir son spectre être creusé en divers endroits. Le fait d'adresser cet utilisateur avec un réseau a pour effet de moyennner ces creux, de façon à ce que le signal reçu présente une bande passante plus uniforme.
- Le retournement temporel est d'autre part, du point de vue du traitement du signal, un filtre adapté. Cette caractéristique permet de profiter également de l'effet de gain d'antenne. Ainsi, à l'émission, le retournement temporel réalise une formation de voie adaptée au milieu dans lequel on désire communiquer. Le rapport signal-sur-bruit est ainsi maximisé.
- Enfin, de la focalisation spatiale découle une autre particularité du retournement temporel appliqué aux télécommunications. Le signal étant plus élevé sur l'antenne visée que dans le reste du milieu, les interférences inter-utilisateurs sont minimisées. Cette propriété est une conséquence directe du gain de réduction d'interférences évoqué précédemment. Ainsi, le retournement temporel offre la possibilité de communiquer à plusieurs utilisateurs en même temps, sur la même bande passante, ce qui augmente l'efficacité spectrale, tout en réduisant les interférences que ceux-ci créent les uns pour les autres.

Nous développerons par la suite les avantages de la technique et en soulignerons également les inconvénients. Mais il nous reste au préalable à montrer l'intérêt du retournement temporel pour les communications large bande, et plus spécifiquement pour l'UWB. Nous avons montré que la réverbération des signaux UWB pose de nombreux problèmes et oblige un traitement en réception, à l'aide d'un récepteur de type RAKE. Ceci est en désaccord avec l'objectif visé quant à la simplicité du récepteur. En utilisant le retournement temporel lors de la phase d'émission, la réponse impulsionnelle est compressée en réception, pour donner une impulsion aussi brève que le signal se propageant dans le vide et à laquelle viennent s'ajouter des lobes secondaires. La détection des symboles s'en trouve simplifiée : le retournement temporel agit tel un précodeur temporel des signaux. Ainsi la tâche du récepteur est réduite : le retournement temporel a pour effet de déplacer la complexité du système de la réception à l'émission. Une telle technique permettrait donc de respecter le cahier des charges en ce qui concerne l'architecture

des récepteurs UWB. Les lobes secondaires créés vont tout de même provoquer des interférences inter-symboles que nous quantifierons par la suite.

Une autre intérêt du retournement temporel pour les communications UWB réside dans sa simplicité de mise en oeuvre. En effet, nous avons vu que 4 des 5 techniques envisagées jusqu'à ce jour utilisent des séquences pseudo-aléatoires afin de réaliser de la communication multi-utilisateurs. Cette approche nécessite un corrélateur en réception, afin de décoder l'information. En utilisant le retournement temporel, c'est le milieu lui même qui fabrique le code car les réponses impulsionnelles vers différents utilisateurs sont pseudo-orthogonales, à condition que le milieu soit suffisamment complexe. Ainsi si une information est émise à destination d'un utilisateur, les autres ne reçoivent pas le message, grâce au codage réalisé par le milieu de propagation. Ceci peut encore considérablement simplifier le type de récepteur à employer en UWB.

Enfin, rappelons que le retournement temporel étant un filtre adapté, celui-ci permet de maximiser le rapport signal-sur-bruit lors de l'émission d'information vers un utilisateur. Cette propriété peut avoir un très grand potentiel dans le domaine de l'UWB. En effet, nous venons de voir que le spectre alloué à l'UWB impose des puissances très limitées. Cette limite risque d'être assez problématique notamment à cause des autres systèmes de communication qui émettent sur la même bande passante à des puissances bien supérieures. Le gain en énergie dû à la focalisation spatiale et à la compression temporelle, conséquences du retournement temporel, va donc être particulièrement utile en UWB. Nous verrons par la suite comment quantifier ce gain en fonction des différents paramètres physiques du système étudié.

Ces différentes raisons nous ont guidés vers une étude des télécommunications large bande par retournement temporel. Nous allons, dans la partie suivante, décrire la première expérience de télécommunications par retournement temporel en ultrasons.

II.1.4 Première expérience en ultrasons

La première expérience de télécommunications par retournement temporel a été réalisée à petite échelle, en utilisant des ultrasons dans l'eau, par Arnaud Derode et ses collaborateurs [52]. En effet, à cette période, aucune expérience de retournement temporel n'avait été réalisée dans le domaine des ondes électromagnétiques car le matériel disponible ne le permettait pas. Le but était de démontrer l'intérêt du retournement temporel en configuration MIMO-Mu, en comparant des communications en milieu homogène et complexe. L'expérience a donc été menée

dans une cuve remplie d'eau qui constitue le milieu homogène de référence. La base, antenne d'émission, est constituée de 23 transducteurs piézo-électriques dont la fréquence centrale vaut 3.2 MHz, et la bande passante 50% à -6 dB. Ce réseau est utilisé en mode retournement temporel pour envoyer simultanément, et sur la même bande passante, cinq séquences de 2000 bits d'information à cinq utilisateurs distants de 4 longueurs d'onde. Ces utilisateurs sont placés en champ lointain de la base d'émission. Afin d'étudier l'effet du désordre sur la communication, l'expérience est réalisée, en plaçant une forêt de tiges métalliques analogue à celles qui ont été décrites dans la première partie de ce manuscrit (fig II.4).

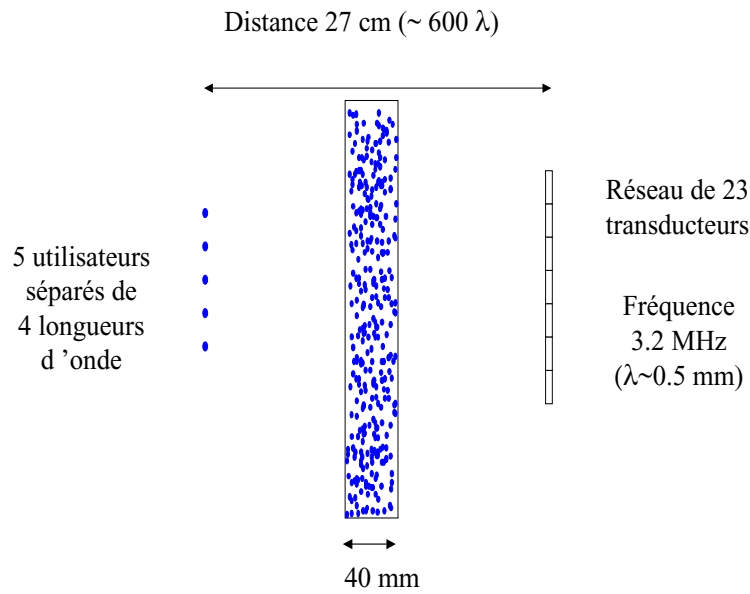


FIG. II.4 – Dispositif expérimental (Source [52]).

La transmission d'information s'effectue en deux temps. Dans une première étape, chaque utilisateur émet successivement une impulsion brève ($3/2$ périodes à 3.2 MHz) à travers le milieu. Chacun des transducteurs qui composent la base reçoit alors un signal temporel qui constitue la réponse impulsionnelle entre celui-ci et l'utilisateur. Lorsque la forêt de tiges n'est pas présente, c'est-à-dire que le milieu est homogène, la réponse est un signal bref (Fig. II.5.a). La réponse à travers la forêt de tiges est en revanche très allongée dans le temps (Fig. II.5.b), car elle résulte de la superposition des contributions des multiples chemins empruntés par l'onde lors de sa propagation. Chacun des signaux temporels est enregistré et cinq collections de 23 réponses sont ainsi numérisées puis retournées temporellement. Si l'une des collections est réémise dans le milieu, l'onde recrée revit son passé pour venir converger vers sa source en une impulsion brève positive (Fig. II.5.c). Ce signal va constituer le symbole élémentaire de la communication : il transmettra un "1". Si l'on souhaite transmettre un zéro, afin de réaliser une transmission

binaire, il suffit de réémettre une collection de signaux de signe opposé, de manière à recevoir un “-1”. Si l’on veut transmettre à un utilisateur un message binaire constitué de ces deux symboles, il suffit de convoluer la collection de réponses impulsionnelles qui lui correspond par la trame de symboles que l’on veut envoyer pour réaliser une modulation de type BPSK. Enfin, pour transmettre aux cinq utilisateurs simultanément, il reste à sommer les signaux fabriqués pour chaque utilisateur et le tout est envoyé par la base.

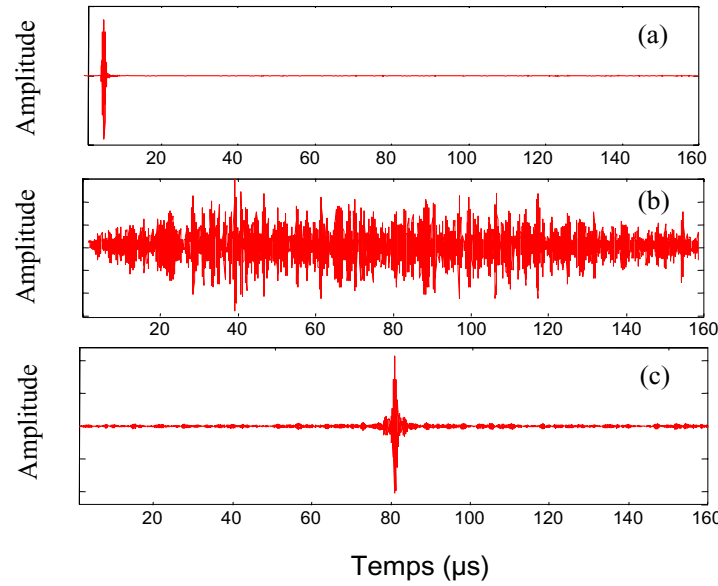


FIG. II.5 – Signal transmis par un utilisateur à une antenne de la base dans l’eau (a), à travers la forêt de tiges (b). Impulsion reçue par l’un des utilisateurs après propagation de l’onde retournée temporellement (c). (Source [52]).

Lorsque cette expérience est réalisée dans le milieu homogène que constitue l’eau, le taux d’erreur de la transmission est énorme : 20%. Au contraire, si l’on transmet ces informations à travers la forêt de tiges, on note une seule erreur sur les 10000 bits envoyés. La raison en est la suivante. Dans l’eau, la qualité de la focalisation spatiale autour d’un utilisateur est conditionnée par l’ouverture angulaire de l’antenne à retournement temporel. La tache focale créée est alors donnée par la figure de diffraction d’une ouverture angulaire de largeur $\lambda F/D$, avec λ la longueur d’onde des signaux, F la distance base-utilisateur et D la largeur de la base. Comme on peut l’apprécier sur la figure II.6.b, la tache focale mesure 16 mm à mi-hauteur. Elle n’est donc pas suffisamment fine pour éviter que les informations transmises aux différents utilisateurs, qui sont distants de 2 mm, ne se recouvrent. Dans le cas du milieu désordonné, lorsque l’onde retournée temporellement est réémise, la forêt se comporte vis-à-vis de utilisateurs comme une source secondaire d’ouverture plus grande que celle de l’antenne proprement dite.

Ainsi, la tache de focalisation est beaucoup plus fine (Fig II.6.a), et les messages envoyés aux divers utilisateurs ne se recouvrent pas.

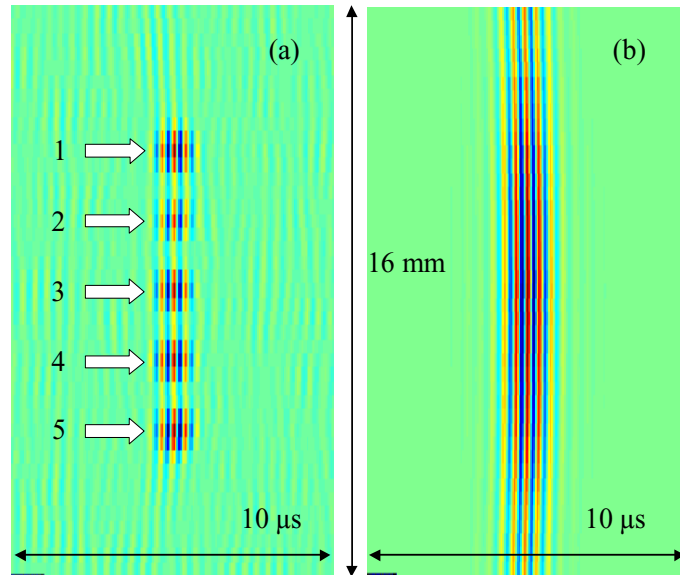


FIG. II.6 – Taches focales obtenues par retournement temporel à travers la forêt de tiges (a), et en milieu homogène (b). (Source [52]).

L'antenne à retournement temporel profite donc du désordre du milieu pour améliorer le débit d'information par rapport à la situation où le milieu est homogène. Le retournement temporel permet d'exploiter les canaux de communication indépendants. Afin de quantifier le nombre de canaux indépendants dont nous pouvons profiter dans un tel milieu, l'expérience suivante a été réalisée. Cette fois, la base est composée de 40 antennes piézo-électriques, et 40 récepteurs sont placés à l'endroit où il y en avait 5. L'espacement entre les éléments d'émission, ainsi qu'entre les divers utilisateurs, est cette fois de l'ordre d'une longueur d'onde. La distance entre la base et les utilisateurs est de 27 cm. Ainsi, une matrice (40 par 40) de réponses impulsionnelles est enregistrée, en milieu homogène et en plaçant la forêt de tiges. Une transformée de Fourier est ensuite appliquée à chaque réponse impulsionnelle. À chaque fréquence, ces matrices de transfert sont alors décomposées en valeur singulières. La figure II.7 représente le nombre de valeurs singulières obtenues, ainsi que leur amplitudes respectives, pour chaque fréquence au sein de la bande passante utilisée.

On observe beaucoup plus de valeurs singulières dans le cas du milieu complexe que dans le cas du milieu homogène. En fixant un seuil arbitraire à -32 dB, on compte 6 valeurs singulières dans l'eau, et 34 à travers la forêt de tiges. Nous avons ici une démonstration expérimentale claire de l'effet du désordre sur le nombre de canaux de communications indépendants dont on peut

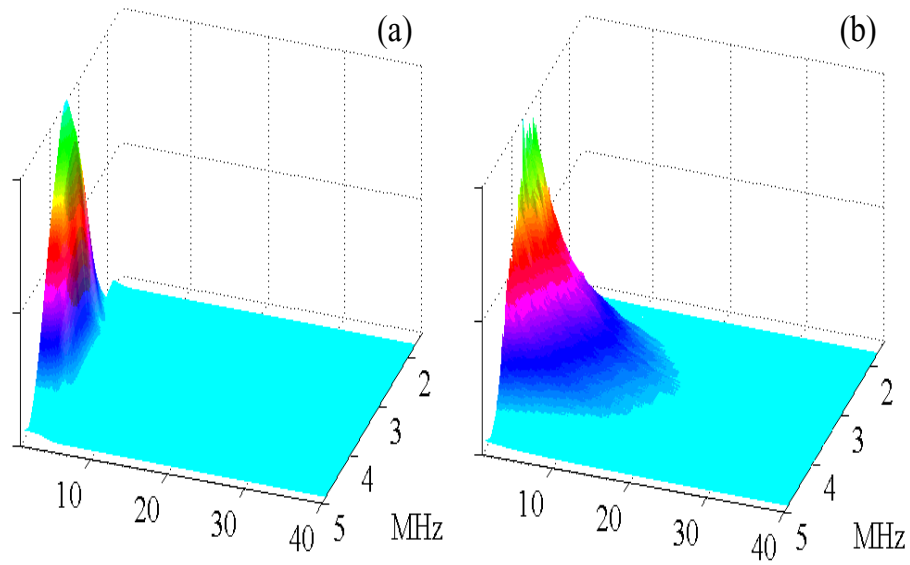


FIG. II.7 – Valeurs singulières obtenues à travers la forêt de tiges (a), et à travers le milieu homogène (b) (Source [52]).

disposer en utilisant un système MIMO. Cette expérience préliminaire, publiée en 2003 dans la revue *Physical Review Letters* par Arnaud Derode et ses collaborateurs [52], prouve que le retournement temporel est un moyen simple de se servir de ces canaux de communication pour augmenter le débit d'information entre une base et un volume donné dans un milieu complexe.

Cette partie s'est attachée à présenter les deux techniques prometteuses qui sont étudiées actuellement afin de faire face à l'augmentation constante des débits d'information. Nous avons vu que l'une de ces techniques est fondée sur les systèmes multi-antennes MIMO, qui permettent de tirer profit de la complexité d'un milieu, afin de créer des canaux de communication indépendants, de manière à multiplexer spatialement l'information. La deuxième technique évoquée s'appuie sur des systèmes à très large bande passante (UWB), grâce auxquels on espère atteindre des débits supérieurs à 100 MB/s. Ces systèmes fonctionnent avec des impulsions très courtes, qui subissent les effets de la propagation comme les réflexions sur les obstacles rencontrés. De plus, nous avons souligné que la dégradation des signaux utilisés dans un environnement complexe est d'autant plus importante que ceux-ci sont brefs. Nous avons alors proposé d'utiliser le retournement temporel qui permet de lutter contre de nombreux problèmes inhérents à ces deux techniques. En effet, nous avons souligné que celui permet de tirer parti de la complexité d'un milieu, sans avoir recours à aucun traitement complexe. D'autre part il permet de compenser la réverbération subie par un signal bref, tout en assurant une énergie maximale au niveau de l'utilisateur visé. En termes de télécommunications, le retournement temporel peut donc être vu comme un précodeur des signaux qui permet de déplacer la complexité du récepteur à l'émetteur. Ces divers aspects en font un candidat idéal pour la réalisation de systèmes MIMO-Mu à très large bande passante. L'expérience réalisée par Arnaud Derode et ses collaborateurs avant le début de cette thèse prouve l'intérêt du retournement temporel dans ce type de communications. Cependant, il est indispensable de compléter cette expérience préliminaire par une étude plus quantitative du retournement temporel appliqué aux communications UWB en environnement "indoor" : c'est le but de la partie suivante. Nous allons voir que si le retournement temporel présente nombre d'avantages, il possède quelques inconvénients que nous quantifierons en fonction des paramètres physiques des expériences réalisées.

II.2 Télécommunications MISO Ultra large bande en environnement indoor à petite échelle

II.2.1 Principe des expériences et dispositif expérimental

Dans la partie précédente, nous avons décrit les premières expériences de télécommunications par retournement temporel réalisées dans le domaine des ultrasons. Cependant, un reproche peut être fait concernant le réalisme du type de milieu utilisé. En effet, la forêt de tiges constitue le milieu complexe de référence étudié au laboratoire, mais elle est assez éloignée des environnements auxquels est confrontée la communauté des télécommunications. Nous avons réalisé des expériences en utilisant un modèle plus réaliste. Limités au domaine ultrasonore pour des raisons matérielles, nous avons procédé à des expériences à échelle réduite, dans un aquarium. Le milieu est une reproduction à échelle centimétrique d'un étage d'un immeuble constitué de pièces, comme le montre la figure II.8.

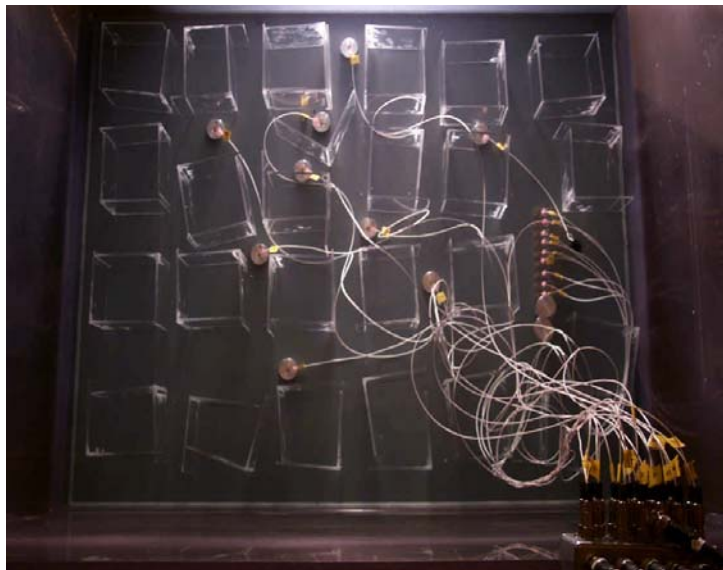


FIG. II.8 – Photo du dispositif expérimental utilisé.

Les pièces sont faites de quatre “murs” en plexiglass de 4 mm d’épaisseur. Chacun de ces carrés mesure 10 cm de côté et 10 cm de haut. Il sont séparés par des couloirs d’une largeur de 4 cm. Cette configuration correspond à un environnement “indoor” réduit à une échelle 1/100. Les ondes électromagnétiques décimétriques sont ainsi remplacées par des ondes millimétriques ultrasonores. De petits piézo-électriques quadripolaires font office d’antennes omnidirectionnelles telles que celles utilisées dans les réseaux locaux. La fréquence centrale de ces

transducteurs est de 1.1 MHz, avec une bande passante de 30%, ce qui correspondrait à un système électromagnétique Ultra Large Bande dont la fréquence porteuse serait de 2 GHz, et la bande passante de 700 MHz. Sur la droite de la photo (Fig. II.8), on voit une base faite de 8 antennes piézo-électriques. Ces dernières sont posées sur un socle en plastique, et séparées de quelques longueurs d’onde. Concernant les utilisateurs, les antennes, collées sur des socles cylindriques en métal, sont distribuées aléatoirement dans les couloirs du bâtiment. Chacun des piézo-électriques est relié à une électronique d’émission et de réception, qui permet de numériser les signaux reçus, ainsi que d’émettre des signaux arbitraires. Dans un premier temps, on réalise l’acquisition d’une collection de 64 réponses impulsionnelles entre les 8 antennes de la base et les 8 utilisateurs. Pour enregistrer ces réponses, chaque utilisateur émet successivement une impulsion d’une durée de $3\ \mu\text{s}$, les 8 antennes de la base enregistrent simultanément les réponses impulsionnelles à une fréquence d’échantillonnage de 32 MHz. La figure II.9 montre un exemple typique d’un tel signal.

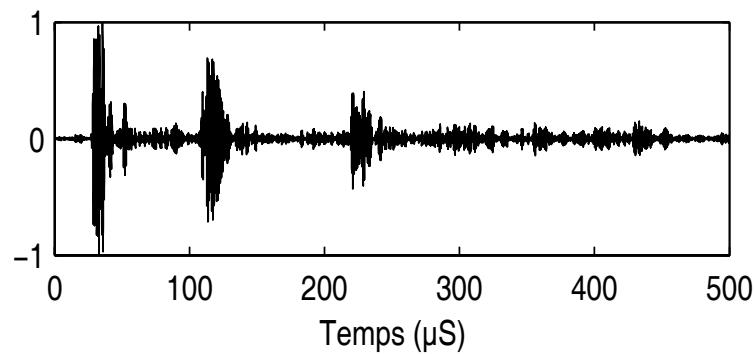


FIG. II.9 – Réponse impulsionnelle typique entre un utilisateur et une antenne de la base dans l’environnement indoor petite échelle.

Cette réponse est très allongée temporellement par rapport au signal initialement émis. Cependant, contrairement au cas de la forêt de tiges, à travers laquelle on enregistre une coda qui ressemble à un bruit aléatoire, on observe ici des paquets d’impulsions assez distincts. La dépendance temporelle de cette réponse s’explique par la géométrie du milieu : on peut y distinguer l’effet de plusieurs murs, sous forme de “cluster”, des tranches de signal plus énergétiques que les autres. Ce type de réponse est couramment modélisé dans la littérature par la modèle de Saleh-Valenzuela [53].

Dans ce type de milieu assez réaliste, les réponses vont cependant être très différentes d’une antenne à l’autre. Les énergies de chaque signaux enregistrés varient en effet en fonction de la position de l’utilisateur par rapport à l’antenne de la base considérée : deux capteurs l’un

en face de l'autre vont créer une réponse d'amplitude très élevée, mais de courte durée. En revanche, la réponse entre deux antennes qui ne se "voient" pas va être très allongée mais d'une amplitude faible. Afin de tirer parti de toute l'information spatiale à laquelle nous avons accès, on normalise alors la collection de signaux enregistrés. Pour ce faire, chaque forme temporelle est divisée par son intensité moyenne. Cette façon de procéder n'est pas optimale en terme énergétique, mais permet de profiter au mieux de tous les trajets multiples.

Durant cette série d'expériences, nous avons voulu, en tout premier lieu, éclaircir un point : la différence entre filtrage adapté à la réception et retournement temporel. En effet, il est connu qu'utiliser un corrélateur qui réalise un filtrage adapté à la réception dans un environnement complexe est un moyen idéal pour "capturer" l'énergie d'une impulsion large bande qui a subi des trajets multiples [54]. Ainsi, si l'on écrit la réponse entre un utilisateur placé en $\mathbf{r}$ et une antenne de la base placée en $\mathbf{a}_i$ sous la forme $h(\mathbf{a}_i \rightarrow \mathbf{r}, t)$, le signal reçu $h(\mathbf{a}_i \rightarrow \mathbf{r}, t) \otimes s(t)$ (où $s(t)$ est le signal émis) peut être décodé par filtrage adapté à la réception :

$$\mathcal{R}_{FA}(t) = [h(\mathbf{a}_i \rightarrow \mathbf{r}, t) \otimes s(t)] \otimes h(\mathbf{a}_i \rightarrow \mathbf{r}, -t) \quad (\text{II.6})$$

Le terme $h(\mathbf{a}_i \rightarrow \mathbf{r}, t) \otimes h(\mathbf{a}_i \rightarrow \mathbf{r}, -t)$ présente un maximum à $t = 0$, et ce maximum est d'autant plus grand que le nombre de trajets multiples est élevé. Dans le cas d'un milieu dont la réponse en fréquence est plate, l'autocorrélation tend vers un Dirac, et le signal est reçu parfaitement : $\mathcal{R}_{FA}(t) = s(t)$.

Si l'on utilise le retournement temporel, c'est le signal $h(\mathbf{a}_i \rightarrow \mathbf{r}, t)$ retourné temporellement qui est convolué par le message $s(t)$ et que l'on envoie dans le milieu. Les ondes produites subissent alors à nouveau les trajets multiples et le signal reçu s'écrit :

$$\mathcal{R}_{RT}(t) = [h(\mathbf{a}_i \rightarrow \mathbf{r}, -t) \otimes s(t)] \otimes h(\mathbf{a}_i \rightarrow \mathbf{r}, t) \quad (\text{II.7})$$

Les signaux reçus en utilisant ces deux techniques sont donc strictement équivalents si le système est linéaire², et cette équivalence est à l'origine d'une certaine confusion entre ces méthodes.

Nous allons voir que dans le cas de systèmes MISO, ou de façon équivalente de systèmes MIMO-Mu, ce dernier va se montrer plus performant. Pour ce faire, dans le cas du filtrage adapté à la réception, nous allons raisonner de façon très simpliste. En effet, nous allons considérer que

²Nous verrons dans la dernière partie que cette remarque peut avoir de l'importance, notamment en ce qui concerne la réception du signal.

chaque antenne de la base émet le même message, c'est-à-dire qu'aucun multiplexage n'est réalisé à l'émission. Dans une telle configuration, si la base cherche à envoyer le signal $s(t)$ à un utilisateur, le signal reçu prendra alors la forme $\sum_i h(\mathbf{a}_i \rightarrow \mathbf{r}, t) \otimes s(t)$, où $\mathbf{a}_i$ est la position de l'antenne i dans la base. Afin de profiter des trajets multiples, et pour réaliser un filtrage adapté, la façon la plus élémentaire de faire est de corrélérer ce signal par $\sum_i h(\mathbf{a}_i \rightarrow \mathbf{r}, t)$ (acquis dans une première étape d'apprentissage au cours de laquelle $s(t)$ est un Dirac). On obtient alors le résultat suivant :

$$\mathcal{R}_{FA}(t) = \left[\sum_i h(\mathbf{a}_i \rightarrow \mathbf{r}, t) \otimes s(t) \right] \otimes \left[\sum_j h(\mathbf{a}_j \rightarrow \mathbf{r}, -t) \right] \quad (\text{II.8})$$

Dans le cas du retournement temporel, le message à transmettre est convolué, pour chaque antenne de la base, par la réponse impulsionnelle retournée temporellement (acquise également lors d'une première phase d'apprentissage), avant émission. Puis, les ondes émises se propagent en sens inverse pour venir converger vers l'utilisateur. Le signal reçu par l'utilisateur prend alors la forme :

$$\mathcal{R}_{RT}(t) = \left[\sum_i h(\mathbf{a}_i \rightarrow \mathbf{r}, -t) \otimes h(\mathbf{a}_i \rightarrow \mathbf{r}, t) \right] \otimes s(t) \quad (\text{II.9})$$

Ainsi le retournement temporel réalise naturellement la somme des autocorrélations temporelles des M réponses impulsionnelles entre chaque antenne de la base et l'utilisateur. Au contraire du cas SISO que nous avons évoqué précédemment, les signaux correspondant à un filtrage adapté en réception et à son analogue à l'émission, le retournement temporel, sont différents. A présent $\mathcal{R}_{FA}(t)$ contient $M(M - 1)$ termes supplémentaires qui correspondent aux inter-corrélations entre les réponses impulsionnelles $h(\mathbf{a}_i \rightarrow \mathbf{r}, t)$ et $h(\mathbf{a}_j \rightarrow \mathbf{r}, t)$, pour des i et j distincts. Ces inter-corrélations augmentent les interférences inter-symboles, ce qui dégrade la communication par rapport au cas du retournement temporel. Il est entendu que le filtrage adapté en réception simpliste que nous avons réalisé n'est pas envisageable : la déviation standard du bruit résultant des $M(M - 1)$ inter-corrélations est approximativement égale à l'amplitude des M autocorrélations. Certaines techniques préconisent alors de réaliser un multiplexage en émettant un flot de données différent par chaque antenne de la base. Des récepteurs de type RAKE en parallèle [55] sont alors utilisés, qui décodent chaque séquence d'information. Cette technique est plus intéressante car, si elle ne permet pas de lutter contre les inter-corrélations, elle augmente la quantité d'information transmise. Mais on se doute que quel que soit l'ajout de

complexité au récepteur, l'effet de ces signaux parasites ne sera jamais annulé complètement. Dès lors, il semble que le retournement temporel soit un outil idéal pour l'utilisation de systèmes MISO, où MIMO-Mu, dans des environnements complexes et réverbérants.

A titre d'exemple, la figure II.10.a montre une impulsion compressée par retournement temporel avec une base faite d'une unique antenne (on obtiendrait le même résultat par filtrage adapté à la réception). Le niveau de lobes est assez élevé, ce qui va causer des erreurs de transmission. Les figures II.10.b et II.10.c montrent la même impulsion pour des miroirs à retournement temporel comportant respectivement 4 et 8 voies. Le niveau de lobes a bien sûr décru proportionnellement à la racine du nombre d'antennes, réduisant ainsi les interférences inter-symboles lors d'un processus de communication. Nous n'avons pas représenté de résultat similaire en utilisant le filtrage adapté à la réception car celui-ci n'a pas d'intérêt : les lobes secondaires ne sont pas atténués, même si l'on augmente le nombre d'antennes de la base.

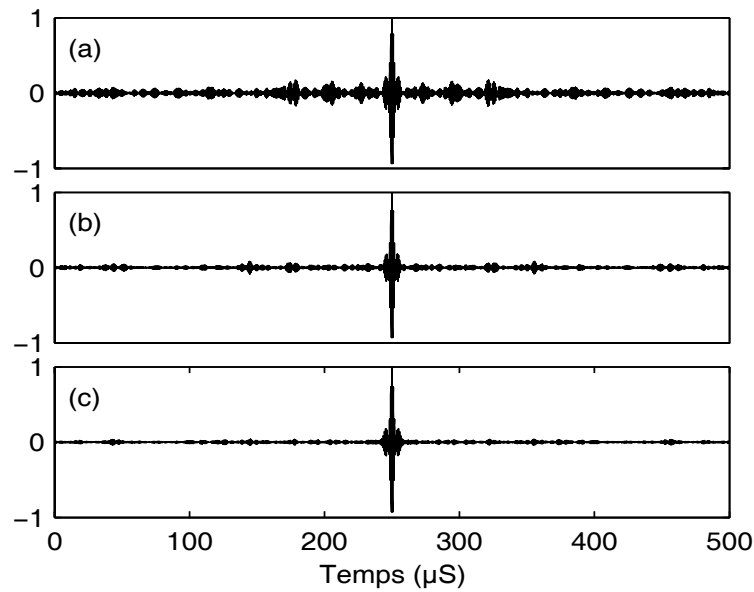


FIG. II.10 – (a) Impulsion compressée sur un utilisateur avec un miroir à 1 voie, (b) avec 4 voies, (c) avec 8 voies.

Dans les sections suivantes nous quantifierons les résultats précédents au travers d'expériences de télécommunication. Nous y verrons que réaliser du retournement temporel, c'est-à-dire du filtrage adapté à l'émission présente de nombreux autres avantages, notamment pour les systèmes MIMO-Mu. En effet, utiliser le multiplexage spatial de l'information est une possibilité qu'offre le retournement temporel de façon naturelle, grâce à la focalisation spatiale qu'il permet.

II.2.2 Résultats expérimentaux

Pour tous les résultats qui vont être présentés par la suite, nous n'avons pas procédé à des communications par retournement temporel "in situ". Au lieu de cela dans un premier temps la matrice des réponses impulsionnelles est enregistrée, puis les signaux $\mathcal{R}_{FA}(t)$ et $\mathcal{R}_{RT}(t)$ sont calculés en utilisant les formules II.8 et II.9, pour chaque configuration étudiée. Nous avons décidé de procéder ainsi en raison du temps de chargement élevé des mémoires des électroniques qui ne permet pas de réaliser ces mesures dans un temps raisonnable. Afin d'estimer le Taux d'Erreur Binaire (TEB) moyen, une séquence de 1000 bits est générée de façon pseudo-aléatoire. Cette séquence est modulée en utilisant une modulation BPSK (Binary Phase Shift Keying), de type antipodale, c'est à dire qu'une impulsion positive représente un "1" et une impulsion négative représente un "0". Puis, $\mathcal{R}_{FA}(t)$ et $\mathcal{R}_{RT}(t)$ sont démodulés numériquement en utilisant une minimisation de la distance quadratique sur le diagramme des constellations. Cette opération est répétée jusqu'à ce que la variance de l'estimation du TEB soit inférieure à une limite fixée à 1%. Durant cette série d'expériences, nous avons fait varier le nombre d'antennes du miroir à retournement temporel, afin d'en étudier l'impact sur la qualité de la communication. Ainsi, pour un nombre d'émetteurs utilisés inférieur au nombre d'antennes total de la base, nous avons moyenné les TEB obtenus pour diverses configurations en choisissant à chaque fois M antennes parmi les 8 utilisables. De même, chaque mesure est un TEB moyen sur les 8 utilisateurs disponibles. Ces diverses estimations permettent d'obtenir un taux d'erreur binaire qui dépend beaucoup moins de la configuration émetteur/récepteur que des caractéristiques du milieu.

Influence du débit d'information

Pour les premières mesures réalisées, nous avons d'abord fait varier le débit d'information envoyé à un utilisateur, en travaillant avec un bruit externe négligeable. Pour ce faire, on change simplement l'intervalle de temps entre deux symboles consécutifs. Ainsi, lorsque l'on utilise le retournement temporel, la trame de symboles, qui sont espacés d'un intervalle de temps δt , est convoluée à la collection de réponses impulsionnelles correspondant à chaque utilisateur. Grâce à l'invariance par translation dans le temps du milieu de propagation, les symboles sont focalisés sur l'utilisateur voulu séparés de l'intervalle de temps δt . La modulation étant de type binaire, le débit d'information est égal à l'inverse de δt . Le taux d'erreur binaire est estimé pour des intervalles de temps allant de $0.5 \mu s$, soit un débit de 2 MBs/s, à $10 \mu s$ ou encore un débit de 100 KBs/s. Sur la figure II.11, nous avons représenté le résultat de ces mesures, en

logarithme décimal, lorsqu'un filtrage adapté est réalisé à la réception, et ce, en utilisant une base faite de 1, 4 et 8 antennes.

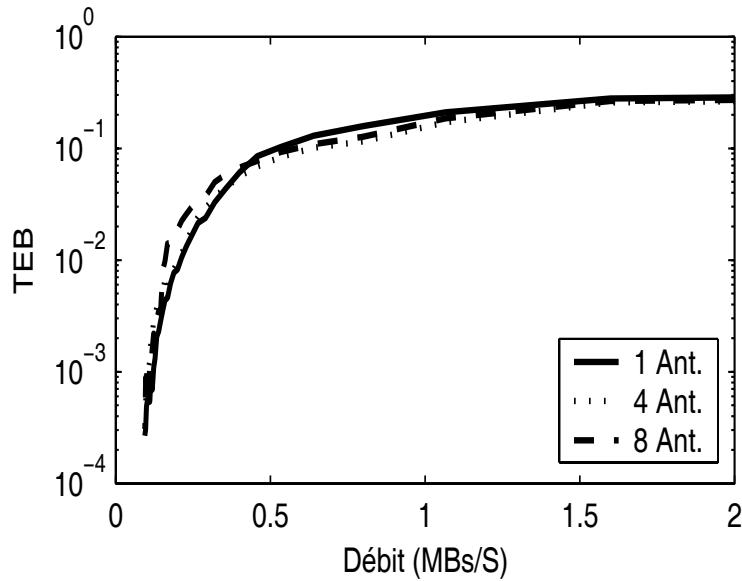


FIG. II.11 – Taux d'erreur binaire en fonction du débit d'information, pour le filtre adapté à la réception. Le miroir est constitué de 1, 4 et 8 antennes.

Il est clair que le fait d'ajouter des capteurs à notre base n'est d'aucune utilité dans le cas présent : cela n'améliore pas la qualité de la transmission. Cette conclusion était aisément envisageable au travers des formules II.8 et II.9, mais peut également se comprendre autrement. En effet, le même signal est émis par toutes les antennes de la base : celle-ci agit donc comme une antenne unique qui serait plus large. Comme de plus l'énergie émise est constante, ajouter des antennes est sans intérêt. Bien que l'addition d'antennes augmente l'amplitude d'un bit d'information, les inter-corrélations entre les signaux des différentes antennes augmentent de la même façon. Ainsi les interférences inter-symbole ne décroissent pas : la qualité de la transmission qui ne dépend que de ces interférences, car le bruit externe est nul, n'est donc pas améliorée. Comme le montre cette courbe, la seule façon de faire diminuer le taux d'erreur de la communication est de séparer les symboles, jusqu'à ce qu'ils ne se perturbent plus mutuellement.

Sur la figure II.12, nous représentons le résultat de la même mesure, réalisée en utilisant le retournement temporel. Comme on pouvait s'y attendre, la courbe obtenue en utilisant une seule antenne de la base est exactement la même que celle obtenue par filtrage adapté à la réception. En revanche, lorsque l'on augmente le nombre d'antennes, les résultats sont améliorés de façon très significative. Ceci est une conséquence du fait que le retournement temporel ne génère pas d'inter-corrélations entre les antennes de la base. Ainsi l'amplitude du pic croît en

fonction de la racine du nombre d'antennes, quand le niveau des interférences inter-symboles reste constant car nous profitons de la diversité spatiale.

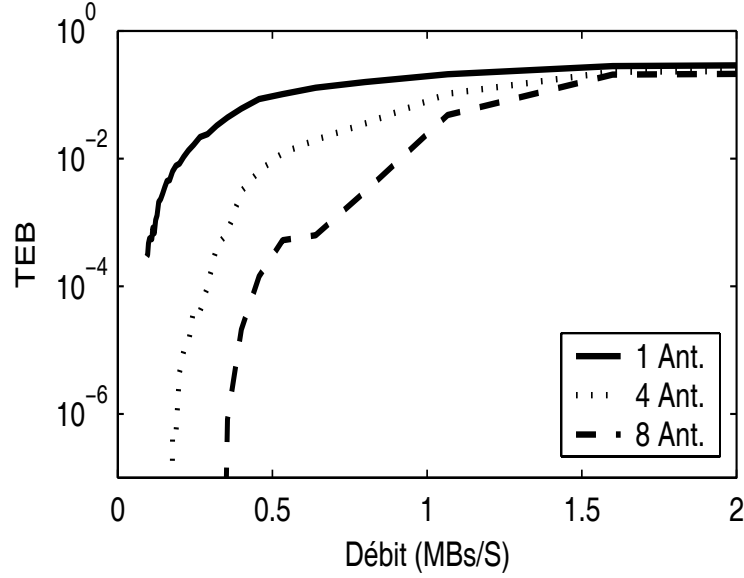


FIG. II.12 – Taux d'erreurs binaire en fonction du débit d'information, pour le retournement temporel. Le miroir est constitué de 1, 4 et 8 antennes.

En procédant de cette manière, il n'est pas nécessaire si l'on veut obtenir un taux d'erreur faible, de séparer les symboles de façon aussi brutale que lorsque l'on réalise un filtrage adapté à la réception.

Influence du bruit externe

Nous avons également étudié la dégradation de la transmission causée par l'ajout d'un bruit externe. Lors de ces mesures, les signaux $\mathcal{R}_{FA}(t)$ et $\mathcal{R}_{RT}(t)$ sont normalisés en énergie. Nous rappelons que l'énergie émise étant constante, l'énergie du pic de retournement temporel reste constante quel que soit le nombre d'antenne utilisé. Ceci permet de s'affranchir du gain d'antennes que procure un réseau d'émission, pour ne se concentrer que sur le gain de diversité spatiale. Le RSB est donc défini indépendamment de la taille du MRT utilisé. Le bruit ajouté est de type blanc gaussien (AWGN). La communication se déroule de la même manière que précédemment, mais cette fois nous fixons le débit à 188 KBs/s, c'est-à-dire que nous nous plaçons au premier quart des courbes représentées dans les figures II.11 et II.12. Les mesures ont été réalisées pour une déviation standard de bruit allant de 0.1 à 0.5. La figure II.13 montre le résultat des estimations du TEB, en logarithme décimal, pour des tailles de base de 1, 4 et 8 antennes.

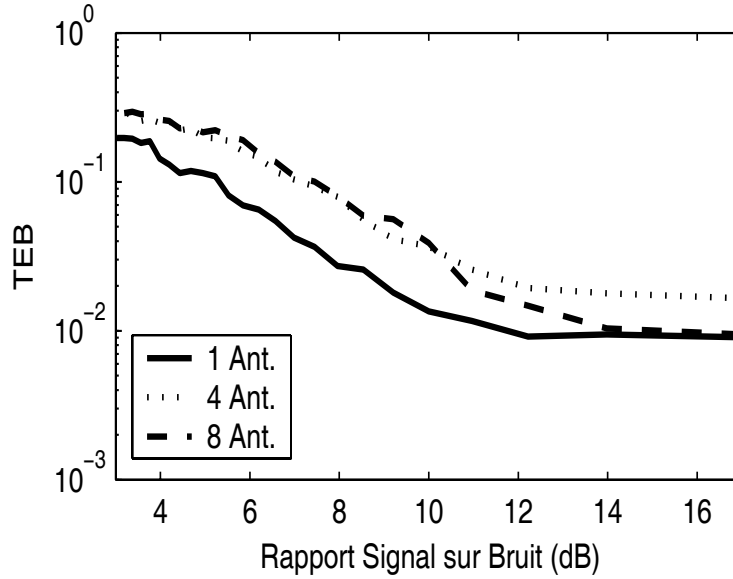


FIG. II.13 – Taux d’erreurs binaire en fonction du bruit externe, pour le filtre adapté à la réception. Le miroir est constitué de 1, 4 et 8 antennes.

Comme on pouvait de nouveau s’y attendre, l’ajout d’antennes ne sert à rien dans ce cas, et la qualité de la transmission est à nouveau limitée par les interférences inter-symboles, qui restent élevées, bien que le débit soit relativement faible. Au contraire, lorsque l’on utilise le retournement temporel, le taux d’erreur binaire est d’autant meilleur que le nombre d’antennes dans le miroir à retournement temporel est grand, même en présence de bruit externe. En effet, comme le montre le figure II.14, dès que le RSB dépasse 6 dB, le TEB décroît de façon très rapide. Ceci est à nouveau très aisé à comprendre. Quand le bruit externe est plus faible que les interférences inter-symboles, la qualité de la transmission est uniquement limitée par ces dernières.

On comprend dès lors que l’ajout d’antennes améliore la robustesse de la communication, celui-ci étant à l’origine d’une diminution des lobes secondaires dus au retournement temporel. Cependant, pour faire ces mesures, nous avons normalisé les signaux $\mathcal{R}_{FA}(t)$ et $\mathcal{R}_{RT}(t)$ en énergie de façon à étudier uniquement le gain en diversité spatiale que procure un ajout d’antenne dans un MRT. Cependant cette normalisation ne prend pas en compte le gain en énergie que procure le retournement temporel et qui est lié à la compression temporelle des réponses impulsionnelles. Ceci est pourtant un paramètre crucial en télécommunications et nous l’étudierons dans la dernière partie de ce manuscrit.

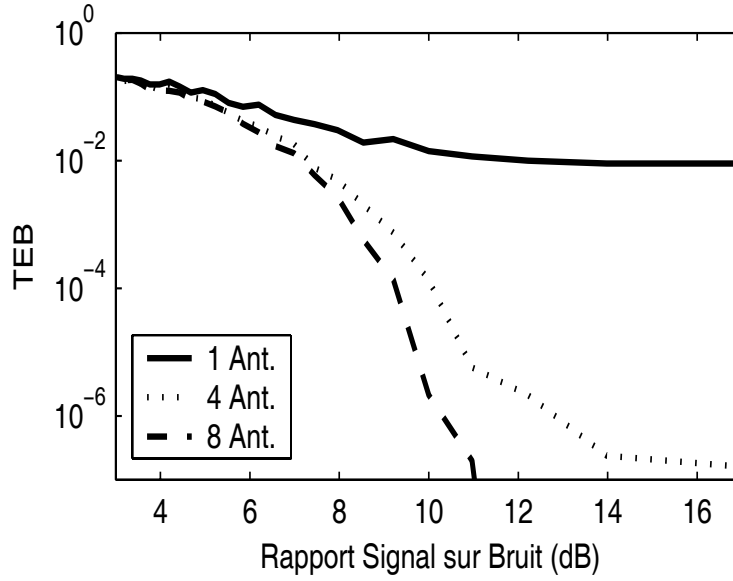


FIG. II.14 – Taux d’erreurs binaire en fonction du bruit externe, pour le retournement temporel. Le miroir est constitué de 1, 4 et 8 antennes.

II.2.3 Modélisation et discussion

Après avoir présenté de façon qualitative les mesures que nous avons réalisées, nous allons à présent nous attacher à comprendre l’effet sur le retournement temporel des divers paramètres physiques importants du système. Pour cela, nous nous limiterons au cas où le bruit externe est nul, celui-ci sera lui pris en compte dans la dernière partie de ce manuscrit. Nous avons pour l’instant représenté le logarithme décimal du taux d’erreur binaire en fonction du débit d’information, qui est un paramètre très utile du point de vue de la communication. Mais cette représentation masque une propriété intéressante, comme le montre la figure II.15, sur laquelle le logarithme décimal du taux d’erreur binaire est tracé en fonction de l’intervalle de temps entre deux symboles consécutifs δt , et ce pour des tailles de base de 1, 2, 4 et 8 antennes.

On remarque sur ces courbes que le logarithme du TEB décroît de façon linéaire en fonction de l’espacement entre les symboles, c’est-à-dire en fonction de l’inverse du débit. Si, de plus, on mesure la pente des droites obtenues pour chacune des bases étudiées, de 1 à 8 antennes, on s’aperçoit que ces coefficients directeurs évoluent eux aussi de façon linéaire en fonction du nombre M de capteurs utilisés pour la transmission de l’information (Fig. II.16). Ainsi, le logarithme du TEB évolue, lorsque l’on utilise le retournement temporel, de façon linéaire en fonction de M et de δt . Au contraire, dans le cas du filtrage adapté, le TEB est indépendant du nombre d’antennes présentes dans la base. Les résultats de ces mesures expérimentales sont

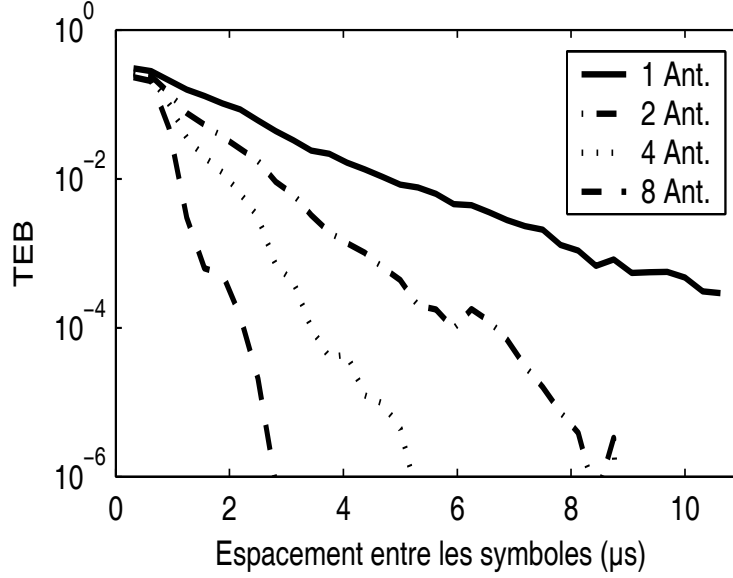


FIG. II.15 – Taux d’erreurs binaire pour un miroir constitué de 1, 2, 4 et 8 antennes.

confirmés par l’approche théorique simple suivante. En théorie du signal, il est connu que la probabilité d’erreur d’une modulation BPSK antipodale, dont le taux d’erreur binaire est un estimateur, est reliée au rapport signal-sur-bruit par la formule suivante [56] :

$$P_e = Q\left(\sqrt{\frac{2E_b}{N_0}}\right) \quad (\text{II.10})$$

où E_b est l’énergie d’un bit d’information, N_0 l’énergie du bruit, et Q la fonction de Marcum, définie en fonction de la fonction d’erreur comme :

$$Q(x) = \frac{1}{2\pi} \int_x^\infty \exp(-u^2/2) du = \frac{1}{2} \left(1 - \operatorname{erf}\left(\frac{x}{\sqrt{2}}\right)\right) \quad (\text{II.11})$$

Ici, le bruit est uniquement du aux interférences inter-symbole créées par le retournement temporel. Par ailleurs, le rapport signal à bruit est supérieur à 1 : on peut donc raisonnablement, en première approximation, écrire le logarithme de la probabilité d’erreur comme :

$$\log(P_e) \approx -\frac{E_b}{N_0} \quad (\text{II.12})$$

Dans le cas du retournement temporel et du filtrage adapté, le niveau d’énergie d’un bit augmente comme le carré du nombre d’antennes M^2 , car il résulte des interférences constructives

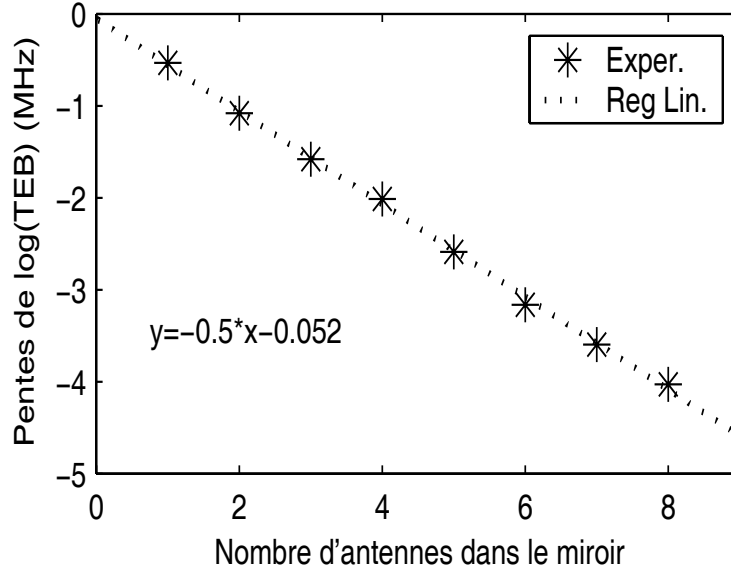


FIG. II.16 – Evolution des coefficients directeurs de $\log(P_e)$ En fonction du nombre d'antennes du miroir.

de chacun des signaux émis par les antennes de la base³. Dans le cas du retournement temporel, le niveau des lobes augmente comme M car il résulte de la somme incohérente des bruits résultants de ces mêmes signaux. On fait ici l'hypothèse que les différents chemins entre les antennes de la base et d'un utilisateur sont parfaitement décorrélés. Au contraire, dans le cas du filtrage adapté à la réception, c'est d'une somme incohérente de M^2 signaux dont résulte le bruit : son énergie évolue elle aussi en M^2 . En résumé, le logarithme de la probabilité d'erreur évolue comme $-E_b(1)/N_0(1)$ et $-ME_b(1)/N_0(1)$ respectivement dans le cas du filtrage adapté à la réception et du retournement temporel, avec $E_b(1)/N_0(1)$ le rapport signal à bruit lorsqu'une seule antenne de la base est utilisée. Or, nous l'avons dit dans le premier chapitre de ce manuscrit, quand une impulsion courte est focalisée à travers un milieu complexe, le carré de l'amplitude de l'impulsion focalisée est proportionnelle à $(\tau\Delta\nu)^2$, où τ est le temps caractéristique d'atténuation dans le milieu et $\Delta\nu$ la bande passante utilisée. L'énergie moyenne des lobes secondaires créés par retournement temporel est quant à elle donnée par $\tau\Delta\nu$, comme schématisé sur la figure II.17.

Si maintenant c'est une séquence de bits qui est envoyée, tandis que l'amplitude d'un symbole reste constante, le RSB est dégradé par l'ajout des bruits (i.e. des lobes secondaires) qui cor-

³Pour nos mesure nous avons normalisé l'énergie d'émission : l'énergie d'un bit et l'énergie du bruit doivent ainsi être divisés par M^2 pour respecter les courbes précédente, ce qui ne change rien au raisonnement qui est exposé.

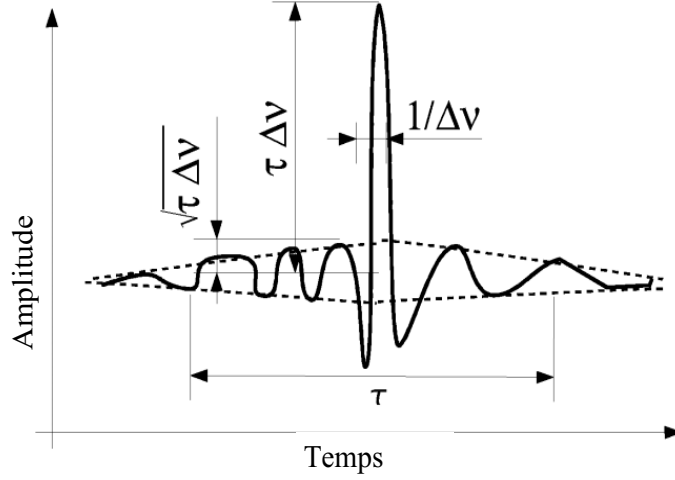


FIG. II.17 – Représentation schématic d'une impulsion compressée par retournement temporel. L'impulsion originale est reconstruite, ainsi que des lobes secondaires. L'amplitude de l'impulsion est proportionnelle à $\tau\Delta\nu$. La déviation standard des lobes secondaire (schématisée par la ligne en pointillé) dure τ et son amplitude est de $\sqrt{\tau\Delta\nu}$.

respondent aux symboles émis avant et après lui. Plus précisément, le niveau des interférences inter-symboles résulte de la somme incohérente des lobes secondaires créés par les n bits envoyés sur la durée τ d'une réponse impulsionnelle, soit $n = \tau/\delta t$. Ainsi l'énergie moyenne des interférences inter-symboles s'écrit, en fonction des paramètres physiques importants du système considéré, $N_0(1) \propto \tau^2\Delta\nu/\delta t$. Cette relation permet enfin d'écrire le rapport signal à bruit créé par retournement temporel pour une antenne :

$$\frac{E_b(1)}{N_0(1)} = \Delta\nu\delta t \quad (\text{II.13})$$

En conclusion, quand le rapport signal à bruit est plus grand que 1, ce qui dans la pratique est toujours vrai dans le cas d'une communication MISO par retournement temporel et sans bruit externe, la probabilité d'erreur respecte la formule suivante :

$$\log(P_e) \approx -\Delta\nu\delta tM \quad (\text{II.14})$$

Cette expression nous permet de justifier l'évolution linéaire du logarithme de la probabilité d'erreur en fonction à la fois du nombre d'antennes M du MRT, et de l'espacement temporel

δt entre chaque symbole. Ainsi, le coefficient directeur de la regression linéaire obtenue sur la figure II.16 n'est autre que la bande passante de notre système, que nous avons estimée à 0.35 MHz à -6 dB, et dont la valeur expérimentale est ici de 0.5 MHz. On observe un bon accord entre les deux valeurs.

Il faut cependant noter que cette discussion n'est valable, comme nous l'avons précisé, que dans le cas où l'on néglige le bruit externe. Nous verrons dans la dernière partie de ce chapitre, comment il est possible de prendre en compte ce bruit externe, et les conséquences que cela implique lors de communications par retournement temporel.

II.2.4 Vers les communications MIMO-Mu Ultra Large Bande

Parmi les intérêts du retournement temporel que nous avons évoqués, la focalisation spatiale qu'il procure n'a pas été prise en compte ici. Cependant, les raisonnements que nous avons exposés concernant les deux techniques étudiées restent valables pour le cas d'une utilisation multi-utilisateurs. En effet, si une base de M antennes s'adresse à un utilisateur, nous avons vu que le bruit temporel créé est dû aux lobes secondaires créés par l'envoi de ces signaux. Ainsi, il va en être de même lorsque l'on regarde le signal créé sur un utilisateur qui n'était pas visé par l'information émise, le bruit spatial étant équivalent au bruit temporel. Lorsque l'on applique le filtrage adapté à la réception, on va donc créer sur un utilisateur à qui l'information n'est pas destinée, des interférences inter-utilisateurs qui résultent de la somme incohérente de M^2 signaux, quand l'énergie d'un pic résulte de la somme cohérente de M signaux. La figure II.18 montre le signal créé par filtrage adapté à la réception sur l'utilisateur choisi et sur un autre utilisateur, et ce dans le cas où une base de 4 antennes est utilisée pour la transmission.

Comme on peut le voir, le niveau de bruit sur l'utilisateur non visé est du même ordre de grandeur que le bruit temporel qui est créé en dehors du pic sur l'utilisateur à qui l'on a envoyé l'impulsion. Une nouvelle fois, ceci est différent dans le cas du retournement temporel. En effet, le signal reçu sur un autre utilisateur résulte de la somme incohérente de M signaux. La figure II.19 montre ainsi le signal créé par la même base que précédemment, sur l'utilisateur choisi et sur un autre. Il est clair que le bruit est beaucoup plus faible. Il va donc être possible de réaliser un multiplexage spatial de l'information de façon bien plus efficace. En effet, si l'on veut à présent s'adresser à N récepteurs simultanément, aux interférences inter-symboles il va falloir ajouter les interférences inter-utilisateurs, afin de calculer le bruit déterministe. Le bruit sur chaque récepteur sera donc la somme incohérente de tous les bruits créés lorsque l'on s'adresse

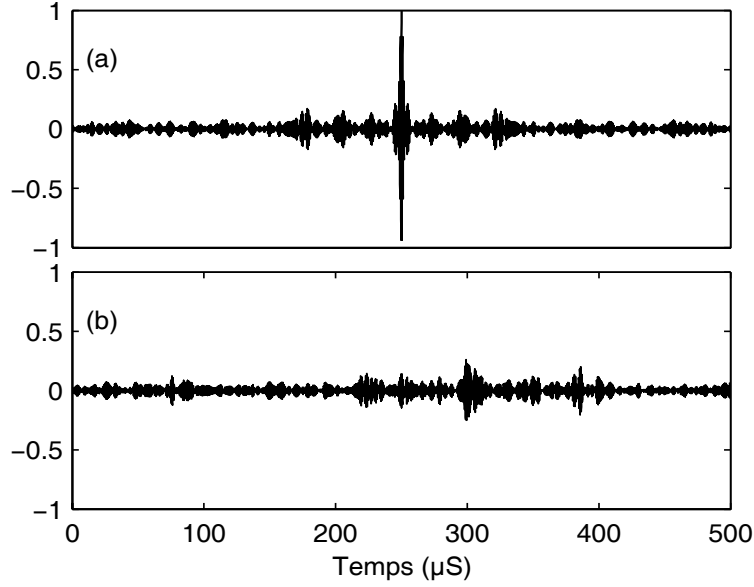


FIG. II.18 – Filtrage adapté à la réception avec une base constituée de 4 antennes, (a) signal sur l'utilisateur visé, (b) signal sur un autre utilisateur.

à chaque utilisateur. Ainsi, pour N récepteurs et M émetteurs, dans le cas du filtrage adapté à la réception, l'énergie moyenne du bruit sera proportionnelle à M^2N . Quant au niveau de l'amplitude d'une impulsion, celui-ci va rester égal à M . Le logarithme de la probabilité d'erreur va donc évoluer pour le filtrage adapté en réception comme :

$$\log(P_e) = -\Delta\nu\delta t \cdot \frac{1}{N} \quad (\text{II.15})$$

Si l'on utilise le retournement temporel, l'énergie moyenne du bruit créé sera proportionnelle à NM , et nous aurons ainsi une probabilité d'erreur respectant la formule :

$$\log(P_e) = -\Delta\nu\delta t \cdot \frac{M}{N} \quad (\text{II.16})$$

Cette étude montre que le retournement temporel permet de réaliser un multiplexage spatial de façon assez simple et efficace. Ceci pourrait se révéler utile dans la conception des systèmes MIMO-Mu ultra large bande passante. Cependant, nous l'avons comparé à une technique très rudimentaire qui, dans le cas de tels systèmes, ne serait jamais utilisée telle quelle. On utiliserait alors des techniques basées sur des séquences pseudo-aléatoires, afin de coder les informations destinées à chacun des utilisateurs. Cependant, il est clair qu'en utilisant uniquement un filtrage à la réception, il est impossible de faire disparaître totalement les effets des inter-corrélations entre les diverses antennes de la base et chaque utilisateur. Le retournement temporel, grâce à

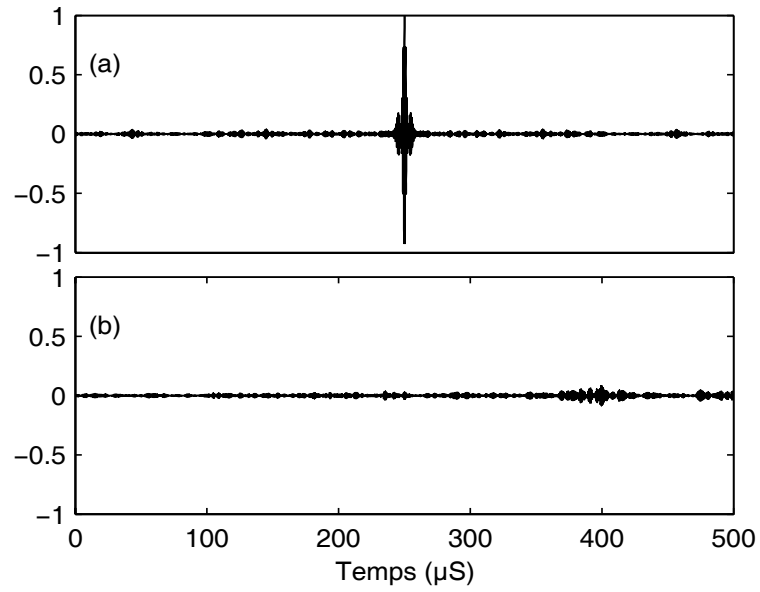


FIG. II.19 – Retournement temporel avec une base constituée de 4 antennes, (a) signal sur l'utilisateur visé, (b) signal sur un autre utilisateur.

la focalisation qu'il procure, diminue ces inter-corrélations de façon naturelle.

Dans cette partie, nous avons voulu nous approcher d'un environnement réel. A cette fin, et toujours limités par le matériel, nous avons réalisé les expériences à petite-échelle avec des ondes acoustiques dans un aquarium. Un modèle centimétrique d'un étage d'immeuble a été réalisé et les ondes décimétriques électromagnétiques ont été remplacées par des ondes ultrasonores d'une longueur d'onde de l'ordre du millimètre. Ceci nous a permis de réaliser des communications de type MISO, et plus spécifiquement de comparer un type de filtrage adapté à la réception simpliste au retournement temporel. Nous avons ainsi démontré que le retournement temporel ne génère pas d'inter-corrélations des signaux provenant des diverses antennes de la base. C'est donc un outil plus efficace qu'un filtrage adapté en réception. Le taux d'erreur binaire a été étudié pour ces deux techniques, en fonction du nombre d'antennes de la base, du débit d'information et du bruit externe. Il a été ainsi mis en évidence que pour les deux techniques, le logarithme de probabilité d'erreur suit une évolution linéaire en fonction de la bande passante et de l'espacement temporel entre deux symboles consécutifs. En revanche, et contrairement au filtrage adapté en réception, le retournement temporel offre de plus un gain de diversité qui est du au nombre d'antennes utilisées dans la base : le logarithme de la probabilité est donc également proportionnel à cette grandeur. Enfin, nous avons vu comment passer d'un système MISO à un système MIMO-Mu. La focalisation spatiale, conséquence naturelle du retournement temporel, permet un multiplexage spatial de l'information aisé. Nous avons montré que le logarithme du taux d'erreur binaire est inversement proportionnel au nombre d'utilisateurs auxquels la base s'adresse. Ces travaux ont en partie fait l'objet d'une publication dans le journal *Radio Science* [57]. Nous verrons dans la dernière partie de ce manuscrit comment il est possible de prendre en compte le bruit externe dans ce formalisme, et nous énumérerons de façon plus systématique les avantages que procure le retournement temporel. Cependant, auparavant, nous allons nous interroger sur le type de filtrage à réaliser en émission : filtrage inverse ou filtrage adapté. C'est le but de la partie suivante

II.3 Filtrage inverse ou filtrage adapté ?

II.3.1 Dispositif expérimental et principe du retournement temporel itératif

Une impulsion compressée par retournement temporel présente des lobes secondaires, conséquences de l'absorption du milieu et de l'impossibilité d'entourer entièrement le milieu avec des capteurs. Ces lobes secondaires sont autant de signaux indésirables lors d'une opération de communication. Ils provoquent des interférences inter-symboles et inter-utilisateurs et se traduisent par une limitation intrinsèque des télécommunications par retournement temporel. Est-il possible d'outrepasser cette limitation inhérentes au retournement temporel classique ? Une méthode originale consiste à mesurer une collection de réponses impulsionnelles entre des antennes d'une base et les utilisateurs, puis à calculer par transformée de Fourier à chaque fréquence la matrice de transfert que l'on notera $\mathbf{H}$. Ensuite il est possible de calculer explicitement la matrice $\mathbf{H}^{-1}$, de manière à obtenir un filtre inverse numérique de la propagation [58]. Cette technique permet de focaliser, en un endroit donné de l'espace, une impulsion dont les lobes temporels et spatiaux sont minimisés, indépendamment du nombre d'antennes de la base et des conditions d'absorption du milieu. Cependant c'est une méthode complexe à implémenter et qui nécessite une puissance de calcul assez élevée.

Dans la référence [59] publiée dans la revue *Waves in Random Media*, nous avons proposé une méthode originale basée sur des itérations du retournement temporel classique et qui converge vers le filtre inverse. Cette technique présente les avantages d'être à la fois aussi précise que le filtre inverse numérique, et presque aussi simple à réaliser que le retournement temporel classique. C'est cette méthode que nous présentons ici. Afin d'en étudier les avantages et désavantages en terme de communication, nous allons réaliser des expériences à petite échelle en comparant retournement temporel classique (RT) et retournement temporel itératif (RTI). Une nouvelle fois, l'information est transportée par des ondes ultrasonores dans une cuve remplie d'eau, à la fréquence de 1.5 MHz. La longueur d'onde associée, dans l'eau, est de 1 mm. Comme le montre la figure II.20, notre milieu complexe de référence est une forêt de tiges de 3 cm d'épaisseur. Chaque tige mesure 0.8 mm de diamètre et leur densité est de 29 tiges/cm². La base est composée de $M = 15$ transducteurs ultrasonores de fréquence centrale 1.5 MHz et dont la bande passante est d'à peu près 60% à -6 dB. Les éléments de la base sont espacés de 3 mm et celle-ci est située à 6 cm de la forêt de tiges. Derrière ce milieu multidiffuseur, à

une dizaine de centimètres, nous avons placé $N = 15$ utilisateurs qui sont des transducteurs de même nature que ceux de la base, et qui sont distants de 4 mm.

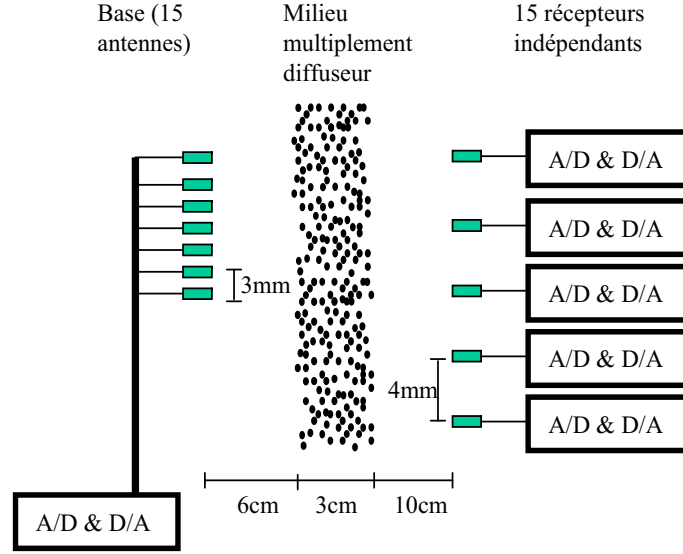


FIG. II.20 – Dispositif expérimental utilisé pour les expériences de communication par retournement temporel et retournement temporel itératif.

A cause des multiples reflexions que subissent les ondes qui traversent le milieu constitué de tiges, l'impulsion émise initialement d'une durée de $1 \mu s$, se transforme en une coda longue de plus d'une milliseconde lorsqu'elle est reçue par un utilisateur. Dans ce long signal, il est impossible de discerner l'arrivée directe de l'onde. On ne pourra donc pas transmettre de l'information sans prendre en compte les reflexions multiples subies par le message envoyé.

Comme dans les précédentes expériences de télécommunications par retournement temporel, une collection de 15×15 réponses impulsionnelles notées $h_{ij}(t)$ est mise en mémoire. Afin d'envoyer une impulsion sur l'utilisateur j , chaque antenne de la base émet alors la version retournée temporellement de la réponse impulsionnelle correspondante, soit $h_{ij}(-t)$. Chaque onde créée se propage en sens inverse dans le milieu pour venir se focaliser sur l'utilisateur visé. Celui-ci reçoit une impulsion d'une durée égale à celle de l'impulsion initiale. Celle-ci est entourée de lobes secondaires, dont la durée caractéristique est de l'ordre de celle de la réponse impulsionnelle. De même, chaque utilisateur à qui l'information n'est pas destinée reçoit un signal composé de lobes secondaires dont la durée est également de l'ordre de celle d'une réponse impulsionnelle. Nous avons vu qu'il est possible de diminuer l'amplitude de ce "bruit" en augmentant le nombre d'éléments de la base. Cependant cette technique nécessite une voie électronique supplémentaire pour chaque capteur ajouté à la base. L'idée simple du retournement temporel

itératif est d'améliorer le processus de retournement temporel en "nettoyant" les lobes secondaires temporels et spatiaux, à l'aide d'un algorithme itératif.

La première étape consiste à créer un objectif spatio-temporel $o_j(t), j \in [1, N]$ pour chaque utilisateur. Celui-ci est constitué d'une impulsion sur l'utilisateur j et de signaux identiquement nuls sur les autres récepteurs. L'impulsion doit bien sûr être réaliste, sa bande passante doit correspondre à celle des transducteurs. On peut par exemple choisir un signal du type $\sin(\pi t B)/\pi t B$, où B est la bande passante des transducteurs. Cette objectif est ensuite retourné temporellement ⁴, et est envoyé de l'utilisateur vers la base (figure II.21.(a)).

Puis, comme dans une procédure de retournement temporel classique, chaque élément de la base d'antennes enregistre et mémorise la réponse impulsionnelle $e_i(t), i \in [1, M]$. Ces signaux sont ensuite retournés temporellement puis renvoyés par la base vers l'utilisateur j (figure II.21.(b)). Tous les utilisateurs reçoivent et enregistrent alors un signal $r_j(t), j \in [1, N]$. Le signal reçu par l'utilisateur j ressemble à l'objectif, mais contient des lobes secondaires temporels. Quant aux autres utilisateurs, ils enregistrent également des lobes temporels, car la focalisation spatiale n'est pas parfaite.

L'idée du retournement temporel itératif est de supprimer les lobes secondaires par RT. On isole les lobes secondaires en en faisant une simple soustraction par l'objectif : $d_j(t) = o_j(t) - r_j(t)$ (figure II.21.(c)). Une fois calculés, ces signaux sont retournés temporellement et émis des utilisateurs vers la base. Celle-ci enregistre alors les signaux résultant $c_i(t)$. Si la base envoie $c_i(-t)$ dans le milieu, une approximation de l'opposé des lobes secondaires calculés précédemment $d_j(t)$ est recrée sur chacun des différents utilisateurs. Il suffit alors à la base d'émettre les signaux $e_i^2(t) = e_i(-t) + c_i(-t)$ afin de créer l'objectif "nettoyé" des lobes secondaires qui avaient été créés par la première opération de retournement temporel. Comme le schématise la figure II.21.(d), les signaux reçus présentent toujours des lobes secondaires, mais ceux-ci sont d'une amplitude moindre. Ce procédé est ensuite itéré à partir de l'étape n°2 afin d'éliminer complètement les lobes secondaires spatiaux et temporels. Lorsque le niveau de lobes secondaires voulu est atteint, on mémorise le dernier signal émis par la base, que l'on nommera $e_{ij}(t)$, et qui assure un objectif sur l'utilisateur j , en émettant avec les M antennes i de la base. Le processus est répété pour les N utilisateurs de sorte que l'on mémorise finalement une collection de $M * N$ signaux, $e_{ij}(t), (i, j) \in [1, M] \times [1, N]$. La méthode ne nécessite qu'un miroir à re-

⁴Ceci est nécessaire si l'objectif temporel est antisymétrique car le retournement temporel crée l'impulsion initiale dans une chronologie inversée.

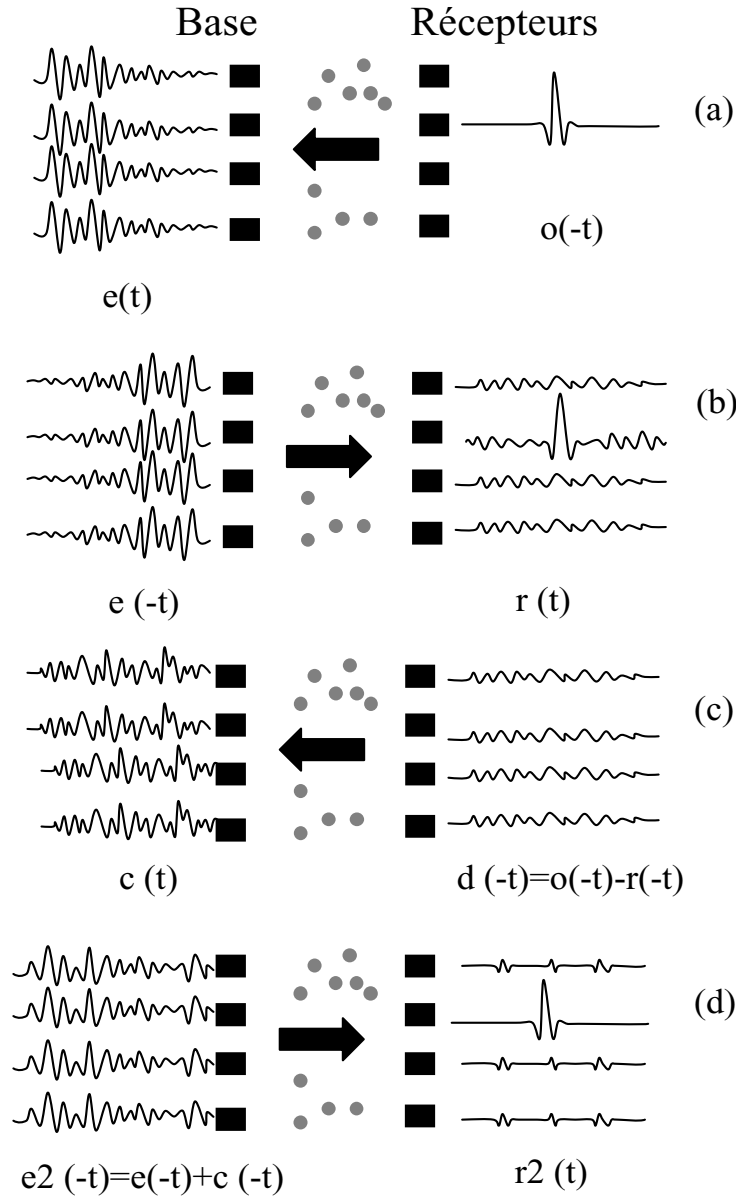


FIG. II.21 – Principe du RT itératif : retournement temporel classique en (a) et (b) ; émission des lobes secondaires des utilisateurs vers la base (c) ; focalisation avec les signaux corrigés (d).

ournement temporel standard et une fonction soustraction, ce qui la rend assez peu coûteuse en calcul. D'autre part, contrairement à une inversion numérique, celle-ci est très stable. Enfin, cette méthode présente un autre avantage : les faibles déformations de l'impulsion occasionnées par les non-linéarités des composant électroniques sont automatiquement corrigées.

La figure II.22 montre une impulsion reçue sur un utilisateur après retournement temporel classique et après 15 itérations du retournement temporel itératif. Le valeur moyenne de l'énergie des lobes secondaires se situe 18 dB en dessous du carré de l'amplitude maximale du pic pour le RT classique. Pour le RTI, ce niveau est 33 dB en dessous du niveau de l'impulsion. Le niveau

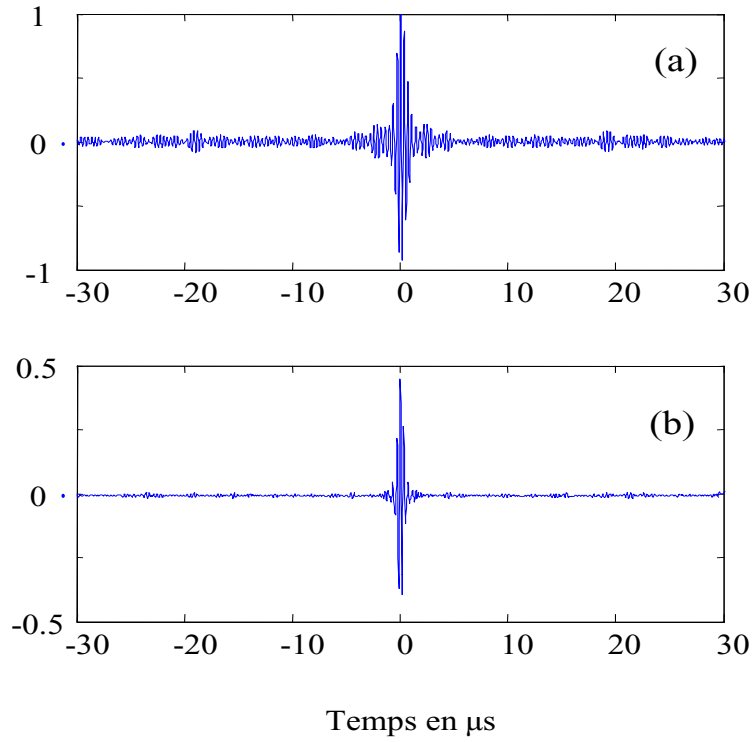


FIG. II.22 – Impulsion reçue sur un utilisateur après retournement temporel (a), et après 15 étapes du retournement temporel itératif.

des lobes temporels, qui comme nous l’avons vu produisent des interférences inter-symboles lors de communications, a donc décru de 15 dB.

Le niveau des lobes spatiaux décroît lui aussi sensiblement dans les mêmes proportions. La figure II.23 montre le rapport entre énergie moyenne des lobes secondaires créés sur chaque utilisateur et le carré de l’amplitude de l’impulsion lorsque celle-ci est émise sur le récepteur 9. On y voit que le niveau de lobes secondaires décroît d’une valeur comprise entre 13 dB et 8 dB selon l’utilisateur sur lequel on le mesure. Ce gain sera utile lors de communications multi-utilisateurs.

Cependant la méthode itérative présente un désavantage : l’amplitude de l’impulsion focalisée décroît de façon significative par rapport au retournement temporel classique. En effet, le retournement temporel est un filtre adapté au sens du traitement du signal. L’énergie d’une impulsion lorsqu’elle est focalisée sur un récepteur donné est donc maximisée. Au contraire, la méthode itérative réalise un filtrage inverse, elle compense les creux dans le spectre du signal en augmentant l’énergie émise aux fréquences auxquelles le milieu est le plus opaque. Ceci a pour conséquence d’accroître la quantité d’énergie “dissipée” dans le milieu, soit de diminuer l’énergie d’une impulsion focalisée sur un utilisateur, pour une puissance d’émission fixée. Sur la figure

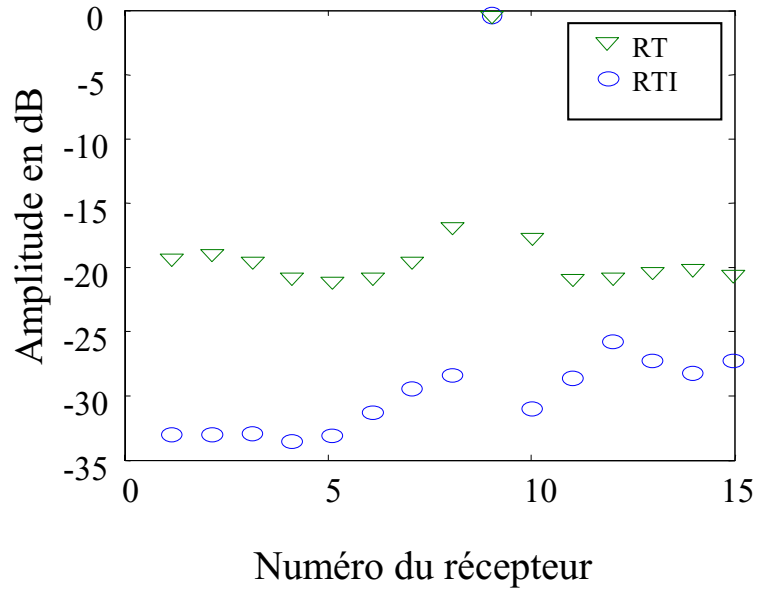


FIG. II.23 – Rapport du niveau des lobes spatiaux au maximum du pic obtenu en focalisant sur l'utilisateur 9 par RT et par RTI en dB.

II.22, on estime que l'amplitude de l'impulsion focalisée par retournement temporel itératif est 6 dB en dessous de celle de l'impulsion focalisée par retournement temporel classique. Ainsi, en présence de bruit externe, le rapport signal-sur-bruit diminue si l'on choisit de communiquer en utilisant la méthode itérative plutôt que la méthode classique. Dans la pratique, nous allons voir qu'un compromis peut être trouvé entre filtrage inverse et filtrage adapté, au travers du nombre d'itérations réalisées. Ceci va permettre de réaliser un équilibre entre une amplitude d'impulsion en adéquation avec le bruit externe et un niveau de lobe secondaires minimisés pour lutter contre les interférences inter-symboles et inter-utilisateurs.

II.3.2 Comparaison des méthodes dans le contexte des télécommunications

Par retournement temporel, dès lors que l'on a connaissance des réponses impulsionnelles, nous avons vu qu'il est possible d'émettre des messages constitués des symboles $s_{j,k}$, avec j un utilisateur et k le numéro du symbole, en envoyant le signal suivant par les i antennes de la base :

$$a_i(t) = \sum_{j,k} h_{ij}(-t + t_k) \otimes s_{j,k}(t) \quad (\text{II.17})$$

où t_k est le temps de focalisation de chaque symbole.

Les symboles pourront être constitués de $+1$ et -1 , afin de réaliser une modulation binaire de

type BPSK, ou pourront être complexes afin de réaliser une modulation en phase à m niveaux. Chaque utilisateur reçoit alors des impulsions séparées du temps t_k , modulées par les symboles $s_{j,k}$.

En utilisant le retournement temporel itératif, il est possible de réaliser exactement la même opération. Dès lors que l'on a connaissance des signaux $e_{ij}(t)$, $(i, j) \in [1, M] \times [1, N]$, qui permettent de créer sur chaque utilisateur l'objectif voulu, il est possible d'envoyer une trame d'information, en émettant par chaque antenne le signal suivant :

$$a_i(t) = \sum_{j,k} e_{ij}(t + t_k) \otimes s_{j,k}(t) \quad (\text{II.18})$$

La figure II.24.(a) montre un extrait d'une suite pseudo-aléatoire de 300 symboles $\{+1, -1\}$ que l'on souhaite émettre vers l'utilisateur 5. De l'information est envoyée vers chacun des utilisateurs simultanément, avec une distance inter-symboles fixe de $1.5 \mu s$. Sur la figure II.24.(b), on représente le message reçu par l'utilisateur 5, par retournement temporel (équation II.17). Les symboles sont détectés par minimisation de la distance quadratique moyenne par rapport à la constellation de référence. Dans ce cas, un symbole reçu sur les 10 est faux. En utilisant la méthode itérative (équation II.18), les interférences inter-symboles sont minimisées, et le message est intégralement bien reçu par l'utilisateur 5, comme le montre la figure II.24.(c).

Pour une communication avec un débit de 0.5 MBs/s, nous avons mesuré expérimentalement un taux d'erreur binaire (défini comme le nombre de bits reçus faux, sur le nombre de bits total envoyés) de 12% par retournement temporel, alors qu'il est de seulement 0.6% lorsque l'on utilise la méthode itérative. Cet exemple montre l'intérêt d'utiliser cette technique afin d'éliminer les interférences inter-symboles et inter-utilisateurs. Nous allons maintenant comparer ces techniques plus en détail, en fonction du débit souhaité et du bruit externe.

Nous avons implémenté une modulation de type PSK (phase shift keying) à 1, 2 et 3 bits par symboles, c'est à dire des modulations de type BPSK, 4-PSK et 8-PSK. La modulation code l'information des n bits grâce à la phase du symbole. Ainsi, pour une modulation binaire, la phase du symbole sera de 0 pour un $+1$ et de π pour un -1 . De même, pour une modulation à 2 bits, c'est-à-dire une 4-PSK, les phases seront 0, $\pi/2$, $2\pi/2$ et $3\pi/2$. Les signaux utilisés $e_{ij}(t)$, $(i, j) \in [1, M] \times [1, N]$ pour la méthode itérative ont été obtenus au bout de 15 itérations. La figure II.25 montre le taux d'erreur obtenu pour différents débits d'information. Celui-ci est calculé comme le nombre de bits erronés sur le nombre de bits total ; pour chaque mesure nous avons transmis autant de bits qu'il était nécessaire afin d'obtenir une estimation de la

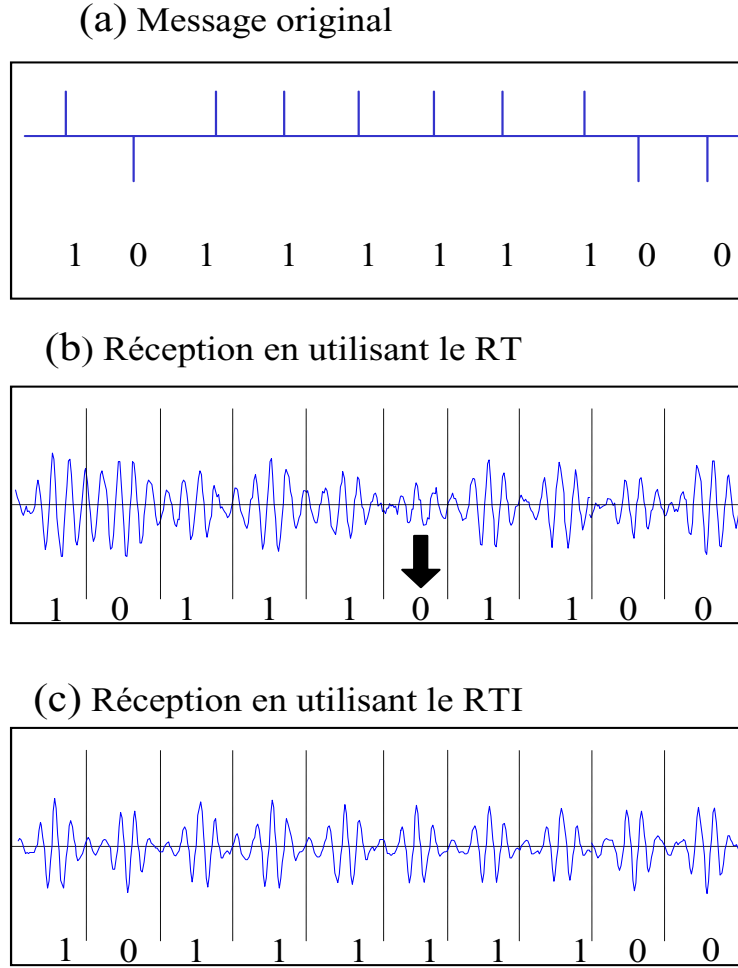


FIG. II.24 – Message à transmettre à l'utilisateur 5 (a), reçu en utilisant le retournement temporel (b), et la procédure itérative (c).

probabilité d'erreur qui soit statistiquement acceptable. La méthode itérative donne de biens meilleurs résultats, spécialement pour les débits d'information assez faibles.

Ces courbes sont assez aisées à interpréter à l'aide d'un modèle simple que nous allons développer ici. Quand nous focalisons un symbole sur un utilisateur, celui-ci est perturbé par les autres symboles qui lui sont destinés, mais également par les symboles qui sont destinés aux autres utilisateurs auxquels on envoie de l'information. Chacun de ces symboles génère un bruit d'interférence η_i , où i est l'indice du symbole considéré. Le bruit total créé sur un symbole donné est simplement la somme de ces interférences :

$$\eta_{int} = \sum_{i=1}^{N_{sym}} \eta_i \quad (\text{II.19})$$

où on note N_{sym} le nombre total de symboles qui créent les interférences.

Les bruits η_i sont indépendants et par conséquent, si l'on suppose qu'ils ont la même variance

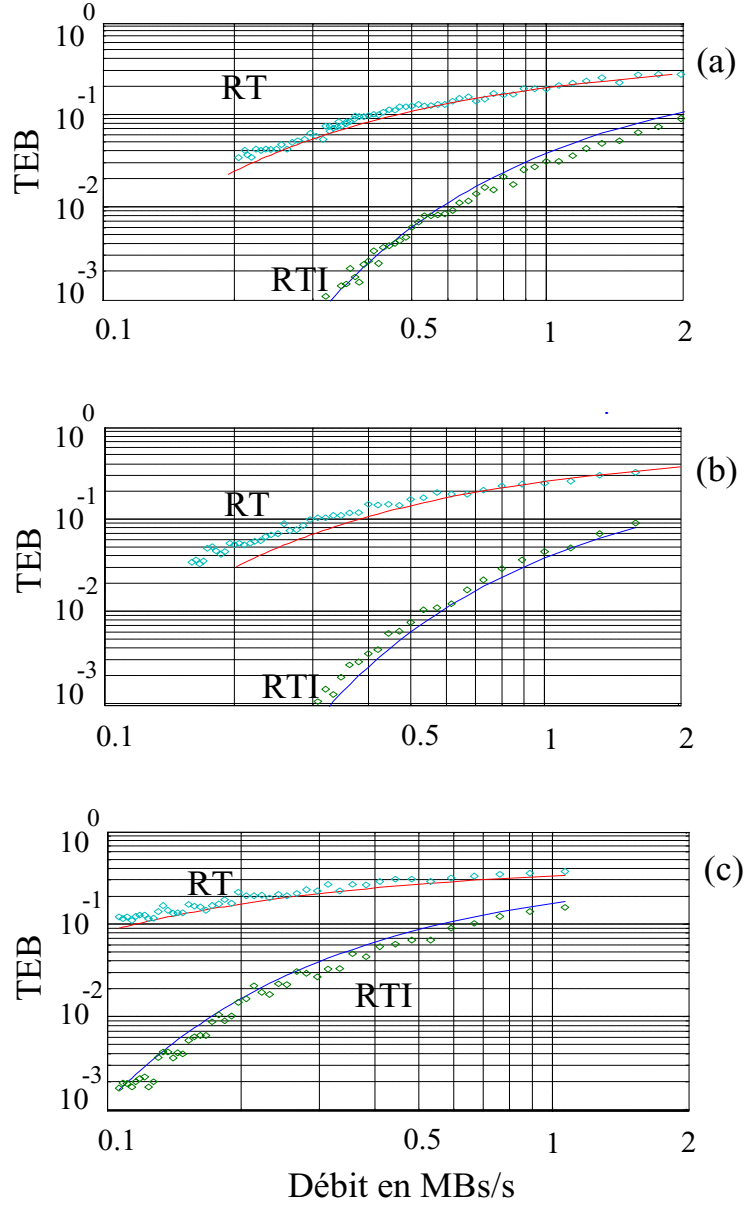


FIG. II.25 – Taux d’erreur binaire en fonction du débit d’information par retournement temporel et par retournement temporel itératif et pour différentes modulations : BPSK (a), 4-PSK (b) et 8-PSK (c).

$\tilde{\sigma}_{int}^2$, la variance totale du bruit η_{int} est : $\sigma_{int}^2 = N_{sym}\tilde{\sigma}_{int}^2$.

On calcule à présent le nombre total N_{sym} de symboles qui interfèrent, comme cela a été fait dans la partie précédente pour un seul récepteur : $N_{sym} = N\tau/\delta t$, avec N le nombre de récepteur, τ le temps de réverbération et δt le laps de temps entre deux symboles consécutifs. En conclusion, on peut écrire la dispersion totale du bruit comme :

$$\sigma_{int} = \tilde{\sigma}_{int}\sqrt{N\tau/\delta t} \quad (\text{II.20})$$

Dans le cas le plus simple de la modulation de type binaire antipodale avec un bruit externe

nul, la probabilité de faire une erreur à la réception d'un symbole est $P_{error} = P(\sigma_{int} > 1)$. Si le nombre de symboles N_{sym} qui interfèrent est grand, on peut faire l'hypothèse que la distribution du bruit suit une loi normale et la probabilité d'erreur est alors donnée par :

$$P_{error} = \frac{1}{2} \operatorname{erfc} \left(\frac{1}{\sigma_{int} \sqrt{2}} \right) \quad (\text{II.21})$$

Concernant les modulations de type m-PSK, une bonne approximation de la probabilité d'erreur est donnée par [56] :

$$P_{error} = \frac{1}{\log_2(m)} \operatorname{erfc} \left(\frac{\sqrt{\log_2(m)} \sin(\frac{\pi}{m})}{\sigma_{int} \sqrt{2}} \right) \quad (\text{II.22})$$

Comme on le voit sur la figure II.25, le taux d'erreur binaire est le même pour les modulations de type BPSK et 4-PSK. En revanche, il est plus élevé dans le cas de la modulation 8-PSK mais il faut noter que cela n'a rien à voir avec les techniques employées, mais avec la modulation elle-même [56]. On a également représenté sur cette figure les courbes théoriques calculées à partir des formules II.21 et II.22. Le paramètre σ_{int} a par ailleurs été mesuré en focalisant une impulsion par chacune des 2 méthodes et en calculant la déviation standard des lobes secondaires sur un intervalle de temps égal au temps de réverbération du milieu τ . Ces courbes théoriques sont en très bon accord avec les points mesurés, ce qui tend à confirmer notre approche.

II.3.3 Comparaison des méthodes en fonction du bruit

Les précédentes mesures ont porté sur les interférences inter-symboles et inter-utilisateurs créés par le retournement temporel et le retournement temporel itératif. Nous avons mis en évidence que celles-ci jouent un rôle déterminant en limitant la qualité de la communication. Le retournement temporel itératif, qui les minimise, est donc à priori plus intéressant que son analogue classique. Cependant, nous n'avons pas pris en compte le bruit qui est toujours présent dans un canal de communication. Il est possible de prendre cette grandeur en compte dans l'équation II.21. Si l'on fait l'hypothèse que le bruit est de type blanc et gaussien, et que l'on note σ_{ext} sa déviation standard supposée indépendante de l'information transmise, on peut écrire la déviation standard du bruit total :

$$\sigma_{total} = \sqrt{\sigma_{int}^2 + \sigma_{ext}^2} \quad (\text{II.23})$$

Pour une modulation de type BPSK, la probabilité d'erreur est alors donnée par la relation suivante :

$$P_{error} = \frac{1}{2} \operatorname{erfc} \left(\frac{1}{\sqrt{2(\sigma_{int} + \sigma_{ext})}} \right) \quad (\text{II.24})$$

Dans le cas présent, comme nous l'avons mentionné précédemment, la comparaison entre le retournement temporel classique et le retournement temporel itératif est plus délicate : un compromis doit être trouvé entre la méthode classique (qui donne une impulsion maximisée en énergie mais des lobes secondaires élevés) et la méthode itérative (qui minimise les lobes secondaires au prix d'une décroissance de l'amplitude de l'impulsion). Afin d'illustrer ceci, il est possible de considérer 2 cas extrêmes :

- Pour une communication dans un canal où le bruit est négligeable, nous utiliserons le retournement temporel itératif jusqu'à converger vers le filtre inverse. En effet, dans ce cas, la qualité de la communication est limitée par les lobes secondaires, et l'on choisira alors une solution qui les minimise.
- Pour une communication dans un canal très bruité, le bruit total ne sera pas dominé par les interférences inter-symboles et inter-utilisateurs mais par le bruit du canal. La solution consistera alors à maximiser l'amplitude des impulsions qui transportent l'information, de façon à travailler avec un rapport signal-sur-bruit le plus élevé possible. On optera alors logiquement pour le retournement temporel classique.

Pour un niveau de bruit intermédiaire, il va être nécessaire de considérer une méthode intermédiaire entre le retournement temporel et le retournement temporel itératif, de manière à assurer un rapport signal-sur-bruit externe optimisé et un niveau de lobes secondaires suffisamment bas. L'avantage de la méthode itérative est qu'elle converge vers le filtre inverse, tout en partant du retournement temporel qui est un filtre adapté. Il va donc être possible de trouver un compromis entre les deux méthodes par la biais du nombre d'itérations que l'on réalise.

Le problème concret que nous avons à résoudre est alors de trouver le nombre d'itérations qui va minimiser la dispersion totale du bruit σ_{tot} pour un niveau de bruit externe donné. Sur la figure II.26.(a) nous avons mesuré l'amplitude de l'impulsion A que l'on focalise par la procédure itérative, et ce en fonction du nombre d'itérations : $A = A(n)$. Le retournement temporel maximisant cette amplitude, nous avons normalisé cette courbe de manière à obtenir une amplitude égale à 1 lorsque l'on ne procède à aucune itération. Sur la figure II.26.(b), nous avons représenté l'écart type $\tilde{\sigma}_{int}$ des lobes secondaires créés, ou encore des interférences inter-

symboles (ISI), en fonction du nombre d'itérations également. Enfin, la courbe sur la figure II.26.(c), qui peut être déduite des 2 courbes précédentes, représente l'écart type des lobes secondaires (i.e des ISI), en fonction de l'amplitude de l'impulsion focalisée : $\tilde{\sigma}_{int} = \tilde{\sigma}_{int}(A)$. Afin de simplifier les calculs qui vont suivre, une interpolation des points expérimentaux de cette dernière courbe a été faite à l'aide d'un polynôme de degrés 3, elle est représentée en vert.

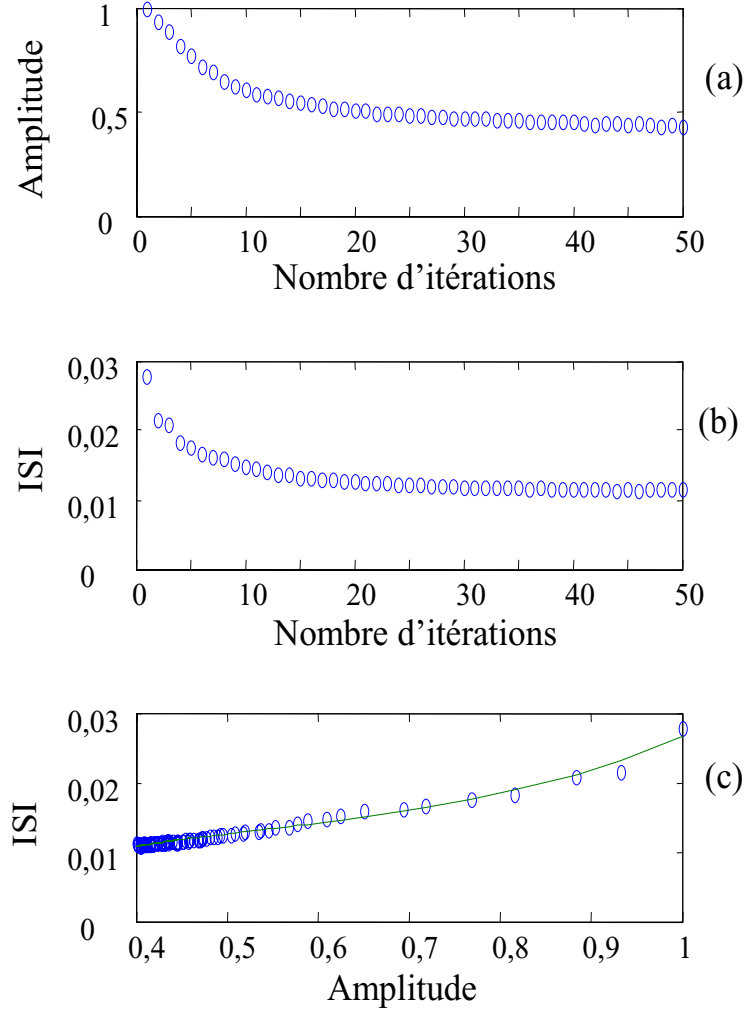


FIG. II.26 – Amplitude d'une impulsion focalisée en fonction du nombre d'itérations (a), Dispersion des lobes secondaires en fonction du nombre d'itérations (b) et Dispersion des lobes secondaires en fonction de l'amplitude de l'impulsion (c).

La dispersion du bruit d'interférences est calculée à partir de l'équation II.20 :

$$\sigma_{int} = \tilde{\sigma}_{int}(A) \sqrt{N\tau/\delta t} \quad (\text{II.25})$$

Si l'on fait l'hypothèse que le bruit est de type blanc et gaussien de variance ν_{ext} , la dispersion

du bruit, normalisée, à la détection est de :

$$\sigma_{ext} = \nu_{ext}/A \quad (\text{II.26})$$

Si l'on substitue à présent les équations II.25 et II.26 dans II.23 on obtient :

$$\sigma_{total}(A) = \sqrt{\tilde{\sigma}_{int}^2(A)N\tau/\delta t + \nu_{ext}^2/A^2} \quad (\text{II.27})$$

Sur la figure II.27 nous avons calculé $\sigma_{total}(A)$, à l'aide de la formule déduite précédemment, pour différentes valeurs de bruit externe ν_{ext} , les autres paramètres étant les suivants : le nombre d'utilisateurs est de $N = 15$, le temps de réverbération τ est de 1.5 ms et les impulsions sont séparées d'une durée δt de 1.5 μ s. Les différentes courbes ont été calculées pour des variances du bruit comprises entre 0.1 et 0.5 de l'amplitude de l'impulsion focalisée par retournement temporel classique.

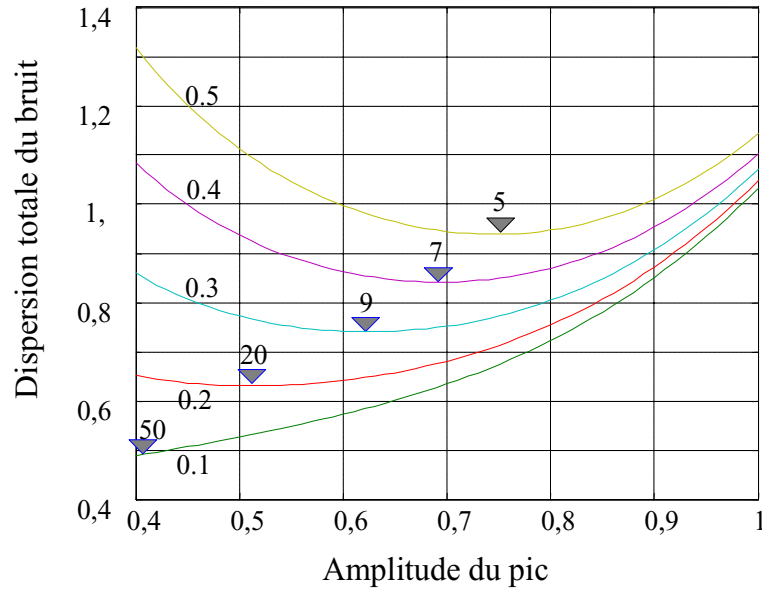


FIG. II.27 – Dispersion totale du bruit en fonction de l'amplitude de l'impulsion, et pour divers niveaux de bruit externe. Les chiffres sur les triangles indiquent le nombre d'itérations optimal à réaliser pour minimiser la dispersion totale du bruit.

La valeur minimale de $\sigma_{total}(A)$ donne la valeur optimale de l'amplitude que l'impulsion focalisée doit avoir afin de lutter contre le bruit externe. Il est alors possible, en regardant sur la figure II.26.(a), d'en déduire le nombre optimal d'itérations à réaliser pour chaque niveau de bruit, de façon à minimiser $\sigma_{total}(A)$.

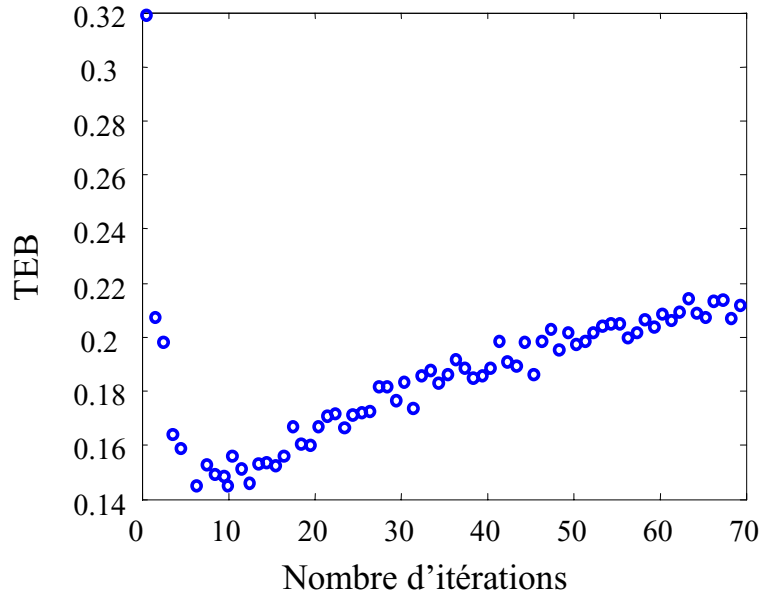


FIG. II.28 – Taux d'erreur binaire mesuré en fonction du nombre d'itérations pour une fraction de bruit de 0.3.

Nous avons voulu vérifier expérimentalement les courbes précédentes en réalisant à nouveau une expérience de communication. En effet, l'équation II.24 nous indique que si la dispersion totale σ_{total} du bruit est minimisée, la probabilité d'erreur est minimisée elle aussi. Nous avons donc réalisé le même type de mesure de taux d'erreur binaire que précédemment, en utilisant une modulation de type BPSK, sur les 15 utilisateurs simultanément et avec un débit de 667 kbs/s, soit un espacement entre les symboles de $1.5 \mu s$. Un bruit externe de type gaussien a été ajouté dont la variance est de 0.3 par rapport à l'amplitude de l'impulsion focalisée par retournement temporel classique. Les résultats de ces mesures sont représentés sur la figure II.28. On voit que le nombre d'itération qui minimise le taux d'erreurs binaire pour la fraction de bruit ajoutée est de 9. Ce chiffre est en parfait accord avec la courbe calculée sur la figure II.27.

Nous avons présenté dans cette partie une méthode qui permet de focaliser des impulsions courtes sur différents utilisateurs à travers un milieu multi-diffuseur. Cette technique, fondée sur des itérations successives du retournement temporel, est capable de créer un objectif arbitraire sur un utilisateur, ce qui minimise les lobes secondaires spatiaux et temporels que produit le retournement temporel classique. D'un point de vue théorique, cette procédure itérative converge vers le filtre inverse de la propagation [59]. Elle suit un schéma simple qui ne nécessite qu'un miroir à retournement temporel et une fonction soustraction, ce qui la rend assez peu coûteuse en calculs. Nous avons présenté des expériences de télécommunications dans un environnement complexe. En utilisant la méthode itérative, nous avons montré que le niveau des lobes secondaires, décroît de 13 dB en moyenne par rapport à la méthode classique. Ceci a pour effet de diminuer les interférences inter-symboles et inter-utilisateurs, ce qui permet d'améliorer considérablement la qualité de la communication en diminuant le nombre d'erreurs commises. Un modèle simple a été développé pour interpréter les résultats obtenus. Cependant, nous avons également souligné que, bien que cette méthode soit optimale si le bruit externe est très faible, le retournement temporel classique est plus efficace dans le cas d'un environnement très bruité. Pour un niveau de bruit intermédiaire, nous avons vu il est possible de trouver un compromis en stoppant la procédure itérative à un certain nombre d'itérations. Ceci permet d'optimiser l'amplitude de l'impulsion que l'on focalise de manière à être plus robuste au bruit externe, tout en réduisant l'amplitude des lobes secondaires afin de limiter les interférences inter-symboles et inter-utilisateurs. Bien que prometteuse, cette technique présente le désavantage d'être difficilement implémentable en temps réel, ce qui est problématique dans les environnements qui changent dans le temps. De plus, se rapprochant du filtrage inverse, elle est très sensible à toute modification du milieu, ce qui la rend peu robuste en pratique. Dans la dernière partie de ce manuscrit, qui traite de télécommunications par retournement temporel électromagnétique à 2.45 GHz, nous nous contenterons donc d'étudier le retournement temporel, qui paraît plus adapté aux problèmes des communications modernes par sa simplicité de réalisation, du moins dans un avenir proche. Notons enfin que cette méthode itérative peut être réalisée numériquement à partir de la base [60]. Les auteurs ont montré que dans ce cas elle nécessite un terme de régularisation afin d'en assurer sa stabilité, et qu'elle converge vers le retournement temporel classique employé avec un récepteur de type minimisation de la l'erreur quadratique moyenne (MMSE) comme cela est présenté dans [61].

II.4 Communications dans le gigahertz en cavité réverbérante

II.4.1 Principe des expériences et dispositif expérimental

Les premières expériences de communication par retournement temporel dans le domaine du gigahertz ont été réalisées par Henty et Stancil [62]. Dans cet article, les auteurs présentent deux séries d'expériences où de l'information est transmise à 4 récepteurs séparés d'une longueur d'onde. Dans un premier temps, la transmission est réalisée dans une salle vide du laboratoire, puis dans la même pièce dans laquelle ils ont placé des diffuseurs. Ils ont montré que le taux d'erreur est très faible lorsque la pièce est "en désordre", alors que celui-ci est élevé lorsque la salle est vide. Cette expérience est une démonstration de ce qui avait d'abord été montré dans [52], à savoir que, par retournement temporel, il est possible de profiter du désordre afin d'augmenter la quantité d'information transmise. En revanche, ces mesures ont été réalisées avec une bande passante très faible et l'aspect compression temporelle n'a donc pas été abordé. Ici, nous proposons d'étudier les communications par retournement temporel "large bande" dans un environnement très réverbérant. Nous avons procédé à des mesures dans la cavité couverte d'aluminium déjà évoquée dans le premier chapitre de ce manuscrit. Nous avons mesuré précédemment un temps de réverbération dans cette cavité de l'ordre de 160 ns, ce qui donne un facteur d'étalement de 16 pour une impulsion initiale d'une durée de 10 ns. Ce type d'étalement temporel est tout à fait de l'ordre de grandeur des étalements que l'on peut observer avec des systèmes ultra large bande passante dans des environnements "indoor" [63]. Ainsi, bien que notre milieu ne soit pas très réaliste, les résultats des études que nous allons présenter seront à priori transposables à des situations réelles.

Les expériences ont été réalisées avec le miroir à retournement temporel présenté dans le chapitre précédent et dont nous rappelons quelques caractéristiques. La fréquence porteuse est de 2.45 GHz et la bande passante de 200 MHz. Nous disposons de 8 voies d'émission et de 8 voies de réception par le biais des multiplexeurs. Les antennes utilisées sont des dipôles achetés dans le commerce qui sont séparés d'une distance de $\lambda/2$. Le miroir à retournement temporel à 8 voies joue le rôle de la base. Les utilisateurs sont également au nombre de 8, et sont placés à une distance de l'ordre de 10 longueurs d'onde du miroir. Comme dans le cas des mesures de tache focale réalisées précédemment, nous utiliserons la linéarité des ondes électromagnétiques pour réaliser nos mesures. En ce sens, lorsque l'on voudra utiliser plusieurs voies du miroir

pour adresser un utilisateur, nous émettrons successivement le message avec chacune des voies, puis nous sommerons les résultats de chaque mesure. De même, lorsque nous souhaiterons nous adresser à plusieurs utilisateurs simultanément, les signaux reçus seront mesurés pour chaque utilisateur séparément à l'aide des multiplexeurs.

Comme lors de chaque expérience de retournement temporel, la première phase consiste à enregistrer les réponses impulsionnelles entre chaque antenne i de la base et les utilisateurs j . Cependant, ici le retournement temporel est réalisé en bande de base : les réponses impulsionnelles sont donc pour chaque couple (i, j) des matrices $(2 * 2)$ notées $H_{ij}(t)$. On rappelle que pour enregistrer une telle matrice, une impulsion d'une durée de 10 ns est émise sur la voie I (en phase) du modulateur, quand le signal sur la voie Q reste nul. Les signaux obtenus en bande de base sont alors les réponses en phase et en quadrature, à partir desquelles on construit la matrice $H_{ij}(t)$, tel que nous l'avons développé dans le premier chapitre de ce manuscrit. Cette opération est réalisée pour les 64 couples constitués par les 8 antennes de la base et les 8 récepteurs. Par la suite, si l'on souhaite transmettre un symbole complexe il suffit d'écrire ce symbole en coordonnées cartésiennes, et le signal à générer en bande de base est simplement la multiplication de la matrice des réponses impulsionnelles par ce vecteur symbole. Ainsi, si l'on veut transmettre une série de vecteurs symboles $s_{jk}, k \in [1, N_{sym}]$ à un utilisateur j avec M antennes de la base, il suffit de calculer en bande de base le signal suivant, qui après modulation sera émis par l'antenne i :

$$e_i(t) = \sum_{k=1}^M H_{ij}(-t + \delta t).s_{jk} \quad (\text{II.28})$$

où δt est l'espacement temporel entre chaque symbole émis.

Si l'on souhaite de plus émettre de l'information à N récepteurs simultanément, il suffit de sommer les divers signaux voulus, de façon à générer en bande de base le signal suivant :

$$e_i(t) = \sum_{j=1}^N \sum_{k=1}^M H_{ij}(-t + k\delta t).s_{jk} \quad (\text{II.29})$$

Du fait de l'invariance par translation dans le temps du signal, le signal reçu après démodulation par l'utilisateur l prendra alors la forme :

$$r_l(t) = \sum_{i=1}^M \left[\left(\sum_{j=1}^N \sum_{k=1}^{N_{sym}} H_{ij}(-t + k\delta t).s_{jk} \right) \otimes H_{il}(t) \right] \quad (\text{II.30})$$

Etant donné que nous utilisons des multiplexeurs plutôt qu'un miroir multivoies, la somme sur les antennes d'émission sera réalisée numériquement après acquisition des signaux.

La modulation employée est de type BPSK antipodale, c'est-à-dire que nous écrirons les symboles s_{jk} , qui représentent des 0 et 1, par les vecteurs $\begin{bmatrix} 1 \\ 0 \end{bmatrix}$ et $\begin{bmatrix} -1 \\ 0 \end{bmatrix}$ ⁵. Le but est d'étudier la qualité de la communication, c'est à dire le taux d'erreur binaire, en fonction de divers paramètres. Pour ce faire, une trame de 500 symboles générés de façon pseudo-aléatoire est émise à un ou plusieurs utilisateurs. Un fois les symboles décodés par minimisation de la distance quadratique, le taux d'erreur binaire est calculé comme le rapport entre le nombre de bits erronés sur le nombre de bits total. L'expérience est réalisée jusqu'à ce que la variance du taux d'erreur binaire soit suffisamment faible. La figure II.29 montre une trame de 500 symboles reçue par un utilisateur après démodulation. On ne représente que le signal reçu en phase (voie I) car le signal en quadrature est composé uniquement de bruit, compte tenu du choix de nos vecteurs symboles. Dans cet exemple, 3 utilisateurs sont adressés simultanément avec un miroir à retournement temporel de 8 voies et un débit d'information de 100 MBs/s.

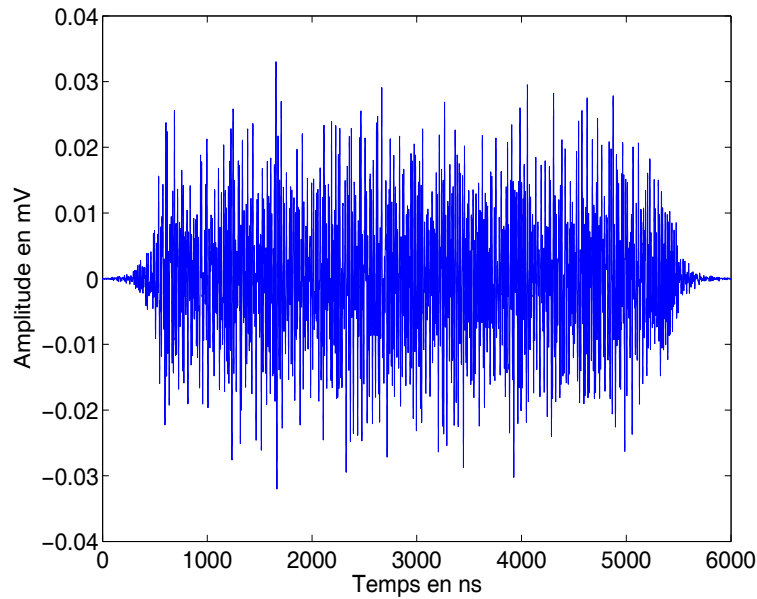


FIG. II.29 – Exemple d'une trame de 500 symboles reçue par un utilisateur après démodulation.

Par la suite, nous allons présenter plusieurs type de mesures. Dans un premier temps, nous avons réalisé des mesures à bruit externe nul. Dans ce contexte, nous présentons des mesures de taux d'erreur binaire en fonction du débit d'information, puis en fonction du nombre d'antennes présent dans le miroir à retournement temporel et du nombre d'utilisateurs. Dans ces mesures, on ne tiendra pas compte de l'énergie émise, les erreurs étant uniquement dues aux lobes secondaires du retournement temporel. Puis, dans la partie suivante, nous verrons comment il

⁵On aurait tout aussi pu choisir de représenter un 0 par une impulsion en I et un 1 par une impulsion en Q, ou toute autre combinaison linéaire, mais ces choix sont en réalité équivalents.

est possible de tenir compte du bruit externe. Nous développerons un formalisme qui permet d'expliquer les résultats et soulignerons alors l'intérêt que procure le retournement temporel en tant que filtre adapté.

II.4.2 Influence du nombre d'antennes dans le miroir et du nombre de récepteurs.

Nous avons en premier lieu étudié le taux d'erreur binaire des communication par retournement temporel en fonction du débit d'information transmis. Nous avons réalisé ces mesures pour des débit par récepteurs allant de 200 MBs/s, c'est-à-dire un espacement entre symboles δt de 5 ns, jusqu'à un débit de 20 MBs/s, soit $\delta t = 50$ ns. Le nombre de récepteurs adressés est égal à 8, ce qui implique que le débit d'information total varie entre 1.6 GBs/s et 160 MBs/s. Ces mesures ont été réalisées pour des tailles de miroir à retournement temporel entre 1 et 8 antennes. Pour chaque taille de miroir étudié, si le nombre d'antennes M était inférieur à 8, nous avons réalisé des moyennes sur les différentes configurations possibles de M antennes parmi les 8 disponibles. Les résultats de ces mesures sont représentés en logarithme décimal sur la figure II.30.

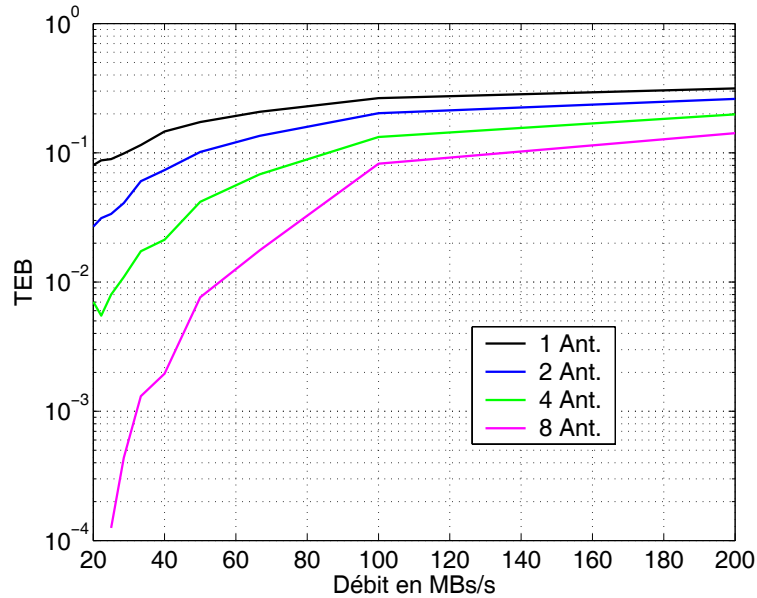


FIG. II.30 – Etude du taux d'erreur binaire en fonction du débit d'information pour différentes tailles de MRT et lorsque 8 récepteurs sont adressés simultanément.

Il est clair que lorsque l'on diminue le débit d'information, la quantité d'erreur commise lors de la transformation d'information est réduite. De même, ajouter des antennes au miroir à retournement temporel a pour effet d'améliorer la qualité de la communication. Une nouvelle fois,

cela s'explique par le fait que les interférences inter-symboles sont diminués dans les deux cas précédents. Nous allons à présent vérifier le formalisme qui avait été développé dans l'expérience à échelle réduite de la deuxième partie de ce chapitre. Pour ce faire, on rappelle lorsque le bruit externe est négligeable, la qualité de la communication est limitée par les interférences inter-symboles et inter-utilisateurs. Le rapport signal-sur-bruit interne du retournement temporel s'écrivant :

$$\text{RSB}_{int} = \Delta\nu\delta t \frac{M}{N} \quad (\text{II.31})$$

où N est le nombre de récepteurs adressés simultanément, M est le nombre d'antennes du MRT, δt l'espacement inter-symboles et $\delta\nu$ est la bande passante utilisée.

Le logarithme de la probabilité d'erreur, comme nous l'avons vu, peut être approché par l'opposé du rapport signal-sur-bruit, soit :

$$\log(P_e) = -\Delta\nu\delta t \frac{M}{N} \quad (\text{II.32})$$

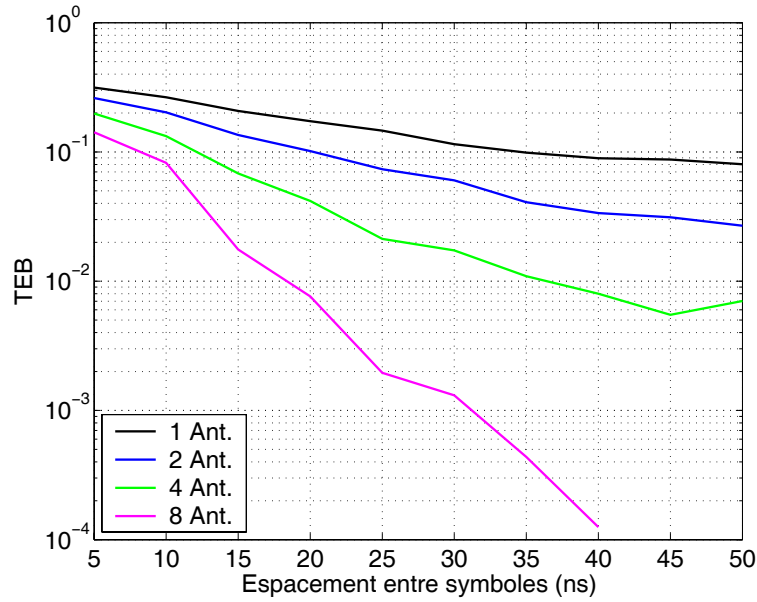


FIG. II.31 – Etude du taux d'erreur binaire en fonction de l'espacement entre les symboles pour différentes tailles de MRT.

La figure II.31 montre l'évolution du logarithme décimal du taux d'erreur binaire en fonction de l'espacement entre les symboles, et ce pour des tailles de miroir à retournement temporel de 1, 2, 4 et 8 antennes. L'évolution est bien linéaire en fonction de l'espacement inter-symboles comme le prévoit la formule II.32.

Comme pour les mesures réalisées à échelle réduite, nous avons par la suite mesuré les coefficients directeurs des droites obtenues précédemment, et ce pour les diverses tailles de miroir à retournement temporel dont nous disposons. Le résultat de ces régression linéaires est représenté sur la figure II.32.

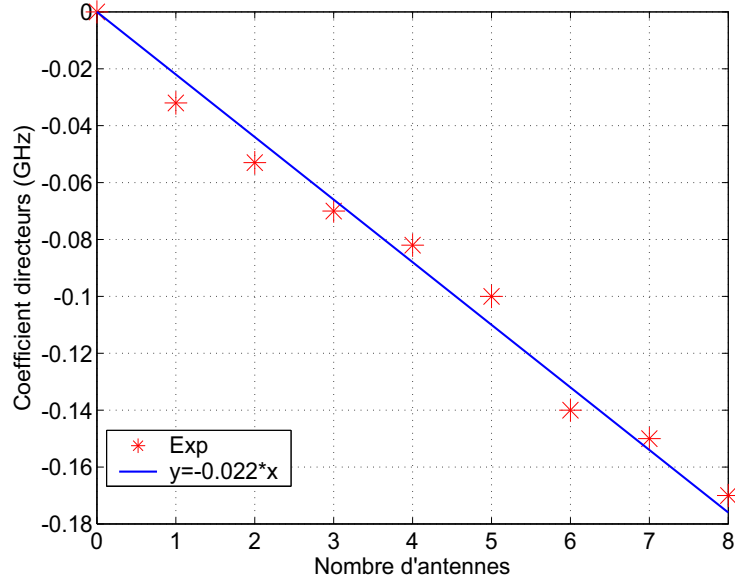


FIG. II.32 – Evolution des coefficients directeurs de $\log(P_e)$ en fonction de la taille du MRT, lorsque 8 utilisateurs sont adressés simultanément.

Comme prédit par la formule II.32, l'évolution des coefficients directeurs en fonction du nombre d'antenne est linéaire. Grâce au coefficient directeur obtenu par la régression linéaire, il est alors possible d'estimer la bande passante utilisée de façon expérimentale : 176 MHz. Cette valeur est tout à fait dans l'ordre de grandeur de la bande passante réelle du système.

Nous avons procédé au même type de mesures, mais en faisant cette fois varier le nombre de récepteurs auxquels on souhaitait envoyer de l'information. La taille du miroir à retournement temporel a été fixée à 2 antennes, et nous avons fait varier le nombre de récepteurs de 1 à 8. Les mesures de taux d'erreur ont été réalisées pour des débits par utilisateur égaux à ceux que nous avons donnés précédemment. La figure II.33 représente l'évolution du logarithme du taux d'erreur binaire en fonction de l'espacement entre les symboles.

On remarque que l'évolution de ces courbes est également linéaire par rapport à l'espacement entre les symboles. Nous avons donc à nouveau réalisé des régressions linéaires de ces droites, et ce pour un nombre d'utilisateurs compris entre 1 et 8. Sur la figure II.34, nous avons tracé l'inverse des coefficients directeurs estimés en fonction du nombre de récepteurs adressés simultanément. Comme prévu par la formule II.32, cette courbe est bien une droite, dont nous avons

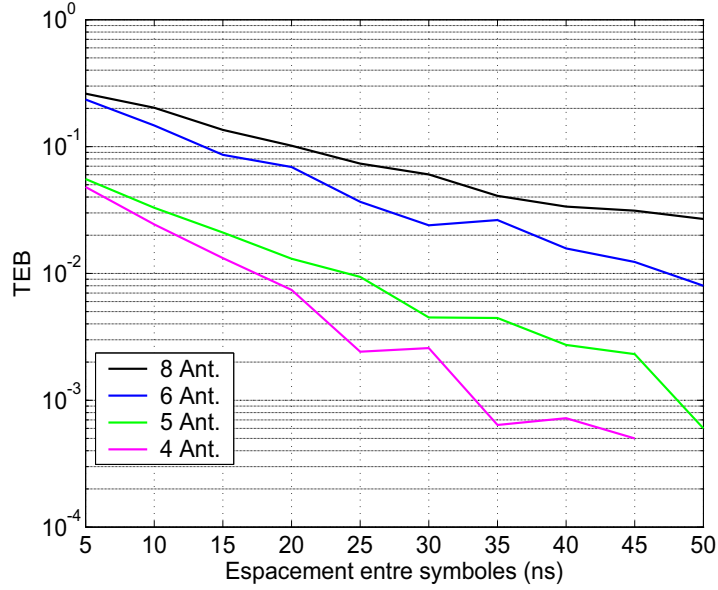


FIG. II.33 – Etude du taux d’erreur binaire en fonction de l’espacement entre les symboles pour un MRT de 2 antennes et divers nombres d’antennes réceptrices.

réalisé une régression, afin d’en calculer le coefficient directeur. Elle est également représentée sur la même figure. A partir de la pente mesurée, nous avons à nouveau accès à une estimation expérimentale de la bande passante : 180 MHz.

La valeur mesurée est elle aussi cohérente avec celle que nous avons estimée précédemment. Ces mesures permettent de valider le formalisme que nous avons établi pour estimer la probabilité d’erreur des communications par retournement temporel lorsque le bruit présent dans le canal est négligeable. Nous allons voir dans la partie suivante comment il est possible de prendre en compte le bruit externe.

II.4.3 Télécommunications dans un canal bruité

Afin de prendre en compte le bruit présent dans le canal de propagation, il est nécessaire de considérer l’énergie qui est émise par le miroir à retournement temporel, lors du processus de communication. Les résultats que nous présentons ici, ainsi que le modèle développé, considèrent un système pour lequel l’amplitude instantanée d’émission est limitée, c’est-à-dire que la puissance instantanée d’émission est limitée. Cette limite en amplitude est fixée à 1, sans que ceci soit limitatif, et ce quel que soit le nombre d’antennes présentes dans le miroir à retournement temporel ainsi que le nombre d’utilisateurs auxquels l’on souhaite s’adresser. Nous allons considérer que lorsque le miroir émet une impulsion à pleine puissance (c’est-à-dire d’ampli-

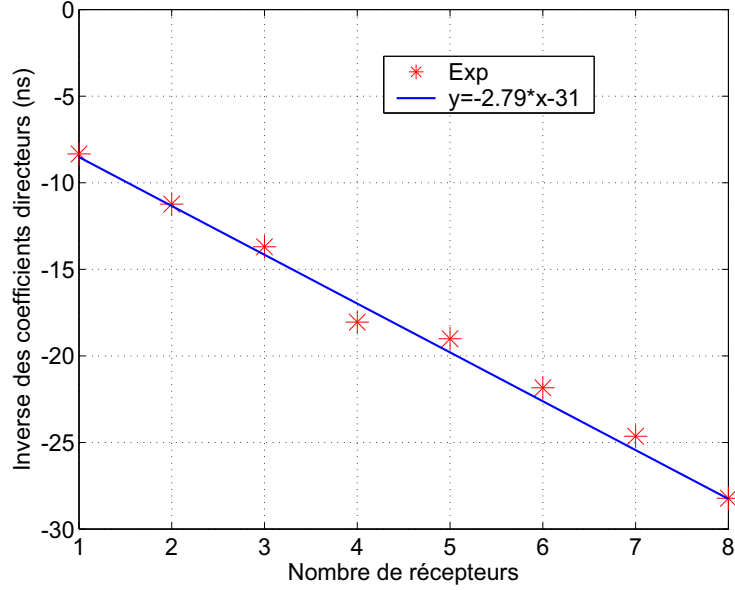


FIG. II.34 – Evolution de l'inverse des coefficients directeurs de $\log(P_e)$ en fonction du nombre de récepteur et pour un MRT de 2 antennes.

tude 1), l'énergie moyenne de la réponse impulsionnelle, considérée équivalente pour tous les utilisateurs, est notée E_s . Le bruit externe dans la bande passante est noté d'une façon usuelle N_0 . Lorsque l'on émet des impulsions successives, c'est à dire sans pratiquer de retournement temporel, le rapport signal-sur-bruit peut ainsi s'écrire :

$$\text{RSB}_0 = \frac{E_s}{N_0} \quad (\text{II.33})$$

Nous allons considérer à présent une communication par retournement temporel avec une base composée de M antennes et vers N utilisateurs. Le bruit total est la somme des interférences inter-symboles et inter-utilisateurs créées par le retournement temporel, et du bruit externe présent dans le canal de communication. L'énergie de chaque symbole, quant à elle, dépend des paramètres physiques qui régissent la qualité de la compression temporelle. Nous allons estimer l'énergie de chaque symbole reçu par un utilisateur quelconque, que l'on assimile à son amplitude maximale mise au carré, ainsi que l'énergie moyenne du bruit total.

Comme précédemment, l'amplitude au carré d'un symbole reçu va dépendre de l'énergie moyenne de la réponse impulsionnelle, ainsi que du gain de compression temporelle, ce qui s'écrit avant normalisation de l'énergie d'émission :

$$E_{RT} = \tau^2 \Delta \nu^2 M^2 E_s \quad (\text{II.34})$$

Le bruit, quant à lui, est la somme du bruit externe et du bruit interne au retournement temporel, que l'on appelle bruit d'interférence. Ce dernier peut s'écrire simplement en utilisant le rapport signal-sur-bruit interne calculé précédemment (equation II.31) :

$$I_{RT} = \frac{E_{RT}}{RSB_{int}} = \frac{\tau^2}{\delta t} \delta \nu M N E_s \quad (\text{II.35})$$

Etant donné que nous avons choisi de travailler à énergie d'émission constante, il est nécessaire de normaliser ces deux grandeurs, afin que l'énergie émise soit constante. Or lorsque l'on fabrique les signaux à émettre par la base, pour chaque antenne, la résultante est la somme de $\tau/\delta t$ réponses impulsionnelles multipliées par les symboles à émettre, et ce pour chaque récepteurs, comme le montre l'équation II.29. Si l'on admet que le signal calculé est une somme de variables aléatoires, on peut alors écrire l'énergie moyenne des signaux fabriqués pour chaque antenne i de la base :

$$E_{em}^i = \frac{\tau}{\delta t} N \quad (\text{II.36})$$

Soit, étant donné que l'émission se fait avec N antennes, une énergie d'émission totale qui s'écrit :

$$E_{em} = \frac{\tau}{\delta t} M N^2 \quad (\text{II.37})$$

Afin de travailler à énergie d'émission constante et égale à l'unité, il est donc nécessaire de normaliser l'émission d'un facteur égal à E_{em} . Par conséquent, le carré de l'amplitude d'un symbole reçu par un utilisateur va être divisé par cette grandeur, ainsi que le bruit créé par les interférences, ce qui donne :

$$E_{RT}^{Norm} = \frac{1}{N} \tau \Delta \nu^2 \delta t E_s \quad (\text{II.38})$$

pour l'énergie d'un symbole, alors que l'énergie des interférences s'écrit :

$$I_{RT}^{Norm} = \frac{1}{M} \tau \Delta \nu E_s \quad (\text{II.39})$$

L'énergie moyenne du bruit total peut alors s'exprimer en fonction de l'énergie du bruit externe et du bruit interne :

$$N_{tot} = N_0 + \frac{1}{M} \tau \Delta \nu E_s \quad (\text{II.40})$$

En se servant des équations II.38 et II.40, que l'on divise par l'énergie moyenne des réponses impulsionnelles E_s , on peut alors définir un nouveau rapport signal-sur-bruit pour les communications par retournement temporel. Celui-ci dépend du rapport signal-sur-bruit initial et des caractéristiques du MRT ainsi que du milieu de propagation et est défini comme :

$$\text{RSB}_{RT} = \frac{\frac{1}{N}\tau\Delta\nu^2\delta t.\text{RSB}_0}{1 + \frac{1}{M}\tau\Delta\nu.\text{RSB}_0} \quad (\text{II.41})$$

Cette nouvelle définition du rapport signal-sur-bruit appelle quelques remarques. En premier lieu, il est utile de vérifier que lorsque le rapport signal-sur-bruit sans retournement temporel SNR_0 est grand, c'est à dire que le bruit externe est négligeable, on retrouve bien le rapport signal-sur-bruit interne RSB_{int} du retournement temporel défini par la formule II.31. La deuxième remarque que l'on peut formuler vis-à-vis de ce dernier résultat concerne le gain en terme d'énergie que peut apporter un système de télécommunications par retournement temporel. En effet, dans divers articles concernant les télécommunications par retournement temporel, dont le plus complet est sans doute [61], les auteurs ne considèrent que les interférences inter-symboles et inter-utilisateurs provoquées par le retournement temporel, sans tenir compte du bruit présent dans le canal de communication. Cependant, le masque de la FCC concernant le spectre en puissance alloué aux communications ultra large bande passante (Fig. II.2), montre à quel point ce facteur va être important, car les puissances d'émission autorisées sont très faibles. Ainsi, si l'on se place dans l'approximation d'un canal très bruité ($\text{RSB}_0 \rightarrow 0$), ou de façon équivalente d'une énergie d'émission très faible, on peut écrire le rapport signal à bruit avec retournement temporel comme :

$$\text{RSB}_{RT} \approx \frac{1}{N}\tau\Delta\nu^2\delta t.\text{RSB}_0 \quad (\text{II.42})$$

Le retournement temporel va donc apporter un gain en puissance sur chaque utilisateur de $10\log_{10}(\frac{1}{N}\tau\Delta\nu^2\delta t)$ dB. Plus la bande passante sera élevée et le temps de réverbération du milieu long, plus ce gain sera élevé. A titre d'exemple, si l'on considère un système dont la bande passante est de 2 GHz, que l'on fait fonctionner à 500 MBs/s ($\delta t = 2$ ns) vers 10 utilisateurs simultanément et dans un milieu dont le temps de réverbération typique est de l'ordre de 50 ns, le retournement temporel offre un gain de 10 dB. Ceci montre un autre intérêt de ce type de systèmes. A nouveau, dans le cas d'un système SISO il est possible de considérer un filtrage adapté à la réception qui donnerait le même gain. Ceci est vrai dans le cas où les amplificateurs de réception sont parfaitement linéaires, c'est-à-dire qu'ils n'ont pas de seuil en dessous duquel

le signal n'est pas amplifié, autrement dit, seulement en théorie. Le retournement temporel, parce qu'il réalise un filtrage adapté physiquement dans le milieu, permet de travailler avec des puissances bien inférieures à celles d'un système conventionnel.

A partir de cette définition du rapport signal à bruit, il est facile de calculer la probabilité d'erreur d'une communication par retournement temporel. Pour une modulation de type BPSK, celle-ci s'écrit :

$$P_e = \frac{1}{2} \operatorname{erfc} \left(\frac{\frac{1}{N} \tau \Delta \nu^2 \delta t \cdot \text{RSB}_0}{1 + \frac{1}{M} \tau \Delta \nu \cdot \text{RSB}_0} \right) \quad (\text{II.43})$$

Afin de valider ce formalisme, nous avons réalisé des mesures de taux d'erreur binaire lorsque du bruit externe est ajouté. Pour des raisons de gain de temps, ces mesures ont été réalisées à partir des réponses impulsionnelles expérimentales, mais en calculant les signaux reçus numériquement (formule II.30). Le nombre d'antennes du miroir à retournement temporel est fixé à 1, ainsi que le nombre de récepteurs. L'espacement entre les symboles est de $\delta t = 18$ ns. Le temps de réverbération de la cavité est toujours à peu près égal à 160 ns, et la bande passante est fixée à 180 MHz, comme mesurée dans la partie précédente. Nous avons fait varier le bruit de -70 dB à +50 dB (une telle dynamique est permise car les mesures sont réalisées à numériquement). Sur la figure II.35 nous représentons le taux d'erreur binaire (qui est égal ici à la probabilité d'erreur) en fonction du niveau de bruit externe en dB.

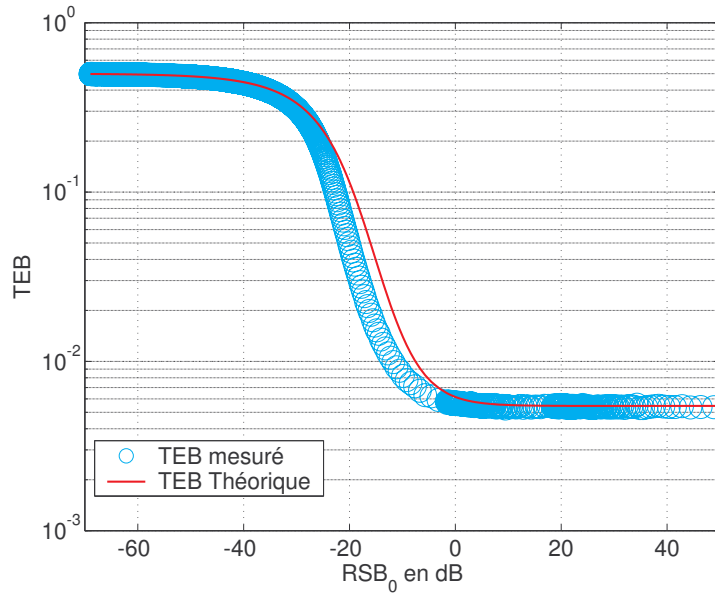


FIG. II.35 – Etude du taux d'erreur binaire en fonction du bruit externe dans le canal.

Nous avons également représenté sur cette courbe les valeurs données par la formule théorique (équation II.43). Les deux courbes se superposent assez bien, ce qui confirme le modèle développé

ici. Sur cette courbe on vérifie que lorsque le rapport signal-sur-bruit externe augmente, la probabilité d'erreur converge vers la limite qui est inhérente au retournement temporel. Il faut également noter que lorsque l'on se place à 0 dB la probabilité d'erreur n'est pas de 15% comme cela devrait être le cas sans aucun filtrage. Celle-ci est déjà limitée par le retournement temporel : grâce au gain en énergie que procure le retournement temporel, même dans un environnement très bruité, il est possible de transmettre de l'information avec un taux d'erreur faible (ici inférieur à 1%).

A l'issue de cette discussion, une étude assez complète des communications par retournement temporel des ondes électromagnétiques a été réalisée. Cependant nous allons voir qu'il est possible de tirer d'autres avantages du retournement temporel, en diminuant la taille des réseaux multi-antennes.

II.4.4 Focalisation sub-longueur d'onde et télécommunications

Nous avons vu que l'augmentation du débit d'information passera nécessairement par l'utilisation de réseaux multi-antennes. L'un des problèmes liés à cette technologie est l'encombrement de ces réseaux. En effet, il est communément admis que pour obtenir une efficacité maximale, des antennes d'un même réseau doivent être séparées d'une distance supérieure à une demi-longueur d'onde [40]. Ceci a pour effet de limiter les applications de ces systèmes : si l'objet "communiquant" est d'une taille petite devant ou comparable à la longueur d'onde (i.e. un ordinateur portable), il ne pourra pas comporter de réseau de réception mais une unique antenne. Une solution contre ce problème pourrait venir des antennes microstructurées que nous avons introduites dans la partie consacrée à la focalisation sub-longueur d'onde. En effet, si nous sommes capables par retournement temporel de focaliser de l'énergie sur des antennes séparées par une distance faible devant la longueur d'onde, comme nous l'avons présenté dans le premier chapitre de ce manuscrit, il est possible par ce même principe de transmettre des informations différentes sur chacune de ces antennes.

Dans le domaine des télécommunications, on considère que deux antennes reçoivent des signaux indépendants si le coefficient de corrélation des signaux reçus par les 2 antennes est inférieur à 0.4. Cette condition étant réalisée, comme nous l'avons montré sur la figure I.39, nous avons voulu prouver l'intérêt d'utiliser de telles antennes dans des systèmes de type MIMO, et c'est naturellement par retournement temporel que nous avons réalisé cette démonstration. Pour ce faire, comme le montre la figure II.36 nous cherchons à transmettre une image en couleur. Celle

ci est séparée en trois flux de données binaires codant les couleurs primaires rouge/vert/bleu (RBG).

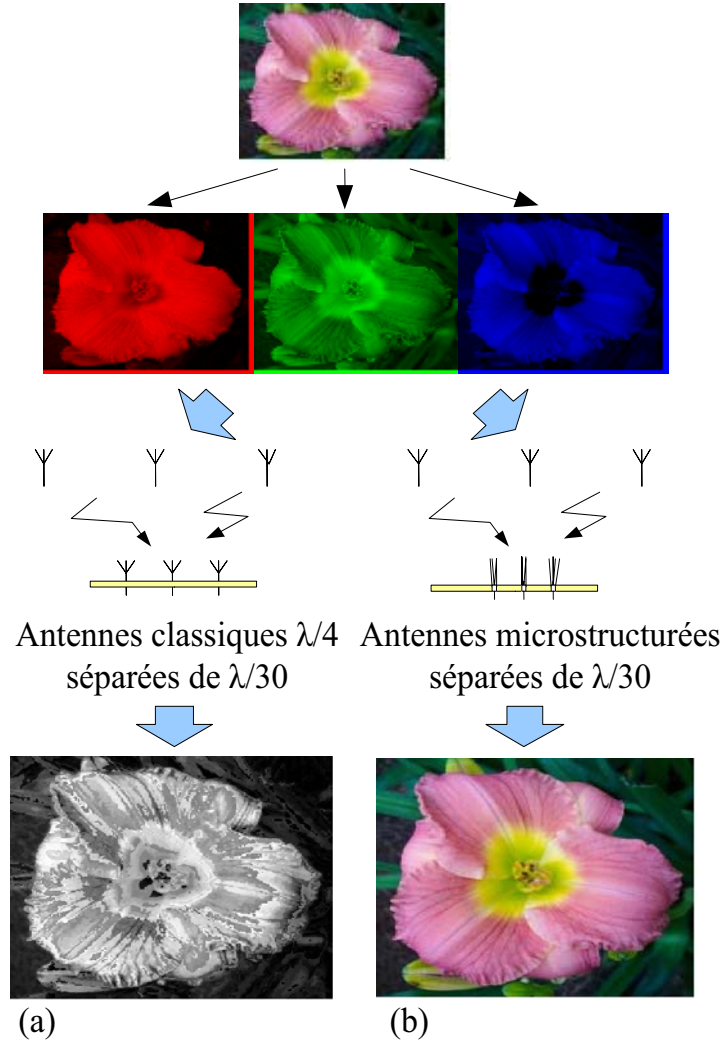


FIG. II.36 – Transmission par retournement temporel dans la cavité d’une photo couleur. Le miroir à retournement temporel émet 3 flux de données binaires (RBG) à 3 antennes séparées de $\lambda/30$. Image reçue avec des antennes ordinaires (a), et avec les antennes microstructurées (b).

Chaque flux de données va être envoyé simultanément, et sur la même bande passante, à trois antennes réceptrices adjacentes d’un réseau de 8 antennes. Le miroir à retournement temporel est constitué de 3 antennes parmi les 8 antennes du MRT. L’expérience est réalisée avec 2 types d’antennes réceptrices : d’abord avec des antennes de type “monopôle sur un plan de masse”, qui sont séparées de $\lambda/30$, puis avec nos antennes microstructurées, décrites dans la partie 5 du premier chapitre de ce manuscrit, et dont l’espacement est également de $\lambda/30$. Dans les deux cas, la modulation est de type BPSK, le débit de 50 MBs/s vers chaque antenne de

réception, soit un débit total de 150 MBs/s. La démodulation est réalisée classiquement et, chaque flux ayant été démodulé, on reconstruit l'image avec la carte RGB reçue. Comme le montre la figure II.36.(a), avec des antennes classiques, l'image reçue est en niveaux de gris. Sur la figure II.36.(b), qui montre l'image reçue avec nos antennes microstructurées, on voit que l'image a été reconstruite sans être dégradée. Ceci s'explique de façon très simple. Lorsque nous utilisons des antennes classiques en réception, les signaux reçus par ces trois antennes sont très corrélés : chaque récepteur reçoit donc le même bit d'information à chaque instant. Ainsi l'image reçue est constituée de 3 cartes RGB identiques, et l'image construite est donc en niveaux de gris. Au contraire, les antennes que nous avons fabriquées reçoivent toutes des flux de données binaires différents car nous sommes capables par retournement temporel de les adresser indépendamment. La carte RGB reçue est donc la même que celle que nous avons émise, à quelques erreurs près. Ceci est une démonstration très visuelle de l'intérêt que pourrait apporter ce type d'antennes dans le domaine des télécommunications.

Dans cette partie nous avons étendu l'étude des communications par retournement temporel au domaine des ondes électromagnétiques. A l'aide du miroir à retournement temporel que nous avons développé à 2.45 GHz, et dans une cavité recouverte d'aluminium, nous avons procédé à des expériences de télécommunication à très haut débit dans un environnement particulièrement réverbérant. Ainsi, nous avons vu que lorsque le bruit extérieur est faible, comme en acoustique, le logarithme de la probabilité d'erreur décroît linéairement avec le nombre d'antennes dont est composé le miroir. De plus, nous avons mis en évidence que celui évolue de façon inversement proportionnelle au nombre d'utilisateurs auxquels nous nous adressons. Puis, nous avons étudié expérimentalement l'influence du bruit externe présent dans le canal de communication. Un modèle simple a été développé pour expliquer les mesures obtenues, qui montre un autre avantage du retournement temporel : celui-ci, en tant que filtre adapté, maximise l'énergie reçue sur un utilisateur. Ceci a pour effet d'augmenter le rapport signal-sur-bruit en réception d'une quantité que nous avons exprimée en fonction des caractéristiques du miroir à retournement temporel et de celles du milieu de propagation. Enfin, nous avons présenté des résultats de communication par retournement temporel vers des antennes microstructurées séparées d'une distance très inférieures à la longueur d'onde. Nous avons montré que contrairement à des antennes classiques qui reçoivent toutes le même signal lorsqu'elles sont placées très proches les unes des autres, il est possible d'adresser ces antennes de façon indépendante. Cette dernière expérience montre l'intérêt d'utiliser de telles antennes dans les systèmes MIMO, particulièrement lorsque l'on procède par retournement temporel.

II.5 Conclusion

Nous avons consacré le deuxième chapitre de cette thèse aux télécommunications par retournement temporel dans les milieux complexes, et plus spécifiquement dans les milieux fortement réverbérants. Cette étude se place dans le contexte des télécommunications de l'avenir. En effet, pour augmenter les débits d'information dans le domaine des communications sans fil, deux techniques sont en train de voir le jour. La première, qui utilise des systèmes multi-antennes de type MIMO, tire profit de la complexité d'un milieu propagation : c'est ce que réalise retournement temporel de façon naturelle [52]. La deuxième technique, basée sur l'utilisation de signaux à très large bande passante, est soumise aux problèmes de l'allongement temporel des signaux et à la faible puissance qui lui est allouée dans les diverses normes internationales. Le retournement temporel, parce qu'il est un filtre adapté à la propagation, permet de lutter efficacement contre ces deux inconvénients.

Après avoir présenté les premières expériences de télécommunications par retournement temporel réalisées au laboratoire, nous nous sommes approchés d'un milieu de propagation réel. Nos études ont été réalisées avec des ondes ultrasonores dans une cuve remplie d'eau. Un modèle décimétrique d'environnement indoor a été fabriqué afin de simuler des communications dans le gigahertz dans un bâtiment. Nous avons ainsi été en mesure de mesurer l'influence des divers paramètres physiques sur le taux d'erreur binaire d'une communication par retournement temporel et un modèle simple a été développé pour interpréter les résultats. Ceci nous a permis de conclure que lorsque le bruit externe est faible, le logarithme de ce taux d'erreur décroît linéairement en fonction de la bande passante du système, du nombre d'antennes d'émission et de l'espacement entre les symboles.

Puis, nous avons comparé le retournement temporel, qui est un filtre adapté, au filtre inverse. Ces expériences ont à nouveau été réalisées avec des ondes ultrasonores. Pour ce faire, une procédure itérative basée sur le retournement temporel a été introduite. Pour un grand nombre d'itérations, celle-ci converge vers le filtre inverse. A travers des mesures de taux d'erreur, nous avons conclu que le retournement temporel doit être appliqué lorsque le canal de propagation est très bruité, alors que le filtre inverse est préférable lorsque le bruit externe est très faible. Dans un cas intermédiaire, il est possible de stopper la procédure itérative afin de minimiser les erreurs commises.

Enfin, des expériences ont été réalisées avec des ondes électromagnétiques à 2.45 GHz, dans une cavité réverbérante. Celles-ci nous ont permis de vérifier les résultats que nous avons d'abord

obtenus grâce à nos études avec des ondes acoustiques. Nous avons également étudié l'influence du nombre de récepteurs auxquels on s'adresse et montré comment celui-ci dégrade la qualité de la transmission d'information. Nous avons ensuite considéré le bruit présent dans le canal et introduit un rapport signal-sur-bruit propre aux communications par retournement temporel. Ce rapport montre que, dans le cas de milieux très bruités, le retournement temporel peut se révéler très utile. Enfin, nous avons exposé une première expérience de communication sub-longueur d'onde par retournement temporel réalisée à l'aide des antenne microstructurées que nous avons décrites dans le premier chapitre de ce manuscrit.

Conclusion et Perspectives

Lors de ce travail de thèse, nous avons montré en quoi le retournement temporel d'ondes électromagnétiques peut se montrer utile pour les générations futures de systèmes de communication sans fil. Les diverses expériences réalisées ont permis de vérifier que celui-ci est à la fois un pré-égaliseur des signaux pour les milieux fortement réverbérants, un multiplexeur spatial pour les environnements complexes et un filtre adapté à la propagation. Ces divers aspects en font un candidat de choix pour les communications à très large bande passante et multi-antennes. Cependant nos mesures ont toujours été réalisées dans des milieux statiques : les futures études devront donc prendre en considération le caractère dynamique des canaux de communication. De même, lors de nos expériences, nous nous sommes toujours limités à des milieux idéaux. Des mesures en environnements réels, dans lesquels certains objets diffusent beaucoup plus que d'autres, sont nécessaires afin de valider la technique.

Au cours de nos travaux, nous avons également obtenu des résultats particulièrement novateurs de focalisation sub-longueur d'onde par retournement temporel. Ces travaux, les premiers sur le sujet, posent de nombreuses questions. Ainsi, la transposition de la méthode employée à d'autres gammes de fréquence du spectre électromagnétique, voir à d'autres types d'ondes, fera certainement l'objet de nombreuses études à venir. De façon plus générale, ces observations font le lien entre le retournement temporel et plusieurs sujets passionnants de recherches actuels (matériaux à indices de réfraction négatifs, cristaux photoniques, plasmons de surface), dont le but est d'obtenir des résolutions indépendantes de la longueur d'onde, que ce soit pour imager un milieu ou pour focaliser de l'énergie. C'est d'ailleurs vers ces domaines que vont s'orienter nos futurs travaux.

Par ailleurs, cette thèse nous a également conduits à nous intéresser à divers sujets qui, n'entrant pas directement dans le cadre du sujet de thèse, ont été omis. Parmi ces sujets l'on peut citer l'étude du cône de rétrodiffusion cohérente dans le domaine du GHz, réalisée avec le même appareillage que notre miroir à retournement temporel. Ces travaux, non terminés, feront un

sujet d'étude intéressant dans le futur. Nous nous sommes également intéressés à des mesures de corrélations de bruits électromagnétiques, comme cela a d'abord été étudié par Weaver et ses collaborateurs en acoustique pour mettre au point un système d'imagerie passive. Le problème semble plus complexe dans le domaine des ondes décimétriques, car comme lors de nos mesures de taches focales, les antennes et câbles utilisés sont autant de diffuseurs des champs mesurés : ceci reste également un problème ouvert. Enfin, à l'aide de notre cavité réverbérante et de notre dispositif impulsionnel large bande, nous avons mis au point une méthode de mesure de section efficace de diffusion d'objets quelconques, fondée sur les travaux de Julien de Rosny en acoustique [13, 64], qui est en cours de publication. Une étude plus complète, donnant également accès à la section efficace d'absorption, pourrait être réalisée avec le même matériel.

Bibliographie

- [1] MR. Andrews, PP. Mitra, and R. deCarvalho. Tripling the capacity of wireless communications using electromagnetic polarization. *Nature*, 409(6818) :316–318, 2001.
- [2] A.L. Moustakas, H.U. Baranger, L. Balents, A.M. Sengupta, and S.H. Simon. Communication through a diffusive medium : Coherence and capacity. *Science*, 287(5451) :287–290, 2000.
- [3] M. Fink. Time reversed acoustics. *Physics Today*, 50 :34, 1997.
- [4] D. Cassereau and M. Fink. Time-reversal of ultrasonic fields. iii. theory of the closed time-reversal cavity. *IEEE Trans. Ultra. Ferr. Freq. Cont.*, 39-5 :579–592, 1992.
- [5] D. Cassereau. Le retournement temporel en acoustique - théorie et modélisation - habilitation à diriger des recherches, université paris 7, 1997.
- [6] M. Fink, D. Cassereau, A. Derode, C. Prada, P. Roux, M. Tanter, J.L Thomas, and F. Wu. Time-reversed acoustics. *Reports on Progress in Physics*, 63-12, 2000.
- [7] A. Derode, P. Roux, and M. Fink. Robust acoustic time reversal with high-order multiple scattering. *Phys. Rev. Lett.*, 75 :4206, 1995.
- [8] A. Derode, A. Tourin, and M. Fink. Limits of time-reversal focusing through multiple scattering : long range correlation. *J. Acoust. Am.*, 107(6) :2987–2998, 2000.
- [9] W.C. Sabine. *Collected Papers on Acoustics*. Harvard University Press, 1922.
- [10] C. Draeger and M. Fink. One-channel time reversal of elastic waves in a chaotic 2d-silicon cavity. *Phys. Rev. Lett.*, 79 :407, 1997.
- [11] D. Royer and E. Dieulesaint. Mesures optiques de déplacement amplitude 10^{-4} à 10^2 angström. *Revue Phys. Appl*, 24 :833–846, 1989.
- [12] A. Derode, A. Tourin, and M. Fink. Random multiple scattering of sound, ii. is time reversal a self averaging process. *Phys. Rev. E*, 64 :036606, 2001.
- [13] Julien de Rosny. *Milieux réverbérants et réversibilité*. PhD thesis, Université Paris 6, 2000.

- [14] R. Carminati and M. Nieto-Vesperinas. Reciprocity of evanescent electromagnetic waves. *J. Opt. Soc. Am. A*, 15-3 :706–712, 1998.
- [15] C. H Papas. *Theory of Electromagnetic Wave Propagation*. Dover Publications, 1988.
- [16] C.A. Balanis. *Antenna Theory, Analysis and Design*. New York, second edition, 1997.
- [17] Arie Voors. Software 4nec2, [http ://home.ict.nl/ arivoors/index.html](http://home.ict.nl/~arivoors/index.html).
- [18] G. Lerosey, J. de Rosny, A. Tourin, A. Derode, G. Montaldo, and M. Fink. Time reversal of electromagnetic waves. *Phys. Rev. Lett.*, 92-19 :193904, 2004.
- [19] *Procédé pour inverser temporellement un onde*. Inventeurs : M.Fink, G. Lerosey, J. de Rosny, A. Derode, and A. Tourin. Brevet N^oFR2.868.894. .Déposants : CNRS.
- [20] G.F. Edelmann, T. Akal, W.S. Hodgkiss, S. Kim, W.A. Kuperman, and H.C. Song. An initial demonstration of underwater acoustic communication using time reversal. *IEEE J. Oceanic. Eng.*, 3 :602–609, 2002.
- [21] G. Montaldo, D. Palacio, M. Tanter, and M. Fink. The time reversal kaleidoscope : a new concept of smart transducers for 3d ultrasonic imaging. *Appl. Phys. Lett.*, 84-19 :3879–3881, 2004.
- [22] J.-L. Thomas, F. Wu, and M. Fink. Time reversal focusing applied to lithotripsy. *Ultrasonic Imaging*, 18 :106–121, 1996.
- [23] A. Derode, A. Tourin, and M. Fink. Ultrasonic pulse compression with one bit time reversal through multiple scattering. *J. App. Pys.*, 85-9 :6343–6352, 1999.
- [24] Ph. de Doncker. Spatial correlation functions for fields in three-dimensionnl rayleigh channels. *PIER*, 40 :55–69, 2003.
- [25] G. Lerosey, J. de Rosny, A. Tourin, and M. Fink. Time reversal of wideband microwaves. *Appl. Phys. Lett.*, 88 :154101–1–3, 2006.
- [26] J.W Goodman. *Introduction to Fourier Optics*. McGraw-Hill, 1996.
- [27] R. Carminati and C Boccara. *Les nanosciences, Chapitre 5 : Champ proche optique, de l'expérience à la théorie*. proofs, 2006.
- [28] J. de Rosny and M. Fink. Overcoming the diffraction limit in wave physics using a time-reversal mirror and a novel acoustic sink. *Phys. Rev. Lett.*, 89 :124301, 2002.
- [29] E. H. Synge. *Philos. Mag.*, 6 :356, 1928.
- [30] Jean-Marie Vigoureux. De l'onde évanescence de fresnel au champ proche optique. *Annales de la Fondation Louis de Broglie*, 28, 3-4 :525–549, 2003.

- [31] C. P. Vlahacos, R. C. Black, S. M. Anlage, and A. Amar F. C. Wellstood. Near-field scanning microwave microscope with $100\mu m$ resolution. *Appl. Phys. Lett.*, 69-21 :3272–3274, 1996.
- [32] A. Copty, F. Sakran, M. Golosovsky, D. Davidov, and A. Frenkel. Low-power near-field microwave applicator for localized heating of soft matter. *Appl. Phys. Lett.*, 84-25 :5109–5112, 2004.
- [33] S. Kim, W.A. Kuperman, W.S. Hodgkiss, H.C. Song, G.F. Edelmann, and T. Akal. Robust time reversal focusing in the ocean. *J. Acoust. Soc. Am.*, 114 :145–157, 2003.
- [34] W.A. Kuperman and J.F Lynch. Shallow-water acoustics. *Physics Today*, 57-10 :55–57, 2004.
- [35] P. Roux, B. Roman, and M. Fink. Time-reversal in an ultrasonic waveguide. *Applied Physics Letters*, 70 :1811–1813, April 1997.
- [36] H.L. Bertoni. *Radio Propagation for Modern Wireless Systems*. Prentice Hall Professional Technical Reference, 1999.
- [37] G.J. Foschini and M.J. Gans. On limits of wireless communications in a fading environment when using multiple antennas. *Wireless Personal Communication*, 6 :311–335, 1998.
- [38] C. Shannon. A mathematical theory of communication. *Bell Labs tech. J.*, 27 :379–423, 623–656, 1948.
- [39] I. Telatar. Capacity of multi antenna gaussian channels. *European Trans. Tel.*, 10-6 :585–595, 1999.
- [40] A.J. Paulraj, R.U. Nabar, and D.A. Gore. *Introduction to space-time wireless communications*. Cambridge university press, 2003.
- [41] W. Jakes. *Microwave mobile communications*. Wiley, 1974.
- [42] S. Alamouti. A simple transmit diversity technique for wireless communications. *IEEE J. Sel. Areas Comm.*, 16-8 :1451–1458, 1998.
- [43] T. Cover and J. Thomas. *Elements of information theory*. Wiley, 1991.
- [44] G. Golden, G. Foschini, R. Valenzuela, and P. Wolniansky. Detection algorithm and initial laboratory results using the v-blast space-time communication architecture. *Electronics letters*, 35-1, 1999.
- [45] Revision of part 15 of the commissions rules regarding ultra-wideband transmission systems, first note and order, federal communications commission, et-

- docket. [http://www.fcc.gov/Bureaus/Engineering Technology/Orders/2002/fcc02048.pdf](http://www.fcc.gov/Bureaus/Engineering%20Technology/Orders/2002/fcc02048.pdf), February 2002.
- [46] R. C. Qiu, R. Scholtz, and X. Shen. Ultra wideband wireless communications, a new horizon. *IEEE Trans. Veh. Technol.*, 54-5, 2005.
 - [47] R. C. Qiu, R. Scholtz, and X. Shen. Ultra wideband for multiple access. *IEEE Commun. Mag.*, 43-2 :80–87, 2005.
 - [48] M.Z. Win and R.A. Scholtz. Impulse radio : how it works. *IEEE Commun. Lett.*, 2-2 :36–38, 1998.
 - [49] Q. Li and L.A. Rusch. Multiuser detection for ds-cdma uwb in the home environment. *IEEE J. Sel. Areas Commun.*, 20-9 :1701–1711, 2002.
 - [50] M.Z. Win and R.A. Scholtz. Ultra-widebandwidth time hopping spread spectrum impulse radio for wireless multiple access communications. *IEEE Trans. Commun.*, 48-4, 2000.
 - [51] M. Tanter, J.-L. Thomas, and M. Fink. Time reversal and the inverse filter . *Acoustical Society of America Journal*, 108 :223–234, July 2000.
 - [52] A. Derode, A. Tourin, J. de Rosny, M. Tanter, S. Yon, and M. Fink. Taking advantage of multiple scattering to communicate with time reversal antennas. *Phys. Rev. Lett.*, 90 :014301, 2003.
 - [53] A.A. Saleh and R.A. Valenzuela. A statistical model for indoor multipath propagation. *IEEE J. Sel. Areas Commun.*, 5-2 :128–137, 1987.
 - [54] M.Z. Win and R.A. Scholtz. On the energy capture of ultra -wide bandwidth signals in dense multipath environments. *IEEE Commun. Lett.*, 2 :245–247, 1998.
 - [55] V.P. Tran and A. Sibille. Uwb spatial multiplexing by multiple antennas and rake decorrelation. COST 273, June 2005. <http://www.ensta.fr/tran/TD-05-091.ppt>.
 - [56] J.G. Proakis. *Digital Communications*. McGraw-Hill Higher Education, 2000.
 - [57] G. Lerosey, J. de Rosny, G. Montaldo, A. Tourin, A. Derode, and M. Fink. Time reversal of electromagnetic waves and telecommunication. *Radio Science*, 40-5 :29–39, 2005.
 - [58] M. Tanter, J-F. Aubry, J. Gerber, J-L. Thomas, and M. Fink. Optimal focusing by spatio-temporal inverse filter : Part i. basic principles. *J. Acoust. Soc. Am.*, 101 :37–47, 2001.
 - [59] G. Montaldo, G. Lerosey, A. Derode, A. Tourin, J. de Rosny, and M. Fink. Telecommunication in a disordered environment with iterative time reversal. *Waves Random Media*, 14 :287–302, 2004.

- [60] W.J. highley, P. Roux, and W.A. Kuperman. Relationship between time reversal and linear equalization in digital communications. *J. Acoust. Soc. Am.*, 120(1) :35–38, 2006.
- [61] T. Strohmer, M. Emami, J. Hansen, G. Papanicolaou, and A. Paulraj. Application of time-reversal with mmse equalizer to uwb communications. *in Proc. IEEE Global Telecommunications Conference*, 5 :3123–3127, 2004.
- [62] B.E. Henty and D.D. Stancil. Multipath-enabled super-resolution for rf and microwave communication using phase-conjugate arrays. *Phys. Rev. Lett.*, 93 :243904–1,243904–4, 2004.
- [63] A. Molisch. Ultrawideband propagation channels - theory, measurement, and modeling. *IEEE Transactions on Vehicular Technolgy*, 54 :1528–1545, 2005.
- [64] Julien de Rosny, Claire Debever, Stephane Conti, and Philippe Roux. Diffuse reverberant acoustic wave spectroscopy with absorbing scatterers. *Applied Physics Letters*, 87(15) :154104, 2005.