

> les caractéristiques des liaisons individuelles entre chaque antenne d'émission et chaque antenne de réception, puis on détermine les signaux à émettre à l'aide d'algorithmes de calcul appropriés. Cela nécessite du temps de calcul si l'on augmente le nombre d'antennes et n'a de sens que pour des bandes passantes étroites.

RÉÉMETTRE LE SIGNAL REÇU, MAIS À L'ENVERS

Une autre solution consiste à exploiter une propriété particulière de la propagation des ondes: son invariance par renversement du temps. L'idée est la suivante. Imaginons une source localisée complètement entourée de capteurs. La source émet un signal, et chacun des capteurs enregistre le signal reçu; si les capteurs réémettent les signaux enregistrés en les retournant temporellement (on les renverse, comme lorsqu'on passe un film à l'envers), les ondes feront le chemin inverse et se focaliseront au niveau de la source.

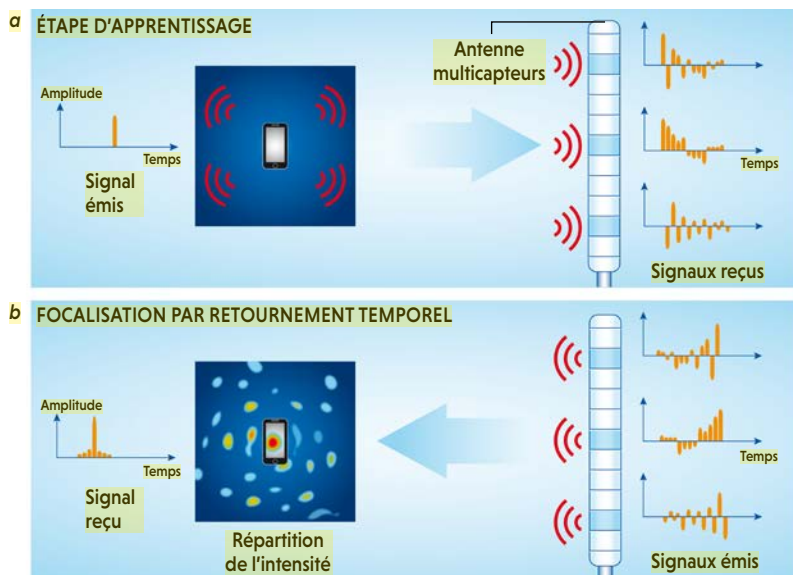
Cette propriété semble nécessiter une surface entière de capteurs. Mais dans les années 1990, l'équipe de Mathias Fink, à l'ESPCI, à Paris, a établi que lorsque le milieu de propagation est assez complexe, un nombre limité d'antennes suffit, sans que la focalisation soit dégradée. Ces physiciens ont ainsi montré que le procédé peut s'appliquer à des ondes électromagnétiques de 2,4 GHz. Et en 2016, les ingénieurs de chez Orange ont fait la démonstration d'un prototype d'une antenne 5G à retournement temporel.

Comment fonctionne-t-elle? Dans une étape d'apprentissage, le mobile émet une brève impulsion; celle-ci arrive au niveau de chaque capteur de l'antenne sous la forme d'une suite d'échos (différente pour chaque capteur), qui est enregistrée (voir l'encadré ci-contre). Cette étape est cruciale et nécessite une bande passante large, pour garantir un bon échantillonnage du signal et une bonne séparation des différents échos. On sait alors que si chaque capteur de l'antenne réémet son signal reçu en le retournant temporellement, on focalisera au niveau du mobile une brève impulsion, presque identique à l'impulsion initiale.

Pour transmettre des données numériques de la station au mobile, on convient d'émettre le signal enregistré retourné temporellement pour coder le bit 1, ou le même multiplié par -1 (son opposé) pour le bit 0. Inutile aussi d'attendre la fin de l'émission d'un enregistrement pour en commencer une autre: il suffit qu'elles soient décalées d'une durée supérieure à

ÉMISSION CIBLÉE GRÂCE AU RETOURNEMENT TEMPOREL

Le principe du retournement temporel des ondes peut être appliqué à la téléphonie afin que l'antenne-relais focalise son signal sur le téléphone destinataire et lui seul. La méthode nécessite une première étape d'« apprentissage » où le destinataire émet un signal bref (a). L'antenne est formée d'une série d'éléments indépendants, dont chacun enregistre le signal reçu. Celui-ci diffère selon la position du capteur, car les chemins suivis par les ondes jusqu'à lui ne sont pas tout à fait les mêmes. Pour émettre de façon ciblée un bit vers le téléphone destinataire, il suffit alors que chaque élément de l'antenne réémette le signal enregistré et retourné temporellement (b). Les ondes qui atteignent alors le téléphone donneront un signal focalisé temporellement, proche de l'original; ailleurs, le signal reçu sera d'intensité faible ou nulle.



l'impulsion originelle pour être discriminées au niveau du mobile. Par ailleurs, on peut profiter de la focalisation spatiale pour émettre simultanément les signaux enregistrés de plusieurs mobiles ou objets connectés. Chacun recevra distinctement le signal qui lui est destiné, et uniquement celui-ci. On peut ainsi augmenter fortement le débit et réduire de beaucoup l'énergie consommée par le réseau, la focalisation spatiale assurant une plus grande efficacité énergétique.

Devant de telles promesses, le retournement temporel sera-t-il choisi pour la 5G? Pas sûr. Ses performances doivent beaucoup à la complexité de l'environnement de propagation et se heurtent à la nécessité de renouveler la phase d'apprentissage lorsque les chemins de propagation ont changé. C'est notamment le cas lorsque l'appareil connecté se déplace: ce n'est pas un problème pour un piéton, ça l'est davantage pour un véhicule en mouvement. Mais les ingénieurs n'ont pas dit leur dernier mot, puisque ceux d'Orange ont conçu un dispositif qui semble fonctionner pour le TGV... ■

BIBLIOGRAPHIE

Y. Chen et al., **Why time reversal for future 5G wireless?**, *IEEE Signal Processing Magazine*, vol. 33(2), pp. 17-26, mars 2016.

T. Dubois, **Application du retournement temporel aux systèmes multi-porteuses : propriétés et performances**, thèse de doctorat de l'INSA de Rennes, 25 mars 2013.

G. Lerosey et al., **Time reversal of wideband microwaves**, *Appl. Phys. Lett.*, vol. 88, article 154101, 2006.

RETOUR SUR « DÉCOLLAGES À CHAUD »

Notre article «Décollages à chaud», paru dans le numéro d'octobre (*Pour la Science* n°480), a suscité plusieurs courriers de lecteurs portant d'une part sur le pilotage des avions et d'autre part sur les effets physiques évoqués. Nous remercions ces lecteurs pour leurs commentaires et les précisions qu'ils nous ont fournies. Nous avons repris certains passages de leurs courriers dans la suite de ce texte afin de clarifier certains points.

Commençons par quelques précisions données par des professionnels sur la phase de décollage.

Lors de l'accélération sur la piste, le pilote ne fait pas décoller l'avion dès que la portance est suffisante, ce qui correspondrait à la vitesse de décrochage. «Au contraire, le pilote empêche l'avion de décoller, lui fait dépasser la vitesse de décrochage pour avoir de la sécurité et tire sur le manche quand est atteinte la "vitesse rotation", déterminée avant le décollage en fonction de la masse, de la température et de la pression.» L'avion commence à s'incliner pour prendre l'assiette de montée, continue à accélérer, puis, lorsque la vitesse décidée pour décoller est atteinte, il quitte le sol. Lors du décollage proprement dit, en tirant sur le manche, «le pilote agit plutôt sur la gouverne de profondeur située sur l'empennage arrière, et non sur les ailerons des ailes, destinés à contrôler le mouvement de roulis [inclinaison sur le côté].»

Un troisième point de pilotage concerne le contrôle de la vitesse de l'avion, en référence à l'encadré de la page 90. Levons l'ambiguïté du texte: lorsque l'avion est en vol, ce que contrôle effectivement le pilote, c'est la vitesse de l'avion par rapport à l'air. Comme l'indiquent les flèches représentant les vitesses dans la figure, la vitesse au sol de l'avion est alors la somme vectorielle de cette vitesse et de la vitesse du vent. Donc, pour une même vitesse par rapport à l'air, la vitesse par rapport au sol sera plus faible si le vent est de face, et plus importante si le vent est de dos, comme écrit dans le texte.

Venons-en à présent aux effets physiques discutés. Un lecteur nous signale que l'un des effets majeurs de l'augmentation de température (ou de la diminution de la densité) de l'air est la réduction de la poussée des moteurs. Nous avons

mentionné cet effet dans notre article, mais nous ne l'avions ni expliqué ni quantifié. L'origine de cet effet est le suivant: «Si l'on augmente le débit de carburant brûlé, la poussée augmente, la température devant la turbine et le régime de rotation augmentent. Lorsque la température devant la turbine ou le régime de rotation atteignent une valeur limite imposée par la résistance des matériaux, on ne peut plus augmenter le débit de carburant, ce qui limite la poussée à une valeur maximale. Or la température devant la turbine ou le régime de rotation augmentent si la densité de l'air alimentant le moteur diminue, puisque le débit massique d'air à chauffer et à brasser par le compresseur diminue. Ainsi, le débit de carburant maximal autorisé, donc la poussée, diminue si la densité diminue.»

Cela signifie que l'accélération de l'avion sera moindre et que le temps et la distance requis pour atteindre la vitesse nécessaire au décollage seront augmentés. Cet effet est dominant par rapport à

par exemple que pour un 737-100, un jour standard, pour les aéroports situés à plus de 6000 pieds d'altitude, le poids maximum au décollage est limité par la vitesse que peuvent supporter les pneus. Dans cette situation, lorsque la température est plus élevée, comme la portance diminue, cette vitesse est atteinte pour une masse plus faible et l'avion doit être moins chargé.

Ces courriers nous donnent l'occasion de rappeler notre démarche: présenter et expliquer ce que peuvent apporter le regard et les outils du physicien dans des situations diverses. Bien souvent, nous choisissons des sujets où nous pouvons expliquer qualitativement et quantitativement le phénomène abordé (par exemple les tubes hurleurs, *Pour la Science* n°479). Mais nous ne nous interdisons pas d'aborder des sujets qui ne sont pas encore totalement compris (comme l'effet Mpemba, *Pour la Science* n°342) ou des situations technologiques complexes, comme dans le cas présent. Après un travail bibliographique substantiel, nous



la diminution de la portance évoquée dans le texte qui, comme nous l'avions écrit, ne permet pas de retrouver les valeurs de longueurs de piste données par les constructeurs.

Ajoutons d'ailleurs que l'une des limites non évoquées pour la vitesse de décollage est la vitesse limite d'utilisation des pneus. Les abaques de vitesse donnés dans la brochure de Boeing sur les caractéristiques des différents modèles de 737 pour la conception des aéroports indique

tentons d'explicitier les phénomènes physiques qui interviennent, d'évaluer quantitativement leur effet et veillons à donner les limites de ces explications. C'est une démarche qui demeure généraliste et nous ne saurions, malgré nos efforts, atteindre la précision du vocabulaire et la richesse de l'analyse que peuvent offrir des années d'expériences dans le domaine concerné. ■

J.-M. COURTY ET É. KIERLIK