



Conception d'interfaces boitiers innovantes pour le radar automobile 77-GHz : Application à la conception optimisée d'une chaine de réception radar en boitier

Charaf-Eddine Souria

► To cite this version:

Charaf-Eddine Souria. Conception d'interfaces boitiers innovantes pour le radar automobile 77-GHz : Application à la conception optimisée d'une chaine de réception radar en boitier. Micro et nanotechnologies/Microélectronique. Université Paul Sabatier - Toulouse III, 2017. Français. NNT : 2017TOU30010 . tel-01653231

HAL Id: tel-01653231

<https://theses.hal.science/tel-01653231>

Submitted on 1 Dec 2017

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

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

THÈSE

En vue de l'obtention du

DOCTORAT DE L'UNIVERSITÉ DE TOULOUSE

Délivré par :

Université Toulouse 3 Paul Sabatier (UT3 Paul Sabatier)

Présentée et soutenue par :

Charaf-Eddine SOURIA

le mercredi 22 février 2017

Titre :

Conception d'interfaces boîtiers innovantes pour le radar automobile 77-GHz.
Application à la conception optimisée d'une chaîne de réception radar en
boîtier

École doctorale et discipline ou spécialité :

ED GEET : Électromagnétisme et Systèmes Haute Fréquence

Unité de recherche :

LAAS-CNRS MOST

Directeur/trice(s) de Thèse :

Thierry Parra, Université Toulouse 3 - LAAS

Gilles Montoriol, NXP Semiconductors

Jury :

Alain Peden, rapporteur

Martine Villegas, rapporteuse

Jean-Baptiste Begueret, examinateur

Ayad Ghannam, examinateur

Christophe Landez, invité

À ma femme Karima

À mes parents Zahra et Hamid

À mon frère Yassine et ma sœur Rim

À mes grands-parents

À tous mes amis et à ceux qu'on oublie

Avant-propos

Le travail présenté dans cette thèse a été effectué dans le cadre d'un contrat CIFRE à l'entreprise Freescale Semiconductors puis NXP Semiconductors à Toulouse, au sein de l'équipe de conception des radars automobiles et au Laboratoire d'Analyse et d'Architecture des Systèmes (LAAS) du Centre National de la Recherche Scientifique (CNRS) de Toulouse, au sein de l'équipe Microondes et Opto-microondes pour systèmes de Télécommunications (MOST).

Je tiens tout d'abord à remercier Monsieur Christophe Landez, responsable du groupe de conception de radars automobiles à NXP Semiconductors pour m'avoir accueilli dans son groupe et également pour la bienveillance qu'il m'a manifestée.

Je tiens à exprimer ma profonde reconnaissance à Monsieur Thierry Parra, Professeur des Universités à l'université Paul Sabatier de Toulouse, ainsi qu'à Monsieur Gilles Montoriol, ingénieur en conception aux fréquences millimétriques à NXP Semiconductors Toulouse, pour la confiance qu'ils m'ont témoignée en acceptant la direction de mes travaux ainsi que pour leurs soutiens, leurs conseils, leurs disponibilités et les échanges scientifiques que nous avons eu.

Je souhaite également remercier chacun des membres de mon jury de thèse d'avoir accepté de consacrer du temps à l'examen de mes travaux de recherche. Je remercie vivement Monsieur Alain Peden, Professeur à l'IMT Atlantique, et Madame Martine Villegas, Professeur à l'ESIEE Paris, pour l'intérêt qu'ils ont porté à ce mémoire en acceptant d'être les rapporteurs de mes travaux. J'exprime également ma reconnaissance à Monsieur Jean-Baptiste Begueret, Professeur à l'Université de Bordeaux, Monsieur Ayad Ghannam, président-directeur général et fondateur de 3DiS Technologies, pour avoir accepté d'examiner mes travaux et de participer au jury de thèse.

Je n'oublie pas également mes amis et collègues de NXP et du LAAS qui m'ont beaucoup apporté sur le plan professionnel et humain et qui m'ont énormément aidé en créant une ambiance agréable et amicale tout au long de ces années de thèse.

Table des matières

Avant-propos.....	5
Table des matières.....	6
Introduction générale.....	9
1 Le radar automobile 77-GHz en boîtier.....	11
1.1 Introduction.....	11
1.1.1 Principe de fonctionnement du radar automobile.....	12
1.1.2 L'utilisation des bandes millimétriques.....	14
1.2 Le système radar automobile.....	18
1.2.1 Spécifications du radar.....	18
1.2.2 Les modulations pour le radar automobile.....	21
1.2.3 Architecture du module radar.....	24
1.3 L'encapsulation en boîtier de l'émetteur-récepteur radar.....	34
1.3.1 De la microsoudure à l'encapsulation en boîtier.....	35
1.3.2 Défis du boîtier radar automobile.....	36
1.3.3 Modélisation du boîtier.....	39
1.4 Conclusion.....	41
2 Modélisation des interconnexions et des composants passifs de l'émetteur-récepteur radar.....	42
2.1 Introduction.....	42
2.2 Caractéristiques du substrat SiGe BiCMOS180.....	43
2.2.1 Propriétés des métaux et des oxydes.....	43
2.2.2 Propriétés du Silicium.....	45
2.2.3 Les niveaux de métallisation (stack-up).....	46

2.3	Validation des modèles d'interconnexions aux fréquences millimétrique	47
2.3.1	Méthodologie de modélisation	48
2.3.2	Modélisation électromagnétique et validation expérimentale	49
2.4	Modélisation de transformateurs intégrés pour applications millimétriques.....	64
2.4.1	Modélisation électromagnétique et principaux paramètres du transformateur.....	64
2.4.2	Validation expérimentale.....	66
2.5	Conclusion.....	70
3	Conception et optimisation de boîtiers pour l'émetteur-récepteur radar automobile 77-GHz	71
3.1	Introduction	71
3.2	Modélisation électromagnétique de composants passifs sur circuit imprimé	71
3.2.1	Description et spécifications des niveaux métalliques du PCB.....	72
3.2.2	Modélisation électromagnétique et validation expérimentale	73
3.3	Modélisation et validation d'interfaces boîtiers « Fan-Out » WLP.....	81
3.3.1	Modèle du boîtier RCP.....	82
3.3.2	Validation expérimentale basée sur une méthodologie d'extraction OSL	85
3.3.3	Application à l'encapsulation de la chaîne de réception d'un radar 77-GHz	89
3.4	Conception d'une interface boîtier « fan-in » WLCSP pour le radar 77 GHz	90
3.4.1	Boîtier WLCSP de départ (version 1) et limitations.....	90
3.4.2	Proposition et conception d'un nouveau boîtier WLCSP (version 2)	93
3.4.3	Validation expérimentale basée sur une nouvelle méthodologie « Half-OSL »	96
3.5	Optimisation de la nouvelle interface boîtier WLCSP	100
3.5.1	Améliorations proposées sur le boîtier WLCSP (version 3)	100
3.5.2	Comparaison de la méthodologie Half-OSL et d'une méthodologie basée sur le principe	
TRL	102	

3.5.3	Validation expérimentale basée sur les méthodologies Half-OSL et TRL.....	105
3.6	Conclusion.....	106
4	Application : Conception de la chaine de réception d'un radar automobile 77 GHz encapsulé .	108
4.1	Introduction	108
4.2	Spécifications et architecture de la chaine de réception radar	108
4.3	Mélangeur de fréquences à 77 GHz	114
4.3.1	Étude des principales topologies de mélangeurs	114
4.3.2	Conception des mélangeurs retenus après la 1 ^{ère} étude	115
4.3.3	Choix du mélangeur et validation expérimentale sur puce nue.....	130
4.4	Conception de la chaine de réception en boîtier WLCSP	137
4.4.1	Conception de la chaine de réception.....	137
4.4.2	Validation expérimentale.....	141
4.5	Conclusion.....	143
	Conclusion générale	145
	Publications	149
	Brevets.....	149
	Conférences internationales	149
	Conférences nationales.....	149
	Références bibliographiques	151
	Bibliographie du chapitre 1	151
	Bibliographie du chapitre 2	154
	Bibliographie du chapitre 3	155
	Bibliographie du chapitre 4	156

Introduction générale

Le développement des systèmes d'aide à la conduite a connu un bond très important durant la dernière décennie. Ces systèmes font partie des options proposées par la majorité des fabricants de voitures et deux voies de développements émergent aujourd'hui. La première voie est celle de la réduction du coût afin d'équiper plus de catégories de voitures de ces systèmes. La deuxième voie est celle de l'amélioration des performances afin de faire face aux exigences croissantes des autorités de sécurité routière et afin d'équiper la voiture autonome. Le radar automobile 76-81 GHz, appelé par abus le radar 77 GHz, se retrouve au cœur de tous les systèmes d'aide à la conduite. Sa principale fonction est la prise en compte de l'environnement de la voiture. Cette fonction est réalisée à travers la détection des cibles pertinentes autour de la voiture, à des distances plus ou moins importantes. Des algorithmes de traitement de signal avancés sont développés par les équipementiers ou les fabricants automobiles pour y parvenir. La qualité du signal exploité par ces algorithmes est critique pour la réussite de leur implémentation. Ce signal est issu de l'émetteur-récepteur radar qui génère une onde en hautes fréquences, l'émet, la reçoit et la transpose en bande de base. Par conséquent, l'élaboration de spécifications précises pour les différents paramètres de cet émetteur-récepteur est indispensable pour le fonctionnement optimal du radar. Il est par conséquent indispensable pour les fournisseurs de circuits intégrés de respecter ces spécifications et suivre leurs différentes évolutions, générées par l'évolution des besoins et des fonctionnalités du module radar. Ces aspects sont détaillés au cours du premier chapitre.

La conséquence de ce lien étroit entre les fonctionnalités du radar et les performances de l'émetteur-récepteur pousse les fournisseurs de semi-conducteurs à développer, à chaque génération, des circuits plus complexes au même coût ou à des coûts plus bas que la génération précédente afin d'assurer à leurs produits une évolution sereine et constante. Cette évolution est principalement réalisée dans le domaine numérique grâce aux avancées technologiques du semi-conducteur. Ces technologies disposent à chaque nouvelle génération de transistors plus petits, permettant l'intégration de plus de fonctionnalités dans des circuits de même taille que la génération précédente. Cette évolution est celle décrite et prédite par les très connues lois de Moore. Dans le domaine analogique, aux fréquences millimétriques en particulier, la contribution de la taille des transistors aux performances et à la taille du circuit est limitée. En effet, les structures passives, très gourmandes en surface, ont un intérêt majeur dans la conception des circuits aux ondes millimétriques. Le choix de ces structures passives et leur modélisation est une activité qui nécessite une connaissance approfondie des principes de l'électromagnétisme et de la propagation des ondes, en plus d'une maîtrise des outils de simulation électriques et électromagnétiques. On s'intéresse à ces aspects au cours du deuxième chapitre de cette thèse.

Parallèlement aux transistors et aux structures passives du circuit intégré, d'autres éléments déterminent les performances de l'émetteur-récepteur comme les interconnexions entre le circuit intégré et le reste du module radar. Ces interconnexions sont réalisées grâce à des microsoudures, ou des contacts dans le

cadre d'une encapsulation en boîtier. Le choix de la technologie d'interconnexion est critique pour les performances du circuit. Plusieurs types d'interconnexions existent et le choix des interconnexions les plus appropriées dépend du coût, des performances électriques et mécaniques ainsi que de la dissipation thermique de ces interconnexions. La multitude des critères à respecter limite le choix et le confine dans le cadre du radar automobile aux microsoudures et plus récemment aux encapsulations au niveau de la plaquette. Deux familles sont présentes dans cette catégorie ; la première utilise une encapsulation externe au circuit « Fan-Out » et la deuxième utilise une encapsulation interne « Fan-In ». Si le boîtier « Fan-Out » présente des avantages significatifs au niveau électrique par rapport au boîtier « Fan-In », le deuxième est très avantageux en termes de dissipation thermique et de coût. Le boîtier « Fan-In » n'a jamais été utilisé dans un produit aux fréquences millimétriques pour les limitations électriques qu'il présente. Nous proposons au cours du troisième chapitre la première conception, et unique à ce jour, d'un boîtier « Fan-In » compatible avec l'encapsulation d'un circuit Silicium aux fréquences millimétriques, en particulier à 77 GHz. Cette conception se base sur un boîtier standard et faible coût de l'industrie, sans modification de processus de fabrication. Ce développement est basé sur la conception des dessins de masques au niveau du boîtier et du circuit intégré.

Après le développement et la validation de ce boîtier, une chaîne de réception utilisant ce boîtier est conçue. La conception de cette chaîne dépasse le cadre de l'application puisqu'il s'agit de complètement revoir l'architecture de la chaîne et de proposer des alternatives à la conception existante afin d'optimiser les performances en bruit du récepteur. Cette étude d'architecture aboutit à la conception de trois mélangeurs. Le mélangeur sélectionné est utilisé pour la conception de la chaîne de réception en boîtier « Fan-In ». De bonnes performances sont obtenues, aussi bien en simulations qu'en mesures. Elles sont présentées en détails dans le quatrième chapitre.

1 Le radar automobile 77-GHz en boîtier

1.1 Introduction

L'introduction du radar embarqué au marché automobile émane de la volonté des constructeurs automobiles d'ajouter plus de confort et de sécurité aux véhicules, en apportant au conducteur plus d'informations sur son environnement, dans un premier temps, et en autorisant le véhicule à manœuvrer sans intervention du conducteur, dans un second temps. D'ailleurs, c'est le développement du radar automobile en association avec d'autres technologies comme la reconnaissance par caméra, appelée vision, et la détection par laser qui ont amené la voiture autonome au monde réel depuis le monde de la science-fiction. En effet, la quantité d'informations que relève le véhicule, en temps réel sur son environnement, devient si conséquente que la présence d'un conducteur « humain » ne sera bientôt plus une nécessité [1.1]. L'image ci-dessous montre la vision de NXP Semiconductor de l'implémentation des capteurs tels le radar, la vision et la communication des véhicules (V2X pour Vehicle-To-X) comme cœur de la voiture autonome.

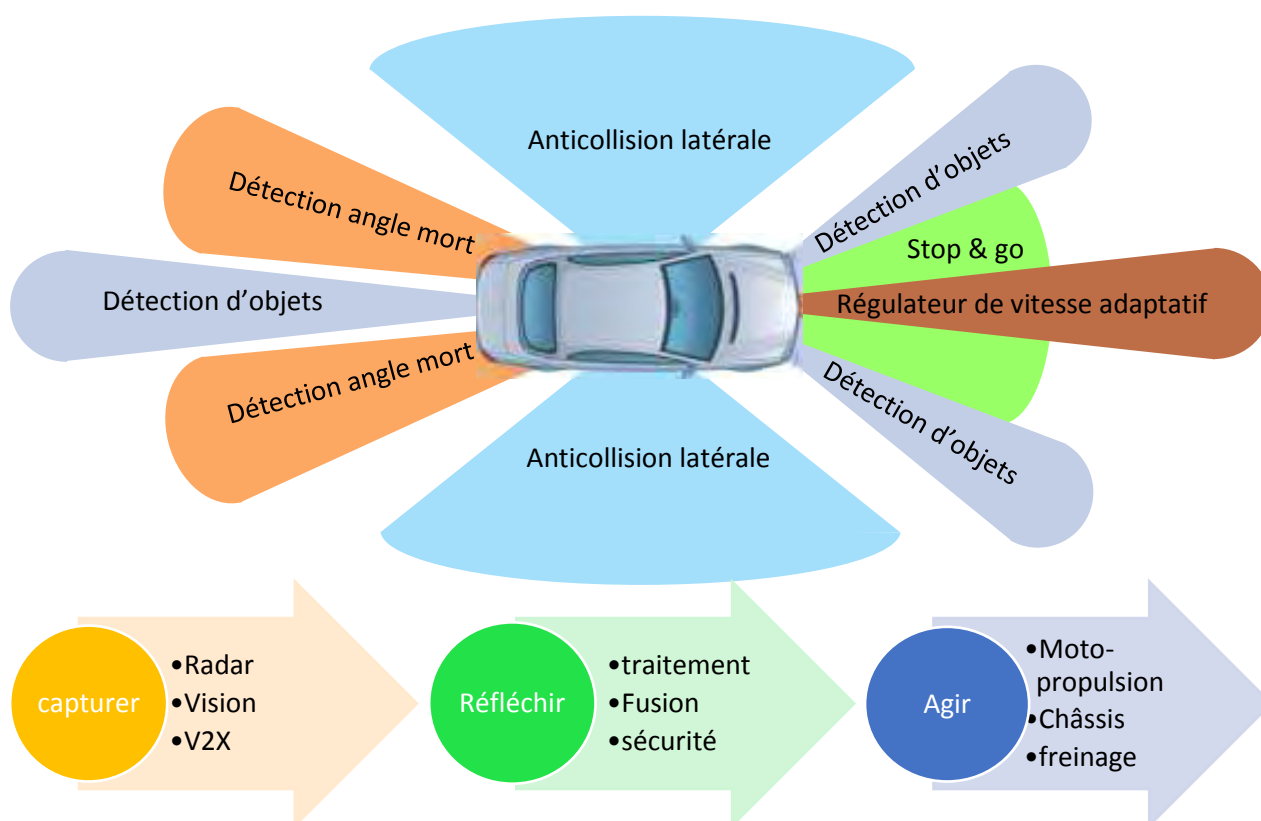


Figure 1 : rôle des capteurs dans la chaîne de développement des voitures autonomes

Le développement du radar automobile s'appuie sur l'existence d'un prédécesseur qui est présent dans la majorité des véhicules depuis de nombreuses années : le sonar de recul appelé par abus de langage « radar » de recul. Le lien entre ces principes est toutefois assez fort malgré la différence entre les

phénomènes physiques intervenant dans chacun des deux. Le radar automobile et le sonar de recul cherchent tous les deux à estimer la distance d'une cible par rapport au véhicule. Cependant, le radar permet d'obtenir bien plus d'informations comme la vitesse et la direction de déplacement de la cible. Le sonar de recul remplit pleinement ses fonctions lorsqu'il s'agit d'assister le conducteur pour un stationnement ou dans le cadre d'un stationnement automatisé. Néanmoins, ce dernier dispositif présente plusieurs inconvénients, notamment sa portée très courte, son incapacité à distinguer les cibles et à estimer les vitesses de rapprochement, en plus de son champs assez restreint nécessitant l'utilisation de nombreux capteurs pour la réalisation de la fonction. Un autre inconvénient du sonar de recul est son manque d'esthétique parce qu'il ne peut pas être dissimulé derrière le pare-chocs et laisse donc apparaître des ouvertures autour du véhicule, contrairement au radar.

1.1.1 Principe de fonctionnement du radar automobile

Le principe général de fonctionnement utilisé par le radar automobile est classique et identique à celui des radars développés dès le début du 20^{ème} siècle. Il s'agit de générer une onde électromagnétique, de l'émettre dans l'environnement que l'on souhaite examiner, puis de recevoir cette onde et d'analyser les modifications qu'elle a subies grâce au traitement du signal. L'onde générée peut être sous forme d'impulsions ou bien d'ondes entretenues. La majorité des radars automobiles sont basés sur des ondes entretenues (CW pour Continuous Wave), généralement modulées, et utilisent l'effet Doppler. Le radar automobile le plus répandu est un radar à ondes entretenues avec modulation linéaire de fréquence.

Les informations que le radar automobile cherche à récolter dépendent de son type. Pour le moment, on distingue principalement deux types de radars : le radar anticollision, appelé radar moyenne portée, et le radar de régulation de distance, appelé radar longue portée. Ces fonctions sont en cours d'évolution avec l'arrivée de la voiture autonome et un nombre de scénarios beaucoup plus important est en cours d'établissement.

1.1.1.1 Le radar moyenne portée

Le radar moyenne portée est un radar anticollision qui a pour objectif de détecter les cibles à proximité immédiate du véhicule. Il prévient le conducteur des dangers l'environnant et déclenche dans certains cas des systèmes de sécurité du véhicule. Ce radar nécessite une portée maximale avoisinant soixante-dix mètres et peut informer de manière très rapide l'ordinateur de bord de l'imminence d'un danger, ainsi que de sa nature. Le véhicule peut alors prendre instantanément des décisions de manière autonome en déclenchant par exemple le système de freinage d'urgence ou la pré-tension des ceintures de sécurité. Sur certaines voitures, il se contente de prévenir le conducteur de l'imminence d'une collision si aucune action n'est prise. Le choix des actions dépend en grande partie du niveau du fonctionnement souhaité pour le radar.

Le radar moyenne portée a pour but principal aujourd'hui la détection des cibles autour du véhicule. Dans un futur proche, ses fonctions se développeront pour également identifier et classer les cibles à travers des algorithmes très puissants. Il devra notamment distinguer les piétons des automobiles et des animaux, ainsi que les cibles fixes des cibles en mouvement. Le radar automobile moyenne portée est calqué sur le radar automobile longue portée à quelques détails près qui sont la bande passante et le champ de vision. La bande passante plus large du radar moyenne portée permet d'atteindre une meilleure résolution en distance. Ce radar subit moins de contraintes concernant l'éloignement des cibles mais il doit disposer d'une très bonne résolution. Cette résolution est définie comme étant la capacité du radar à détecter séparément deux cibles très proches l'une de l'autre. La résolution du radar moyenne portée est de l'ordre de quelques dizaines de centimètre.

La détection des piétons est un des axes d'étude les plus importants aujourd'hui en termes d'algorithmes de traitement du signal, que ce soit pour le radar automobile ou pour la vision [1.2]-[1.4]. Cette activité est motivée par la criticité de la détection de piétons en termes de sécurité routière, d'engagements gouvernementaux et de préparation de la voiture autonome. L'organisme Euro NCAP (pour European New Car Assessment Program), en charge de l'évaluation de la sûreté des voitures en Europe, note les voitures en leur attribuant un nombre d'étoiles allant de 0 à 5. Cet organisme donne une grande importance à la détection des piétons dans sa notation. Une des difficultés de la détection des piétons est que les spécifications matérielles ne sont pas encore basées sur des scénarios avancés de détection de piétons. C'est donc à la partie logicielle de relever ce défi. Des avancées importantes ont été réalisées dans ce domaine et plusieurs algorithmes ont été confrontés avec succès à des scénarios pratiques.

1.1.1.2 Le radar longue portée

Le radar longue portée est principalement utilisé aujourd'hui pour la régulation de distance. Le radar détecte les véhicules circulant devant lui et calcule leurs vitesses, leurs distances et leurs positions angulaires. Ces informations sont alors exploitées pour définir la vitesse à laquelle devrait rouler le véhicule pour garder une distance de sécurité correcte par rapport aux véhicules le précédant. Ce dispositif est appelé « Adaptive Cruise Control » (ACC).

Le principal défi pour ce radar est la portée à atteindre qui doit être de l'ordre de 200 mètres. Cette contrainte impose à l'émetteur la génération d'une onde à un niveau de puissance isotrope rayonnée équivalente (PIRE) élevé, et impose au récepteur une grande sensibilité, soit un faible bruit afin de détecter correctement l'onde reçue après sa réflexion par la cible, et après son atténuation par son parcours aller-retour. Le paramètre clé qui reflète la qualité de la détection longue portée est le rapport signal sur bruit ou SNR pour "Signal to Noise Ratio". Ce dernier caractérise le rapport entre la puissance du signal d'intérêt et la puissance du bruit. C'est un paramètre présent dans tous les systèmes radar et de télécommunication. Sa formulation est très intéressante dans le cadre du radar automobile puisqu'elle permet d'identifier immédiatement les principaux contributeurs à la portée du radar. En effet, le SNR

est proportionnel au produit du facteur de bruit du récepteur et de la puissance de sortie de l'émetteur. D'autres paramètres interviennent dans la portée comme le gain des antennes et la résolution de bande passante (RBW pour Resolution BandWidth) de la transformée de Fourier Rapide (FFT pour Fast Fourier Transform). Cette dernière transforme le signal temporel en spectre fréquentiel. Ces notions seront étudiées plus en détails par la suite.

1.1.2 L'utilisation des bandes millimétriques

La fréquence de fonctionnement du radar automobile intervient dans la spécification de nombreux paramètres du module radar comme la résolution et la portée. Le choix de la bande de fréquence dépend notamment des normes d'attribution de fréquences et de l'atténuation de la puissance de l'onde dans l'espace libre aux fréquences de fonctionnement.

1.1.2.1 Attribution des fréquences

Une agence des Nations Unies appelée Union Internationale des Télécommunications (UIT) est en charge de l'élaboration et la planification de recommandations pour les Télécommunications à l'échelle mondiale. Son principal rôle est de permettre à tous les équipements transmettant des ondes électromagnétiques de fonctionner dans de bonnes conditions en évitant les interférences entre les équipements ou en les régulant quand elles sont inévitables. Une table reprenant toutes les fréquences allouées en Europe est régulièrement mise à jour [1.5]. Un organisme européen nommé ETSI (pour European Telecommunications Standards Institute) est en charge de ces réglementations. L'organisme en charge de cette réglementation aux États-Unis est le FCC (pour Federal Communications Commission). Les principaux standards européens relatifs aux radars automobiles sont les suivants :

- EN 302 858 : équipement radar automobile fonctionnant dans la bande de fréquence de 24.05 GHz à 24.25 GHz ou 24.50 GHz
- EN 301 091 : équipement radar automobile fonctionnant dans la bande de fréquence de 76 GHz à 77 GHz
- EN 302 264 : équipement radar automobile courte portée fonctionnant dans la bande de fréquence de 77 GHz à 81 GHz
- EN 302 288 : équipement radar automobile courte portée fonctionnant dans la bande de fréquence de 21.65 GHz à 26.65 GHz. Ces bandes ne doivent pas être utilisées par les nouveaux équipements depuis le premier juillet 2013 au sein des pays de la Conférence Européenne des Administrations des Postes et Télécommunications (CEPT).

Les standards ci-dessus spécifient les bandes de fréquence allouées aux radars automobiles, les puissances maximales à émettre en plus d'autres spécifications. On distingue deux bandes de fréquence allouées aux radars automobiles. La première bande se situe entre 24 et 24.50 GHz et la deuxième entre 76 GHz et 81 GHz. Pour la deuxième bande, il est fait une distinction entre la bande 76-77 GHz et la

bande 77-81 GHz, la première étant réservée au radar longue portée (LRR) et la deuxième au radar courte portée (SRR). Une autre bande 21.65-26.65 GHz avait été provisoirement allouée au radar courte portée afin de permettre aux équipementiers développant des modules radar à 24 GHz d'étendre rapidement leurs modules à des applications de courte portée tout en développant en parallèle de nouveaux modules dans la bande de fréquence permanente entre 77 et 81 GHz. Le passage par ces fréquences temporaires était motivé par les contraintes associées à l'investissement et au temps de développement importants, qui résultent des difficultés technologiques et de conception liées à la bande 77-81 GHz.

On remarque que la bande de fréquence du radar courte portée est plus importante que celle du radar longue portée et cela est lié à la notion de résolution en distance discutée plus haut. Une modulation sur une plus large bande de fréquence permet d'obtenir une meilleure résolution. Cette notion sera abordée plus loin.

1.1.2.2 L'atténuation de l'onde

De manière générale, plus la fréquence est élevée plus l'atténuation est importante. Nous nous intéressons ici en particulier à l'atténuation en espace libre dans le cas d'un module radar automobile.

Absorption atmosphérique

Un des facteurs à prendre en compte pour les systèmes basés sur la propagation d'une onde en espace libre terrestre est l'absorption atmosphérique. Cette absorption peut être étudiée en fonction de la fréquence ou de la longueur d'onde, ces deux dernières étant inversement proportionnelles et liées par l'équation :

$$F \times \lambda = c$$

Avec $c = 299\,792\,458$ m/s, la vitesse de propagation de la lumière dans le vide, F la fréquence et λ la longueur d'onde.

La courbe suivante, présentée dans [1.6], montre l'atténuation des ondes en fonction de la fréquence ou de la longueur d'onde :

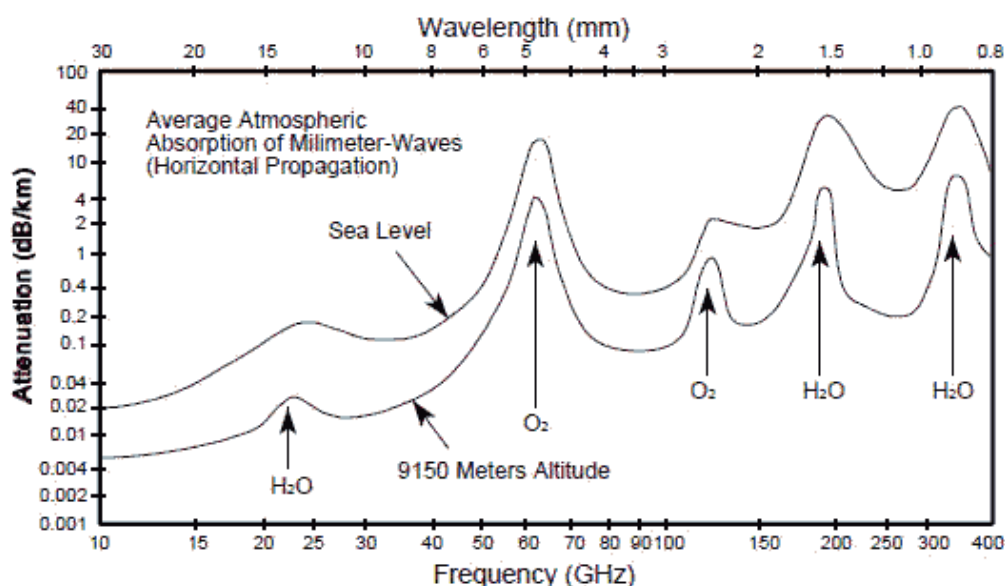


Figure 2 : absorption atmosphérique en propagation horizontale

Étudions le cas du radar longue portée décrit plus haut. La portée de ce radar devant avoisiner les 200 mètres, l'onde parcourt alors 400 mètres aller-retour avant de revenir au radar. Si on considère les deux fréquences de fonctionnement de 24 GHz et 77 GHz, les pertes respectives pour ces 2 fréquences sont de 0.2 et 0.4 dB/km. Pour la distance de 400 mètres considérée, ces pertes par kilomètre correspondent à 0.08 et 0.16 dB de pertes, respectivement. En d'autres termes, une onde radar à 77 GHz qui est réfléchi par une cible se trouvant à 200 mètres ne perd que 3,75% de sa puissance par l'absorption atmosphérique, et cette perte n'est même que de 1.85% pour une onde à 24 GHz.

On déduit de ce cas pratique que l'absorption atmosphérique n'a pratiquement pas d'impact sur les performances du radar que ce soit à 24 GHz ou à 77 GHz. Ceci est principalement lié aux portées relativement courtes du radar automobile comparées aux systèmes de télécommunications mobiles et satellitaires où les émetteurs et récepteurs peuvent être espacés de centaines voire de milliers de kilomètres.

Affaiblissement en espace libre

L'absorption atmosphérique n'est pas l'unique source d'atténuation lors d'une propagation en espace libre. L'onde électromagnétique, par sa nature, perd de l'énergie en se propageant dans une direction donnée puisque cette énergie se disperse quand l'onde s'éloigne de l'émetteur. Cette atténuation dépend principalement de la distance entre l'émetteur et le récepteur ainsi que de la fréquence. Dans le cadre du radar, le niveau du signal réfléchi dépend également de la surface équivalente radar (SER), ou « Radar Cross Section » (RCS) en anglais, de la cible. Cette grandeur est une propriété physique inhérente aux objets et reflète l'importance de la surface de réflexion de ces objets. Elle dépend, entre autres, de la nature des matériaux constituant l'objet, de sa forme et de la fréquence du signal réfléchi par l'objet. Ci-dessous les SER de quelques cibles particulièrement intéressantes dans le cas du système radar automobile, exprimées en dBsm [dB(m²)] et où R est la distance entre le radar et la cible [1.7]:

Piéton :	-10 dBsm
Motocycle :	7 dBsm
Voiture :	$\min \{ 10\text{Log}(R)+5 \text{ dBsm}, 20 \text{ dBsm} \}$
Camion :	$\min \{ 20\text{Log}(R)+5 \text{ dBsm}, 45 \text{ dBsm} \}$

L'équation (1), recommandée par l'UIT-R P.525-2 [1.8], donne l'affaiblissement (A_{0r}) de la propagation en espace libre de l'onde d'un système radar.

$$A_{0r} = 103,4 + 20 \log f + 40 \log R - 10 \log \sigma \quad (\text{dB}) \quad (1)$$

Où :

σ : surface équivalente radar de la cible (m^2)

R : distance entre le radar et la cible (km)

f : fréquence de fonctionnement du système (MHz).

On remarque à partir de cette équation que pour une SER fixe donnée, l'affaiblissement est proportionnel au carré de la fréquence f et à la puissance 4 de la distance d . En d'autres termes, si la fréquence double, l'atténuation est 4 fois plus importante (6 dB d'atténuation supplémentaires). Si la distance double, l'atténuation est 16 fois plus importante (12 dB d'atténuation supplémentaires).

La puissance reçue par le récepteur (P_r) est décrite en fonction de l'affaiblissement en espace libre (A_{0r}), de la puissance de l'émetteur (P_t) et du gain des antennes de réception (G_{rx}) et d'émission (G_{tx}). Elle s'écrit :

$$P_r = P_t + G_{tx} + G_{rx} - A_{0r} \quad (\text{dBm}) \quad (2)$$

En considérant : $G_{tx} = 23 \text{ dBi}$, $G_{rx} = 23 \text{ dBi}$, $P_t = 10 \text{ dBm}$, $f = 76.5 \text{ GHz}$ et $\sigma = \{-20, 10\} \text{ dBsm}$, on obtient la courbe P_r (dBm) en fonction de R (m) de la Figure 3.

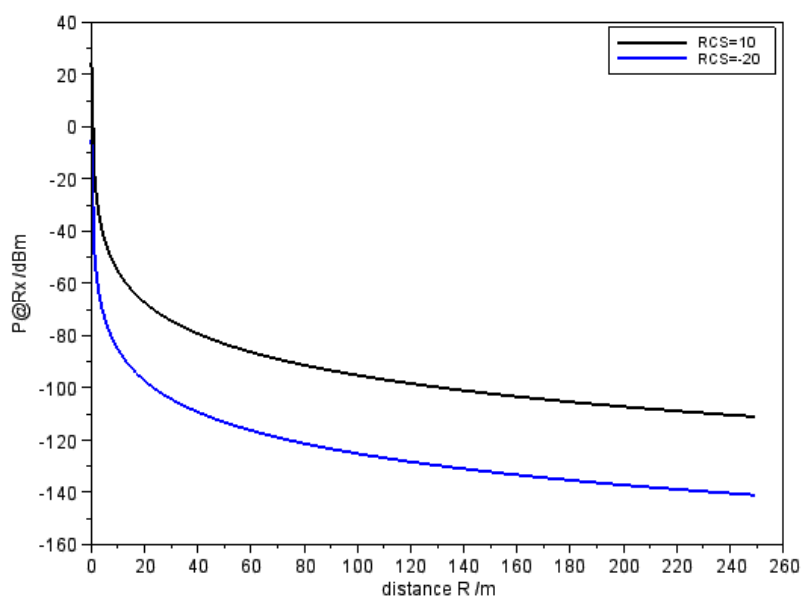


Figure 3 : évolution de la puissance reçue en fonction de la distance pour 2 valeurs de RCS {-20, 10}

Dans les conditions d'une SER (ou RCS) de -20 dBsm, correspondant à un piéton de petite taille, la puissance reçue par le radar après la réflexion de l'onde par ce piéton à une distance R de 80 m avoisine -120 dBm. Ce niveau de puissance est particulièrement faible puisque -174 dBm correspondent au plancher de bruit absolu à la température ambiante. Dans le contexte du radar, le bruit de l'émetteur et le bruit du récepteur viennent s'additionner à cette valeur, d'où la criticité du dimensionnement des bruits de l'émetteur et du récepteur radar.

1.2 Le système radar automobile

Le système radar automobile est un ensemble de moyens matériels et logiciels qui communiquent entre eux afin de réaliser la fonction de détection de la distance, de la forme, de la vitesse et de la position angulaire des cibles. Ces fonctions sont traduites en spécifications pour le produit, spécifications qui répondent aux besoins spécifiques grâce à des choix effectués au niveau de l'architecture du système radar. Cette architecture est conçue de sorte que toutes les spécifications qui en découlent soient réalisables et optimisent l'utilisation des technologies auxquelles elles font appel.

1.2.1 Spécifications du radar

Après avoir décrit de manière globale les fonctions du radar automobile, il s'agit donc maintenant de traduire ces fonctions en spécifications.

Le radar automobile, à l'instar de tous les radars, opère en réalisant six fonctions, essentiellement (Figure 4). La première est la génération de l'onde électromagnétique à une fréquence ou sur une plage de fréquences données. La deuxième est l'émission de cette onde pour une propagation en espace libre à une puissance et dans des conditions de rayonnement données. La troisième est la réception du signal

émis après sa réflexion par des cibles et sa transposition en fréquence plus basse pour permettre son traitement. La quatrième fonction est l'analyse et le traitement du signal réfléchi afin d'en déduire toutes les informations utiles. La cinquième fonction est la gestion de l'environnement pour préparer la sixième fonction qui est la prise de décision. Cette dernière peut être sous forme d'action d'accélération, de freinage, ou de changement de direction mais la décision de ne pas réagir peut également être prise. Ces fonctions sont réalisées, sur le plan matériel, grâce à un ou plusieurs circuits intégrés, un circuit imprimé sur lequel ces circuits intégrés sont reportés, des antennes ainsi que d'autres composants, notamment ceux qui sont en charge de la gestion de l'alimentation du module radar. Sur le plan logiciel, le système requiert deux types, à savoir des programmes pilotant le matériel et des programmes traitant les informations.

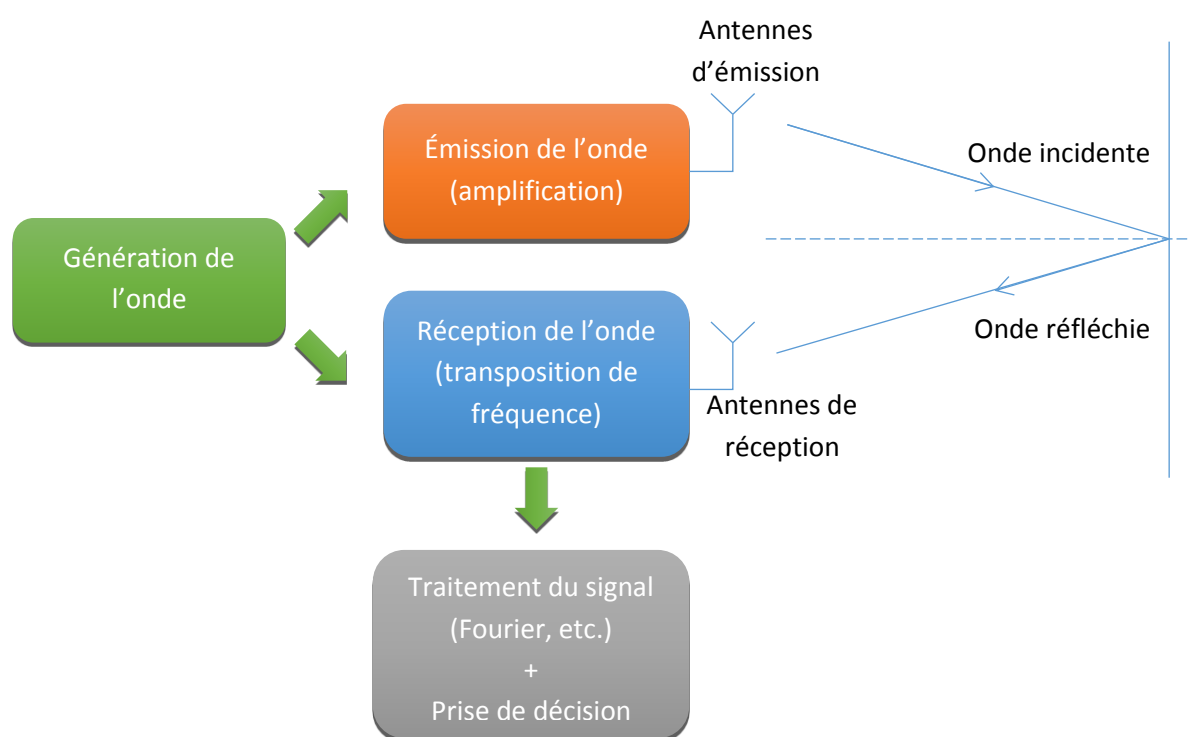


Figure 4 : diagramme de blocs simplifié du radar

De façon plus complète que pour certains radars militaires et civils, pour lesquels on s'intéresse uniquement aux informations relatives à la distance ou à la vitesse des cibles, le radar automobile requiert trois informations principales qui sont la détection de la vitesse de la cible ainsi que sa position à travers sa distance et son azimuth. L'azimut est défini comme étant l'angle d'incidence de l'onde dans un plan horizontal parallèle au plan de la chaussée. La détection de l'azimut de la cible est possible grâce à l'utilisation de plusieurs antennes recevant toutes le signal réfléchi. La différence de phase entre les signaux reçus par ces antennes permet de déterminer l'angle d'incidence du signal. Par ailleurs, le développement de radars réalisant, également, la fonction de détection de l'élévation est en cours. La distance et la vitesse de la cible sont étroitement liées au type de signaux émis par le radar. En effet, les ondes électromagnétiques peuvent être émises sous forme d'impulsions ou d'ondes entretenues. Les

impulsions donnent généralement des informations sur la distance mais pas sur la vitesse. Réciproquement, les ondes entretenues donnent des informations sur la vitesse mais pas sur la distance quand elles ne sont pas modulées. Pour résoudre cette limitation, plusieurs types de modulations ont été investigués pour disposer d'un système d'équations à deux variables : distance et vitesse. La modulation qui s'est imposée comme référence dans le domaine des radars automobiles récemment est la modulation linéaire de fréquence d'ondes entretenues (FMCW ou LMCW pour Frequency ou Linear Modulated Continuous Wave). Les radars à impulsion modulée restent très présents à 24 GHz. La modulation FMCW permet d'obtenir des informations sur la distance et la vitesse des cibles simultanément, comme on le verra dans la partie qui suit. Le module de Continental, un des pionniers du radar automobile 77 GHz, en est un excellent exemple. Il se base sur une modulation FMCW rapide sur la bande 76-77 GHz et peut atteindre une portée de 250 m avec son module Premium [1.9].

Une autre spécification critique pour le fonctionnement du radar concerne ses caractéristiques en bruit. Le signal émis par le radar perd une grande partie de son énergie lors de sa propagation et surtout lors de sa réflexion. Le niveau du signal reçu se situe souvent à de très faibles niveaux comme le montre l'équation (2) décrite plus haut. La limitation de détection des cibles est liée au bruit associé au spectre du signal reçu en plus du décalage fixé par les algorithmes de taux de détection de fausses alarmes. Les deux principales caractéristiques en bruit d'un radar sont le facteur de bruit du récepteur (NF ou Noise Figure) et le bruit de phase du signal émis ($\mathcal{L}(f)$ ou PN pour Phase Noise). Le facteur de bruit permet de déterminer le plancher de bruit du radar et il est très important même en présence d'une cible unique. Le bruit de phase est critique quand il y a deux cibles ou plus puisque les bruit de phase des signaux réfléchis sont transposés en bande de base. Le problème devient extrêmement contraignant pour l'application radar quand une des cibles est proche, et dispose donc d'un niveau de bruit de phase très élevé par rapport aux amplitudes des signaux utiles provenant de cibles plus éloignées.

La Figure 5 représente un scénario multi-cibles où le bruit de phase du signal réfléchi par un camion proche du radar (en rouge) masque le signal réfléchi par un cyclomoteur plus éloigné sur la chaussée (en vert) [1.10]. On peut voir à travers les courbes simplifiées de spectres des signaux qu'avec un bruit de phase important, le camion masque complètement le signal du motocycle.

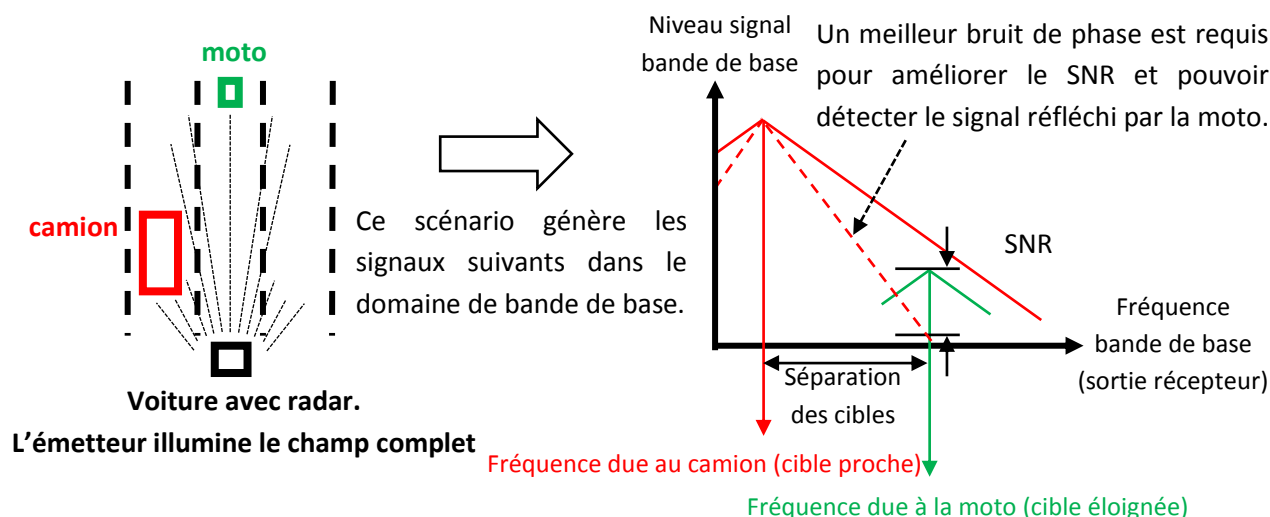


Figure 5 : scénario radar multi-cibles

Le niveau de bruit fait donc partie des spécifications les plus importantes du système et mérite une attention particulière lors de la conception du module radar. Une diminution du niveau de bruit améliore directement des paramètres critiques du radar comme la portée et la capacité de détection dans un contexte multi-cibles.

1.2.2 Les modulations pour le radar automobile

Deux principales modulations sont présentes dans les radars automobiles. La première est la modulation FMCW, présentée plus haut, et la deuxième est le décalage à plusieurs fréquences (MFSK pour Multiple Frequency Shift Keying).

1.2.2.1 Modulation FMCW

La modulation FMCW est une modulation assez simple, comparée aux modulations numériques, et apporte des avantages importants aussi bien au niveau de la génération de l'onde qu'au niveau du traitement du signal en bande de base. Elle est basée sur une onde entretenue dont la fréquence varie linéairement en fonction du temps (Figure 6). Les signaux d'écho sont comparés et la différence de fréquence (Δf) permet de déterminer le délai temporel entre les signaux émis et réfléchis (Δt). Si un décalage de fréquence Doppler (f_D) apparaît, il peut fausser la mesure s'il n'est pas dissociable de la différence de fréquence générée par le retard.

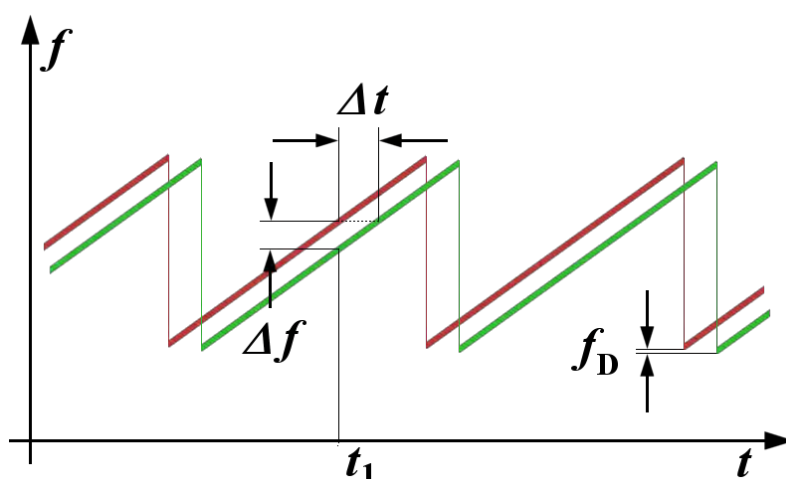


Figure 6 : principe de la modulation FMCW [1.11]

Une modulation FMCW classique, dite lente, est basée sur une variation linéaire de la fréquence dans la bande allouée sur un intervalle de temps de l'ordre de la milliseconde. La modulation FMCW lente permet de générer des ondes sous plusieurs formes. Le Tableau 1 présente les formes les plus connues et l'intérêt de chacune d'elle. Une génération de rampe en forme de dents de scie permet de détecter la distance d'une cible se déplaçant à la même vitesse et dans la même direction que le radar. Afin de prendre en compte l'effet Doppler et de l'utiliser pour déterminer la vitesse de la cible, la rampe ascendante est suivie d'une rampe descendante. Cette opération permet d'obtenir deux équations à deux variables, qui sont la distance et la vitesse. Il est ainsi possible de déterminer la distance et la vitesse d'une cible grâce à une rampe ascendante qui est suivie d'une rampe descendante quand le radar détecte une seule cible. Cependant, le radar fait souvent face à une multitude de cibles. Afin de distinguer dans ce cas les vraies cibles des cibles dites « ambiguës », il faut rajouter d'autres rampes montantes et descendantes à la séquence, avec des pentes différentes, afin d'éliminer ces cibles ambiguës. Ces dernières cibles sont des solutions données par les premières rampes et que les rampes supplémentaires permettent d'éliminer grâce à la différence des pentes d'une rampe à l'autre. Cette opération complexifie l'acquisition et le traitement du signal puisqu'on ne connaît pas à l'avance le nombre de cibles qu'il faut détecter. Dans certains cas, il est nécessaire d'implémenter, après l'extraction, un traitement avancé de suivi de cibles, appelé tracking, pour reconnaître les cibles réelles [1.12].

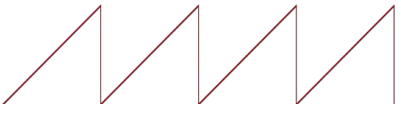
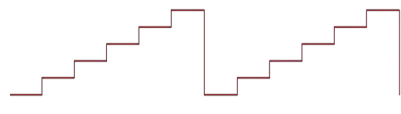
Type de modulation FLCW lente	intérêt	Schéma
dents de scie	utilisée pour les longues portées et lorsque l'influence de la fréquence Doppler est négligeable	
triangulaire	permet une séparation aisée de la différence de fréquence (Δf) et la fréquence Doppler (f_D)	
onde carrée	utilisée pour une mesure de distance très précise à très courte portée. Elle présente l'inconvénient de ne pas distinguer plusieurs cibles	
en escalier	utilisée pour des mesures interférométriques et élargit la plage de mesure sans ambiguïté	

Tableau 1 : différentes formes de modulation FMCW

Depuis quelques années, grâce à des efforts de recherche importants au niveau des principaux équipementiers automobiles et des fournisseurs de circuits intégrés pour le radar, une nouvelle modulation FMCW a été développée. Cette modulation FMCW rapide, appelée « Fast chirp FMCW modulation » ou « Chirp Sequence », récupère l'information de distance à partir d'une seule rampe rapide puisque l'effet Doppler est négligeable devant la période de la rampe. Cette dernière est passée de l'ordre de la milliseconde en modulation lente à l'ordre de la dizaine de microsecondes en FMCW rapide.

Dans une modulation FMCW, la résolution en distance dépend directement de la bande de fréquence de la modulation. Pour le radar moyenne portée, la résolution est très importante, d'où la large bande de fréquence allouée à cette application. La résolution du radar (ΔR) et la bande de fréquence de modulation sont inversement proportionnelles, comme le montre l'équation (3).

$$\Delta R = \frac{c}{2 \times (f_1 - f_0)} \quad , \text{ avec } c = 299\,792\,458 \text{ m/s} \quad (3)$$

Par conséquent, la seule limitation de la résolution est la largeur de la bande de fréquence. C'est d'ailleurs la raison pour laquelle les standards de radar ont alloué une bande de 4 GHz centrée à 79 GHz

pour le radar moyenne portée. Cependant, la bande passante peut être limitée par la technologie utilisée pour la fabrication du circuit intégré radar, puisqu'il est souvent difficile de concevoir un générateur de signaux avec une bande passante dans un rapport supérieur à 5% par rapport à la fréquence d'opération, c'est-à-dire une bande qu'il est possible de qualifier de très large.

L'application radar automobile moyenne portée demande une résolution de l'ordre de dizaines de centimètre. Un calcul rapide à partir de l'équation (3) permet de vérifier si la bande de fréquence allouée par l'IUT permet de passer les spécifications de résolution du radar. La largeur de bande de fréquence maximale est $f_1 - f_0 = 4 \text{ GHz}$. En arrondissant la vitesse de la lumière à $c = 3 \times 10^8 \text{ m/s}$, on trouve une résolution minimale de $\Delta R = 3,75 \text{ cm}$. Un radar automobile moyenne portée peut donc atteindre une résolution de 3,75 cm s'il utilise une rampe qui parcourt la totalité de la bande de fréquence allouée, [77-81] GHz, en une fois. Cependant, les générateurs d'onde ne sont pas toujours capables de fournir une telle rampe ou préfèrent utiliser des bandes de fréquences moins larges et plus rapides. Ces considérations limitent la résolution minimale réelle du module à des dizaines de centimètres.

1.2.2.2 Modulation MFSK

Quelques équipementiers utilisent la modulation MFSK pour se différencier de leurs concurrents. Cette modulation permet d'améliorer la capacité de détection des radars automobiles par rapport à une modulation FMCW [1.13]. Elle utilise une logique de commande de modulation MFSK (MMCL pour MFSK Modulation Control Logic) pour suivre le balayage linéaire de fréquence tout en générant le décalage de fréquence (FSK pour Frequency Shift Keying) à chaque pas dans le temps (Figure 7).

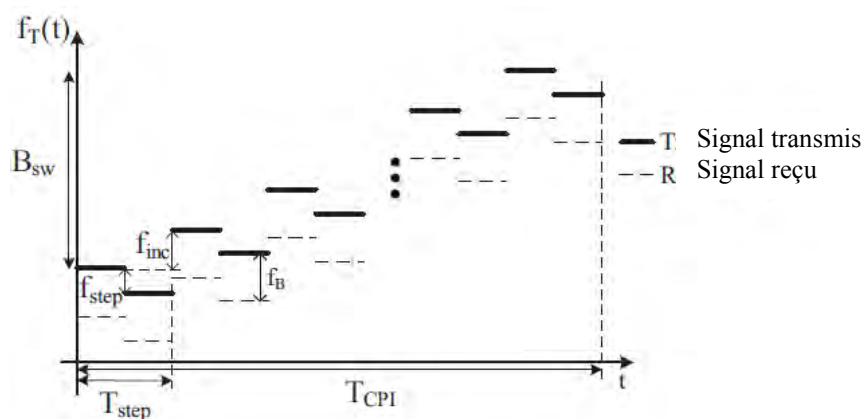


Figure 7 : principe de la modulation MFSK

1.2.3 Architecture du module radar

Le module radar automobile est dérivé d'une architecture radar classique avec un générateur de fréquence, un émetteur, un récepteur et un module de traitement du signal. Le radar bistatique utilise plusieurs antennes pour l'émission et la réception, généralement de 2 à 4 pour l'émission et de 2 à 12 pour la réception [1.9]. Une architecture de radar bistatique est choisie pour améliorer l'isolation entre

l'émetteur et le récepteur. Contrairement au radar monostatique, ce radar utilise des antennes séparées pour l'émission et pour la réception. La mise en forme du signal est basée sur une modulation FMCW et nécessite un convertisseur analogique numérique, qui peut être intégré dans l'émetteur-récepteur ou dans le microcontrôleur, pour le traitement numérique du signal reçu. Le module est également constitué d'un processeur de traitement de signal pour la réalisation de la transformée de Fourier et d'autres opérations mathématiques dans les domaines temporel et fréquentiel.

La génération des signaux radar automobile modulés en FMCW est basée sur un oscillateur contrôlé en tension (VCO) utilisé en boucle ouverte ou en boucle fermé.

L'émission de l'onde consiste à amplifier le signal à travers une chaîne d'amplification et de lui adjoindre, dans certains cas, un contrôle de phase grâce à un déphaseur (phase shifter). Le dernier bloc de cette chaîne d'émission (TX Channel), constitué par l'amplificateur de puissance (PA), est le plus critique de la chaîne parce qu'il doit produire la plus forte puissance dans la chaîne avec une consommation limitée, mais aussi pour des contraintes de dissipation thermique.

Le signal sortant de l'émetteur est envoyé à des antennes qui le transmettent dans l'espace libre dans des directions imposées par leur diagramme de rayonnement.

La réception de l'onde consiste à capter le signal réfléchi au niveau de la chaîne de réception (RX Channel) grâce à des antennes de réception et de transposer ce signal en bande de base grâce à un mélangeur de fréquences (Mixer). Les fréquences des signaux utiles en bande de base sont principalement entre 100 kHz et 10 MHz en modulation FMCW rapide. Ces fréquences sont entre le continu et des dizaines de kHz dans un radar à impulsions. Généralement, plus la fréquence dans cette bande est élevée plus la cible est éloignée. Les deux principales contraintes associées à cette chaîne sont le niveau de bruit qui doit être le plus faible possible et à la puissance de compression en entrée qui doit être la plus élevée possible.

Une fois le signal transposé en fréquence, le convertisseur analogique-numérique (ADC pour Analog to Digital Converter) le transforme en signal numérique, ce dernier étant transmis au microcontrôleur (MCU pour Microcontroller Unit) qui réalise une transformée de Fourier (FFT) sur ce signal puis traite les données.

La Figure 8 présente le diagramme de blocs d'un module Radar 77-GHz à base de circuits intégrés développés par NXP Semiconductors.

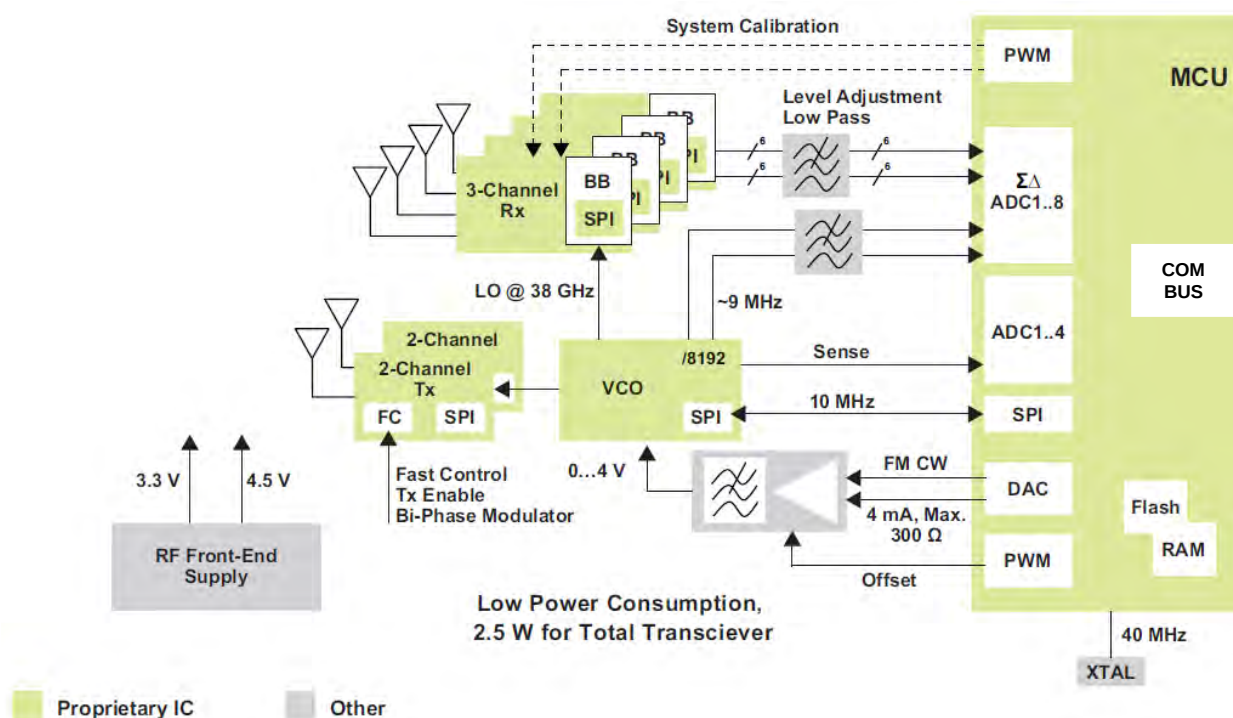


Figure 8 : diagramme de blocs d'un module radar basé sur des circuits intégrés NXP

Le système décrit par la figure 8 est un module radar composé d'un microcontrôleur, 4 récepteurs à 3 canaux, 2 émetteurs à 2 canaux, 1 VCO, plusieurs antennes et un circuit de gestion d'alimentation en plus de quelques filtres et autres composants. La génération de fréquence se fait grâce au VCO en boucle ouverte et des DAC et ADC présents dans le microcontrôleur. Le signal généré est émis en espace libre grâce aux 4 canaux d'émission. Le signal réfléchi est capté par les antennes de réception puis transmis via les 12 canaux des récepteurs. La conversion analogique-numérique est réalisée au niveau du microcontrôleur puis le signal est transformé en spectre fréquentiel grâce à une FFT. Des algorithmes permettent par la suite d'identifier les cibles.

1.2.3.1 L'émetteur-récepteur radar

L'émetteur-récepteur radar est fourni aux équipementiers sous forme d'un ou plusieurs circuits intégrés qui incluent la fonction de génération de fréquence, d'émission et de réception des ondes dans les bandes de fréquences millimétriques. Les puces les plus basiques incluent uniquement les blocs millimétriques de base (PA, Mixer, VCO) et des entrées-sorties numériques permettant de les contrôler. Le microcontrôleur avec d'autres composants, est en charge de piloter la génération de fréquence et des rampes de la modulation FMCW, ainsi que le traitement du signal à la réception. Les récentes générations d'émetteur-récepteur radar incluent de plus en plus de fonctions analogiques et numériques afin de minimiser la communication avec le microcontrôleur et simplifier le développement du module radar complet. Cela a été possible grâce aux évolutions importantes des technologies silicium au niveau transistor permettant de garder d'excellentes performances aux ondes millimétriques et d'être capables en même temps d'intégrer un nombre important de fonctions numériques. Avec l'importance croissante

de l'intégration numérique aujourd'hui deux technologies sont en concurrence dans le domaine des émetteurs-récepteurs radars 77-GHz. Ces technologies sont, d'une part, les technologies BiCMOS dont la longueur de grille est inférieure à 180 nm et, d'autre part, les technologies CMOS avec des longueurs de grille inférieures à 50 nm. Les technologies CMOS présentent l'avantage d'une très bonne intégration des fonctions numériques grâce à la taille des MOS, mais elles souffrent encore de certaines limitations aux ondes millimétriques, spécialement à des températures dépassant 100° C. De leur côté, les technologies BiCMOS disposent de transistors bipolaires très performants aux fréquences millimétriques, même à 150° C, mais les transistors MOS de ces technologies sont toujours en retard par rapport aux évolutions des technologies dites « full CMOS », ce qui limite l'intégration numérique et la réduction de consommation d'énergie du circuit. Il convient de noter qu'une autre technologie à base d'arséniure de gallium (GaAs) était très présente sur le marché radar automobile à ses débuts. Les deux inconvénients de cette technologie sont sa très faible capacité d'intégration numérique et son coût beaucoup plus élevé que les technologies BiCMOS. Cette limitation l'a rendu obsolète pour le développement des récentes générations de radars automobiles.

Le choix de l'architecture du circuit intégré et de ses sous-blocs est en partie lié à la technologie utilisée. On se focalisera sur des architectures basées sur des technologies BiCMOS, mais la majorité des aspects traités sont communs aux technologies CMOS et BiCMOS.

La génération de l'onde millimétrique, son émission ainsi que sa réception sont les fonctions fondamentales d'un émetteur-récepteur radar. Cependant, d'autres fonctions analogiques et numériques de plus en plus importantes sont incluses dans l'architecture des récents émetteurs-récepteurs radars, comme le contrôle et la sûreté de fonctionnement (norme ISO26262) et la gestion d'alimentation. On traitera ici ces principales fonctions.

La génération de fréquence :

La génération de la fréquence s'effectue grâce à un oscillateur opérant à la fréquence utile du radar ou à une de ses fréquences sous-harmoniques. Un multiplicateur est alors utilisé dans ce dernier cas. Aux radiofréquences ou aux fréquences millimétriques, on utilise principalement un oscillateur contrôlé en tension (VCO) pour réaliser la génération de l'onde. Afin de maîtriser la fréquence d'oscillation, deux architectures existent (Figure 9). La première approche est le VCO en boucle ouverte (Open-loop) (Figure 9(a)). Cette architecture permet de réaliser une modulation FMCW rapide sans contraintes importantes au niveau de l'émetteur-récepteur [1.14]. La modulation est principalement réalisée par logiciel en passant par des convertisseurs analogique-numérique et numérique-analogique. La deuxième architecture est la boucle à verrouillage de phase (PLL). Elle consiste à verrouiller la phase à laquelle une division par N du signal du VCO opère. Cette opération est réalisée en comparant ce signal divisé à un signal de référence (Figure 9(b)). Avant l'arrivée de la modulation FMCW rapide, les PLL étaient utilisées pour générer la fréquence modulée. Comme la période de modulation était de l'ordre de la milliseconde, la PLL convenait parfaitement. Avec l'utilisation de la modulation FMCW rapide, les PLL

classiques n'ont plus permis de répondre aux spécifications, ce qui a poussé au développement de l'architecture Open-loop, pour laquelle la modulation est effectuée par programmation au niveau du microcontrôleur [1.15]. Depuis, des activités de recherche sur les architectures de PLL ont permis de réaliser des modulations FMCW rapides. Cette nouvelle architecture a fait ses preuves et s'est même montrée supérieure à l'architecture Open-loop, en termes de linéarité notamment [1.16].

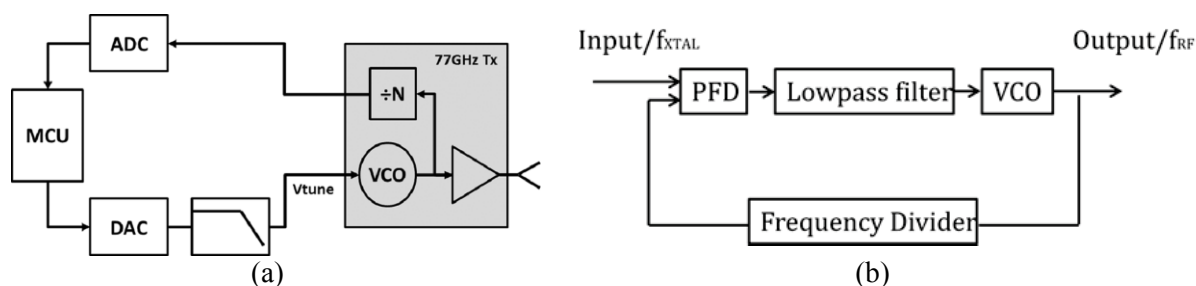


Figure 9 : génération de fréquence en boucle ouverte (a) et en boucle à verrouillage de phase (b)

La chaîne d'émission :

Après sa génération et sa mise en forme, le signal aux fréquences millimétriques est amplifié avant de sortir de la puce et d'être envoyé vers l'antenne d'émission. Le principal bloc en charge de cette opération est le PA. Il est généralement précédé d'autres amplificateurs et d'un déphaseur. Le PA doit fournir un certain niveau de puissance issu des spécifications techniques de l'émetteur-récepteur radar [1.17]. Les principales contraintes de la conception du PA sont la consommation électrique, la variation de puissance en sortie et le bruit d'amplitude. Garantir la consommation la plus faible possible passe par le choix de l'architecture la plus appropriée, en fonction des transistors utilisés, mais également par le choix des conditions de polarisation des transistors et les impédances à présenter en sortie de l'amplificateur. Ces dernières incluent les réseaux d'adaptation de sortie au niveau puce ainsi que la transition boîtier.

La Figure 10 donne un exemple de chaîne d'émission d'un module radar.

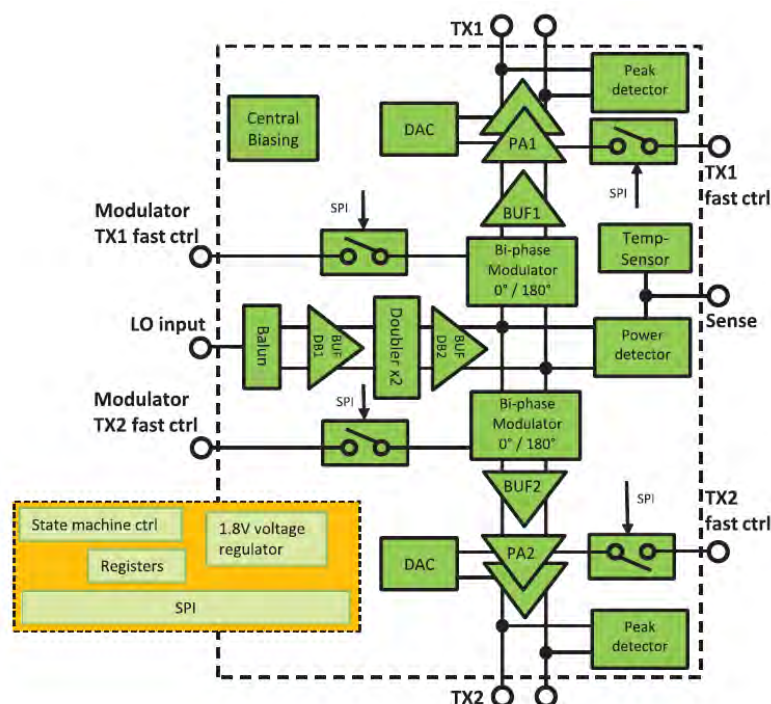


Figure 10 : schéma de blocs typique d'une chaîne d'émission [1.17]

La chaîne de réception :

La chaîne de réception est en charge de transposer en bande de base le signal reçu après réflexion par une cible. Le signal reçu par l'antenne de réception arrive au récepteur où il est transposé, à l'aide d'un mélangeur de fréquences, en bande de base. Cette transposition de fréquence du signal reçu nécessite le mélange de ce dernier avec un signal d'oscillateur local. L'écart en fréquence entre ces deux signaux dépend de l'éloignement de la cible et de sa vitesse relative. Cet écart se traduit, en bande de base après le mélange, par un signal en très basses fréquences. En effet, la bande de base d'intérêt pour le radar automobile se situe entre 100 kHz et 10 MHz [1.18]. En dehors de cette bande, les cibles potentielles sont soit trop proches, soit trop lointaines, et donc sans intérêt pour le fonctionnement du radar.

Comme discuté plus haut, le niveau de bruit dans cette chaîne doit rester le plus faible possible. En application du théorème de Friis, toutes les pertes de puissance subies par le signal reçu depuis l'antenne jusqu'au cœur du premier étage d'amplification se répercutent directement sur la sensibilité du récepteur. Chaque dB de pertes introduit un dB de dégradation de la sensibilité. Cela fait du chemin de l'antenne jusqu'au cœur de l'amplificateur faible bruit, ou du mélangeur si le premier n'est pas présent, un chemin extrêmement critique. Ce chemin comporte l'antenne, une ligne de transmission sur circuit imprimé reliant l'antenne à l'émetteur-récepteur et des microsoudures ou une interface boîtier reliant le circuit imprimé aux plots de la puce. Pour cette raison, la qualité de l'encapsulation en boîtier est d'un intérêt majeur pour les performances de la chaîne de réception.

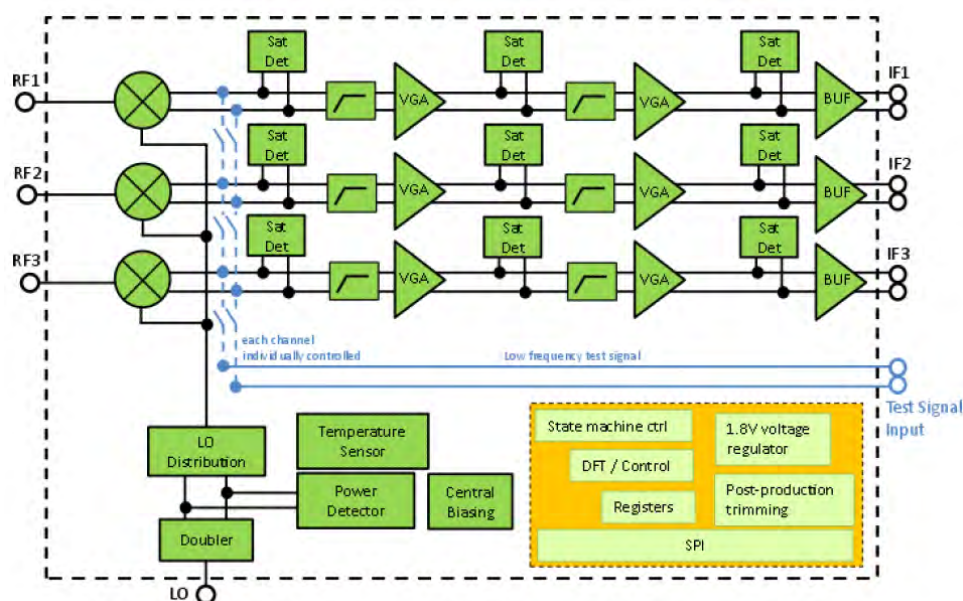


Figure 11 : diagramme de blocs typique d'une chaîne de réception [1.18]

1.2.3.2 Le microcontrôleur

Le microcontrôleur est le circuit intégré le plus important du module radar après l'émetteur-récepteur. Il contient une multitude de fonctions numériques grâce à la présence d'une mémoire vive (RAM pour Random Access Memory) et d'une grande capacité de programmation. Ces fonctions lui permettent de contrôler l'émetteur-récepteur radar et de traiter les signaux générés par le radar en bande de base. La communication avec l'émetteur-récepteur radar se fait principalement grâce au protocole SPI (pour Serial Peripheral Interface). La transmission du signal en bande de base, du récepteur vers le microcontrôleur, est généralement réalisée en mode analogique. C'est le microcontrôleur qui numérise ces signaux grâce à un ADC. Cependant, les émetteurs-récepteurs radars les plus récents incluent des ADC. Les signaux transmis au microcontrôleur sont donc déjà numérisés et on utilise souvent le protocole propriétaire MIPI (pour Mobile Industry Processor Interface) pour la transmission des données [1.19].

Une fois le signal numérisé, l'une des tâches les plus critiques d'un radar en bande de base reste à réaliser : il s'agit de convertir le signal du domaine temporel au domaine fréquentiel. Cette opération mathématique est réalisée grâce à une transformée de Fourier rapide (FFT). Nous allons nous focaliser sur la FFT d'un signal issu d'une modulation FMCW rapide [1.20]. En effet, deux ou trois FFT sont réalisées dans ce cas avant l'application des algorithmes de détection de cibles. Ces opérations sont appelées 3D-FFT par abus de langage (Figure 12).

La Figure 12 montre une des approches, les plus suivies, pour la réalisation de la FFT des signaux radar. On dispose au départ, à la sortie des ADC, d'un signal temporel en fonction des échantillons (Samples), des rampes (Ramp) et des canaux de réception (Ch.). Il s'agit, lors de la première FFT, d'intégrer le signal en fonction du nombre d'échantillons (Sample) pour donner un signal en fonction de la portée

(Rang). La deuxième FFT est réalisée en fonction des canaux de réception (Ch.) et permet d'obtenir l'angle d'incidence (Angle). La dernière FFT est réalisée en fonction des rampes de la modulation lente (Ramp) et donne le signal en fonction de la fréquence Doppler du signal.

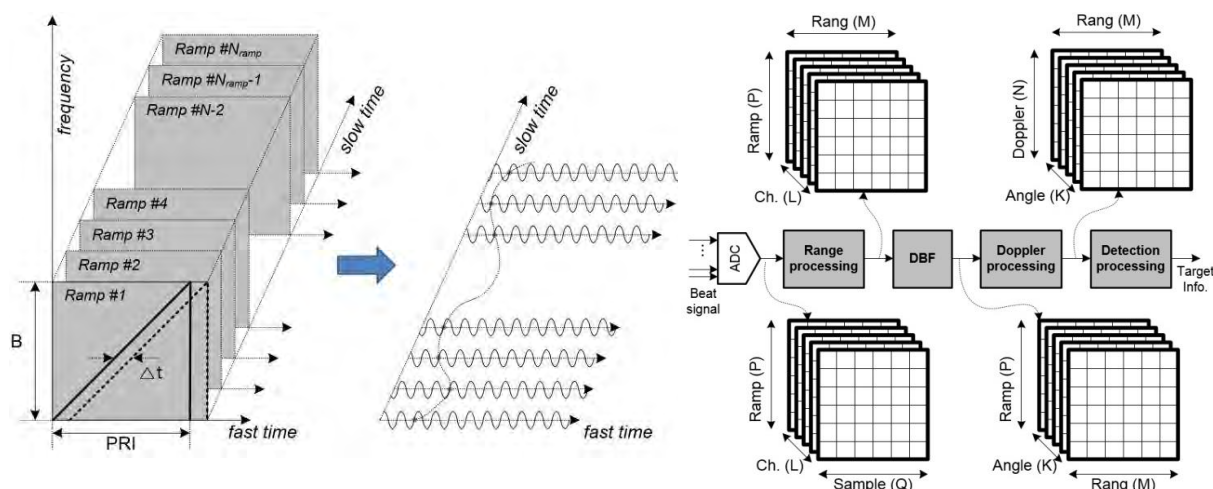


Figure 12 : principe de la FFT 3D

Une fois la FFT réalisée, le microcontrôleur est en charge d'exploiter le signal temporel converti en une matrice en fonction de l'angle d'incidence, la portée et la fréquence Doppler. C'est le traitement de cette matrice qui permet d'identifier les différentes cibles ainsi que leurs propriétés. La détection des cibles est généralement basée sur des algorithmes adaptatifs sur le principe des taux constants de fausses alertes ou « Constant False alarm rate » (CFAR). Les principaux algorithmes basés sur le CFAR sont CA-CFAR (pour Cell-Averaging CFAR) et OS-CFAR (pour Ordered Statistics CFAR) [1.21]. Des travaux récents ont donné naissance à un nouvel algorithme CFAR appelé AND-OR CFAR (pour « AND » et « OR » CFAR) [1.22].

La Figure 13 présente le schéma de blocs d'un microcontrôleur radar automobile de NXP Semiconductors incluant les fonctions nécessaires pour s'interfacer avec l'émetteur-récepteur radar et d'autres capteurs et systèmes automobiles [1.23].

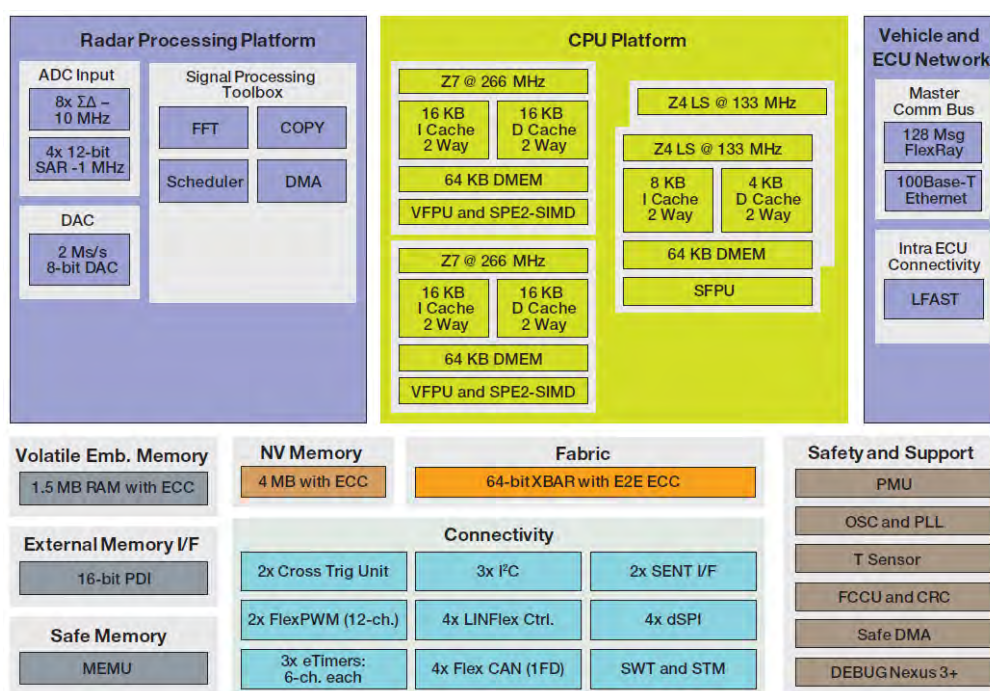


Figure 13 : fonctions d'un microcontrôleur NXP pour le radar automobile

Ce microcontrôleur inclut des fonctions de traitement du signal radar comme la conversion analogique-numérique, la FFT et d'autres fonctions de traitement des signaux temporels ou fréquentiels. Il inclut également des fonctions de communication avec l'émetteur-récepteur radar comme le protocole SPI et la conversion numérique-analogique. D'autres fonctions présentes dans le microcontrôleur permettent au microcontrôleur de communiquer avec différents capteurs et composants électroniques du véhicule.

1.2.3.3 Le circuit imprimé et les antennes

Le circuit imprimé (PCB pour Printed Circuit Board) du module radar est la structure physique sur laquelle sont disposés l'émetteur-récepteur radar, le microcontrôleur et tous les composants montés en surface (SMD pour Surface Mounted Device) du module. Il permet de réaliser les interconnexions entre les composants du module radar et intègre également les antennes d'émission et de réception ainsi que les pistes qui les connectent à l'émetteur-récepteur. Le choix du PCB est critique pour les performances du radar puisqu'il impacte les pertes des lignes microruban réalisant les interconnexions pour les signaux millimétriques, ainsi que les performances des antennes. Tous ces éléments dépendent directement des propriétés des métaux et des diélectriques mis en œuvre sur le PCB ainsi que des tolérances avec lesquelles leur fabrication est conduite.

Les tolérances de fabrication des PCB sont définies par la norme NFC 93-713 sous forme de plusieurs classes. De manière générale, plus la classe est élevée plus les dimensions minimales autorisées pour les pistes et les trous (vias) sont petites. Le Tableau 2 résume les principales classes utilisées dans l'industrie. Les modules radar sont généralement basés sur des PCB multicouches de classe 6. Ce choix émane du compromis à trouver entre les performances du PCB et son coût.

Critère	Classe						
	1	2	3	4	5	6	7
Largeur minimale des conducteurs (mm)	0.7	0.45	0.28	0.19	0.13	0.09	0.1
Espacement minimal entre conducteurs (mm)	0.6	0.45	0.28	0.19	0.13	0.09	0.1
Largeur des pastilles (mm)	-	1.65	1.25	1.05	0.85	0.65	0.55
Diamètre du trou traversant une pastille (mm)	-	0.8	0.7	0.6	0.45	0.35	0.3

Tableau 2 : spécifications de fabrication des circuits imprimés en fonction de la classe

Un autre paramètre très important pour le choix du PCB réside dans le type de diélectrique utilisé pour la couche supérieure (la plus haute). Ce diélectrique impacte directement les performances des antennes et des lignes véhiculant le signal millimétrique. Les premiers modules radar automobiles utilisaient un substrat Rogers RT5880 [1.24], conçu à la base pour des applications militaires et aérospatiales, et donc très cher. La réduction des coûts de production a poussé les fabricants à se diriger vers des substrats performants et à plus faible coût comme Le Rogers RO3003 [1.25]. Ce dernier est utilisé pour la couche supérieure et du FR4 est utilisé pour les autres couches. Ces diélectriques ont fait leurs preuves en termes de performances électriques en fonction de la fréquence et de la température, notamment pour les bandes de fréquence radar 24 et 77 GHz. Ils ont à la fois de très faibles pertes et une très faible variation de propriétés électriques en fonction de la température, comme présenté sur la Figure 14.

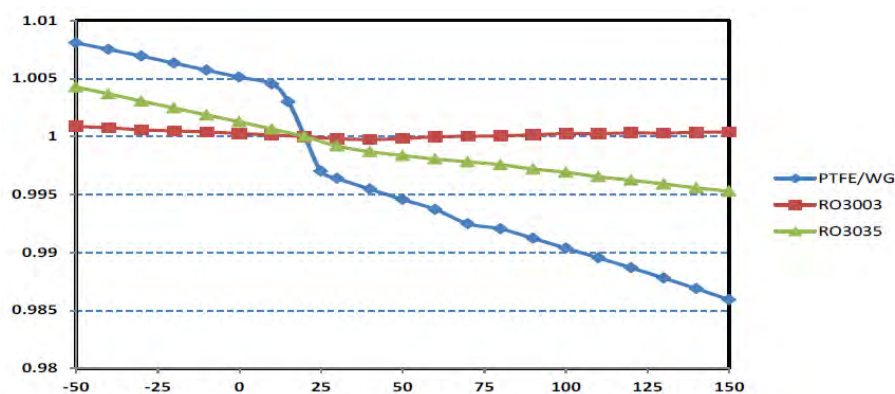


Figure 14 : variation de la constante diélectrique de quelques substrats Rogers en fonction de la température [1.20]

Après l'identification des principales propriétés du PCB, on s'intéresse maintenant à la conception des antennes. Ces dernières peuvent être conçues sous différentes formes. Les antennes utilisées dans le radar automobile sont les antennes paraboliques de type Cassegrain [1.26] et les antennes patch sur PCB [1.27]-[1.28]. La plupart des constructeurs de modules radar ont opté pour des antennes patch sur PCB pour minimiser les coûts et la complexité du module radar. Certes, ce choix limite le rendement de l'antenne, à cause des contraintes imposées par les caractéristiques du PCB, mais il présente un très bon compromis performances/coût. Les travaux de recherche menés sur les antennes des modules radar se

sont intéressés à différentes architectures pour connecter les antennes patch entre elles afin de minimiser les pertes et proposer des diagrammes de rayonnement optimaux pour le radar (Figures 15 et 16).

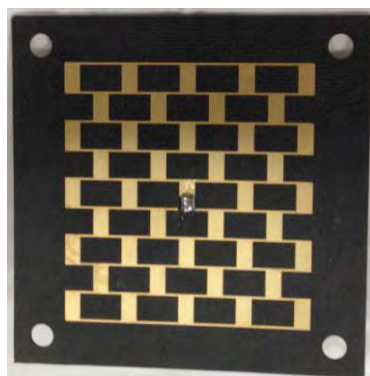


Figure 15 : Antenne patch pour radar automobile 77-GHz en RO3003 [1.27]

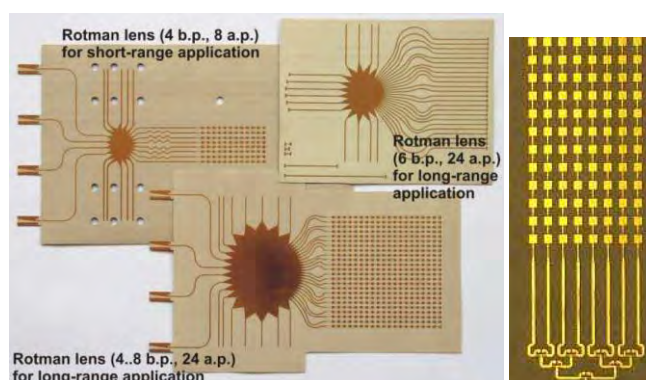


Figure 16: Antenne patch pour radar automobile 24-GHz en RT5880 [1.28]

1.3 L'encapsulation en boîtier de l'émetteur-récepteur radar

L'une des principales contraintes du module radar est la gestion de l'interconnexion entre l'émetteur-récepteur radar intégré sur semi-conducteur et le circuit imprimé. Cette contrainte émane de la fréquence très élevée des signaux d'émission et de réception, autour de 77 GHz, et de la criticité des pertes des interconnexions à ce niveau. En effet, il est possible de perdre plus de 3 dB au niveau de cette interconnexion uniquement, c'est-à-dire la moitié de la puissance. Ces pertes se traduisent par une perte de la moitié de la puissance émise par l'émetteur, impliquant généralement de doubler la consommation du PA, et une multiplication par deux du facteur de bruit du récepteur. Si cette perte de puissance n'est pas compensée par une augmentation de la consommation au niveau de l'émetteur, elle se traduira par une dégradation de la portée du radar de 30 mètres. Au niveau du récepteur, cette augmentation du facteur de bruit, dégrade la portée du radar longue portée d'environ 30 mètres également, ce qui détériore les performances du radar de façon significative. En d'autres termes, si la transition de l'émetteur-

récepteur radar vers le PCB génère des pertes de 3 dB au niveau de l'émetteur et du récepteur, la portée du radar diminue d'environ 50 à 60 mètres.

1.3.1 De la microsoudure à l'encapsulation en boîtier

Historiquement, l'interconnexion entre l'émetteur-récepteur radar et le circuit imprimé était réalisée par câblage de microfils, appelé également microsoudures. Les fabricants de circuits intégrés fournissaient alors des puces provenant directement de plaquettes découpées, appelées puces nues. C'étaient alors les équipementiers qui étaient en charge du câblage sur le PCB. Comme illustré sur la Figure 17, des travaux présentent des solutions où les entrées/sorties millimétriques du circuit intégré sont directement reliées aux antennes par microfils. Les autres signaux sont reliés par microfils aux plots du boîtier QFN (pour Quad Flat No-leads) [1.29]. Cette architecture permet d'enlever quelques contraintes associées au PCB mais n'apporte pas d'améliorations majeures dans le chemin reliant le circuit aux antennes.

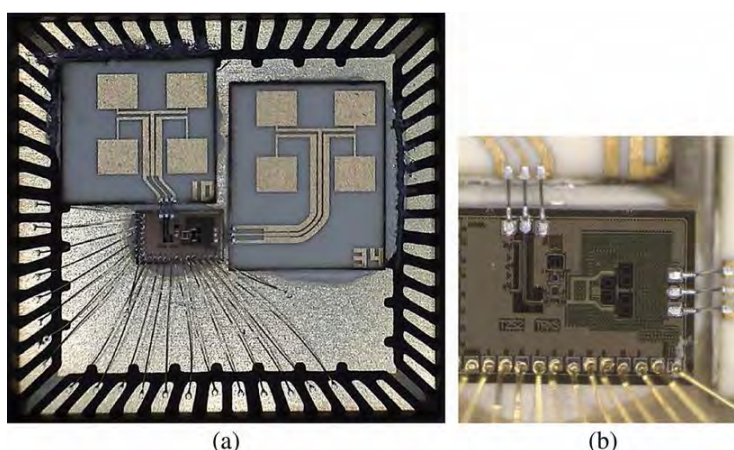


Figure 17 : photos de puces radars intégrées cablées par microfils [1.29]

Un des avantages du câblage par microfils, par rapport à la mise en boîtier, réside dans la simplicité du cycle de production et de test pour le circuit intégré. En effet, la mise en boîtier alourdit la chaîne de production puisque les plaquettes sont envoyées à un fabricant de boîtiers pour l'encapsulation. Compte tenu de la simplicité de la géométrie et des matériaux intervenant dans le câblage par microfils, il est également plus facile de modéliser des microfils qu'une transition boîtier. De plus, les microsoudures ne présentent pas de problème de fiabilité de joints de brasage ou « Solder Joint Reliability » (SJR) nécessitant, dans certains cas, l'application d'une couche appelée « underfill » entre la puce et le PCB. L'underfill est également utilisé dans certains cas pour améliorer la dissipation thermique quand cette dernière est limitée par le boîtier [1.30].

Cependant, le câblage par microfils présente deux inconvénients majeurs par rapport à l'encapsulation. Tout d'abord, la transition réalisée par les microsoudures apporte une désadaptation importante, qui peut être compensée mais au détriment d'une perte de bande passante du circuit. Le deuxième inconvénient est que le câblage par microfils présente de plus grandes variations de fabrication qu'une interface boîtier [1.31]. Des travaux ont été menés afin de réduire la longueur du câblage en enterrant la puce dans le

PCB, alignant ainsi horizontalement les plots de la puce à la couche supérieure du circuit imprimé [1.32]. Ces innovations posent cependant plusieurs contraintes liées à la production en grande échelle et au coût.

Compte tenu des avantages que présente l'encapsulation en boîtier par rapport au câblage par microfils, l'industrie du radar automobile se dirige de plus en plus vers des circuits intégrés en boîtier. Cette solution présente aussi un avantage significatif pour les équipementiers automobiles. En effet, les spécifications du circuit intégré ne sont plus données par le fournisseur au niveau de la puce nue mais au niveau des contacts du boîtier. Par conséquent, il bascule au fournisseur du circuit intégré la charge de choisir le type de boîtier à utiliser, puis de s'assurer que le produit passe toutes les spécifications électriques et mécaniques après son encapsulation.

1.3.2 Défis du boîtier radar automobile

Le boîtier du circuit intégré radar doit satisfaire, en plus des nombreuses contraintes relatives à la réglementation de sécurité pour les véhicules routiers, des contraintes relatives aux fréquences de fonctionnement très élevées du radar. En effet, les entrées des récepteurs et les sorties des émetteurs du radar automobile véhiculent des signaux dont les fréquences se situent autour de 77-GHz. Les matériaux conducteurs et diélectriques utilisés dans la réalisation du boîtier doivent être soigneusement choisis pour minimiser les pertes à ces fréquences. De plus, le boîtier peut apporter une désadaptation importante à ces fréquences de fonctionnement.

Par conséquent, le boîtier figure parmi les éléments les plus critiques dans la fabrication du circuit radar. Il est très important de choisir le bon type de boîtier afin de garantir la conception d'un circuit radar de fonctionnement optimal. Certains paramètres du circuit, comme le facteur de bruit du récepteur et la puissance de sortie de l'émetteur, sont spécifiés au niveau des contacts du boîtier comme précisé plus haut. Ainsi, la conception optimisée des blocs situés au voisinage immédiat de l'interface d'entrée-sortie du boîtier, comme le mélangeur ou l'amplificateur de puissance, nécessite une connaissance précise des caractéristiques du boîtier. La co-conception du circuit intégré et de l'interface boîtier peut aussi s'avérer nécessaire.

1.3.2.1 Types de boîtiers pour le Radar automobile

Il existe de très nombreuses familles de boîtiers de circuits intégrés sur le marché des semi-conducteurs. Les principales différences entre les boîtiers sont les matériaux utilisés et le type d'interconnexions. On distingue, entre autres, des boîtiers plastiques, céramiques, métalliques et à base de polymères comme le polyimide (PI) et le polybenzoxazole (PBO). L'interconnexion entre le circuit intégré et le PCB se fait avec des métaux sous forme de broches, de pattes ou de boules.

L'une des principales contraintes de la mise en boîtier des circuits intégrés industriels, en général, et le radar automobile en particulier, résulte du volume de la puce avec son encapsulation. Ce volume doit rester raisonnable, ce qui restreint la mise en boîtier des circuits à des encapsulations au niveau plaquette ou « Wafer Level Packaging » (WLP). Ces boîtiers ont l'avantage d'avoir une surface très proche de la taille de la puce nue et un volume qui dépend directement de la taille de la puce et du type d'interconnexions utilisé. Pour les boîtiers WLP, on privilégie généralement les interconnexions sous forme de boules de diamètre avoisinant 300 μm pour le marché automobile.

Les boîtiers WLP sont présents sous forme de boîtier « fan-in » et « fan-out » [1.33]. Les boîtiers « fan-out » englobent la puce, générant un circuit en boîtier d'une taille de boîtier supérieure à la taille de puce (Figure 18 (a)). Cette configuration a l'avantage d'utiliser un nombre plus important de contacts dont une partie située à l'extérieur de la région de la puce. Le boîtier « fan-in » restreint l'encapsulation à la région définie par la puce et tous les contacts sont, par conséquent, réalisés dans cette région (Figure 18 (b)). Le choix entre ces deux boîtiers est basé sur plusieurs critères dont le coût, la fiabilité, la dissipation thermique, les possibilités de routage, les performances électriques et les dimensions du boîtier. Si la concurrence entre les deux boîtiers est forte pour les applications faibles fréquences, inférieures à 5GHz, grâce au très bon compromis que proposent les deux solutions, elle l'est beaucoup moins en très hautes fréquences, supérieures à 40 GHz. En effet, à ces fréquences le « fan-out » présente des avantages électriques importants par rapport au « fan-in » en termes de désadaptation et de couplage au niveau des entrées/sorties hautes fréquences. Néanmoins, de récents travaux, notamment lors de cette thèse, proposent des solutions permettant de s'affranchir en partie des limitations électriques du « fan-in », le rendant exploitable aux fréquences millimétriques avec des performances rivalisant avec celles du « fan-out ».

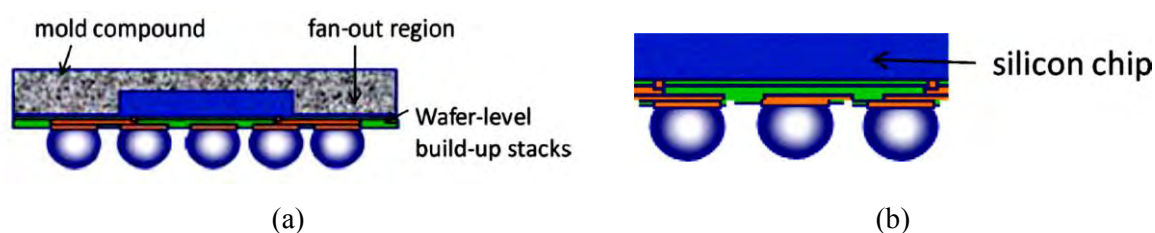


Figure 18 : boîtiers de type "Fan-Out" (a) et de type "Fan-In" (b)

1.3.2.2 Rôle du boîtier WLP dans les performances du circuit

La mise en boîtier de la puce impacte directement les paramètres mécaniques, thermiques et électriques du circuit intégré [1.34]. Mécaniquement, la puce est sujette à un stress mécanique en permanence à cause des variations importantes de température subies par le produit. Pour s'assurer que cela est sans danger pour le circuit, un test de fiabilité appelé fiabilité des joints de brasage ou « Solder Joint Reliability » (SJR) est réalisé. Le type de boîtier et sa taille sont déterminants dans les résultats de ce test. Concernant les contraintes thermiques, la dissipation permise par le boîtier est très importante parce

que la chaleur générée par les transistors du circuit passe en grande partie à travers le boîtier avant d'être dissipée par le circuit imprimé. Enfin, électriquement, la contribution du boîtier WLP est souvent négligeable en très basse fréquence où l'ordre de grandeur des inductances, résistances et capacités introduites par le boîtier sont négligeables devant les éléments du circuit. Cependant, les dimensions du boîtier aux fréquences millimétriques deviennent critiques parce qu'elles se rapprochent de la longueur d'onde des signaux. La longueur des pistes métalliques connectant le plot de la puce à la boule avoisine les 500 μm , assimilable qualitativement à une inductance série d'environ 500 pH. Cette longueur introduit par conséquent un déphasage autour de 70° à 77-GHz.

1.3.2.3 Défi de la conception d'un boîtier radar 77 GHz en Fan-In WLP

Il est connu que les boîtiers « fan-out » WLP sont parmi les boîtiers qui satisfont la plus grande partie des contraintes techniques et industrielles de l'encapsulation des radars automobiles 77 GHz. C'est par ailleurs le seul boîtier utilisé par les fournisseurs de circuits intégrés radar aujourd'hui. Deux noms, eWLB (pour embedded Wafer Level Ball grid array) et RCP (pour Redistributed Chip Packaging), sont présents dans la littérature mais les différences entre les deux sont marginales et n'impactent pas les performances électriques de l'encapsulation [1.35]-[1.36]. Si toutes les considérations électriques font pencher la balance vers le « fan-out » WLP, d'autres considérations liées au coût et à la dissipation thermique ont poussé des responsables de produits à envisager l'utilisation de boîtiers « fan-in » WLP et notamment le boîtier WLCSP qui est très utilisé pour les applications basses fréquences nécessitant un boîtier très compact et à faible coût. De plus, le boîtier « fan-in » permet une dissipation thermique plus importante, relâchant ainsi les contraintes de dissipation d'énergie du circuit par rapport au « fan-out ».

Le défi de faire du WLCSP un boîtier pour le radar automobile 77 GHz s'est donc retrouvé au cœur de ces travaux de thèse. Une des nombreuses difficultés associées à ce défi est qu'il ne s'agit pas d'améliorer un boîtier en optimisant ses dimensions mais d'apporter des innovations majeures permettant de réduire considérablement ses pertes avec la contrainte de n'impacter en aucun cas le coût ou la capacité de production en grand volume de ce boîtier. Il faut par ailleurs noter qu'aucun article n'a été publié concernant cette problématique et la littérature ne présente qu'un seul circuit encapsulé en WLCSP fonctionnant aux fréquences millimétriques [1.37]. Cependant, deux particularités de ces travaux les éloignent du contexte de cette thèse. Premièrement, ils utilisent une technologie GaAs très différente de la technologie Silicium que nous devons envisager ici. Deuxièmement, le boîtier utilisé est de type 3D-WLCSP, ce qui implique l'utilisation de plusieurs couches métalliques et plusieurs épaisseurs de diélectriques au niveau du boîtier, facilitant ainsi la conception du circuit grâce à l'éloignement des signaux et des masses. L'utilisation de ce type de boîtier ne peut pas être envisageable dans notre étude à cause de l'augmentation de coût qu'elle génère.

1.3.3 Modélisation du boîtier

La prise en considération de l'effet du boîtier sur les performances électriques est nécessaire pour garantir une conception optimale du circuit intégré. Il faut par conséquent disposer d'un modèle fidèle au comportement électrique de l'interface boîtier. Par ailleurs, la modélisation du boîtier permet de mieux comprendre son fonctionnement et d'optimiser ses performances. Historiquement, les interfaces boîtiers étaient associées à des circuits électriques équivalents de type RLCG, c'est à dire un modèle classique de ligne constitué d'une résistance (R) et d'une inductance (L) en série, et d'une capacité (C) et d'une conductance (G) en parallèle. Ces modèles purement électriques sont toujours utilisés pour des applications basses fréquences. Ils sont basés sur la décomposition de la structure en plusieurs éléments simples qui peuvent, chacun, être assimilé à un ou plusieurs composants électriques RLCG [1.38]-[1.40].

Les modèles de boîtiers pour les applications millimétriques sont aujourd'hui basés sur des simulations électromagnétiques. Il s'agit de dessiner des structures géométriques, en 2 ou 3 dimensions, qui reprennent fidèlement les caractéristiques géométriques et électriques des éléments du boîtier (boules, couches de redistribution, diélectriques, etc.). Les dessins sont basés sur des dessins de masques de la puce, du boîtier et du circuit imprimé. Ils sont certes simplifiés, mais gardent une représentation électriquement fidèle à l'original. Les caractéristiques électriques sont, quant à elles, extraites de la documentation existante et en particulier des manuels de conception quand cela est possible. La principale difficulté à ce niveau est de pouvoir disposer des caractéristiques électriques des matériaux mis en œuvre extraites ou applicables aux fréquences millimétriques. Dans ce cas, la première approche consiste à extrapoler les valeurs des caractéristiques selon des équations théoriques ou empiriques à partir des valeurs disponibles en basses fréquences. La deuxième approche consiste à mesurer les caractéristiques électriques des matériaux aux fréquences d'intérêt en utilisant des méthodologies présentes dans la littérature. Cette deuxième approche est plus coûteuse et délicate mais permet d'avoir un niveau de confiance plus élevé dans les caractéristiques électriques utilisées. Généralement, le choix est pondéré par la criticité des paramètres recherchés et le degré de difficulté de l'extraction par la mesure.

Une fois les caractéristiques géométriques et électriques des éléments de l'interface boîtier déterminées, le simulateur électromagnétique le plus adapté aux besoins est choisi. À ce niveau, si la théorie derrière les algorithmes de modélisation électromagnétique est rigoureuse et bien connue, décrite par les équations de Maxwell, les méthodes de résolution de ces équations sont différentes d'un simulateur à l'autre, chacun ayant sa propre implémentation.

Les équations de Maxwell font partie des lois fondamentales de la physique et constituent la base de l'électromagnétisme. Elles décrivent de manière précise l'ensemble des phénomènes électromagnétiques régissant les champs sous forme d'équations intégrales [1.41]-[1.42]. Nous

reportons ci-dessous quelques-unes des équations de Maxwell reliant les champs électriques et magnétiques.

$$\begin{aligned}
 \vec{E}(\vec{r}, t) &= \vec{E}_0 e^{j(\omega t - k_x x - k_y y - k_z z)} & \nabla^2 \vec{E} - \frac{1}{c^2} \frac{\partial^2}{\partial t^2} \vec{E} &= 0 \\
 \nabla \times \vec{H} &= \vec{S} + \frac{\partial(\epsilon \vec{E})}{\partial t} & c &= \frac{1}{\sqrt{\epsilon \mu}} \\
 \nabla \times \vec{E} &= -\frac{\partial(\mu \vec{H})}{\partial t} & \vec{S} &= 0 \\
 \nabla \cdot (\epsilon \vec{E}) &= \rho & \rho &= 0 \\
 \nabla \cdot (\mu \vec{H}) &= 0 & \Delta E - \frac{1}{c^2} \frac{\partial^2}{\partial t^2} E &= 0
 \end{aligned}$$

La résolution des équations de Maxwell pour des formes simples (domaine d'étude régulier) est assez facile, surtout quand la forme présente des symétries électriques. Cependant, il est impossible de trouver des solutions exactes aux problèmes faisant intervenir des géométries complexes, comme une interface boîtier par exemple. Plusieurs méthodes numériques ont alors été introduites afin de résoudre les équations de Maxwell de manière approximative mais suffisamment précise pour le type d'applications concerné. Les principales méthodes de résolution des équations de Maxwell sont la méthode des éléments finis (FEM), la méthode des moments (MoM) et la méthode des différences finies dans le domaine temporel (FDTD). Les deux premières sont des résolutions fréquentielles et la troisième est une résolution temporelle. Chaque méthode présente des avantages et des inconvénients qui la distinguent des autres, et le choix entre ces différentes méthodes dépend des besoins de l'utilisateur en termes de géométrie du problème à traiter et de ses dimensions. La méthode des moments a l'avantage de résoudre très rapidement et de manière efficace des problèmes de géométrie plane comme cela est très souvent le cas des circuits intégrés sans boîtier. Ensuite, les dimensions peuvent se situer de l'ordre du micromètre, pour les circuits intégrés, jusqu'à des centaines de mètres dans l'analyse de la propagation d'ondes autour de navires ou d'avions, par exemple. La méthode des différences finies dans le domaine temporel présente des avantages significatifs en termes de rapidité de calcul lorsqu'il s'agit d'analyser des volumes importants. La méthode des éléments finis permet de résoudre des problèmes de toutes les tailles et formes mais le temps de calcul est directement proportionnel au volume de la structure à analyser et à la fréquence de fonctionnement. Elle est préconisée pour des structures comme les boîtiers de circuits intégrés.

La Figure 19 décrit de manière qualitative la correspondance entre les méthodes de calcul numériques et la complexité du problème.

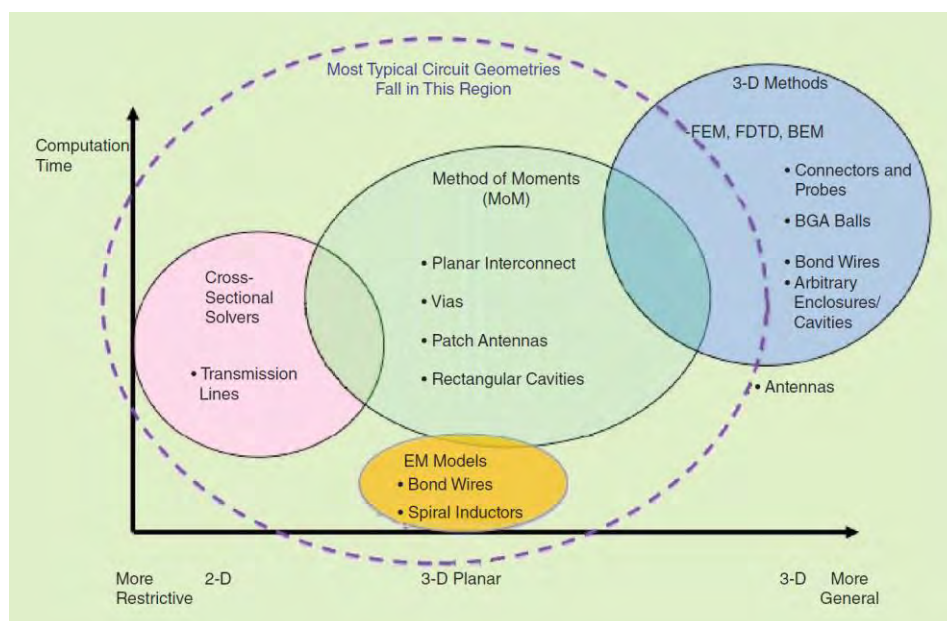


Figure 19 : comparaison des temps de simulation de différents types de simulateurs EM selon la nature du problème [1.43]

1.4 Conclusion

Ce premier chapitre décrit le contexte dans lequel s'est située notre étude. Nous sommes partis des spécifications systèmes très haut niveau du radar automobile et des contraintes liées à cette application, pour ensuite descendre progressivement vers les composants plus élémentaires. Les principales fonctions du radar ont été abordées et en particulier l'émetteur-récepteur radar 77 GHz. Nous nous sommes attardés sur l'encapsulation en boîtier de ce dernier. En effet, comme on a pu le montrer, compte tenu de la criticité des interactions entre le silicium, les couches métalliques du circuit intégré et les éléments du boîtier, cette encapsulation doit être considérée comme une composante à part entière de la puce d'émission-réception.

L'importance de l'interaction entre la puce et son boîtier disqualifie toute conception n'en tenant pas compte, en particulier lors de la conception du mélangeur de fréquence et de l'amplificateur de puissance. L'impossibilité d'effectuer des mesures directes de la transition du PCB vers la puce, à cause de l'absence de ports de mesures au niveau de la puce, rend la modélisation du boîtier encore plus importante.

En outre, une des grandes difficultés de la co-conception de la puce et de son boîtier réside dans l'utilisation d'un boîtier « fan-in » WLP, en l'occurrence le WLCSP. L'absence totale d'implémentation antérieure de ce boîtier aux fréquences millimétriques ainsi que les très nombreuses contraintes de couplage associées à son utilisation à ces fréquences donnent une opportunité d'innovation unique et un goût particulier au travail à réaliser au cours de cette thèse.

2 Modélisation des interconnexions et des composants passifs de l'émetteur-récepteur radar

2.1 Introduction

Après avoir évoqué les différents composants du module radar automobile et leurs fonctions, on se focalise maintenant sur l'émetteur-récepteur et les défis qui lui sont associés. Les travaux présentés ici sont principalement axés sur la conception aux fréquences millimétriques. A ces fréquences, on doit disposer, d'une part, de transistors très performants, et d'autre part, d'une technologie permettant la conception d'éléments passifs à faibles pertes. De plus, cette conception requiert des outils de modélisation électrique et électromagnétique adaptés à ces fréquences, ainsi qu'une très bonne connaissance des caractéristiques du substrat utilisé.

Les principaux éléments passifs utilisés pour les conceptions en bande millimétrique que nous présenterons par la suite sont les lignes de transmission, les inductances, les transformateurs d'impédance, les baluns, les diviseurs de puissance et les capacités. La modélisation des lignes de transmission, notamment microrubans, est souvent incluse dans les kits de conception, « Process Design Kit » (PDK). Ces modèles électriques sont généralement sous forme d'équations ou de circuits RLCG équivalents, basés sur des ajustements aux mesures. Ils représentent assez fidèlement le comportement électrique réel des lignes de transmission mais ne sont pas en mesure de rendre compte d'un couplage éventuel avec d'autres éléments du circuit. Pour certaines technologies, il peut être fourni un modèle électrique de lignes couplées mais qui se limite à deux ou plusieurs tronçons de lignes planaires et parallèles. Par ailleurs, aucun modèle ne sera fourni pour un couplage 3D, compte tenu de la trop grande variété des cas de figures. Ces limitations empêchent l'utilisation de modèles électriques pour simuler des structures impliquant des interactions entre différentes lignes, comme les transformateurs ou les capacités inter-digitées. Le meilleur moyen de prendre en compte les interactions entre ces lignes est la modélisation électromagnétique. Les simulateurs planaires basés sur la méthode des moments (MoM) et les simulateurs 3D basés sur la méthode FEM sont privilégiés pour ce type de structures. Néanmoins, au préalable, il est indispensable de valider les résultats de simulations électromagnétiques par quelques mesures. Même avec une confiance assez élevée par rapport au simulateur EM et aux caractéristiques fournies pour le substrat, il n'est pas rare d'être surpris par des comportements inattendus n'ayant pas été anticipés. En outre, plus la fréquence est élevée, plus l'importance de cette validation augmente. Il est donc indispensable pour la conception d'éléments passifs pour un radar automobile 77-GHz d'inclure le plus tôt possible la validation expérimentale dans le flot de conception.

2.2 Caractéristiques du substrat SiGe BiCMOS180

Avant d'entamer la conception d'éléments passifs, il est important de connaître les propriétés du substrat utilisé et les performances qu'il peut autoriser. On s'intéresse particulièrement aux métallisations, aux matériaux diélectriques les séparant, des oxydes principalement, et au substrat qui supporte l'ensemble du circuit qui est ici le Silicium. Ces paramètres définissent en partie les architectures susceptibles d'être implémentées, et les facteurs de qualité qui pourront être obtenus pour les éléments passifs.

2.2.1 Propriétés des métaux et des oxydes

Au niveau des métaux et des oxydes, il faut connaître le nombre de couches métalliques à disposition, l'épaisseur et la conductivité de chacune d'elles, ainsi que les pertes diélectriques et les permittivités des couches d'oxydes séparant ces couches métalliques. Ces paramètres sont essentiels pour la modélisation des éléments et la détermination de leurs performances.

2.2.1.1 Les métallisations

Les dimensions et les propriétés physiques des métaux définissent en partie leurs caractéristiques électriques. Les deux propriétés principales des métaux impactant les performances électriques du circuit sont l'épaisseur de la métallisation et sa conductivité (ou sa résistivité). Si la conductivité impacte directement les pertes du matériau, l'impact de l'épaisseur de la métallisation sur les pertes peut être plus limité en raison de l'effet de peau.

Conductivité électrique :

La conductivité électrique caractérise l'aptitude des matériaux à transmettre un courant électrique grâce au déplacement des charges électroniques libres. Son unité est le siemens par mètre ($S.m^{-1}$). Son inverse est la résistivité et se mesure en ohm-mètre ($\Omega.m$). Les métallisations du circuit intégré sont généralement réalisées avec un alliage à base d'aluminium. La conductivité électrique de l'aluminium se situe autour de $3.50 \times 10^7 S.m^{-1}$. Il n'est pas rare que des couches de métallisation du substrat contiennent également du cuivre, dont la conductivité est d'environ $5.96 \times 10^7 S.m^{-1}$, afin d'augmenter la conductivité de la métallisation.

Effet de peau :

L'effet de peau est un phénomène électromagnétique qui conduit à une réduction de la pénétration des lignes de courant dans le conducteur à mesure que la fréquence augmente. Les lignes de courant ont tendance à rester proches de la surface aux très hautes fréquences. Cette propriété est très intéressante pour des circuits fonctionnant à hautes fréquences, comme le radar 77-GHz. En effet, les contraintes sur l'épaisseur des lignes métalliques intervenant dans la propagation des ondes millimétriques sont réduites.

La formule (4) permet de calculer l'épaisseur de peau :

$$\delta = \frac{1}{\sqrt{\sigma \times \mu \times \pi \times f}} \quad (\text{m}) \quad (4)$$

Avec :

δ : épaisseur de peau en mètre [m]

f : fréquence du courant en hertz [Hz]

μ : perméabilité magnétique en henry par mètre [H/m]

σ : conductivité électrique en siemens par mètre [S/m]

Pour un signal se propageant à 77 GHz sur un support métallique dont la conductivité se situe autour de 5×10^7 S/m et la perméabilité relative autour de 1, on obtient une épaisseur de peau d'environ 0.25 μm . Ce résultat est très intéressant parce qu'il permet de conclure qu'il suffit de disposer de couches métalliques dont l'épaisseur est légèrement supérieure à 0.25 μm pour que l'épaisseur du métal porteur du signal n'ait aucune influence sur les pertes.

2.2.1.2 Les oxydes

Les oxydes utilisés sont des matériaux diélectriques. Ils sont donc caractérisés par une permittivité, une épaisseur et des pertes.

La permittivité diélectrique :

La permittivité diélectrique ε est un paramètre important puisqu'il décrit la réponse du diélectrique à un champ électrique appliqué. On parle souvent de permittivité relative ou constante diélectrique ε_r qui est une valeur normalisée par rapport à la permittivité du vide ε_0 , de sorte que :

$$\varepsilon = \varepsilon_r \times \varepsilon_0 \quad , \quad \text{avec } \varepsilon_0 = 8.854187 \times 10^{-12} \text{ F.m}^{-1}$$

Les valeurs de constantes diélectriques des oxydes inter-métaux présents dans les technologies monolithiques silicium sont globalement entre 3 et 7. Il peut exister plusieurs oxydes, avec des permittivités relatives différentes, entre deux couches métalliques superposées. Une permittivité équivalente est extraite dans ces cas-là pour simplifier les modèles électromagnétiques.

L'épaisseur :

L'épaisseur des diélectriques inter-métaux est définie par la technologie. Combinée aux permittivités, elle permet de déterminer la capacité entre les couches de métallisation. C'est un paramètre important pour la conception des éléments passifs et leur dimensionnement.

Les pertes diélectriques :

La tangente de l'angle de perte diélectrique, notée $\tan \delta$, est un autre paramètre à prendre en compte au niveau du diélectrique. Sa formulation, prenant en compte les pertes par conduction, est la suivante [2.1]:

$$\tan \delta = \frac{\epsilon_r'' + (\sigma_c / \omega \epsilon_0)}{\epsilon_r'} \quad (5)$$

Avec ϵ_r' et ϵ_r'' les composantes réelle et imaginaire de la permittivité diélectrique complexe ϵ_r du milieu.

Les valeurs de $\tan \delta$ sont comprises entre 10^{-2} et 10^{-3} à 77 GHz pour les diélectriques mis en œuvre dans les technologies.

2.2.2 Propriétés du Silicium

Les substrats utilisés en BiCMOS SiGe sont faiblement dopés. La résistivité de ces substrats est généralement comprise entre 10 et 20 $\Omega \cdot \text{cm}$, ce qui permet d'obtenir d'assez bonnes performances pour les éléments passifs ne disposant pas d'un écran métallique entre le signal et le substrat. Cependant, des couches fortement dopées (couches enterrées P^+/N^+) sont implantées à la surface du substrat P^- . Ces couches servent à prévenir des phénomènes de verrouillage ("latch-up") dans les circuits logiques CMOS. Leur résistivité étant beaucoup plus faible qu'un substrat faiblement dopé, elles génèrent des pertes importantes dans les éléments passifs placés au-dessus et ne disposant pas d'écran métallique. Pour éviter de telles situations, des masques négatifs permettent d'éviter le dépôt de ces couches enterrées fortement dopées dans les zones où l'intégration de composants passifs est prévue.

Les pertes dans le substrat sont minimales pour une technologie BiCMOS SiGe quand aucun dopage supplémentaire n'est présent. La valeur de résistivité du substrat dans ce cas est autour de 18 $\Omega \cdot \text{cm}$. Cette résistivité reste très faible comparée à un substrat isolant comme le GaAs. Les plages de résistivité généralement rencontrées pour les substrats silicium le situent dans une catégorie entre l'isolant et le conducteur. Un tel substrat génère donc des pertes aux fréquences millimétriques. On distingue principalement des pertes par induction électrique et des pertes par induction magnétique [2.2]. La Figure 20 ci-dessous représente les formes des pertes associées à une inductance intégrée au-dessus du substrat.

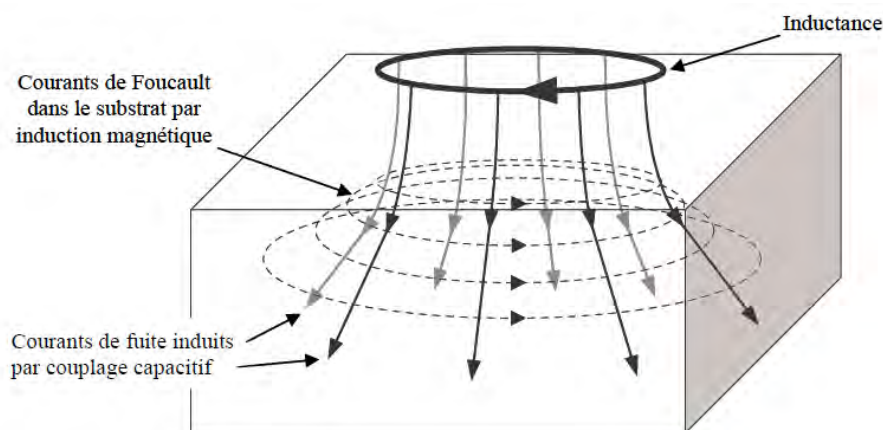


Figure 20 : pertes associées à une inductance sur Silicium

2.2.2.1 Pertes par induction électrique

Un couplage existe entre le composant passif et le substrat. Ce couplage est modélisé par une capacité. Il résulte de cette capacité des courants de conduction et de déplacement dans le substrat vers la masse la plus proche. Cette énergie électrique est finalement dissipée par effet joule dans le substrat [2.3].

2.2.2.2 Pertes par induction magnétique

Le champ magnétique créé par les courants qui circulent dans les éléments passifs au-dessus du substrat induit des courants « images » dans le substrat, appelés courants de Foucault. Ces courants génèrent un champ magnétique qui s'oppose au champ principal. De l'énergie magnétique est ainsi convertie en chaleur ici encore par effet joule dans le volume du substrat. Plus la résistivité du substrat est faible, plus les conséquences de ce phénomène sont importantes.

2.2.3 Les niveaux de métallisation (stack-up)

Les métaux implémentés dans les technologies monolithiques ont des épaisseurs de l'ordre du micromètre, comprises généralement entre $0.2\ \mu\text{m}$ et $3\ \mu\text{m}$. Dans les technologies destinées aux applications RF, les métallisations supérieures sont souvent différentes des autres. Elles sont plus épaisses et leur alliage est de composition différente afin de présenter une meilleure conductivité. Du cuivre est souvent rajouté à l'alliage à base d'aluminium pour améliorer la conductivité électrique.

La technologie utilisée dans le cadre des travaux de cette thèse dispose de six niveaux de métallisations. Les métaux minces vont de M1 à M4 et 2 métaux épais LM et LM2 sont utilisés au niveau des couches supérieures.

La Figure 21 illustre les six niveaux de métallisation de cette technologie BiCMOS SiGe [2.4].

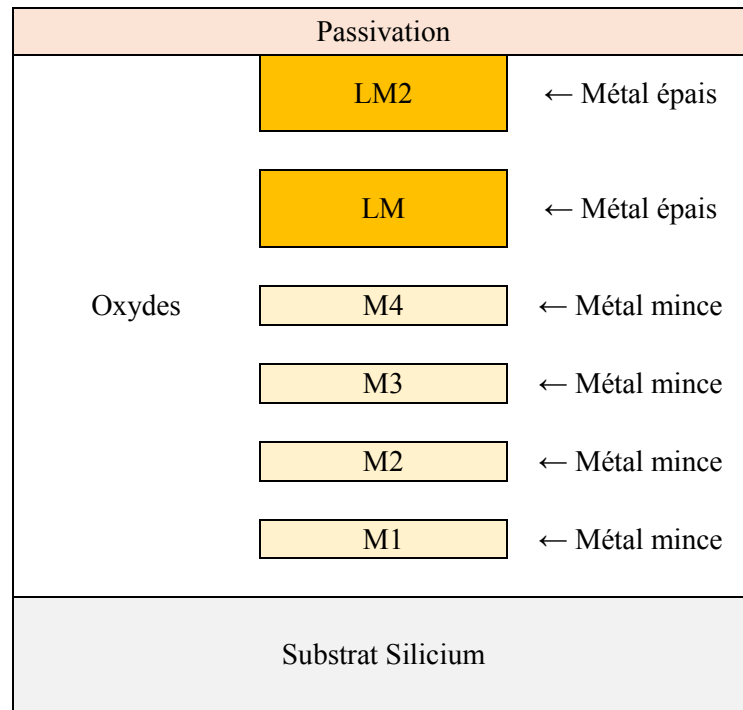


Figure 21 : niveaux de métallisation de la technologie BiCMOS utilisée

Il n'est pas courant d'avoir deux couches métalliques épaisses et cet avantage est très important en termes de conception. Effectivement, il donne plus de libertés dans le choix des architectures des éléments passifs et notamment en permettant la conception d'inductances et de transformateurs sur deux niveaux sans être pénalisé par l'épaisseur des métaux.

Cette technologie contient un nombre important d'oxydes avec, certaines fois, plusieurs oxydes différents entre deux métaux. Les permittivités varient d'un oxyde à l'autre mais on considère un seul oxyde équivalent entre deux métaux. La permittivité et la hauteur (ε_{eq}, h_{eq}) de cet oxyde, lorsqu'il remplace deux oxydes superposés dont les couples permittivité et hauteur sont respectivement (ε_1, h_1) et (ε_2, h_2) , s'écrivent :

$$\varepsilon_{eq} = \frac{\varepsilon_1 \varepsilon_2 (h_1 + h_2)}{\varepsilon_1 h_2 + \varepsilon_2 h_1} \text{ et } h_{eq} = h_1 + h_2, \quad (6)$$

2.3 Validation des modèles d'interconnexions aux fréquences millimétrique

Après avoir défini le substrat et les niveaux métalliques de la technologie, la modélisation électromagnétique des structures passives peut commencer. Des ajustements des propriétés des matériaux pourront éventuellement être réalisés au cours de l'étape de validation expérimentale des simulations. La première phase de la modélisation d'éléments passifs intégrés consiste à concevoir des

éléments assez simples. Cela permet de limiter les contraintes associées à la structure, afin d'identifier facilement les éventuelles différences entre le modèle et la mesure.

2.3.1 Méthodologie de modélisation

Nous nous intéresserons lors de cette partie à la méthodologie de modélisations des structures passives. En premier lieu, des outils de simulation EM sont utilisés pour modéliser ces structures. Par la suite, ces structures sont fabriquées en plaquettes et elles sont validées expérimentalement.

2.3.1.1 Outils de simulation électromagnétique

L'une des questions les plus courantes qui se pose au moment de la mise en place d'une modélisation électromagnétique est la convenance de l'outil de simulations EM utilisé et sa précision par rapport au sujet traité. Un flux de simulation, basé sur l'interopérabilité entre les logiciels de simulation Virtuoso de Cadence et ADS de Keysight a été développé lors de cette thèse. Ce flux permet la réalisation de simulations électriques et électromagnétiques en utilisant les mêmes bibliothèques [2.5]. Le simulateur EM du logiciel ADS de Keysight est donc préféré aux autres simulateurs EM.

Néanmoins, les performances de ce logiciel seront comparées à celles du simulateur HFSS et les résultats obtenus seront confrontés à des mesures afin de s'assurer de son bon fonctionnement jusqu'à 110 GHz et de la convenance des temps de calcul par rapport à nos besoins. Ce simulateur propose deux méthodes de résolution des équations de Maxwell : la méthode des moments (MoM) et la méthode des éléments finis (FEM). Le simulateur FEM est préféré puisqu'il permet de simuler également des structures 3D, ce que ne permet pas la méthode MoM. En effet, il pourra être utilisé pour la modélisation de l'interface boîtier sans devoir valider de nouveau les éléments modélisés précédemment. En plus d'une bonne précision, on s'attend à un temps de simulation raisonnable afin de valider cette méthode.

2.3.1.2 Banc de validation expérimentale

Nous venons de présenter la méthodologie de modélisation électromagnétique des structures passives. Afin de confirmer la validité des modèles EM et les ajuster si besoin, une validation expérimentale est requise. Des mesures sont réalisées en paramètres S à l'aide d'un VNA (pour Vector Network Analyzer) 110 GHz et d'une station sous pointes. Les pointes utilisées sont de type GSG avec un espacement de 150 μm entre la pointe signal et chacune des deux pointes de masse. Le système est calibré pour fixer le plan de référence au niveau des pointes. Nous utilisons les algorithmes SOLT et TRL et un kit de calibrage fourni par le constructeur des pointes, selon les besoins de caractérisation.

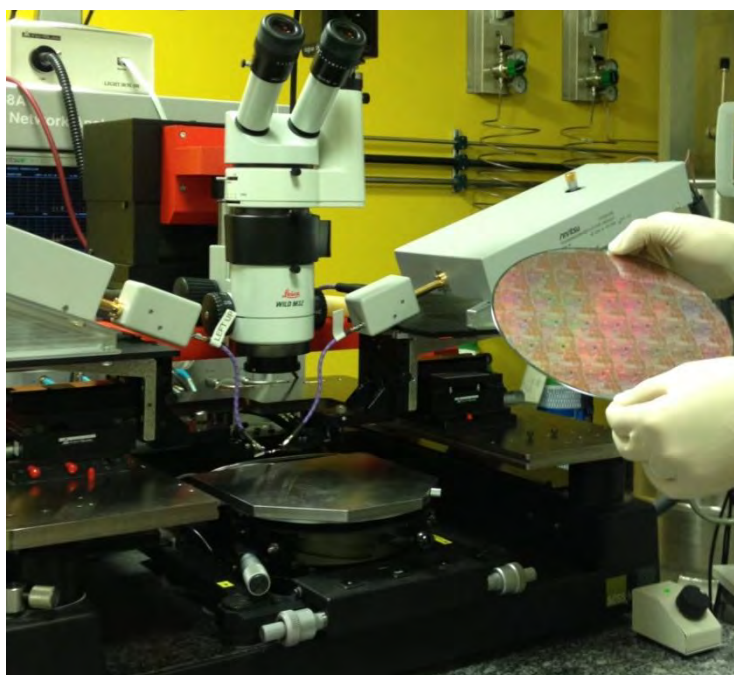


Figure 22 : photo du banc de test sous pointes et d'une plaquette à mesurer

La photo de la Figure 22 présente la station sous pointes, des mélangeurs externes permettant au VNA d'atteindre 110 GHz et une plaquette sous test contenant, entre autres, les structures de lignes et plots à caractériser.

2.3.2 Modélisation électromagnétique et validation expérimentale

La modélisation d'interconnexions simples correctement choisies permet de valider la simulation électromagnétique en comparant les résultats obtenus à la mesure, et de lever ainsi le doute sur de nombreux paramètres. Pour ce travail préliminaire, on se focalisera donc sur les structures de plots et de lignes de transmissions. Ces composants sont généralement assez simples à simuler et à caractériser alors qu'ils permettent d'apporter des informations importantes sur la qualité de la simulation.

2.3.2.1 Plots standards pour des signaux de fréquences millimétrique

Les plots pour la connexion de la puce sont la plus part du temps présents dans la conception et dans la caractérisation des blocs millimétriques. Ils sont intégrés dans l'adaptation du bloc pour présenter des performances optimales en aval du plot, ou épluchés de la mesure du bloc si on s'intéresse aux performances du bloc en amont du plot. Dans les deux cas, une caractérisation précise du plot est requise. Il est souhaitable d'utiliser la taille de plots la plus petite possible mais cette réduction des dimensions est limitée par la nature des interconnexions entre la puce et son environnement, comme par exemple le type des pointes de mesure. Dans le cas où la puce est encapsulée, la taille minimale des plots varie généralement entre 0.0025 mm^2 et 0.01 mm^2 selon le type de boîtier. La capacité parallèle présente peut ainsi dépasser 100 fF, ce qui impacte significativement les performances de circuits opérant autour de

77-GHz. En effet, la formule (7) calculant la capacité (C) d'un condensateur plan, basé sur des valeurs typiques pour les technologies monolithiques, donne une valeur de 78 fF.

$$C = \varepsilon_0 \times \varepsilon_r \times \frac{A}{d} \quad (7)$$

Avec :

$\varepsilon_0 = 8,84 \times 10^{-12} \text{ F.m}^{-1}$, la permittivité électrique du vide

$\varepsilon_r = 4$, la permittivité relative du diélectrique

$A = 0.01 \text{ mm}^2$, la surface du plan métallique du plot

$d = 4.7 \text{ }\mu\text{m}$, la distance entre le plan métallique du plot et la masse

Dans un premier temps, nous procédons à la modélisation électromagnétique d'un plot intégrant un écran métallique compte tenu que cette structure est plus facile à modéliser qu'un plot au-dessus d'un substrat. Les plots des technologies Silicium utilisent une couche métallique supplémentaire plus épaisse que les autres métallisations afin de garantir un bon contact et une bonne résistance mécanique. Le métal est généralement un alliage de cuivre et d'aluminium (AlCu). Le motif du dernier niveau métallique du plot déborde des motifs métalliques inférieurs de manière plus ou moins significative selon les technologies et le dessin des masques. Sur la Figure 23 ci-dessous une vue en perspective présente le modèle électromagnétique du plot à caractériser où le métal en violet correspond à l'AlCu :

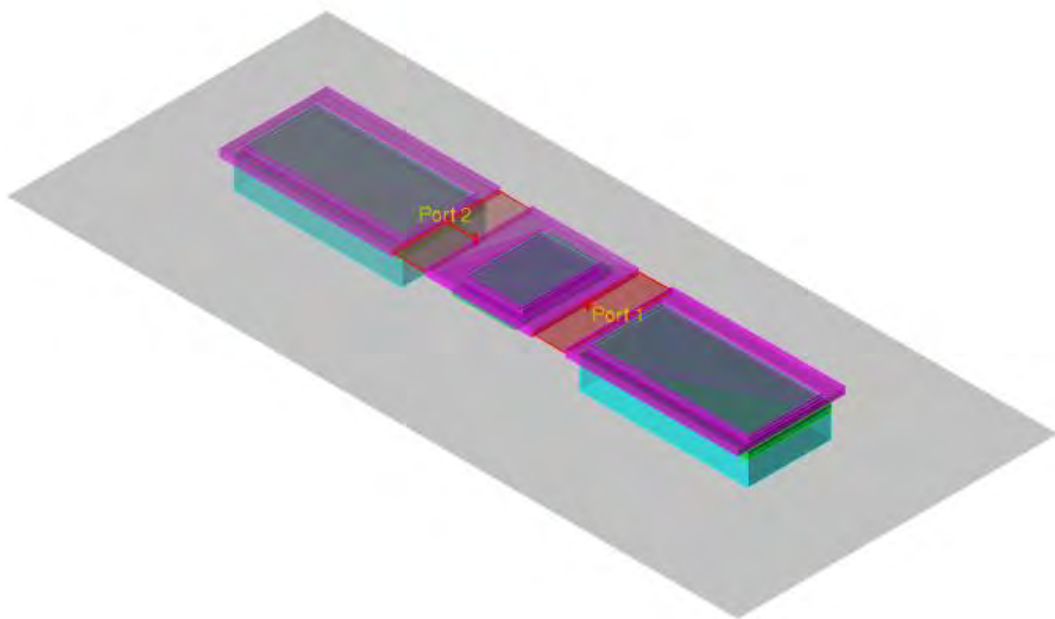


Figure 23 : vue 3D du plot de puce

Ce plot métallique est un plot de type masse-signal-masse (GSG pour Ground-Signal-Ground) et la couche métallique supérieure est un carré d'environ 100 μm par 100 μm . Il a été simulé à l'aide des outils ADS et HFSS des sociétés Keysight et ANSYS respectivement. Les deux méthodes de résolution numérique Momentum et FEM ont été utilisées pour ADS et la méthode FEM a été utilisée pour HFSS.

Par la suite, on identifiera les simulations réalisées à l'aide de l'outil ADS Momentum par MoM, les simulations réalisées à l'aide de l'outil ADS FEM par FEM et les simulations réalisées à l'aide de l'outil HFSS FEM par HFSS. Les modèles EM utilisés dans les deux simulateurs sont présentés sur la Figure 24. Ils ont des géométries identiques et utilisent les mêmes propriétés de matériaux.

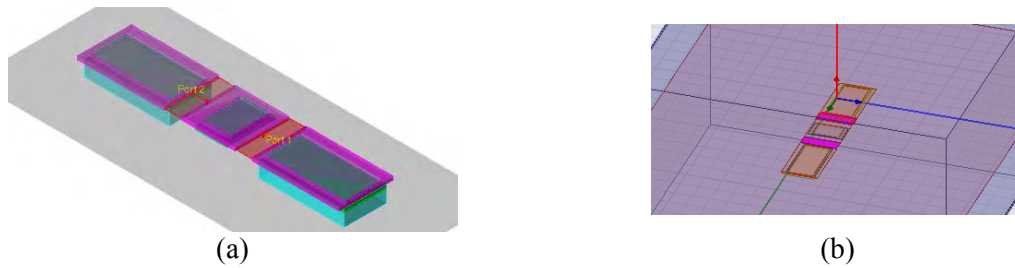


Figure 24 : comparaison entre les vues 3D ADS (a) ET 3D HFSS (b) utilisant les mêmes types de port

La Figure 25 présente les résultats des différentes simulations en paramètres S du plot métallique en dipole. Ces résultats de simulations sont accompagnées de mesures réalisées dans les conditions les plus proches de la simulation en utilisant des paramètres S 1-port du VNA 110 GHz sous pointes.

Comme le plot dispose d'une composante capacitive dominante, on peut négliger son inductance et on considère que la capacité parallèle du plot est égale à:

$$C_{plot} = \frac{-1}{\text{Im}(Z_{11}) \times \omega} \quad (8)$$

Avec :

$$\omega = 2\pi \times \text{fréquence}$$

Z la matrice des paramètres impédances

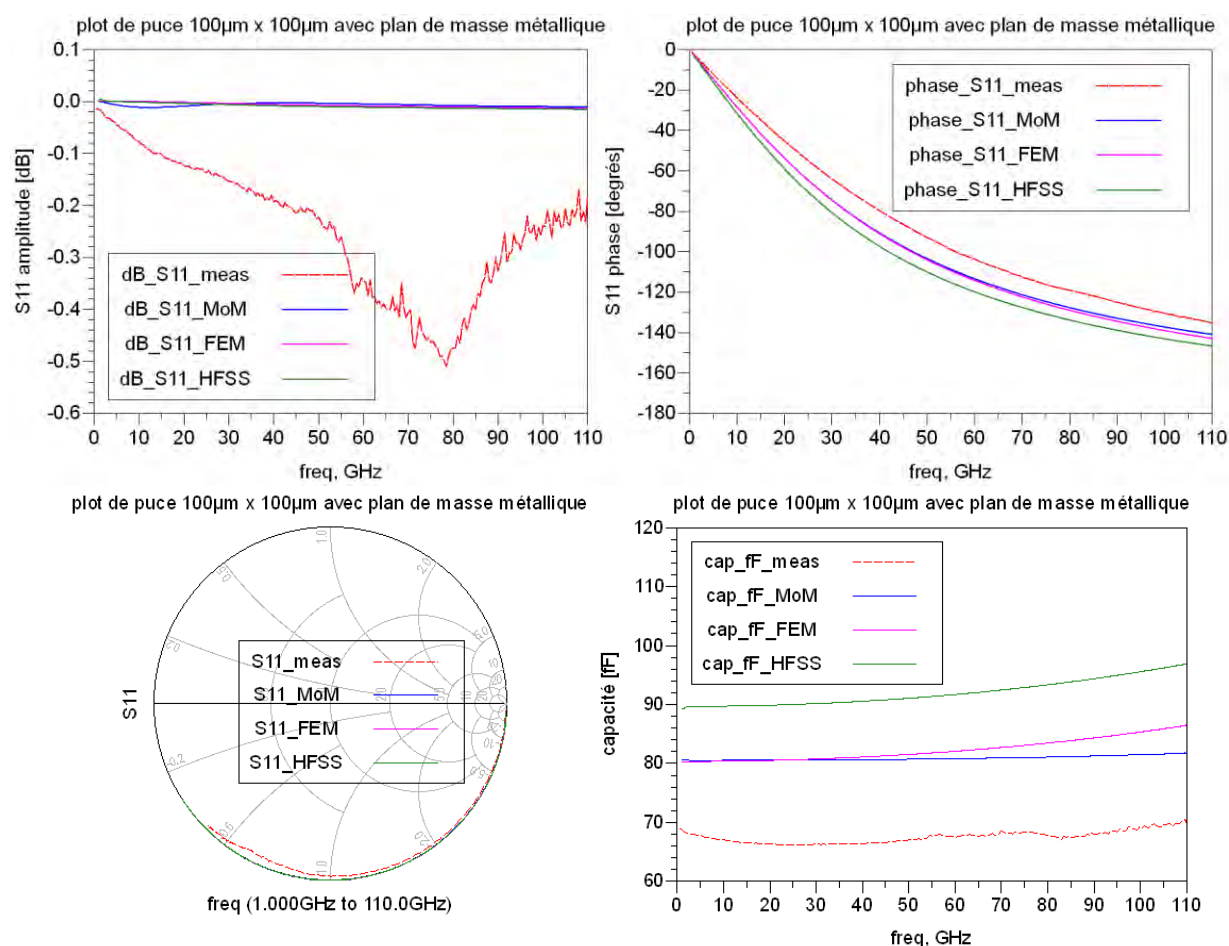


Figure 25 : résultats de simulations ADS MoM, ADS FEM et HFSS et de mesures des plots d'accès avec plan de masse

Les trois outils de simulations donnent des résultats assez proches en amplitude et en phase, mais l'extraction de la capacité à partir des paramètres S obtenus montre un écart de 10 fF (11 %) entre HFSS et ADS. Cet écart ne dépend pas de la méthode de calcul ou de la géométrie et des propriétés de la structure. En effet, MoM et FEM sur ADS donnent des résultats identiques jusqu'à 40 GHz et la structure utilisée dans HFSS a été directement importée au format « SAT » depuis ADS. La principale différence entre les deux outils est la manière avec laquelle les ports sont définis. En effet, même lorsqu'on choisit des ports non calibrés dans les simulateurs, les outils concernés effectuent des opérations de calibrage qui leur sont propres pour compenser l'inductance virtuelle du port. Des informations sur la théorie des ports non calibrés peuvent être trouvées dans l'aide des outils. Des travaux [2.6] confirment la grande dépendance entre les résultats de la simulation du plot GSG et les dimensions du port non calibré ainsi que son emplacement. Des simulations basées sur une variation de l'ordre d'une dizaine de microns des dimensions du rectangle constituant le port localisé montrent une variation de phase avoisinant 30° (19%) à 67 GHz. La différence entre les simulations ADS et HFSS est de l'ordre de 5° (4%) à la même fréquence.

On remarque que les simulations et les mesures donnent des résultats assez similaires. Ceci dit, on s'est fixé comme cible de précision une différence de valeur de capacité ne dépassant pas 10%, or on a ici un écart d'environ 20%.

À l'aide du simulateur ADS FEM, nous avons recherché l'origine de l'écart entre les simulations et la mesure. Le choix de se focaliser sur un simulateur FEM est principalement lié à l'intérêt de cet outil pour la caractérisation du boîtier comme déjà expliqué. Pour la préférence d'ADS par rapport à HFSS, le principal argument est l'interopérabilité naturelle offerte par cet outil concernant l'exploitation des dessins en provenance de Cadence Virtuoso. Les deux outils d'aide à la conception partagent la même base de données appelée « OpenAccess ».

Après des investigations autour de nombreux paramètres du modèle électromagnétique tels le maillage, les ports, les dimensions et les propriétés des matériaux, il s'est avéré que la définition des ports explique une grande partie de cet écart entre les simulations et les mesures, comme le montre le Tableau 3.

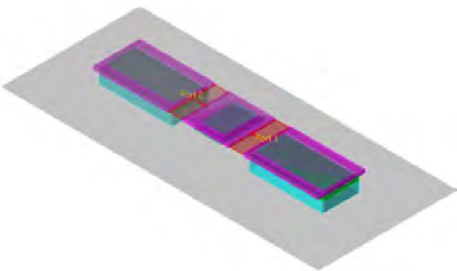
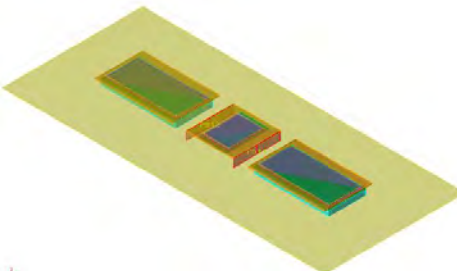
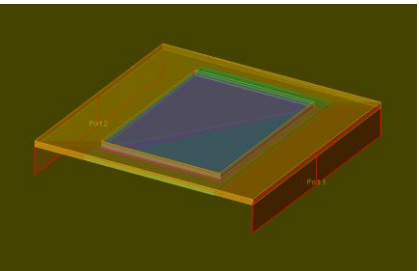
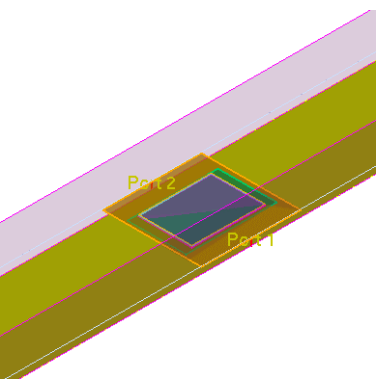
Type de ports	Dessin 3D du modèle	Phase à 80 GHz (degrés)	Capacité du plot (fF)	
Non calibré : ports définis entre plot de signal et plots de masse		-128	81	Point de départ
Non calibré : ports définis entre plot de signal et plan métallique de masse		-129	82	Étape 1
Non calibré : ports définis entre plot de signal et plan métallique de masse après suppression des plots de masse		-129	82	Étape 2
Calibré en TML (ou Wave Port dans HFSS) : ports définis entre plot de signal et plan métallique de masse		-125	75	Étape Finale

Tableau 3 : évolution de la capacité parallèle simulée en fonction de la définition des ports

Les différentes étapes ont permis de confirmer l'impact de la définition des ports d'excitation sur les résultats de simulation des paramètres S des plots. On constate que les écarts en phase sont très faibles, de l'ordre de 4° à 80 GHz, mais que la différence qui en résulte sur la valeur de capacité n'est pas négligeable.

Les graphes de la Figure 26 comparent les mesures aux simulations ADS FEM avant et après ajustement des ports.

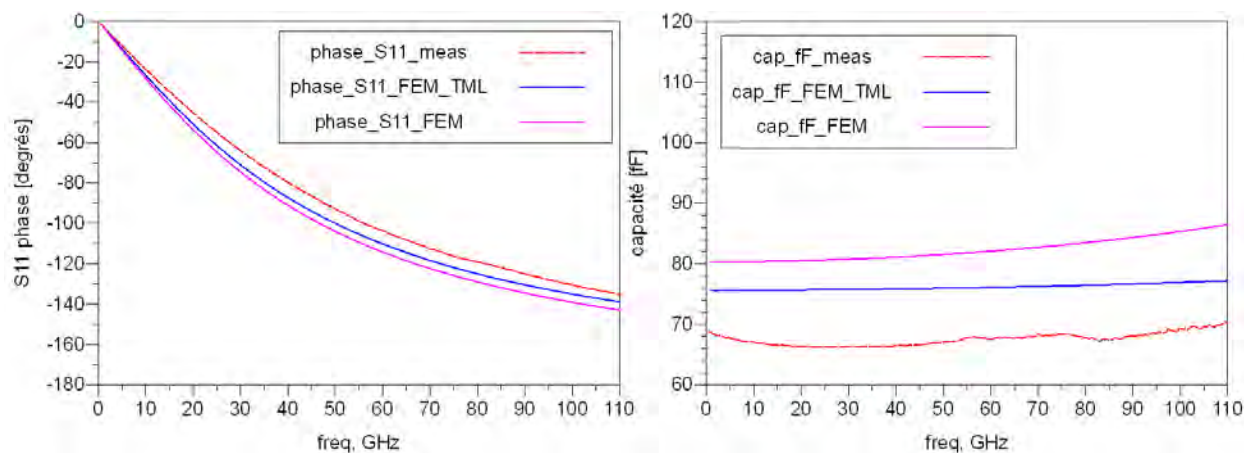


Figure 26 : comparaison des simulations et des mesures de la capacité et de la phase S_{11} du plot de puce $100\ \mu\text{m} \times 100\ \mu\text{m}$ avec plan de masse métallique pour des ports localisés et des ports de type guide d'onde

La différence entre la capacité mesurée et simulée est maintenant d'environ 10%, ce qui donne une assurance quant à la qualité de la modélisation EM.

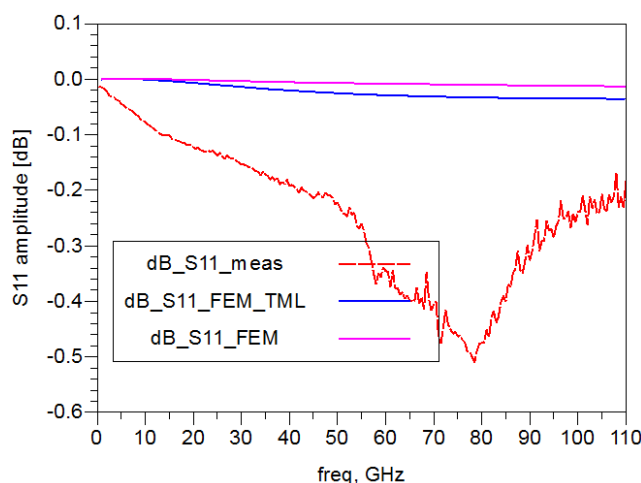


Figure 27 : comparaison des simulations et des mesures de l'amplitude S_{11} du plot de puce $100\ \mu\text{m} \times 100\ \mu\text{m}$ avec plan de masse métallique pour des ports localisés et des ports de type guide d'onde

Outre les différences constatées au niveau de la phase, on remarque une légère différence entre les simulations et les mesures en amplitude (Figure 27). La simulation ne reflète pas les pertes d'environ 0.3 dB observées sur la mesure de S_{11} . Ces pertes sont équivalentes aux pertes générées par une résistance d'environ $2\ \Omega$ en série avec la capacité. Une partie de ces pertes s'explique au niveau de la mesure par le contact de la pointe avec la métallisation AlCu qui ne présente pas un très bon contact et génère une résistance parasite en plus des effets de rayonnement entre le plot et les pointes. Une autre partie s'explique par la représentation des vias au niveau des plots de masse par des rectangles parfaits. La raison derrière cette approximation est la limitation du temps de calcul. Cet écart n'a pas été

investigé d'avantage dans le cadre de ce travail vu la très faible contribution de ces pertes sur la transmission, environ 0.15 dB.

2.3.2.2 Plots optimisés pour minimiser la capacité parallèle

On s'intéresse aux possibilités de réduction de la capacité parallèle introduite par le plot métallique qui impacte les performances du circuit aux ondes millimétriques. Par conséquent, nous cherchons à concevoir un plot minimisant la capacité parallèle tout en restant compatible avec les processus de fabrication et d'interconnexion de la puce. Le plot a donc été réduit à $100\mu\text{m} \times 60\mu\text{m}$. Ces dimensions correspondent au minimum autorisé par les règles de mise en boîtier « fan-out » WLP des puces. Cette opération permet de réduire d'environ 40% la capacité parallèle totale. La deuxième modification consiste à créer une ouverture au niveau du plan métallique de masse M1 à la verticale du plot. Cela réduit encore la capacité mais le signal « voit » alors le substrat Silicium. Le principal défi associé à cette architecture est donc lié à l'intervention des pertes du Silicium dans les caractéristiques du plot : la réduction de la capacité ne doit pas être réalisée au détriment d'une dégradation importante de son coefficient de qualité. Le plot métallique étudié, avec ouverture du plan de masse, est présenté sur la Figure 28.

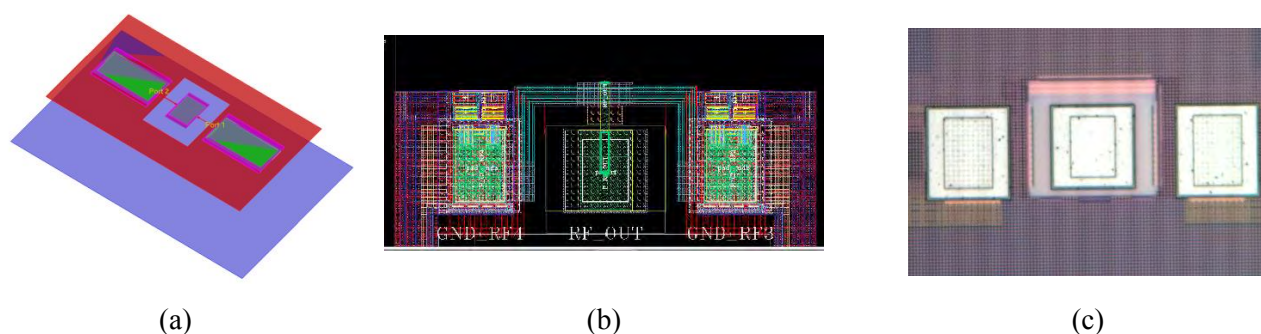


Figure 28 : modèle en vue 3D (a), dessin de masque (b) et photo (c) du plot GSG sans masse sous le plot signal

Le plot « signal » qui se retrouve en regard du Silicium après ouverture du plan de masse présente plus de pertes que le plot standard intégré au-dessus d'un plan de masse métallique complet. En revanche, cette configuration de plot possède l'avantage d'être jusqu'à quatre fois moins capacitive qu'un plot standard si la résistivité du Silicium est assez élevée, c'est-à-dire supérieure à $5\ \Omega\cdot\text{cm}$. En effet, le Silicium présente un comportement qui dépend de sa résistivité et de la fréquence des signaux. Des études réalisées sur la propagation d'ondes pour des structures sur Silicium ont permis d'établir en 1971 un abaque résistivité-fréquence qui dégage l'existence de trois régions de propagation [2.7]. La Figure 29 illustre ces trois régions pour des lignes de différentes largeurs.

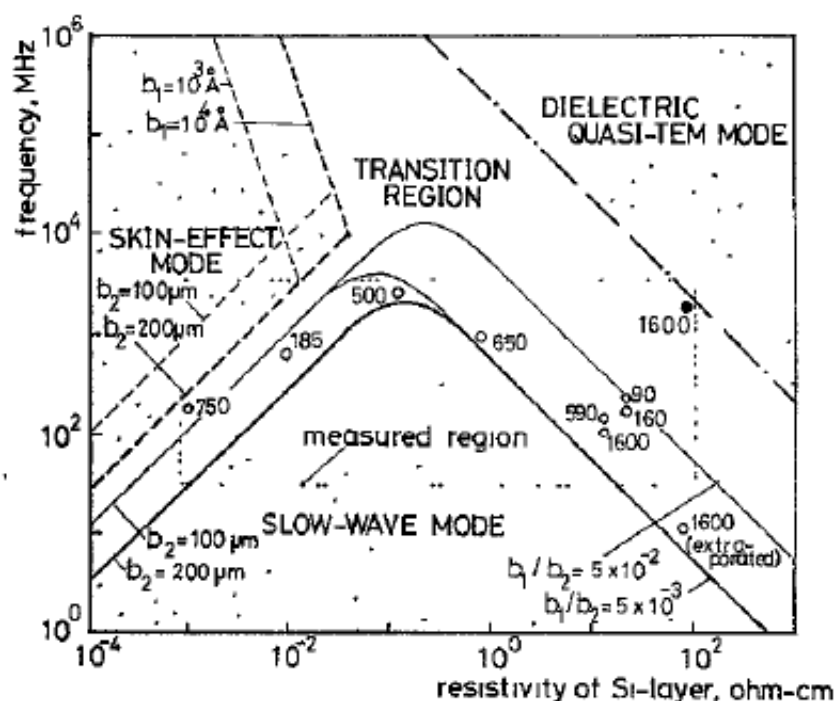


Figure 29 : modes de propagation d'une onde en fonction de la résistivité du Silicium et de la fréquence [2.7]

Le mode de propagation « onde lente » confère des caractéristiques particulières aux lignes de transmissions qui peuvent être intéressantes pour certaines applications. Cependant, ici, ce mode de propagation n'est pas pertinent vu les fréquences mises en jeu. On peut ainsi se focaliser sur les modes de propagation « effet de peau » et quasi-TEM. À des fréquences supérieures à 10 GHz, le mode de propagation « effet de peau » se produit principalement pour des résistivités du Silicium assez faibles, inférieures à 0.1 $\Omega\cdot\text{cm}$. Le Silicium se comporte dans ce mode comme un conducteur à pertes. Pour ces mêmes fréquences, le mode de propagation quasi-TEM est lié à des résistivités du Silicium généralement supérieures à 1 $\Omega\cdot\text{cm}$. Le Silicium se comporte alors dans ce cas comme un diélectrique que l'énergie traverse.

Afin d'identifier l'impact de la résistivité du Silicium sur les caractéristiques du plot, nous avons réalisé des simulations du plot pour différentes valeurs de résistivité du Silicium. La Figure 30 montre les résultats obtenus à partir des simulations EM des paramètres du dipôle. Nous nous intéressons en premier lieu à la capacité générée par le plot de la puce en fonction de la fréquence et de la résistivité du Silicium en $\Omega\cdot\text{cm}$. L'étude est axée sur la bande de fréquences d'intérêt pour le radar, [76-81] GHz.

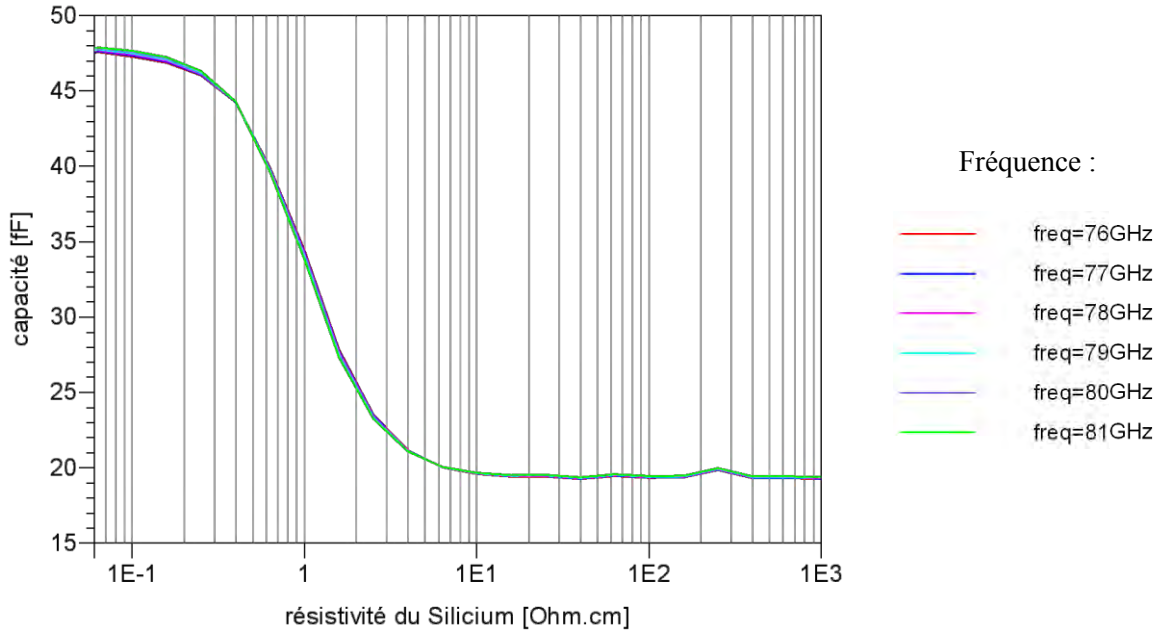


Figure 30 : variation de la capacité du plot en fonction de la résistivité du Silicium dans la bande [76-81] GHz

Les résultats obtenus à partir des simulations EM sont en parfait accord avec les modes de propagations décrits par [2.7]. Le comportement du Silicium est très similaire à un conducteur jusqu'à une résistivité de 0.5 $\Omega\cdot\text{cm}$ avec une capacité parallèle d'environ 45 fF qui résulte des épaisseurs de diélectriques déposées en face supérieure du silicium. Le calcul de capacité à partir de l'équation (7) donne une capacité d'environ 40 fF, en considérant le Silicium assimilable à un conducteur. À partir de 10 $\Omega\cdot\text{cm}$, le Silicium se comporte comme un diélectrique, et il présente une très faible capacité, avoisinant 19 fF. Cela représente une réduction de 60% de la capacité du plot. On constate que la région entre 0.5 et 10 $\Omega\cdot\text{cm}$ est une région de transition entre les deux modes avec une valeur de tangente pour la capacité assez importante, rendant la conception de circuits dans cette zone assez risquée. En effet, les variations liées à la fabrication auront un impact significatif sur les performances des circuits.

La technologie BiCMOS utilisée pour les puces radars présente une résistivité de substrat autour de 18 $\Omega\cdot\text{cm}$, avec une variation de fabrication affichée autour de 20%. Néanmoins, le plot de la technologie BiCMOS utilisée devrait permettre de rester dans la région où la capacité est à sa valeur minimum. La validation expérimentale permettra de confirmer ou infirmer cela.

Si on s'intéresse maintenant aux pertes générées par ce plot en fonction de la résistivité du Silicium et de la fréquence. La Figure 31 représente le coefficient de qualité du le plot, défini comme étant :

$$Q_{plot} = -\frac{Im(Z_{11})}{Re(Z_{11})} \quad (9)$$

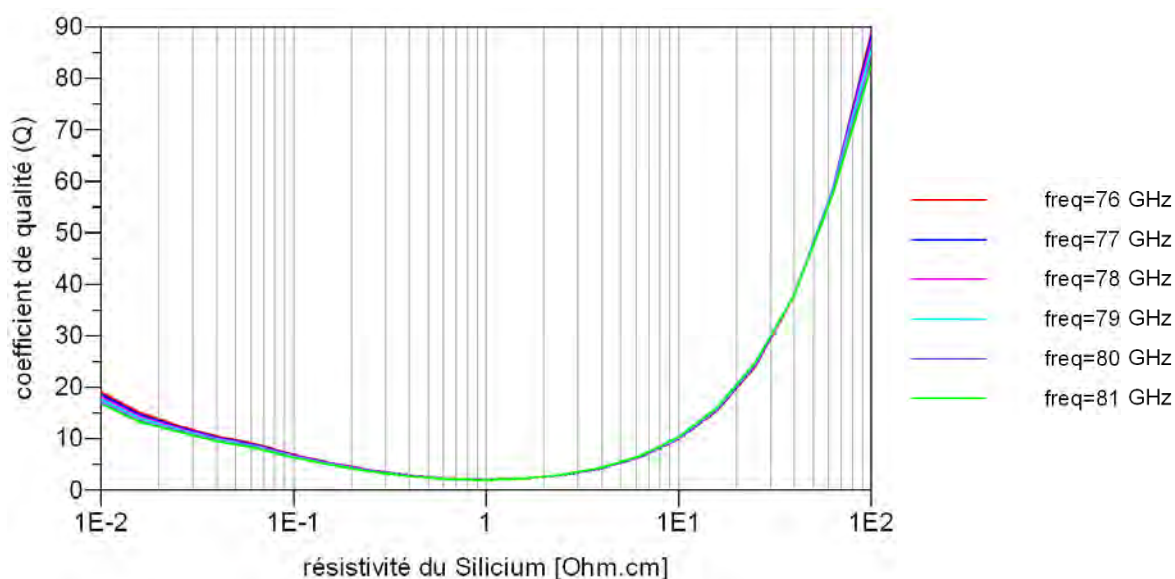


Figure 31 : variation du coefficient de qualité du plot en fonction de la résistivité du Silicium dans la bande [76-81] GHz

Si la capacité du plot devient quasiment constante à partir de 10 $\Omega\cdot\text{cm}$, les pertes du plot, quant à elles, s'améliorent significativement entre 10 et 1000 $\Omega\cdot\text{cm}$. La pire région d'opération se situe autour de 1 $\Omega\cdot\text{cm}$ où les pertes du plot avoisinent 2 dB en transmission pour un coefficient de qualité autour de 2. D'après les données disponibles sur la résistivité du Silicium non dopé de la technologie utilisée, les pertes en transmission du plot devraient être comprises entre 0.4 et 0.2 dB.

Ces résultats montrent que l'utilisation d'un plot pour lequel le substrat n'est pas masqué par le plan de masse aura un facteur de qualité moins bon qu'un plot classique. Cependant, la contrainte concernant les dimensions minimales du plot fait de cette configuration un bon candidat si la résistivité du Silicium est supérieure à 10 $\Omega\cdot\text{cm}$. Il permettra de réduire la capacité parallèle de 60% en gardant les mêmes dimensions. Les pertes additionnelles qu'il générera ne devraient pas dépasser 0.3 dB.

2.3.2.3 Lignes microrubans

Après l'étude des deux types de plots, nous abordons maintenant la modélisation des lignes de type microruban. La théorie derrière ce type de lignes est bien maîtrisée depuis son invention en 1952 et elle reste la base de toute conception en ondes millimétriques [2.8]. D'autres types de lignes comme les lignes coplanaires ou les lignes à ondes lentes ont largement été étudiés. Ces dernières ont brièvement été évoquées dans la partie précédente lors de l'étude des modes de propagations dans le silicium. Un travail complet sur les ondes lentes a été réalisé [2.9]. Il démontre qu'il est possible de concevoir des lignes coplanaires basées sur le principe d'ondes lentes en remplaçant le plan de masse inférieur par des barreaux flottants. Ces lignes permettent d'obtenir un meilleur coefficient de qualité et une miniaturisation supérieure, comparées aux lignes microrubans. Ces lignes coplanaires à ondes lentes sont effectivement intéressantes pour la conception aux fréquences millimétriques mais on privilégiera

lors de ces travaux des lignes microrubans puisque la miniaturisation ne fait pas partie des contraintes du projet.

En premier lieu, on s'assure lors de cette étude de la qualité de la modélisation EM des lignes de transmission en vérifiant que les modèles reflètent le comportement attendu pour ces lignes. L'étude porte principalement sur les paramètres S des lignes de transmission et les paramètres secondaires qui en découlent comme l'impédance caractéristique et la constante de propagation γ .

Les paramètres Z_0 et γ sont extraits à partir des paramètres ABCD, selon (10) et (11). Les paramètres ABCD sont eux-mêmes dérivés des paramètres S et les formules intervenant dans cette conversion sont bien connues.

$$Z_0 = \sqrt{\frac{B}{C}} \quad (10)$$

$$\gamma = \tanh^{-1} \sqrt{\frac{B \times C}{A \times D}} \quad (11)$$

Nous commençons l'étude avec la ligne la plus utilisée dans les circuits de cette technologie. C'est une ligne microruban dont le ruban signal de $8.5 \mu\text{m}$ de largeur est réalisé sur le dernier niveau métallique de la technologie (LM2), au-dessus d'un plan de masse implémenté sur le premier niveau métallique de la technologie (M1). La particularité de cette ligne est qu'elle présente une impédance caractéristique de 50Ω ainsi que les pertes les plus faibles possibles pour une ligne 50Ω dans cette technologie, puisque l'espacement entre le signal et le plan de masse est le plus grand possible.

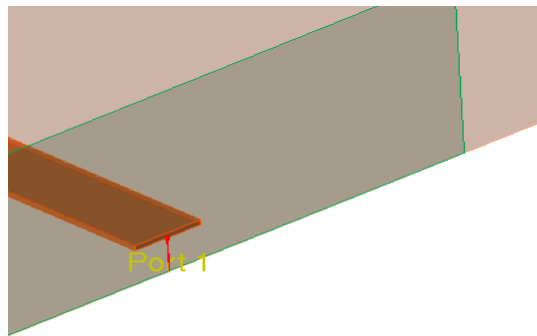
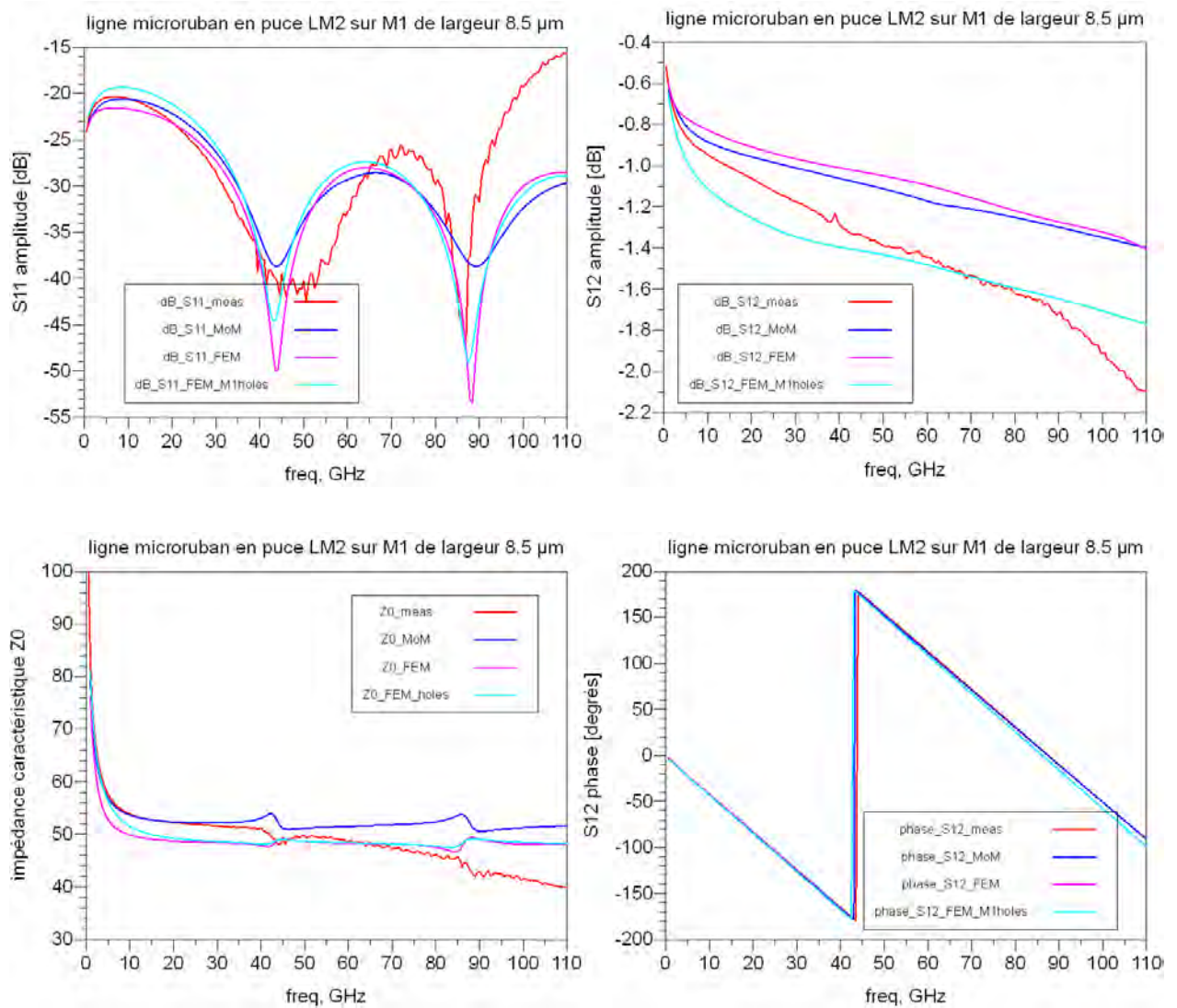


Figure 32 : vue 3D de la ligne microruban modélisée

Pour des questions de densité fixée par les règles de dessin, les plans de masse doivent intégrer des petites ouvertures sous forme de carrés, de l'ordre du micron, espacés d'une dizaine de microns. Ces motifs sont négligés lors des simulations EM compte tenu que leur simulation se révèle impossible à cause de la taille du maillage trop faible quand ils sont inclus. Cependant, la résistance additionnelle due

au plan métallique avec trous est quantifiée grâce à une division par 2 de la résistivité du métal M1. Les simulations correspondant à cette configuration sont nommées « M1holes ».

Les résultats de simulation ADS MoM et FEM de la ligne microruban que nous venons de décrire, en utilisant une longueur de ligne avoisinant 1 millimètre, sont reportés sur la Figure 33 et comparés aux mesures. La caractérisation des ces lignes est réalisée de la même manière que précédemment pour les plots, mais cette fois avec des paramètres S de quadripôle. Un calibrage TRL multi-ligne est appliqué à la place du calibrage SOLT parce qu'il s'agit ici de retirer également la contribution des plots.



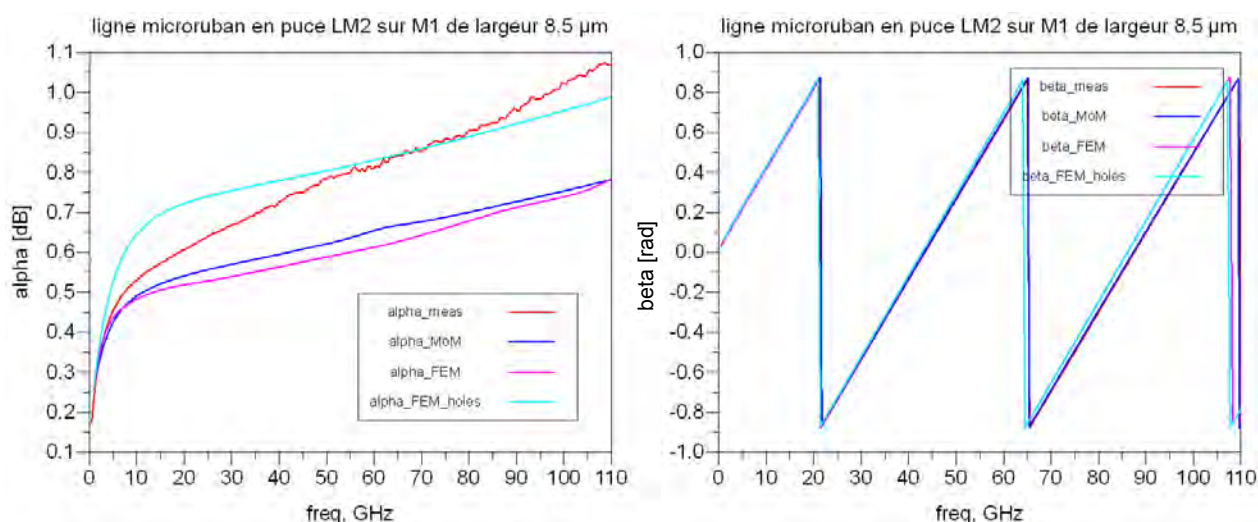


Figure 33 : résultats de simulations ADS MoM et FEM et de mesures d'une ligne microruban 50 Ω de longueur 1mm

Les trois principales informations qui ressortent des simulations sont les suivantes :

- Les deux simulateurs MoM et FEM donnent des résultats très proches. Il faut noter que les résultats initiaux en FEM étaient assez différents des simulations MoM et des investigations au niveau des deux simulateurs ont permis de découvrir la source de cette différence. Malgré la simplicité de la structure, FEM n'arrive pas à une très bonne convergence même après l'atteinte du nombre maximum d'affinages défini. Il a donc fallu étendre ce nombre à 50 au lieu de 20 et augmenter le nombre minimal de cas de convergence successifs à 5 au lieu de 2. Toutes les simulations FEM ultérieures utiliseront ces dernières options.
- Les pertes de 0.8 dB/mm à 80 GHz pour la ligne microruban sont comparables aux performances des technologies BiCMOS similaires et notamment la technologie 0,13 μm de ST Microelectronics [2.10]. Cependant, ces valeurs risquent d'être légèrement optimistes étant donné que les ouvertures sur le plan de masse ont été négligées. Ces ouvertures augmentent la résistivité de M1 et peuvent impacter les pertes de transmission.
- L'impédance caractéristique simulée est très proche de 50 Ω pour la ligne microruban, comme attendu, et les légères différences entre les deux simulateurs proviennent des écarts de définition de maillage ou de ports, qui sont par construction différents dans les deux simulateurs. La vitesse de propagation donnée par les deux simulateurs est comparable.

Par la suite, nous étudions une autre ligne utilisant les mêmes couches métalliques mais plus large et donc d'impédance plus faible. Le but est de confronter les modèles aux mesures dans des conditions d'adaptation différentes de 50 Ω . La Figure 34 reporte les résultats obtenus pour une ligne d'une largeur de 15 μm .

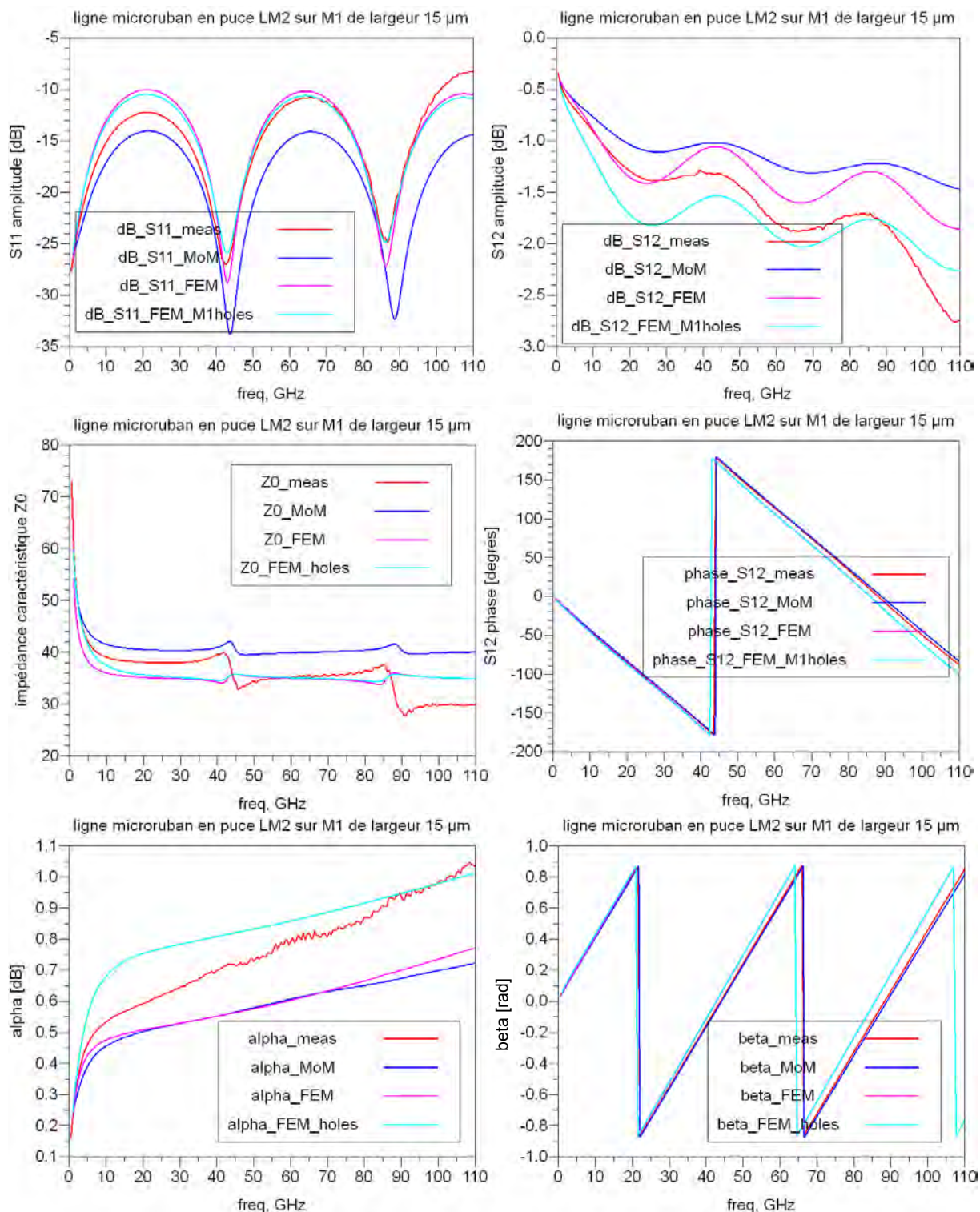


Figure 34 : résultats de simulation et de mesures d'une ligne d'une largeur de 15 μm

Le modèle de la ligne microruban de largeur 15 μm présente une impédance caractéristique entre 35 et 40 Ω . Les pertes sont légèrement inférieures par rapport à la ligne 50 Ω ce qui est conforme avec une augmentation de la largeur de la ligne qui réduit la résistance.

On voit une très bonne correspondance entre les simulations et les mesures pour les deux lignes de transmission en termes d'impédance caractéristique et de vitesse de propagation ou de phase. Néanmoins, les pertes données par la simulation se révèlent inférieures. Les pertes additionnelles

obtenues en mesures ne se retrouvent pas en simulations en raison vraisemblablement de la simplification du plan de masse métallique. La simulation « M1holes » a permis de retrouver des pertes très proches des pertes mesurées en hautes fréquences. Les pertes en basses fréquences sont pessimistes parce que la résistivité appliquée pour émuler les trous dans le plan M1 est indépendante de la fréquence alors que la contribution des trous aux pertes réelles dépend de la fréquence.

2.4 Modélisation de transformateurs intégrés pour applications millimétriques

Les transformateurs d'impédance sont des éléments passifs dotés de propriétés très intéressantes qui les rendent extrêmement utiles dans la conception des circuits. Parmi les propriétés d'un transformateur, on peut citer la transformation d'impédance qu'il peut apporter par des nombres de tours différents entre primaire et secondaire, sa taille compacte, la conversion de signaux « single-ended » en signaux différentiels et réciproquement, en plus de son intérêt pour la polarisation des transistors.

Les travaux décrits en [2.11] ont porté sur les transformateurs intégrés aux fréquences millimétriques. Nous les avons exploités pour la compréhension des contributeurs aux pertes d'un transformateur et le choix de la meilleure topologie par rapport aux spécifications à respecter lors de notre étude. Le transformateur qu'on cherche à concevoir doit s'accompagner du minimum de pertes possible et doit réaliser la conversion des signaux du mode « single-ended » vers le mode différentiel. De plus, nous l'utiliserons dans la gestion de la polarisation des transistors du circuit. Par conséquent, après l'étude de quelques transformateurs qui paraissaient pertinents, un seul transformateur a été retenu et on se focalisera uniquement sur celui-ci durant ces travaux. A ce stade, il ne s'agit pas d'obtenir le transformateur le mieux optimisé possible pour un circuit donné, mais de disposer d'un transformateur générique qui soit bien caractérisé, à très faibles pertes et qu'on puisse utiliser, par exemple, entre des étages d'amplification ou aux entrées de mélangeurs et sorties d'amplificateurs de puissance. Le transformateur choisi est un transformateur octogonal à un tour par enroulement et disposant d'un ruban central du côté différentiel (Figure 35). Son diamètre et la largeur de ses lignes ont été basés sur une optimisation des pertes.

2.4.1 Modélisation électromagnétique et principaux paramètres du transformateur

Le modèle EM du transformateur décrit précédemment est conçu à l'aide de l'outil ADS à partir des deux dernières couches métalliques de la technologie, LM et LM2, toutes les deux épaisses. Le transformateur n'intègre pas d'écran entre les lignes et le silicium. Néanmoins, des couches spécifiques sont rajoutées au niveau du dessin du masque pour disposer d'un silicium le moins dopé possible. Pour garder le maximum de flexibilité, la conception du transformateur est basée sur 5 accès afin de pouvoir

l'utiliser pour une large gamme d'applications. Dans certains cas, il est judicieux de remplacer la connexion directe vers la masse par un autre type de connexion, comme par exemple par une capacité de découplage et une ligne pour une amenée de polarisation. Le transformateur à 5 accès conçu (Figure 35) le permet.

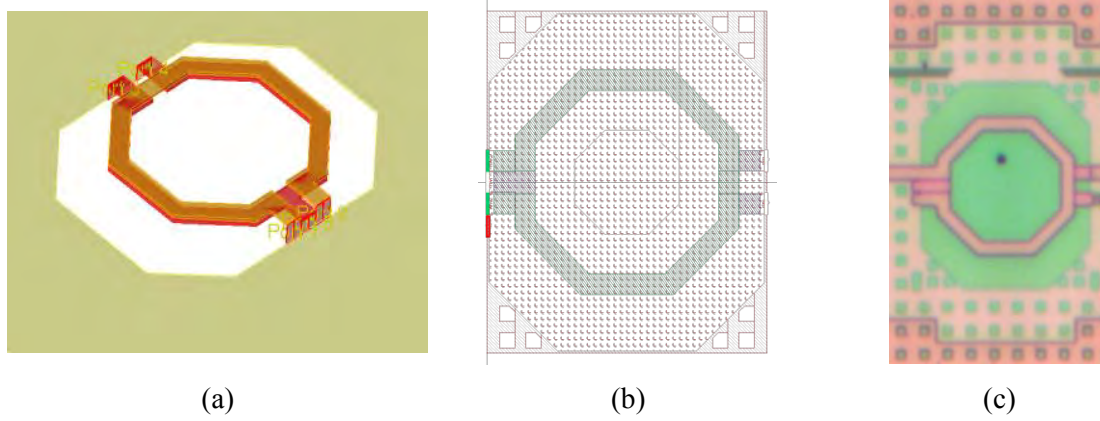


Figure 35 : modèle 3D (a), dessin de masque (b) et photo (c) du transformateur d'impédance conçu

Pour caractériser un transformateur, on peut considérer le schéma du quadripole simplifié ci-dessous, avec L_p et L_s les inductances respectives des lignes primaires (LM2) et secondaires (LM) et M l'inductance mutuelle :

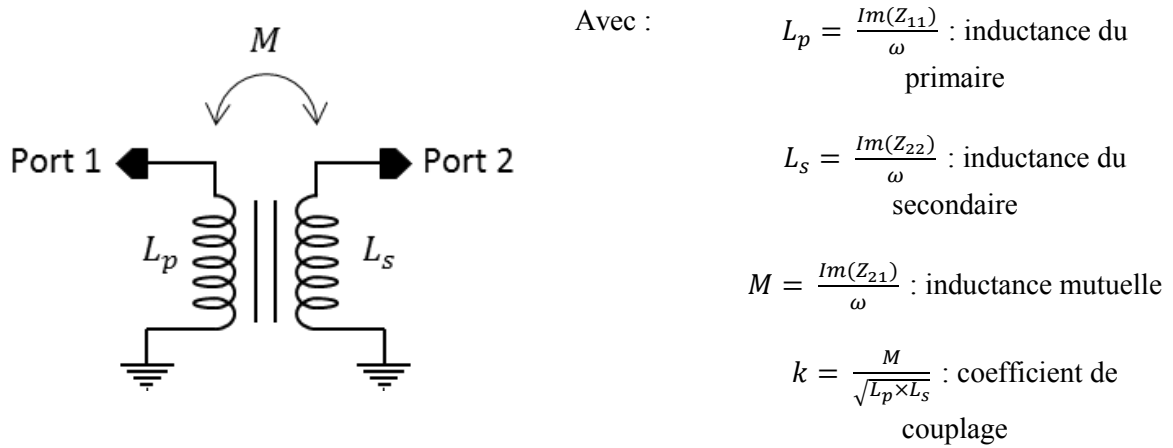


Figure 36 : schéma simplifié et principaux paramètres d'un transformateur

En plus de ces paramètres caractérisant les inductances en jeu, on définit des paramètres liés aux pertes dans le transformateur qui sont le coefficient de qualité des inductances primaire (Q_p) et secondaire (Q_s) ainsi que les pertes d'insertion minimales (IL_m).

$$Q_p = \frac{Im(Z_{11})}{Re(Z_{11})} \quad (12)$$

$$Q_s = \frac{Im(Z_{22})}{Re(Z_{22})} \quad (13)$$

$$IL_m = \frac{1}{1 + 2 \times (x - \sqrt{x^2 + x})} \quad (14)$$

Avec :

$$x = \frac{Re(Z_{11}) \times Re(Z_{22}) - [Re(Z_{12})]^2}{[Re(Z_{12})]^2 + [Im(Z_{12})]^2}$$

Les sources des pertes dans ce transformateur sont les mêmes que celles des lignes de transmissions pour lesquelles aucun masquage du Silicium n'est réalisé (cf. II.3.b).

Ce transformateur inclut également la fonction balun. En effet, il est en mesure de convertir un signal asymétrique (« single-ended ») en signal différentiel. On s'intéresse donc également au déphasage entre les deux voies du signal différentiel, qui doit par construction être très proche de 180°.

2.4.2 Validation expérimentale

La caractérisation expérimentale de ce transformateur devrait idéalement être réalisée à l'aide d'un VNA 110 GHz 4-ports. Cependant, ce matériel est pour l'instant assez rare et n'était pas présent dans les laboratoires impliqués dans ces travaux. Par conséquent, il a été décidé de se baser sur des mesures en quadripôle réalisées par un VNA 110 GHz standard, et de concevoir deux transformateurs, l'un intégrant une résistance 50 Ω au niveau de la deuxième voie de l'accès différentiel (transformateur 1), l'autre au niveau de la première voie (transformateur 2).

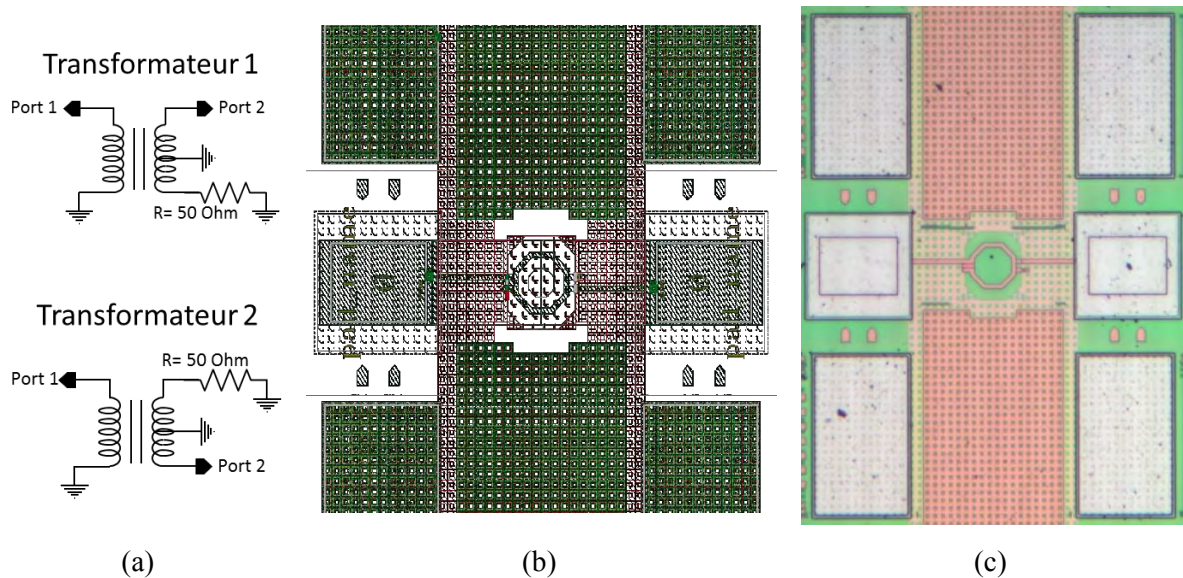


Figure 37 : schéma (a), dessin de masque (b) et photo (c) du transformateur incluant les plots

Ces deux configurations permettent de confronter les simulations aux mesures et d'extraire le déphasage des voies à partir de la mesure. Les résultats de simulations et de mesures sont comparés sur la Figure 38. Pour les mesures, un calibrage TRL multi-lignes est réalisé afin de s'affranchir de l'influence des plots et de disposer des mêmes plans de référence que pour la simulation. Les mesures correspondant

au transformateur 1 sont nommées « meas1 » et celles correspondant au transformateur 2 sont nommées « meas2 ».

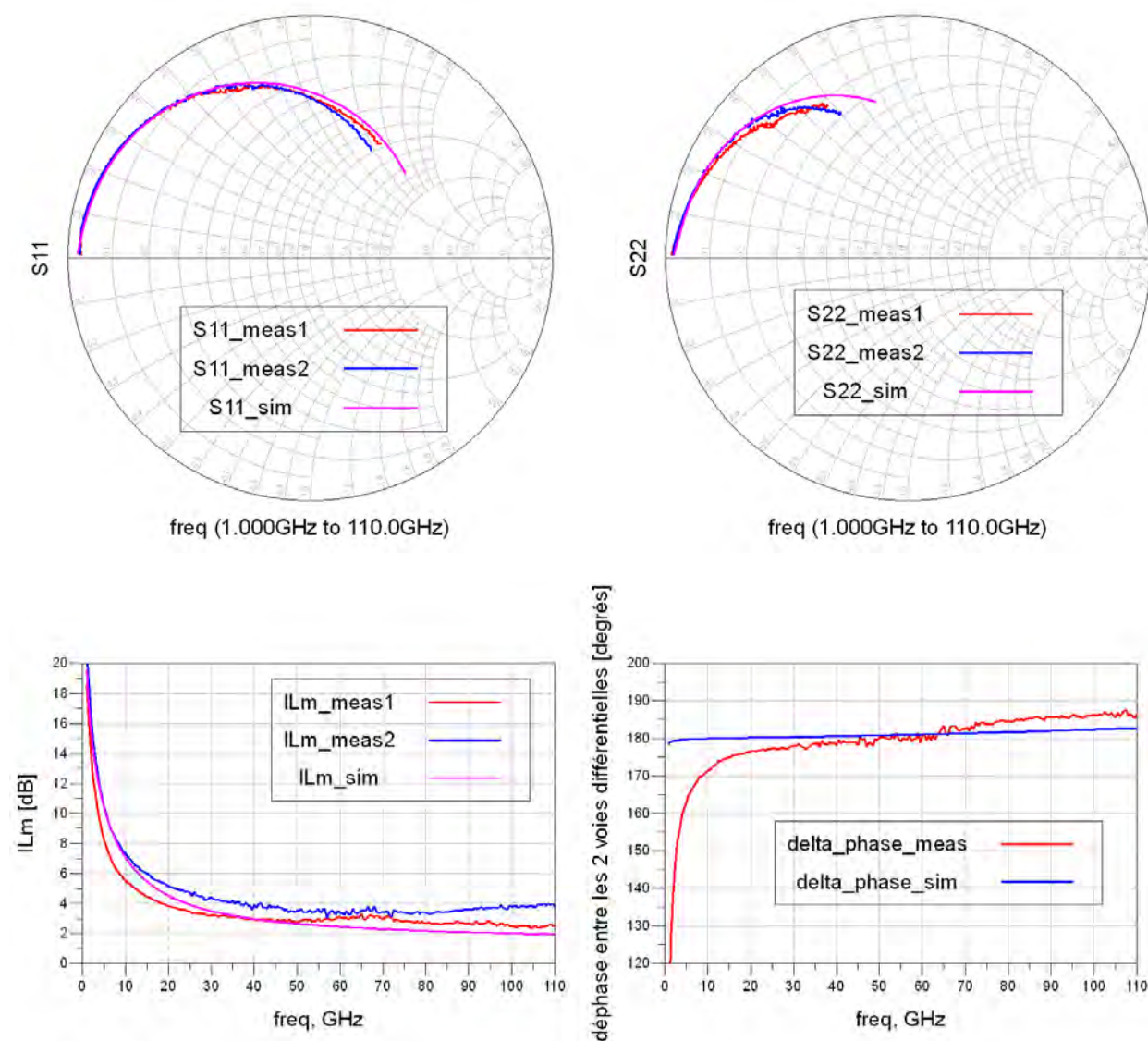


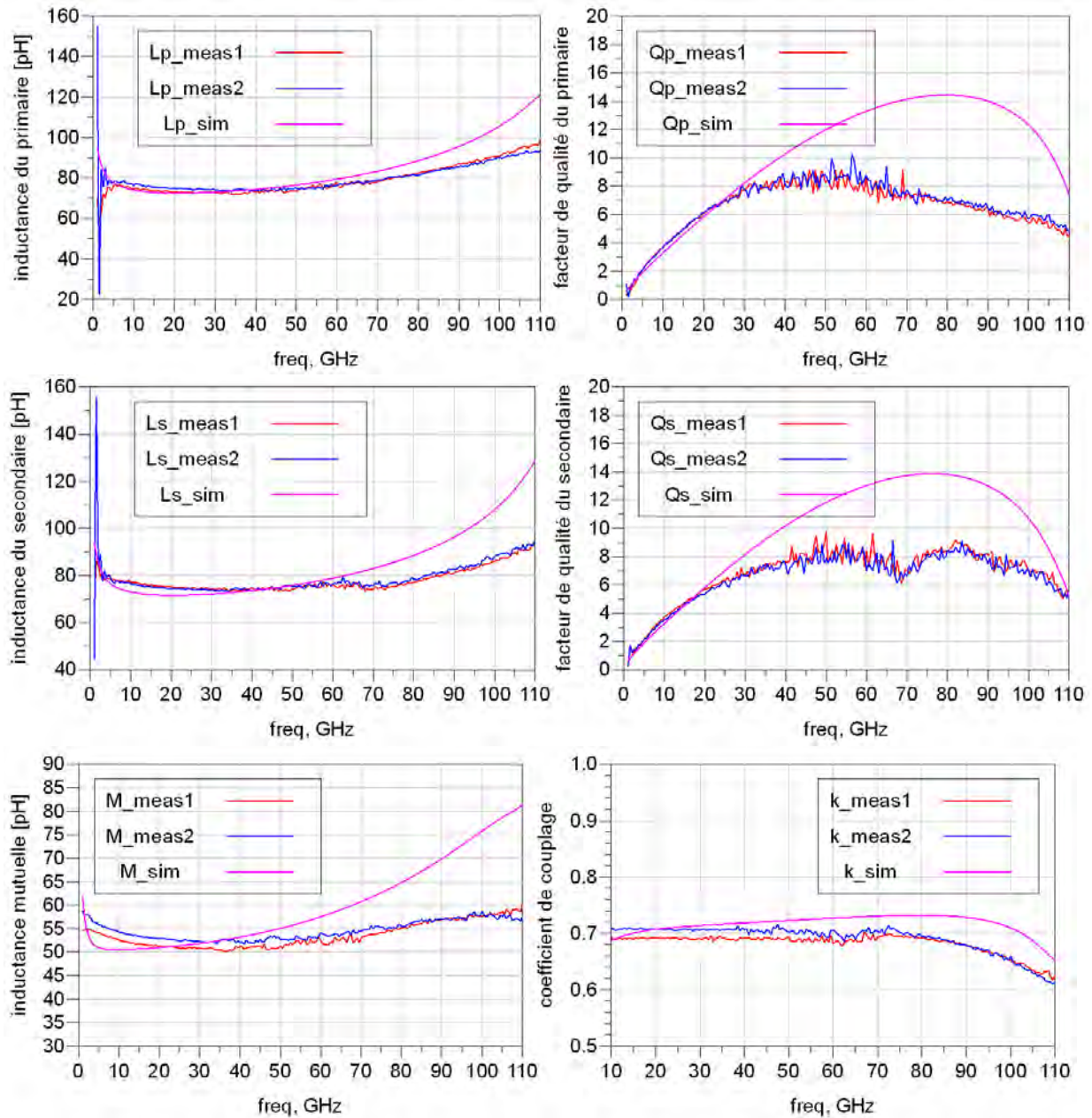
Figure 38 : résultats de simulation et de mesure du transformateur conçu

Le déphasage est très proche de 180° dès 30 GHz comme attendu par la simulation. La conversion de signaux asymétriques « Single Ended » en un signal différentiel est donc correctement réalisée. Les pertes mesurées sont légèrement supérieures à celles obtenues par la simulation, mais l'écart relevé reste très faible sur l'ensemble de la bande de fréquences et se situe au même niveau que la précision du calibrage et de la mesure. Les impédances des voies asymétrique et différentielle sont également très proches entre simulation et mesure.

Afin de compléter la caractérisation du transformateur et d'extraire les inductances des enroulements primaires et secondaires ainsi que le coefficient de qualité, un nouveau transformateur est intégré. Il reprend les mêmes dimensions que le transformateur de la Figure 37, mais le ruban central est supprimé et un port sur chacune des lignes est directement connecté à la masse. Cette configuration permet de

disposer de tours complets sur les 2 niveaux du transformateur, comme le montre la Figure 36, présentée précédemment.

Nous allons nous concentrer ici sur les paramètres qui n'ont pas pu être extraits de la caractérisation de la configuration hexapode du transformateur. La Figure 39 présente les résultats des simulations et des mesures de deux puces identiques de ce transformateur appelées « meas1 » et « meas2 ».



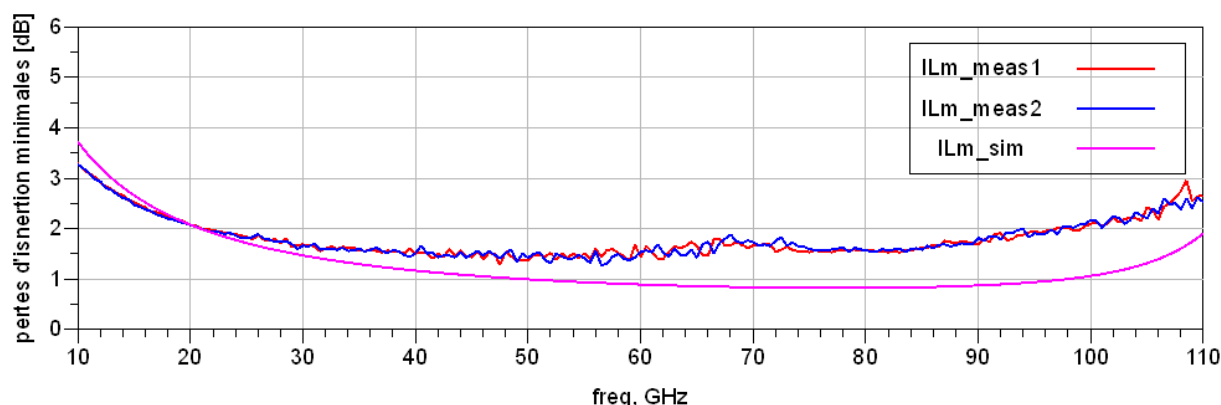


Figure 39 : résultats de simulation et de mesure du transformateur en quadripôle

Ces résultats montrent une bonne correspondance entre la simulation et la mesure, en particulier pour les inductances et le coefficient de couplage qui correspondent bien à la simulation jusqu'à 60 GHz, avec un léger décalage au-delà. Les pertes du transformateur sont plus faibles en simulation qu'en mesures comme déjà remarqué sur la modélisation du transformateur en hexapode. Cette différence impacte directement le facteur de qualité, inversement proportionnel aux pertes. Les investigations ont permis d'éloigner les pistes relatives à la structure du plan de masse, aux pertes diélectriques et dans le Silicium. Une dégradation importante dans le modèle de ces trois éléments cumulés n'aurait impacter que de 0.2 dB les pertes du transformateur, d'après les simulations de pires cas.

D'autres travaux menés au LAAS sur des inductances à très fort coefficient de qualité, ont trouvé l'origine d'un problème présentant beaucoup de similitudes avec celui-ci, même si les fréquences de travail et les ordres de grandeur des inductances sont différents. Ces travaux publiés récemment [2.12] lient la dégradation du facteur de qualité à la mesure elle-même et non à la simulation. En effet, une interaction a lieu entre le transformateur et les pointes de mesure résultant à une augmentation de la radiation. Cette interaction dégrade ainsi les performances des inductances en rajoutant des pertes importantes non prises en compte par la simulation. Le ratio entre le facteur de qualité mesuré et simulé dans les travaux de Bushueva est autour de 0.6. Ce ratio est le même dans les mesures réalisées ici comme on peut le voir sur le graphe de la Figure 40.

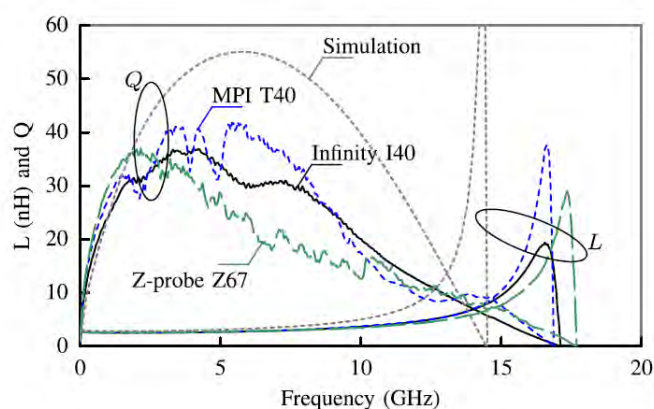


Figure 40 : coefficient de qualité simulé et mesuré en fonction du type de pointes [2.11]

Les contraintes associées aux travaux de cette thèse ne permettent pas de démontrer rigoureusement que cette explication s'applique ici et si elle permet de justifier la totalité de l'écart observé entre les simulations et les mesures. Néanmoins, il a été décidé de considérer le modèle EM du transformateur comme étant correct et de revenir éventuellement dessus si on s'aperçoit par la suite, lors de mesures expérimentales de blocs incluant ce transformateur, qu'il y a une différence importante entre les simulations et les mesures.

2.5 Conclusion

La modélisation électromagnétique a été au cœur de ce chapitre. Après une description de la méthodologie, de l'environnement, des propriétés et des caractéristiques des technologies devant être prises en considération dans ce travail de modélisation, quelques éléments passifs de base ont été modélisés. Les composants ayant fait l'objet de modélisations EM ont été les plots d'accès coplanaires, des lignes microrubans d'impédance caractéristique 50Ω et faible impédance, et enfin d'un transformateur. Les caractéristiques simulées pour ces éléments ont été confrontés à des mesures afin d'affiner les modèles et comprendre les éventuelles différences. Il ne s'agissait pas d'apporter des innovations majeures au niveau de la modélisation de structures telles les lignes de transmission ou les transformateurs, mais de se référer plutôt à des travaux jugés pertinents et à les adapter au contexte de ce projet. Ces structures ne peuvent être intégrées dans des circuits plus complexes qu'une fois des preuves tangibles apportées pour leur validité et leur robustesse. Chacune des structures étudiées a permis d'affiner des aspects de la modélisation EM et de comprendre plus finement les sources potentielles d'erreur permettant de corriger les modèles. Si les plots d'accès et les lignes microrubans sont des éléments très utilisés dans toutes les technologies et par tous les concepteurs, le transformateur quant à lui est une structure moins utilisée, spécialement dans l'industrie où on préfère privilégier les topologies à base de lignes microrubans et de capacités Métal-Isolant-Métal (MIM) pour réaliser les fonctions de conversion des signaux du mode asymétrique au mode différentiel, de même que pour réaliser les structures d'adaptation et/ou de polarisation. Ce transformateur entrera ensuite dans la conception d'un mélangeur pour le radar automobile 77 GHz, qui remplacera un mélangeur basé sur des éléments passifs plus élémentaires, comme nous le décrirons dans le dernier chapitre.

3 Conception et optimisation de boîtiers pour l'émetteur-récepteur radar automobile 77-GHz

3.1 Introduction

La mise en boîtier des circuits intégrés fait partie des grands défis des systèmes fonctionnant en ondes millimétriques et en particulier les radars automobiles 77 GHz. Les contraintes de qualification automobile et de production en masse viennent s'ajouter aux contraintes liées aux fréquences de fonctionnement très élevées. Par conséquent, le choix du boîtier optimal est très restreint et impose un compromis délicat entre performances et coût. Ces contraintes s'appliquent également au choix des matériaux et de la classe de fabrication du PCB sur lequel est reporté le boîtier.

Tous les fournisseurs de circuits intégrés radar se sont tournés vers des solutions de mise en boîtier au niveau de la plaquette (WLP pour Wafer Level Packaging) de type « fan-out » quand il a fallu remplacer le câblage par microfils. Le coût relativement faible de cette technologie et les pertes associées à la transition du PCB vers la puce confèrent à ce boîtier un bon rapport performances/coût. Cependant, les limitations de ce boîtier ont poussé les industriels à s'intéresser à des boîtiers « fan-in », en particulier le WLCSP. En effet, les limitations du boîtier « fan-out » sont principalement liées à la dissipation thermique permise par cette technologie, à aux dimensions du boîtier bien supérieures à celles de la puce, ainsi qu'à son coût plus élevé qu'un boîtier « fan-in » standard. Par conséquent, le principal enjeu des travaux présentés ici est de réussir à concevoir, modéliser et valider expérimentalement un boîtier WLCSP qui conviendrait à la transmission de signaux radar à 77 GHz. Afin d'y parvenir, l'étude commencera par la caractérisation du PCB, inséparable du boîtier, puis nous mènerons la caractérisation du boîtier « fan-out » RCP, pour nous intéresser finalement au boîtier « fan-in » WLCSP. En parallèle de la conception et de la modélisation des interfaces boîtier, l'étude se penchera également sur le développement de méthodologies permettant de caractériser rigoureusement la transition. Ce dernier travail nous permettra de confirmer la précision des modèles électromagnétiques conçus.

3.2 Modélisation électromagnétique de composants passifs sur circuit imprimé

La majeure partie des composants électroniques du système radar est reportée sur un circuit imprimé (PCB). De manière générale, les contraintes imposées pour le PCB par un boîtier de type WLP impactent

la conception PCB puisqu'elles restreignent, entre autres, le choix du substrat et des largeurs de lignes [3.1]. Pour un module radar automobile, le PCB est également utilisé pour le routage des signaux de fréquences millimétriques et pour la réalisation des antennes patch d'émission et de réception. Le choix de commencer par la caractérisation du PCB avant d'aborder l'interface boîtier est lié à la structure de la transition du PCB vers la puce. En effet, on ne peut pas séparer la modélisation du boîtier de celle du PCB étant donné que les couplages entre les deux sont importants et qu'on ne peut connaître le mode de propagation de l'onde parcourant le boîtier qu'une fois le signal relevé sur les lignes de transmission microruban du PCB. De plus, l'optimisation du PCB a autant d'intérêt que l'optimisation du boîtier lui-même comme on le verra plus loin.

3.2.1 Description et spécifications des niveaux métalliques du PCB

Plusieurs couches métalliques sont utilisées pour router le nombre important de signaux au niveau du PCB. Néanmoins, pour ce qui concerne les signaux millimétriques, on ne s'intéresse qu'à la couche supérieure. Dans ces conditions, un plan de masse métallique plein peut séparer ce niveau des autres couches (Figure 41). Le matériau diélectrique de ce dernier niveau doit répondre à des spécifications électriques bien définies, que ce soit pour les pertes diélectriques ou la variation des caractéristiques électriques en fonction de la température. Les couches métalliques utilisées sont des couches de cuivre d'une épaisseur de l'ordre de la dizaine de microns pour une épaisseur du diélectrique RF de l'ordre de la centaine de microns.

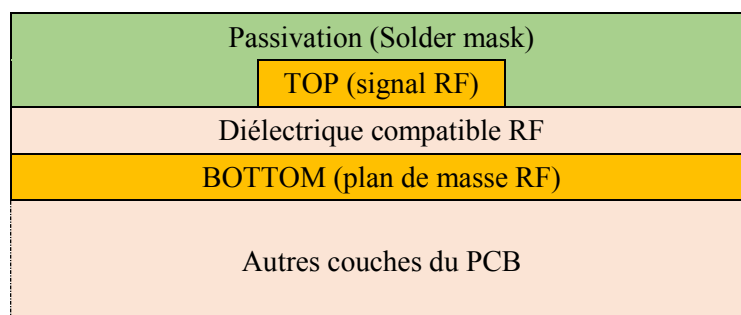


Figure 41 : niveaux métalliques (stack-up) du PCB utilisé dans ces travaux

La passivation est la couche verte qu'on aperçoit sur la grande majorité des PCB. Elle permet de protéger ce dernier mais ses propriétés électriques ne sont pas maîtrisées en hautes fréquences. Par conséquent, des masques négatifs sont appliqués au niveau des pistes transmettant des signaux de fréquences millimétriques pour éviter que le vernis dégrade les performances. Un autre élément important entrant dans la conception des structures passives microondes intégrées sur le PCB est le trou métallisé (via). Les vias ne sont pas utilisés explicitement pour le routage du signal de fréquence millimétrique, mais ils connectent les deux derniers niveaux métalliques pour les contacts à la masse. La gestion des vias se révèle être une source d'amélioration importante pour les performances RF de la transition. La Figure

42 présente une coupe transversale d'une section de PCB où on peut distinguer des vias RF s'arrêtant au niveau du plan métallique de masse:



Figure 42 : coupe transversale du PCB au niveau des vias

La description des couches du PCB et les coupes réalisées permettent d'acquérir de nombreuses informations utiles pour la conception d'éléments passifs sur PCB pour les fréquences millimétriques. Ces informations seront complétées par la validation expérimentale des performances des structures modélisées. Ces travaux permettront ainsi de caractériser rigoureusement le PCB utilisé dans la constitution de la transition boîtier. Le résonateur en anneau est particulièrement intéressant pour cette caractérisation comme on le verra par la suite.

3.2.2 Modélisation électromagnétique et validation expérimentale

Différentes structures sur PCB sont étudiées dans cette partie pour caractériser le plus rigoureusement possible ce PCB. Il s'agit dans un premier temps de s'assurer de la validité de ses propriétés électriques, puis de confirmer la qualité des modèles avant l'introduction du modèle du boîtier. Tout le savoir-faire acquis à travers la modélisation des lignes et des plots sur puce sera exploité et appliqué au PCB. Nous commençons par la modélisation de résonateurs en anneau afin de valider la précision de la simulation et d'extraire la permittivité effective grâce à la mesure. Nous étudions aussi par la suite des lignes microrubans $50\ \Omega$ largement utilisées pour connecter les entrée/sorties du boîtier aux antennes et pour interconnecter, combiner ou diviser des signaux RF au niveau du PCB.

3.2.2.1 Résonateurs en anneau

Les résonateurs en anneau font partie des structures préférées des concepteurs PCB pour la caractérisation des constantes diélectriques des matériaux [3.2]- [3.3]. Ces résonateurs sur PCB ont également d'autres utilisations plus exotiques comme la détection de fissures dans les métaux [3.4] ou l'évaluation de la qualité d'aliments comme la viande grâce à l'évolution de la constante diélectrique de cette dernière [3.5].

On utilise ici un résonateur à base d'un seul anneau et dans lequel les lignes de transmission microrubans sont finies par des arcs, comme il peut être relevé sur la Figure 43.

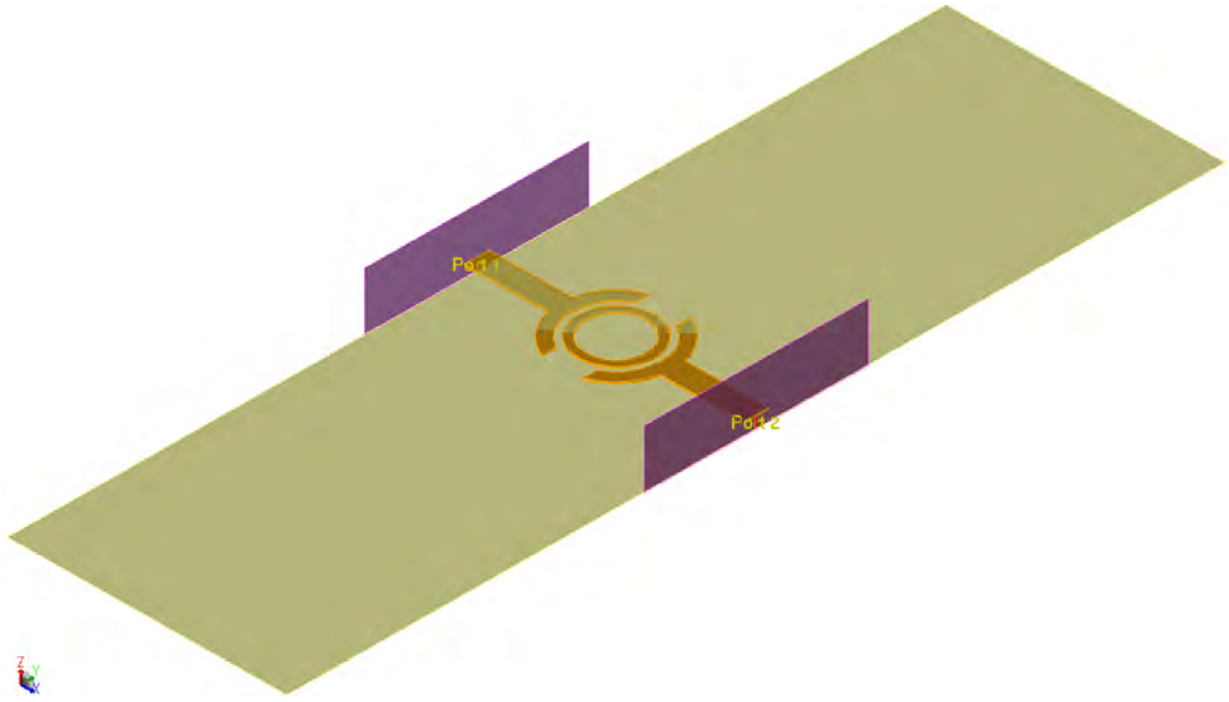


Figure 43 : vue 3D du modèle EM du résonateur en anneau

Une propriété très intéressante du résonateur en anneau est la relation qui lie son rayon à sa fréquence de résonance et à la permittivité effective du diélectrique. L'équation (15) permet de caractériser facilement un matériau diélectrique à partir de la mesure de la fréquence de résonance puisque le rayon de l'anneau est imposé par la conception. La résonance se produit quand la circonférence de l'anneau est égale à un multiple de la longueur d'onde guidée (λ_g) [3.6].

$$2 \pi r = n \lambda_g \quad \text{avec } n = 1, 2, 3 \dots \quad (15)$$

Où r est le rayon moyen de l'anneau et n le numéro de la fréquence harmonique.

En outre, λ_g exprimé en fonction de la fréquence de résonance (f_0) permet d'écrire la relation (15) qui lie cette fréquence au rayon de l'anneau :

$$2 \pi r = \frac{n c}{f_0 \sqrt{\epsilon_{eff}}} \quad (16)$$

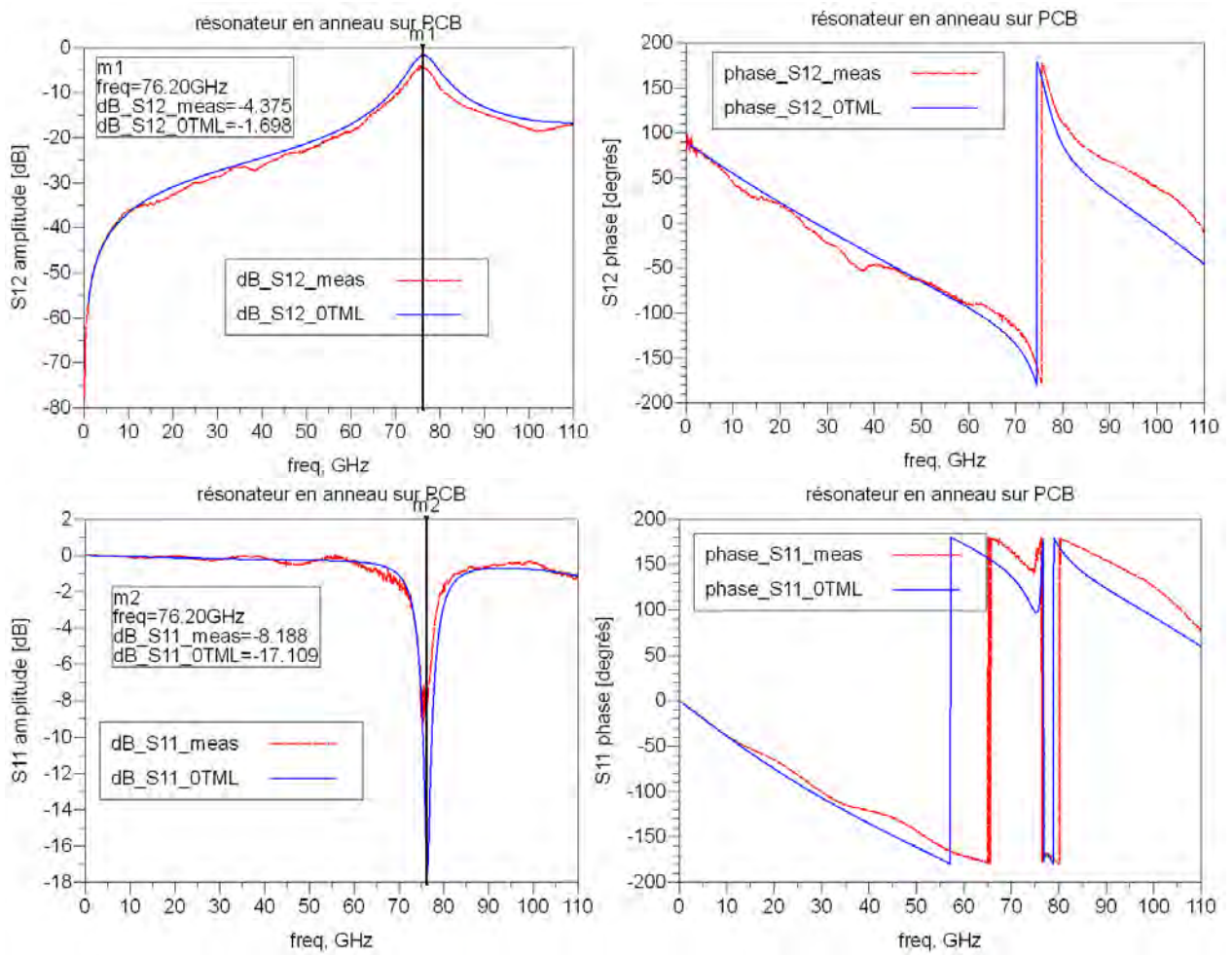
Où c est la vitesse de la lumière et ϵ_{eff} la constante diélectrique effective.

Dans notre étude, on s'intéresse principalement à la permittivité effective pour des fréquences autour de 76 GHz. Le résonateur en anneau est donc conçu de sorte à résonner à cette fréquence. La valeur de la permittivité relative (ϵ_r) du substrat PCB est calculée à partir de la permittivité effective (ϵ_{eff}). Les deux grandeurs sont liées par la formule (17) suivante, validée pour les structures microrubans dont la largeur des pistes est supérieure ou égale à la hauteur du substrat [3.7] :

$$\varepsilon_{eff} = \frac{\varepsilon_r + 1}{2} + \left(\frac{\varepsilon_r - 1}{2} \right) \left[\frac{1}{\sqrt{1 + \frac{12h}{w}}} \right] \quad (17)$$

Où h est la hauteur du substrat et w la largeur de la piste métallique véhiculant le signal.

Sur la Figure 44, nous reportons les résultats de simulations et de mesures obtenus pour ce résonateur en anneau en considérant pour la simulation un port guide d'onde défini au plus près du résonateur, nommé « 0TML ». Au niveau de la mesure, un calibrage TRL est réalisé pour ramener le plan de référence aussi au plus près du résonateur.



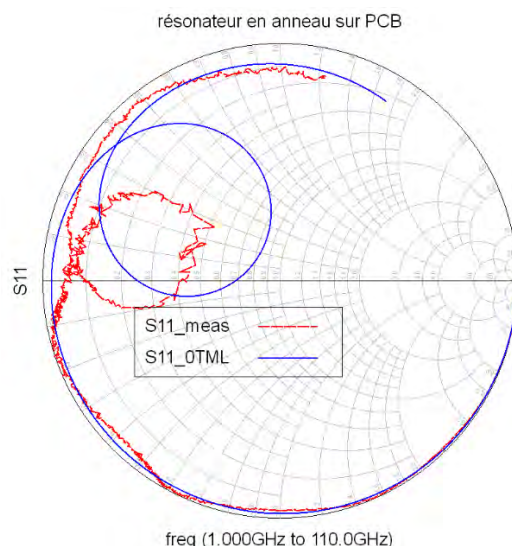


Figure 44 : comparaison des résultats de simulations et de mesures du résonateur en anneau sur PCB

Les mesures et les simulations montrent la même fréquence de résonance autour de 76.2 GHz. Cependant, on observe une différence dans la réponse, particulièrement pour les courbes tracées sur l'abaque de Smith. De plus, les pertes à la fréquence de résonance sont optimistes au niveau de la simulation. Cette constatation laisse penser que le rayonnement de l'anneau est sous-estimé, ce qui lie directement le problème aux conditions de simulations, notamment aux conditions aux limites.

En effet, il s'est avéré après des investigations que l'utilisation du port en guide d'onde dans ces conditions est la source du problème. Ce port, en violet sur la figure 43, délimite énormément le volume de simulation sur deux des six faces de la boîte. Afin de contourner cette limitation, deux approches ont été considérées : la première consiste à utiliser un port non calibré pour avoir un volume important sur tous les axes, nommé « noTML », alors que la deuxième se base sur le rallongement de l'accès de 500 μm de chaque côté en gardant un port guide d'onde et sur l'épluchage de cet accès par calcul. Cela permet de conserver une excitation en guide d'onde tout en élargissant de 500 μm la boîte de simulation dans les deux directions contraintes par les ports. La Figure 45 montre les résultats obtenus. On constate que les deux solutions proposées présentent des niveaux de pertes plus proches de la mesure et que la dernière solution est la plus proche de la mesure au niveau de la phase. Cela est tout à fait conforme à ce qui est attendu étant donné que la solution sans calibrage ne permet pas l'établissement correct d'un mode de propagation quasi-TEM avant l'excitation du résonateur, perturbant ainsi la phase.

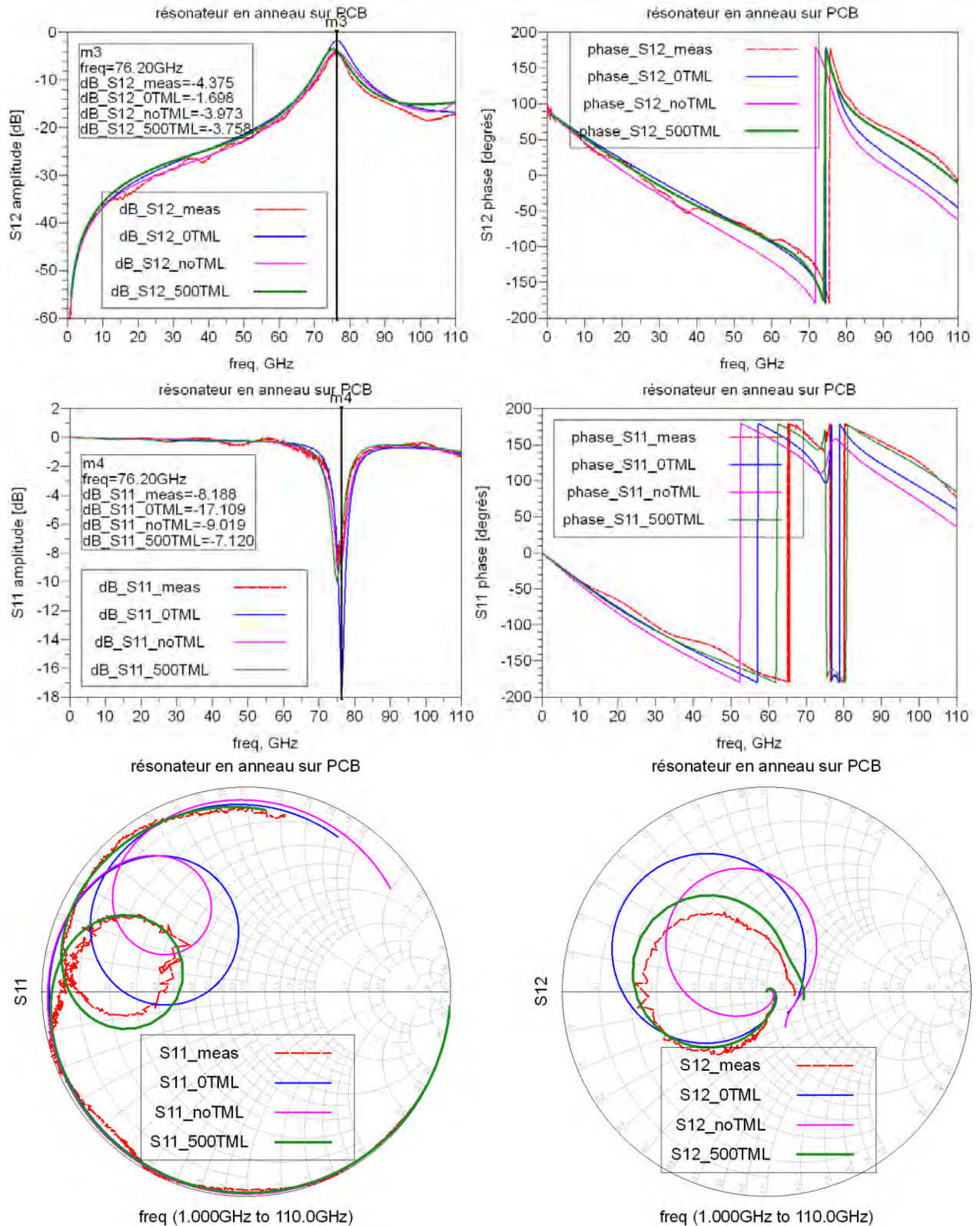


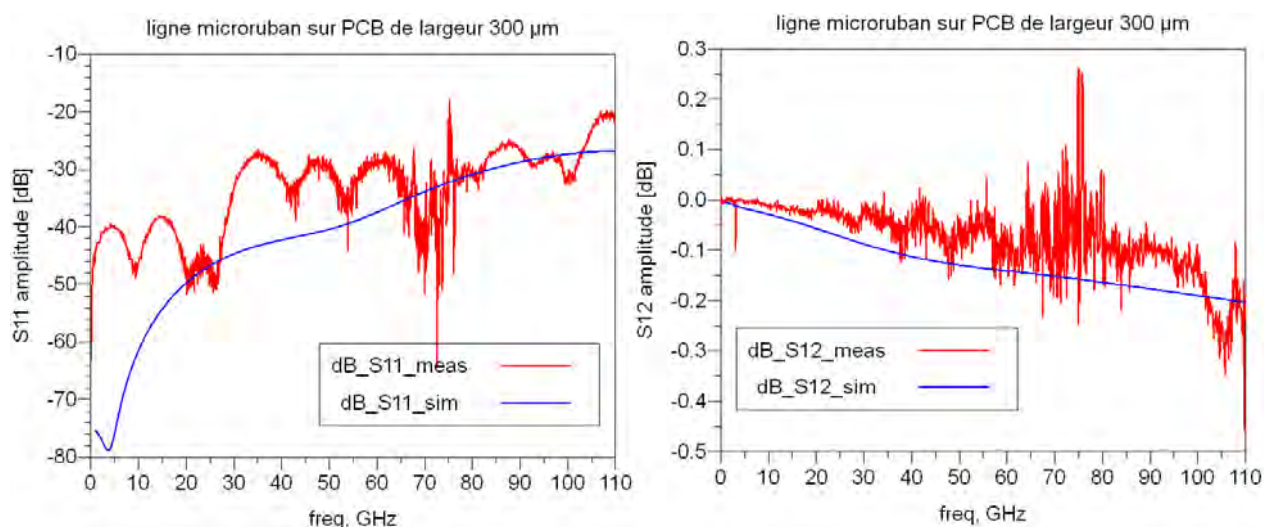
Figure 45 : comparaison des résultats de simulations et de mesures du résonateur en anneau pour différentes définitions de ports

Nous pouvons conclure que cette étude confirme la valeur de la constante diélectrique donnée pour le PCB. Par ailleurs, cette étude valide le modèle électromagnétique du résonateur en anneau. Cela permet de passer aux étapes suivantes avec certaines assurances.

3.2.2.2 Lignes microrubans

À la suite de la modélisation des résonateurs en anneau présentée précédemment, on s'intéresse aux lignes microrubans sur PCB, et en particulier aux lignes microrubans 50 Ω , largement présentes dans les circuits imprimés du module radar automobile. Cette étude est assez similaire à celle réalisée au niveau du circuit intégré et présentée au chapitre précédent. Les principales différences se situent pour les couches de métallisation et pour les ordres de grandeur des dimensions en jeu. On se base cependant sur les mêmes paramètres pour caractériser les lignes et confronter les simulations aux mesures.

Pour cette technologie PCB, les lignes d'impédance 50 Ω ont une largeur de 300 μm et leurs pertes ne dépassent pas 0.2 dB par millimètre de longueur. On vérifie cela grâce à la modélisation d'une ligne sur PCB et à la comparaison des performances obtenues à celles issues des mesures. Sur la Figure 46, les mesures montrent une adaptation quasiment parfaite vu la méthode d'extraction des performances de cette ligne. En effet, on se base sur un calibrage TRL utilisant des retards de la même largeur que la ligne sous test. Par conséquent, l'impédance présentée par le retard est celle utilisée pour normaliser l'impédance de la ligne sous test. Comme on peut le voir, cette mesure reste cependant pertinente parce qu'elle permet bien de vérifier le déphasage et les pertes apportées par une ligne 50 Ω longue d'un millimètre.



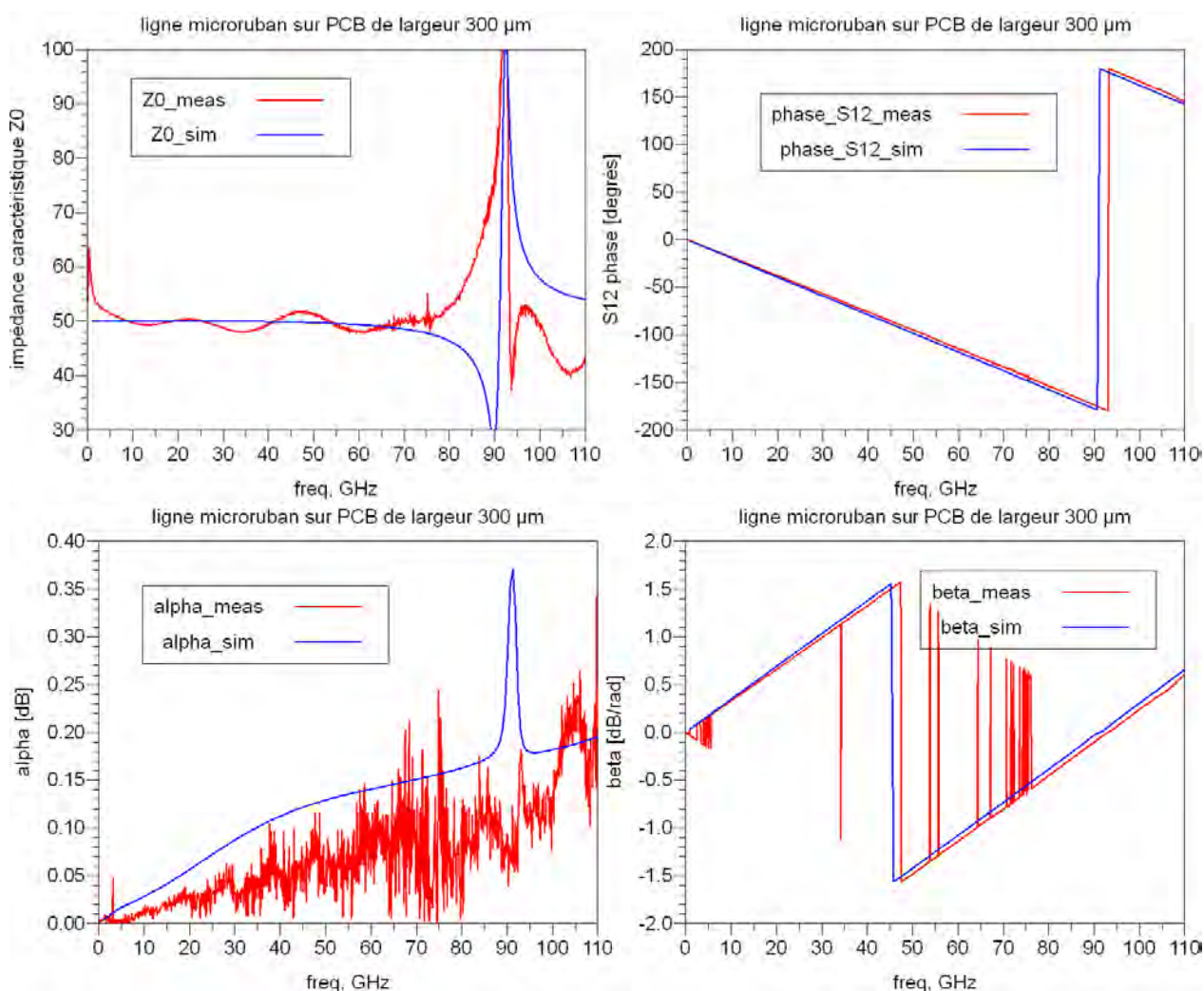


Figure 46 : comparaison des simulations et mesures d'une ligne microruban 50 Ω sur PCB

Cette comparaison des performances simulées et mesurées d'une ligne 50 Ω permet de confirmer la précision du modèle et en particulier la vitesse de propagation de l'onde. Le très léger écart observé sur la phase et la constante de phase (beta) est équivalent à un écart de 20 μm entre la simulation et la mesure d'une ligne d'un millimètre de longueur. Cet écart peut être expliqué par les incertitudes de mesure liées au calibrage et à la répétabilité du poser de pointes.

Par ailleurs, on remarque que l'impédance caractéristique de la ligne est exactement 50 Ω , que ce soit pour la simulation ou la mesure.

Afin de compléter cette analyse, nous nous intéressons à la modélisation de la ligne utilisée comme un « thru » dans le calibrage TRL, en incluant les plots sur PCB. Cet exercice présente une double difficulté. D'une part, la définition des ports sera approximative, comme souvent constaté avec les pointes GSG. Cet aspect a déjà été étudié lors de l'analyse des pots au niveau du circuit intégré, et présenté au chapitre précédent. D'autre part, comme évoqué plus haut, il est difficile d'effectuer un posé de pointes précis sur la couche d'or présente sur le PCB étant donné que les pointes ne laissent pas de traces sur ce

matériau. Il est par conséquent normal de s'attendre à de légères différences, ne remettant pas en cause l'utilité de cette comparaison. La structure modélisée est illustrée par la Figure 47.

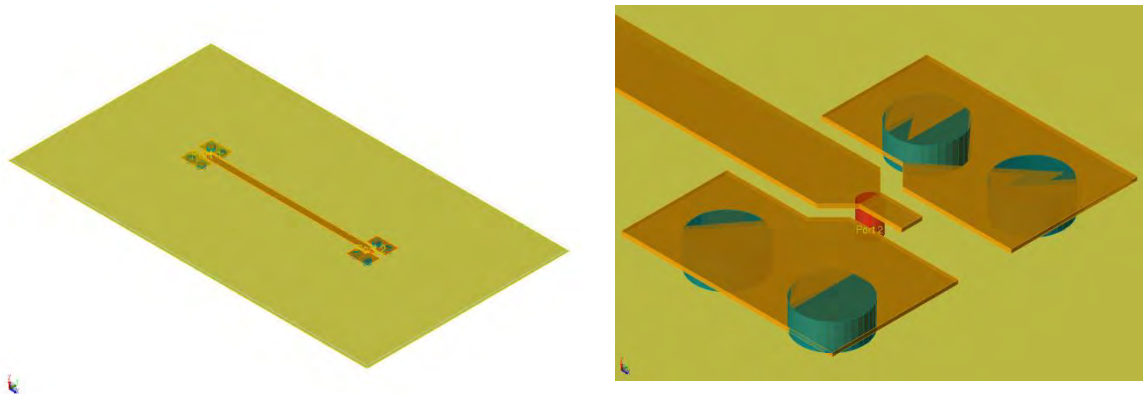
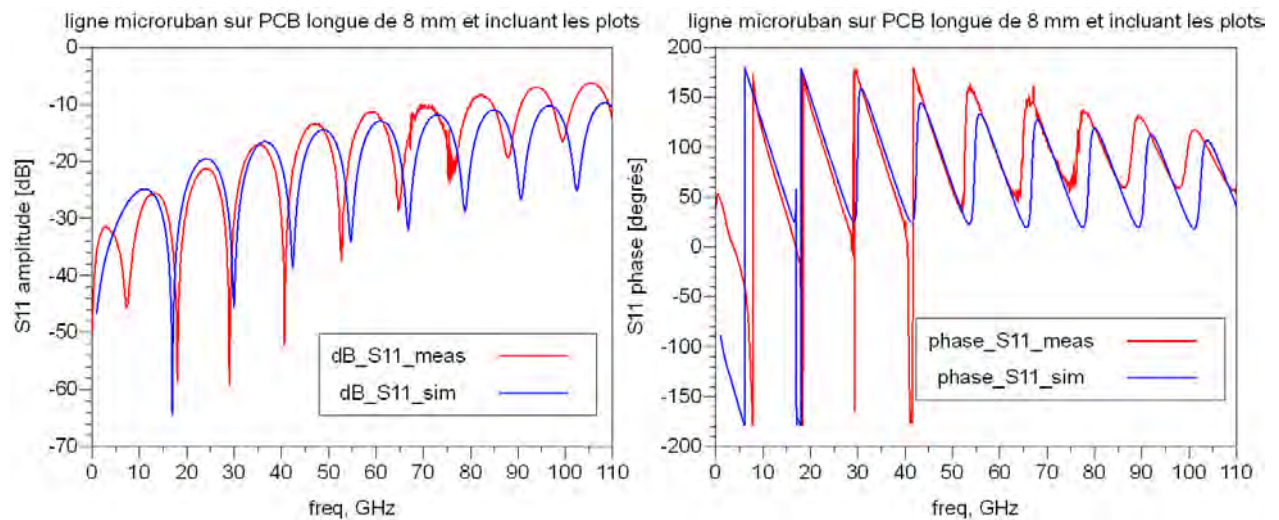


Figure 47 : modèle de la ligne PCB incluant les plots

La Figure 47 permet de voir le port utilisé pour l'excitation de la ligne par les pointes GSG. Après le test de plusieurs types de port, le port représenté ici s'est révélé le plus représentatif de la mesure, même s'il est référencé à la masse sur le plan métallique inférieur et non par rapport aux plots de masse situés de part et d'autre du plot signal.

La Figure 48 présente les simulations et les mesures des paramètres-S de cette ligne.



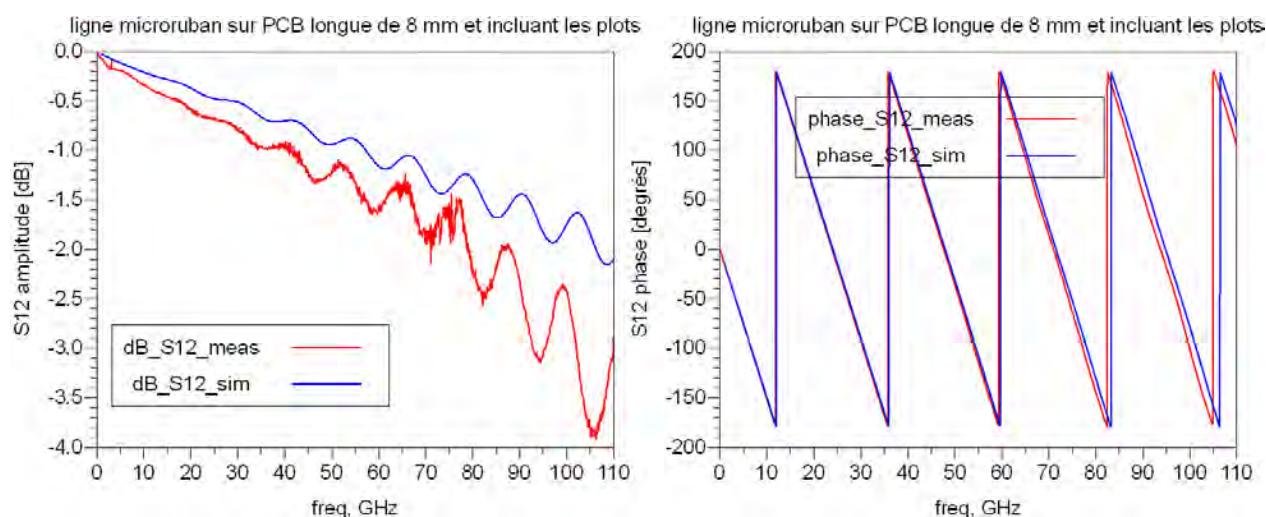


Figure 48 : comparaison simulations / mesures des performances d'une ligne microruban 50 Ω sur PCB incluant les plots

On observe une bonne correspondance entre les simulations et les mesures de la ligne malgré son importante longueur. La différence de pertes entre les simulations et les mesures émane de la sous-estimation des pertes des plots, puisque pour des lignes sans plots les valeurs de ces pertes simulées et mesurées sont très proches.

Cette caractérisation permet de conclure que la modélisation des lignes sur PCB est suffisamment fidèle. Néanmoins, il sera préférable d'éplucher les plots nécessaires aux mesures afin de modéliser les pertes avec plus de précision.

3.3 Modélisation et validation d'interfaces boîtiers « Fan-Out » WLP

Rappelons qu'il est essentiel de disposer d'un modèle du boîtier fidèle à la réalité afin de concevoir correctement la fonction sur la puce de semi-conducteur, et cette fidélité étant d'autant plus critique que les fréquences sont élevées. La contribution de l'interface boîtier au facteur de bruit et à la puissance de sortie fait partie des aspects qui ne peuvent être bien prédits qu'avec une modélisation correcte du boîtier.

Les boîtiers « fan-out » WLP sont soumis à moins de contraintes que les boîtiers « fan-in », étant donné que les boules servant à connecter les accès millimétriques de la puce au PCB sont situées à l'extérieur de la puce. Par conséquent, ces boules sont entourées d'un diélectrique à faibles pertes, optimisant ainsi les performances du circuit en hautes fréquences. Néanmoins, la modélisation de ce boîtier n'est pas pour autant simple aux fréquences millimétriques. En effet, prendre correctement en compte les nombreuses interactions entre le PCB, le boîtier et la puce, produites sous différentes formes, n'est pas trivial [3.8].

De plus, il ressort de la littérature que la conception de circuits encapsulés, basée sur des modèles électromagnétiques de l'interface boîtier, souffre d'un manque de précision aux fréquences millimétriques lorsque ces modèles EM ne sont pas validés et affinés à partir de mesures. Ces imprécisions peuvent mener, par exemple, à une surestimation de la puissance de sortie d'émetteurs radar fonctionnant à la bande de fréquence [76-81] GHz [3.9]. L'imprécision des méthodes basées uniquement sur des simulations EM oblige les concepteurs à garder des marges importantes sur la puissance de sortie de l'amplificateur, s'accompagnant d'une augmentation importante de la consommation causée par le surdimensionnement des transistors. La chaîne de réception est soumise à des contraintes similaires concernant l'adaptation en entrée et les performances en bruit.

3.3.1 Modèle du boîtier RCP

Le boîtier « fan-out » WLP utilisé ici est le RCP (Redistributed Chip Package). Il est similaire au boîtier « fan-out » eWLB à quelques différences près, notamment la présence dans le RCP d'un anneau métallique entourant le Silicium, appelé EGP (Embedded Ground Plane).

Des modèles de boîtier RCP pour le radar avaient été développés par le passé dans le cadre de ce projet. Cependant, leur précision ne permettait pas de prédire suffisamment correctement le comportement du boîtier. Par conséquent, les performances du circuit encapsulé étaient dégradées. Afin d'améliorer les performances affichées par le radar, il s'est donc avéré nécessaire de disposer d'un modèle d'interface boîtier plus précis. Une nouvelle évaluation des différentes géométries intervenant dans la transition du PCB vers la puce est effectuée à base de coupes transversales et images aux rayons X au laboratoire d'analyses NXP à Toulouse. Ces travaux d'inspection des géométries viennent s'ajouter au savoir-faire acquis au niveau de la modélisation EM grâce aux travaux réalisés au niveau du circuit intégré. Nous présentons sur la Figure 49 quelques-unes des coupes transversales et des images rayons X réalisées.

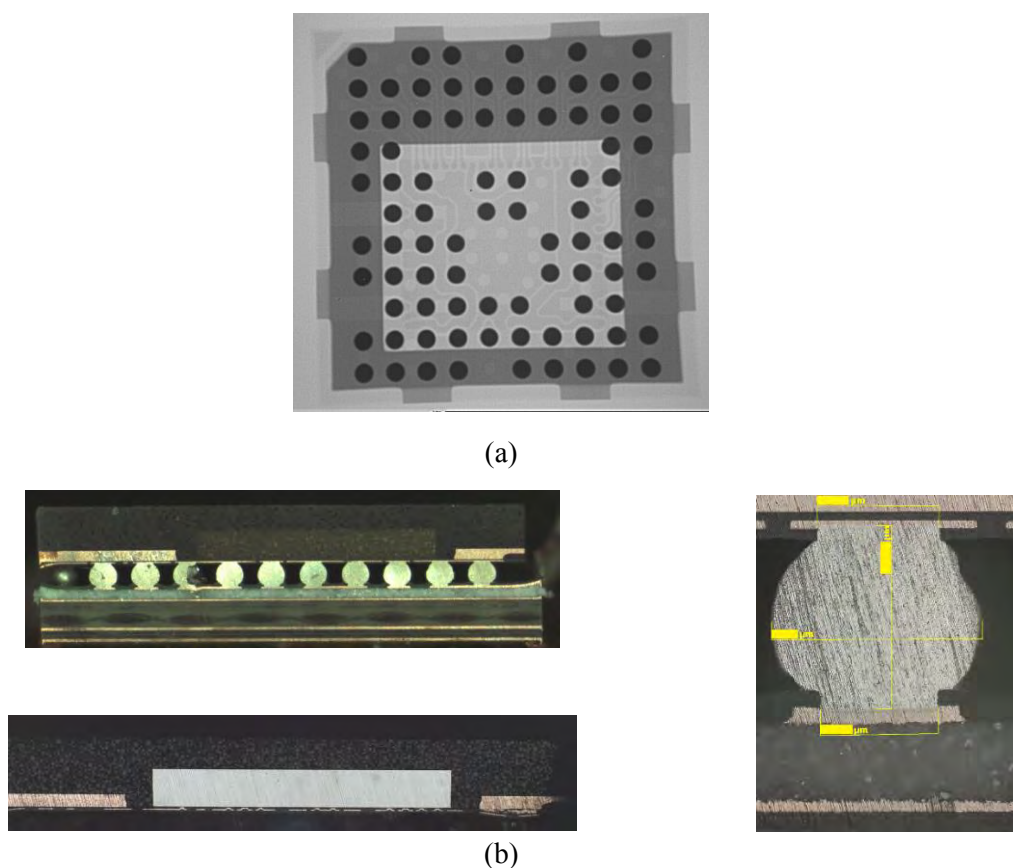


Figure 49 : image aux rayons X (a) et coupes transversales (b) du boîtier RCP

Les mesures réalisées sur différentes géométries de l'interface boîtier ont permis d'ajuster le modèle électromagnétique de cette interface. La hauteur des boules et l'espacement entre l'EGP et la puce font partie des paramètres réajustés suite à ce travail. La Figure 50 présente une vue en perspective du boîtier.

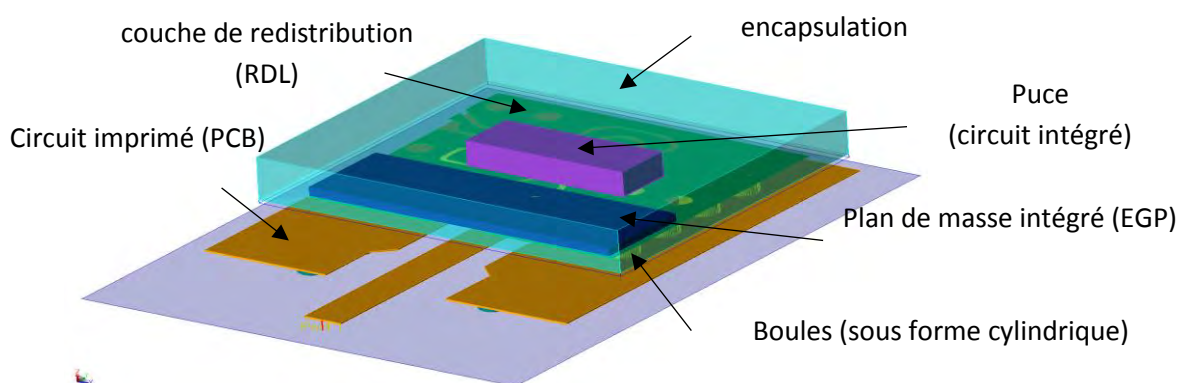


Figure 50 : Vue en perspective de l'interface boîtier modélisée

On remarque que le modèle conçu sous ADS FEM utilise des boules sous forme de cylindres pour des raisons de simplification. Ce modèle a été comparé à un modèle basé sur les mêmes géométries excepté les boules sous forme de cylindres « sim_cylinder » (Figure 52 (a)) qui ont été remplacées par des formes plus sphériques « sim_ball » (Figure 52 (b)) en utilisant l'outil EMPro de Keysight, le même éditeur de l'outil ADS.

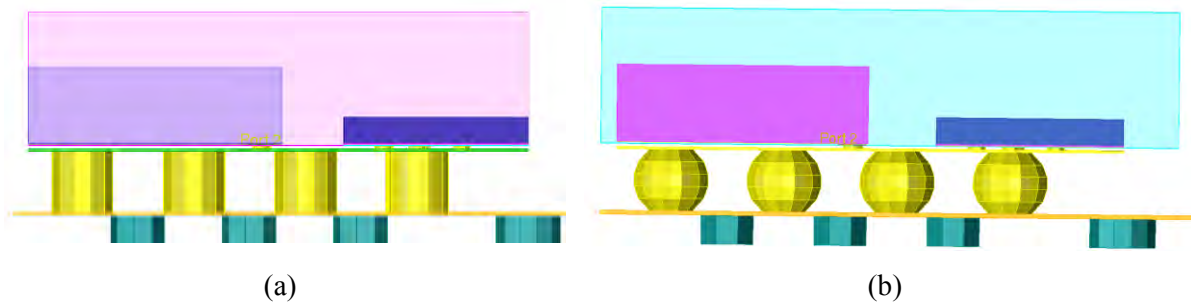
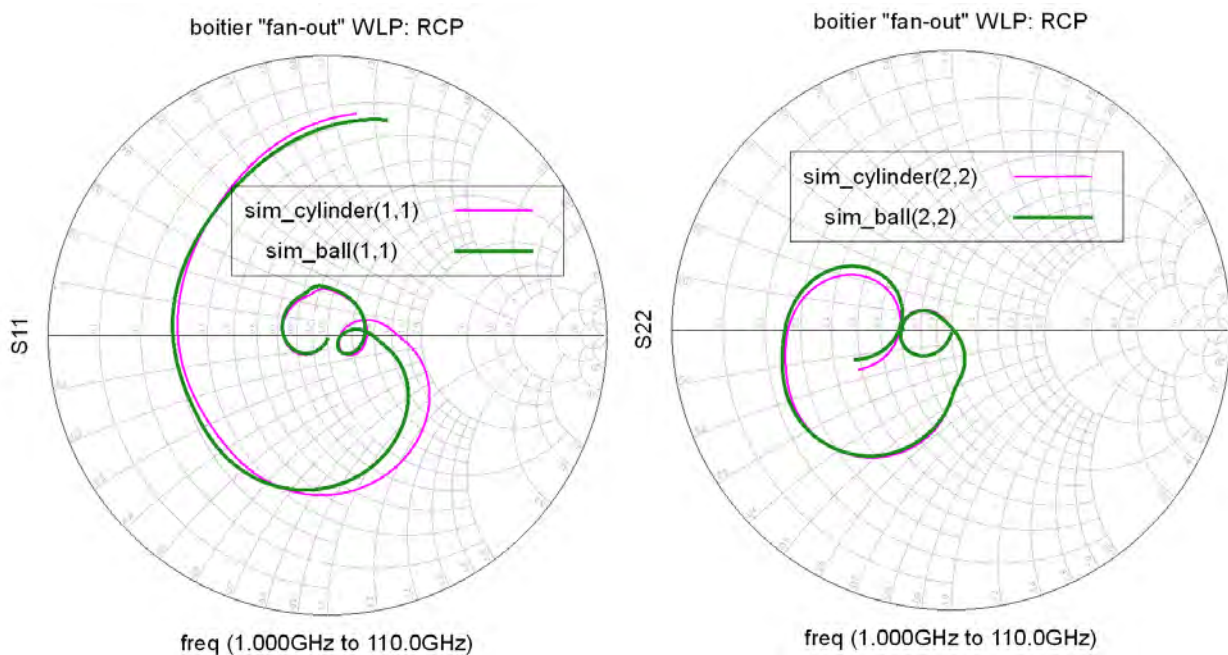


Figure 51 : coupes transversales du boîtier RCP avec des boules sous forme cylindrique (a) et sous forme sphérique (b)

Les deux modèles donnent des résultats très similaires comme on peut le constater sur la Figure 52. Un léger décalage est observé pour le retard (delay). Il provient de la différence entre les chemins de propagation du signal au niveau des boules. En effet, aux fréquences élevées les champs sont principalement localisés à la surface des boules. Par conséquent, les boules sous forme de sphères présentent un retard supplémentaire, ce qui correspond aux résultats de la simulation de la phase de S_{12} dans le modèle intégrant des boules sphériques. Cette différence reste tout de même faible et n'empêche pas d'utiliser le modèle ADS basé sur des cylindres pour la comparaison entre les simulations et la mesure et la compréhension des phénomènes régissant la propagation des ondes entre le PCB et la puce.



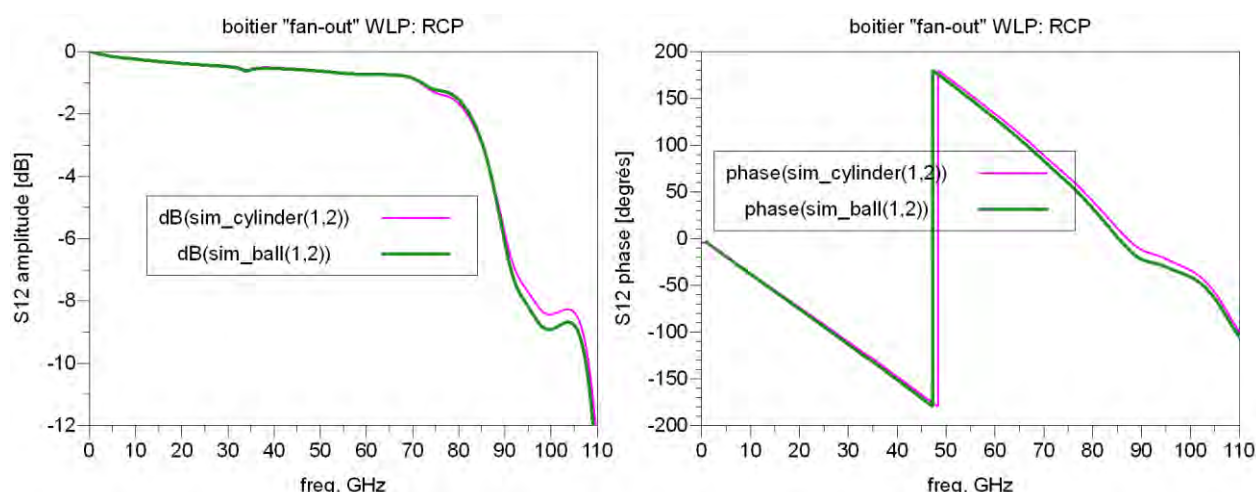


Figure 52 : comparaison des modèles RCP avec des boules cylindriques et sphériques

3.3.2 Validation expérimentale basée sur une méthodologie d'extraction OSL

La validation expérimentale de l'interface boîtier pose un problème fondamental : l'absence d'un deuxième accès au niveau de la puce permettant la réalisation d'une mesure paramètres S en quadripôle. Des travaux se sont intéressés à cette problématique mais sans présenter de validation expérimentale de toute l'interface du PCB à la puce. Une des réalisations les plus complètes a proposé de considérer uniquement les lignes au niveau de la puce et les boules du boîtier [3.10]. Il a été possible de mesurer en quadripôle cette transition en posant les pointes directement sur les boules. Cette approche est intéressante mais ne peut pas être appliquée dans le contexte de ces travaux à cause de la complexité de la structure du boîtier. En effet, la structure dans [3.10] dispose de boules posées directement sur les plots de la puce, ce qui limite la désadaptation apportée par le boîtier et permet de raisonner uniquement sur le délai de propagation. Dans notre cas, on a pu constater à travers les simulations EM qu'une désadaptation importante est apportée par le boîtier. De plus, nous devons absolument considérer le boîtier et le PCB comme étant une seule transition vu les fortes interactions présentes entre les deux. Nous aborderons par conséquent une approche spécifiquement développée pour pallier toutes ces limitations. Cette approche est basée sur des mesures en dipôle et l'algorithme de calibrage OSL (pour Open-Short-Load).

3.3.2.1 Description de la méthodologie OSL (Open-Short-Load)

Si l'algorithme OSL est couramment utilisé pour retirer la contribution des interconnexions d'un dispositif sous test, il n'a encore jamais été utilisé à notre connaissance pour l'extraction d'interfaces boîtier. Cet algorithme a été préféré à l'algorithme TRL (pour Thru-Reflect-Line) puisqu'il permet de librement choisir le plan de référence au niveau de la puce vu la nature dipolaire des standards. L'algorithme TRL aurait nécessité l'utilisation de lignes sur puce de différentes longueurs, ces lignes

nécessitant en cascade des modifications au niveau du boîtier afin de relier leurs deux extrémités à l'interface boîtier.

Dans la méthodologie développée ici, la transition (S_P), joue le rôle des interconnexions de test. Le modèle de l'interface boîtier est alors extrait en appliquant l'algorithme OSL à partir de mesures de terminaisons connus, comme le montre la Figure 53.

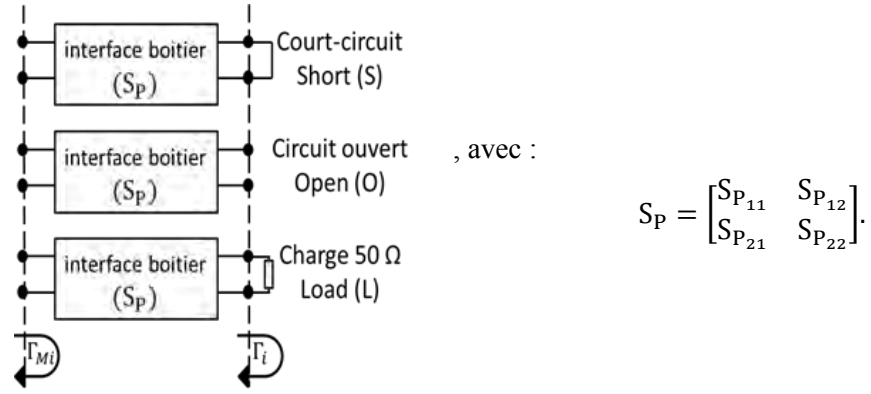


Figure 53: Le principe de l'algorithme OSL appliqué à l'extraction du modèle de la transition

Trois standards Open, Short et Load sont conçus et caractérisés sur puce nue avant d'être encapsulés dans le boîtier RCP. De cette façon, l'interface boîtier se comporte exactement comme des interconnexions de test décalant l'impédance vue par le port de l'analyseur de réseau vectoriel (VNA) de Γ_i à Γ_{Mi} , selon (18) :

$$\Gamma_{Mi} = S_{11} + \frac{S_{21}S_{12}\Gamma_i}{1 - S_{22}\Gamma_i}, \quad \text{avec } i = O, S, L \quad (18)$$

Les paramètres S des 3 standards sont mesurés dans les 2 configurations en puce nue et en boîtier, et donnent 6 coefficients de réflexion Γ_i et Γ_{Mi} , avec $i = O, S, L$. La boîte des paramètres S (S_P) du modèle est alors extraite en utilisant les équations (19), (20) et (21), dérivés de (18) :

$$S_{P11} = \Gamma_{ML} \quad (19)$$

$$S_{P21}S_{P12} = \frac{(\Gamma_S - \Gamma_O)(\Gamma_{MS} - \Gamma_{ML})(\Gamma_{MO} - \Gamma_{ML})}{\Gamma_S\Gamma_O(\Gamma_{MS} - \Gamma_{MO})} \quad (20)$$

$$S_{P22} = \frac{\Gamma_S(\Gamma_{ML} - \Gamma_{MO}) + \Gamma_O(\Gamma_{MS} - \Gamma_{ML})}{\Gamma_S\Gamma_O(\Gamma_{MS} - \Gamma_{MO})} \quad (21)$$

Pour l'extraction de S_{P21} et S_{P12} à partir de (20), la réciprocity de l'interface boîtier permet de considérer $S_{P21} = S_{P12}$. L'amplitude et la phase de S_{P21} (ou S_{P12}) sont ensuite calculées séparément. L'amplitude est calculée en utilisant directement (22), alors que pour la phase, il est nécessaire de

regarder la valeur de la phase à basse fréquence pour décider si la phase commence à 0 ou à π , comme présenté dans la méthode OSL étendue [3.11].

$$|S_{P_{21}}| = |S_{P_{12}}| = \sqrt{|S_{P_{12}} S_{P_{21}}|}. \quad (22)$$

3.3.2.2 Validation expérimentale

Sur la Figure 54, nous procédons à la comparaison des résultats fournis par le modèle extrait à partir de cette méthodologie à ceux issus des simulations EM présentées plus haut.

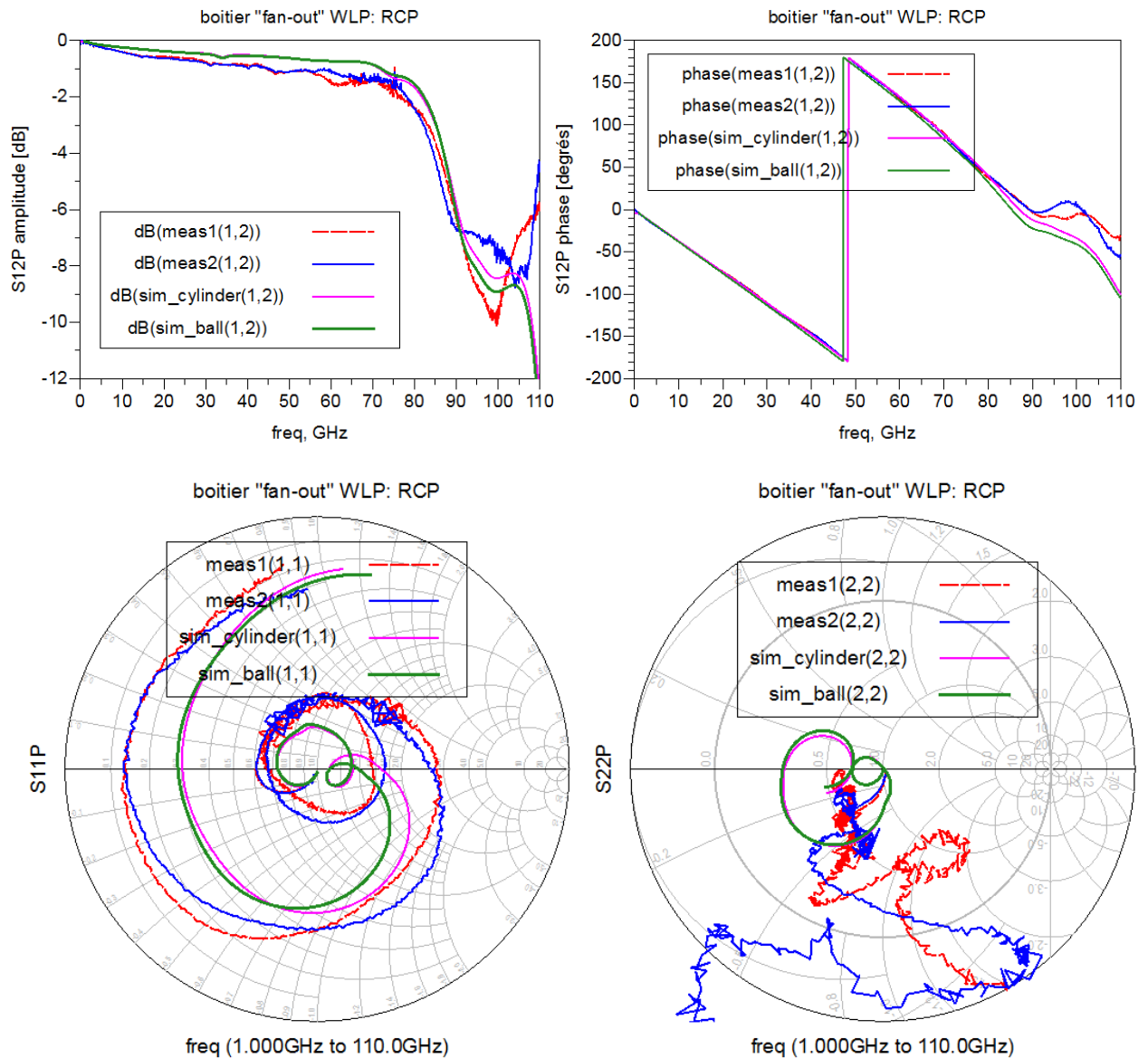


Figure 54 : comparaison des simulations et des extractions par la mesure des paramètres S du boitier RCP

On constate que les paramètres S simulés et mesurés sont assez comparables, en particulier les paramètres S11P et S12P (ou S21P). Les pertes de la transition ainsi que le déphasage qu'elle introduit sont bien modélisés et les mesures ne présentent pas d'aberrations. Cependant, on remarque que la phase du paramètre S22P est différente de ce qui est attendu de cette structure. Par ailleurs, l'amplitude est

supérieure à l'unité, comme s'il s'agissait d'un système actif. À cause de ces limitations, il n'est pas possible de directement comparer le modèle basé sur les simulations à celui basé sur les mesures de la transition. La source des erreurs de ce modèle issu des mesures a été investiguée et présentée à la conférence internationale IEEE RWS2016 [3.12]. En effet, un problème de plan de référence a été identifié. Le standard « Load » inclut le plot de la puce, connectant le circuit au boîtier. Dans les standards « Open » et « Short », ce plot est mesuré avec les standards et sa contribution à la transition est retirée. On se retrouve ainsi avec un algorithme qui enlève la contribution du plot pour certains standards et pas pour d'autres, d'où l'aberration du comportement du S22. Néanmoins, le même article propose une méthodologie permettant d'exploiter les résultats de cette approche afin de corriger rigoureusement la simulation EM. Cette correction est basée sur la comparaison des 3 standards OSL en boîtier à des simulations EM de ces 3 standards. Il s'agit de terminer le modèle EM du boîtier par des modèles des standards Open, Short et Load, identiques à ceux conçus et mesurés, et les comparer aux trois mesures respectives de ces standards en boîtier. Cette comparaison a permis d'accorder correctement les simulations et les mesures. En effet, elle a permis de constater la nécessité de rajouter une capacité parallèle d'une dizaine de fF, au niveau du plot de la puce, dans le modèle EM, pour tenir compte de l'effet du substrat semi-conducteur. Après la correction de cet écart grâce à l'ajout de la capacité sous forme d'élément localisé au niveau du plot de la puce, les simulations des standards correspondent parfaitement aux mesures comme on peut le voir sur la Figure 55.

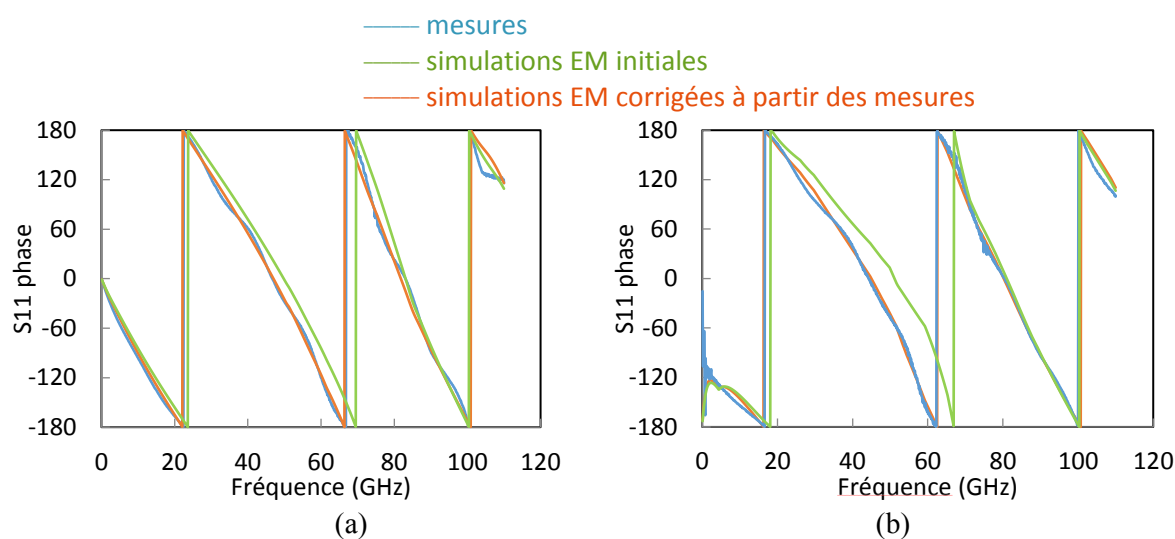


Figure 55 : comparaison des simulations et des mesures des structures « circuit ouvert (O) » (a) et « charge 50 Ω (L) » (b) en boîtier avant et après correction

Grâce à ces travaux, on dispose maintenant d'un modèle EM de l'interface boîtier RCP, fidèle à la mesure, qui permet de comprendre le comportement de la transition en identifiant les contributeurs aux pertes et les sources de désadaptation.

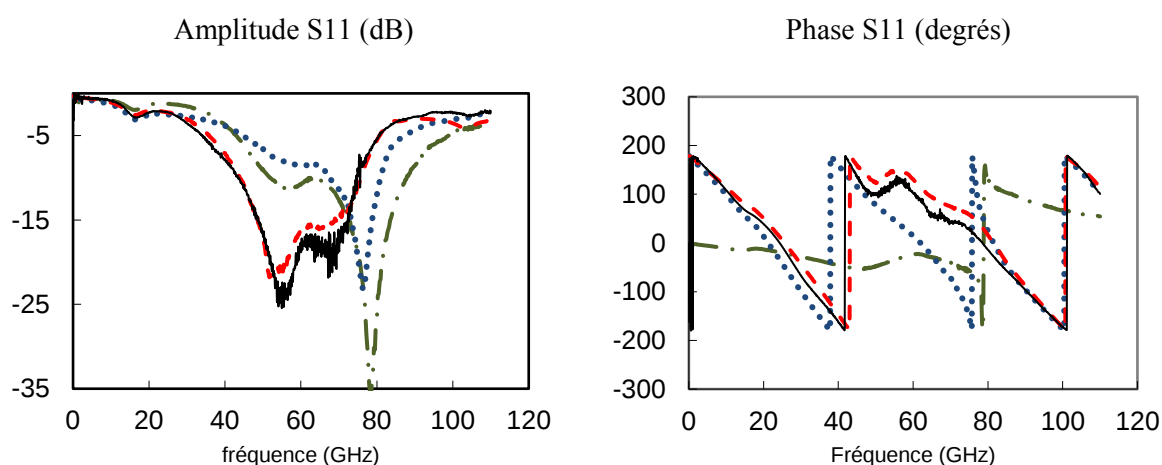
La Figure 54 montre que l'adaptation est bonne et les pertes ne dépassent pas 1 dB des basses fréquences jusqu'à une fréquence d'environ 60 GHz. Ces pertes s'expliquent, en partie, par la contribution du plot

et des lignes sur PCB, avec environ 0.3 dB de pertes introduites par chacun de ces éléments. Les pertes diélectriques et celles liées à la résistivité des métaux interviennent également. Dans la bande d'intérêt pour le radar automobile, [76-81] GHz, des pertes par désadaptation en plus d'un début de coupure en fréquence viennent s'additionner aux contributeurs cités précédemment, portant les pertes à environ 1.7 dB. Ces pertes restent très correctes pour ces fréquences. Elles sont équivalentes aux pertes introduites par une ligne 50 Ω sur puce d'une longueur de 1.5 millimètres.

Ces bonnes performances électriques étaient attendues mais elles ne suffisent pas, à elles seules, à juger des performances globales du boîtier « Fan-Out » WLP. En effet, les performances thermiques du Fan-Out sont moins bonnes que celles du « Fan-In » WLP à cause de l'encapsulation, entourant la puce, qui limite la dissipation thermique. Le coût du boîtier « Fan-Out » est également supérieur à celui du boîtier « Fan-In ».

3.3.3 Application à l'encapsulation de la chaîne de réception d'un radar 77-GHz

Le modèle EM validé et corrigé est utilisé pour déterminer l'adaptation en entrée d'un récepteur radar 77 GHz encapsulé en boîtier RCP. Les canaux de réception de ce circuit présentent, sur puce nue, des coefficients de réflexion meilleurs que -30 dB à 77 GHz. Une fois la puce encapsulée en boîtier WLP, les mesures montrent le déplacement de l'adaptation vers les basses fréquences d'environ 10 GHz. Le modèle développé est alors associé aux caractérisations du récepteur en puce nue pour modéliser la transformation d'impédance apporté par le boîtier. Le résultat de ce travail est comparé par la suite aux mesures et aux simulations impliquant le modèle EM initial dont disposait l'équipe (Figure 56). Comme attendu, les résultats fournis par le modèle généré en utilisant la méthodologie développée correspondent très bien aux mesures au niveau du boîtier, et sa validité jusqu'à 110 GHz est confirmée. Ces résultats démontrent que ce modèle de transition prédit avec précision la transformation d'impédance introduite par le boîtier et permettra ainsi une conception optimisée des futures générations du radar automobile 77 GHz en boîtier.



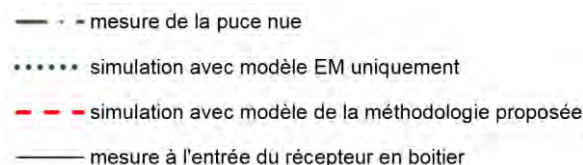


Figure 56 : Coefficient de réflexion du récepteur radar encapsulé en boîtier

Les bonnes performances électriques du boîtier « Fan-Out » sont confirmées mais l'amélioration de la dissipation thermique et du coût ne peut être réalisée qu'en utilisant un boîtier de type « Fan-In ».

3.4 Conception d'une interface boîtier « fan-in » WLCSP pour le radar 77 GHz

Après la modélisation réussie d'un boîtier « fan-out », nous nous intéressons maintenant à la conception et à la modélisation d'un boîtier « fan-in » WLP qui sera le premier de ce type à être utilisé pour une technologie Silicium à des fréquences aussi élevées. Comme discuté au premier chapitre, les rares travaux de recherche impliquant l'utilisation de ce boîtier aux fréquences millimétriques sont basés sur un boîtier 3D au coût élevé et reposent sur une technologie GaAs. Jusqu'à aujourd'hui, les radars présents sur le marché ou présentés par les fournisseurs de circuits intégrés sont câblés en microfils ou utilisent des boîtiers de type « fan-out » WLP. Une première conception du boîtier « fan-in », calquée sur la conception du boîtier RCP, a été réalisée. Elle a permis d'évaluer les performances de ce boîtier et de confirmer, comme évoqué dans le premier chapitre, les difficultés d'utiliser ce type de boîtier pour la bande de fréquence visée [76-81] GHz. Néanmoins, nous avons entrepris de modéliser ce boîtier pour ouvrir la voie vers la conception d'un boîtier WLCSP performant et surtout utilisable pour un radar 77 GHz. Les pertes de cette interface boîtier ne devront pas dépasser de 0.3 dB celles du RCP, soit 2 dB de pertes d'insertion au maximum. La désadaptation des deux côtés de l'interface boîtier devra rester meilleure que -7 dB afin de faciliter l'adaptation au niveau de la puce, et d'éviter l'utilisation de réseaux d'adaptation au niveau PCB. Des innovations importantes au niveau de la conception sont attendues afin d'atteindre ces objectifs.

Dans la suite de ce paragraphe, nous nous intéressons aux réalisations ayant permis d'atteindre les objectifs énumérés précédemment. Nous présentons aussi la validation expérimentale, basée sur une nouvelle méthodologie, de ce nouveau boîtier.

3.4.1 Boîtier WLCSP de départ (version 1) et limitations

La différence fondamentale entre un boîtier « fan-in » et un boîtier « fan-out » WLP est la disposition des boules par rapport à la puce (figure 57). Dans le cas du « fan-in », toutes les boules sont situées au niveau de la puce, alors que dans le cas du « fan-out », certaines sont situées au niveau de la puce mais la plupart se trouve placée sur une région au-delà de bords de la puce. Les boules pour les signaux à

haute fréquence sont par conséquent disposées dans cette région afin de minimiser leur couplage avec le circuit. Dans le cas du « fan-in », les couplages entre les boules et la puce doivent être gérés avec une grande attention.

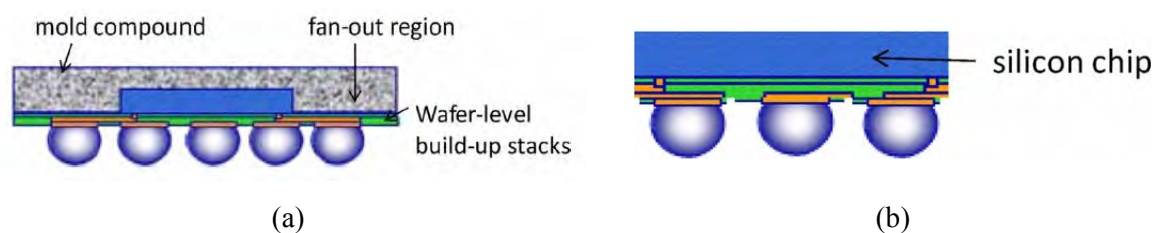


Figure 57 : boîtiers de type "fan-out" (a) et "fan-in" (b)

La superposition des métallisations du boîtier aux lignes de signal des blocs millimétriques crée des couplages importants. Ces couplages peuvent générer des résonances et compromettre la fonctionnalité même du circuit. Afin d'éviter une telle situation, la conception de la puce a été réalisée de sorte à garder un espace entre l'anneau des plots de puce et l'extrémité du Silicium. Seul un plan métallique de masse est présent dans cette zone nue comme le montre la figure 58.

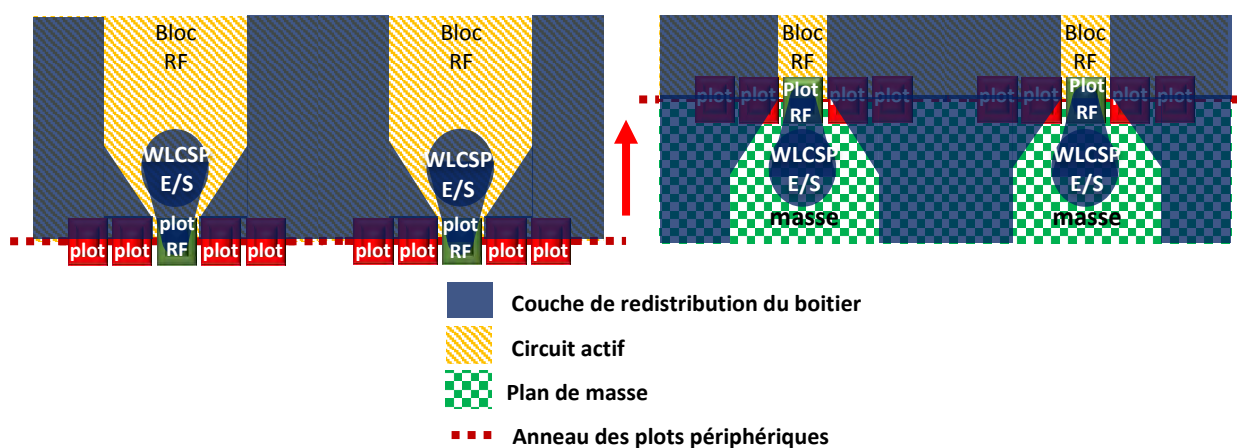


Figure 58 : dessin 2D de la disposition des couches du boîtier par rapport à la puce du boîtier

Cette configuration a un coût en termes de surface de Silicium vu le rajout d'une surface supplémentaire. Cependant, elle permet de s'affranchir des problèmes de contre-réaction entrée-sortie au niveau des blocs d'amplification millimétriques et de mélange. Un modèle préliminaire correspondant à ce boîtier a été développé et comparé aux mesures extraites de la méthodologie OSL présentée précédemment (Figure 59).

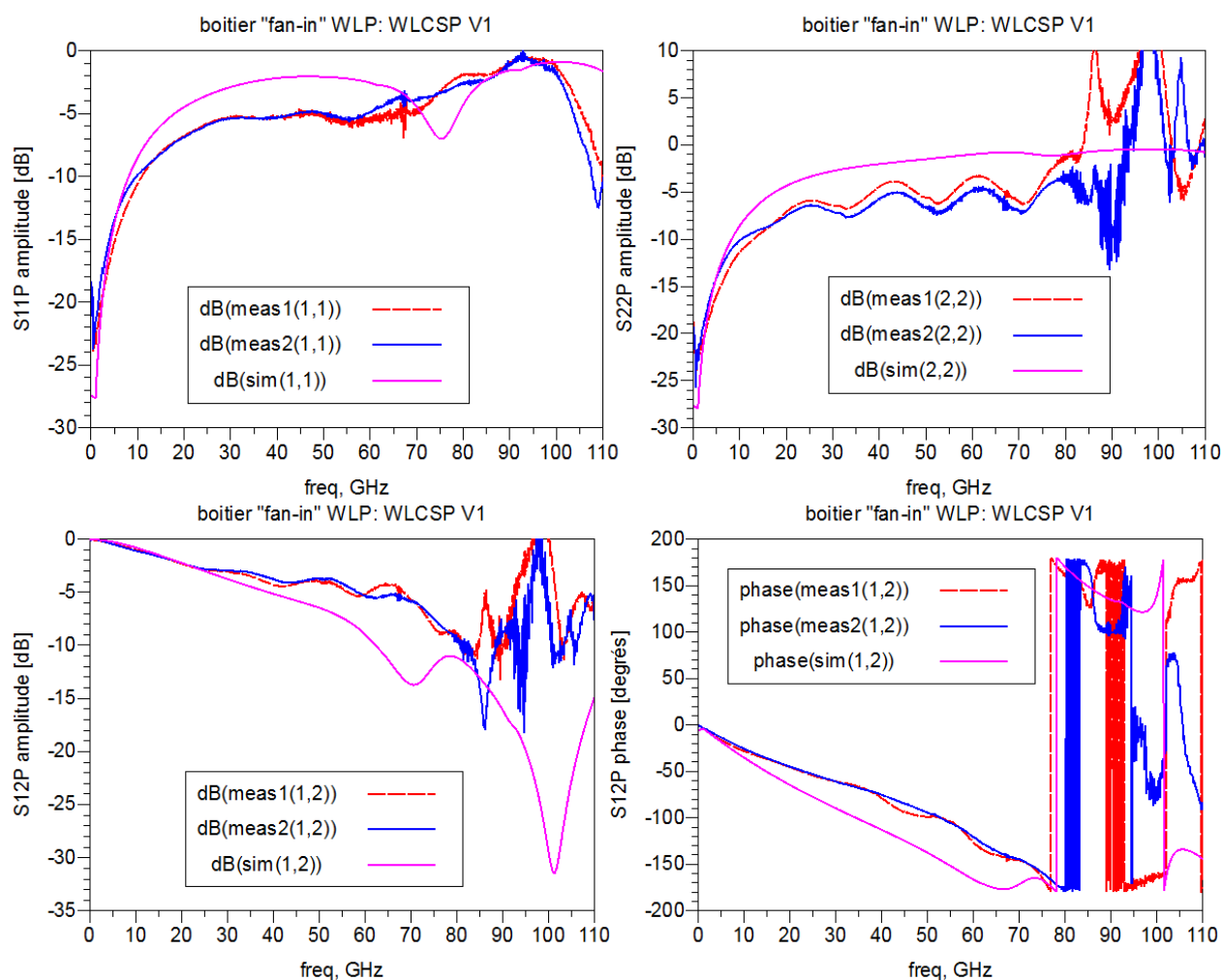


Figure 59 : résultats de simulations et de mesures de l'interface boîtier WLCSP V1

Les résultats des simulations et des mesures sont légèrement différents mais montrent tous deux des performances médiocres pour le boîtier WLCSP, que ce soit pour l'adaptation que pour les pertes d'insertion. Ce boîtier tel qu'il est conçu ne peut en aucun cas être utilisé pour un radar opérant à 77 GHz. En effet, il présente des pertes d'insertion d'environ 10 dB dans la bande d'intérêt. Même parfaitement et idéalement adapté des côtés PCB et puce, ce qui est impossible, les pertes dans la bande se situent autour de 5 dB.

La principale source de désadaptation et de pertes résulte de la capacité parallèle, constituée entre le niveau RDL et le plan métallique de masse sur la puce, dont la valeur est importante. La très faible impédance présentée par cette capacité ainsi que l'inductance de la boule viennent s'ajouter à la désadaptation introduite par le PCB créant au final une structure s'accompagnant de niveaux de pertes élevés. La conception d'un boîtier utilisable en bande [76-81] GHz sera basée, entre autres, sur ces informations. Cette conception fait l'objet du paragraphe suivant.

3.4.2 Proposition et conception d'un nouveau boîtier WLCSP (version 2)

Les résultats obtenus pour le boîtier WLCSP initial ont permis de distinguer d'ores et déjà les principaux contributeurs aux pertes et aux désadaptations de l'interface boîtier. Il s'agit maintenant de proposer des architectures pour réduire leur contribution et disposer d'une interface utilisable pour le radar dans la bande de fréquence 77 GHz. Le principal de ces contributeurs est la capacité parallèle entre le niveau RDL et le plan métallique de masse intégré sur la puce, d'un côté, et la structure sur le PCB, de l'autre.

Les moyens mis en œuvre pour la réduction de la capacité parallèle entre le niveau RDL et la masse de la puce ont fait l'objet d'une demande de brevet [3.13], déposée en fin d'année 2015. L'approche suivie s'inspire des travaux réalisés sur la conception des plots de la puce avec et sans plan métallique inférieur de masse, présentés dans le deuxième chapitre. Ces travaux avaient montré qu'il était plus intéressant, pour certaines dimensions de plots, d'utiliser des plots en regard du Silicium au lieu de plots protégés par l'écran métallique. Les dimensions des plots du niveau RDL sont déjà fixées aux valeurs minimales pour la technologie de boîtier, compte tenu de la taille des boules utilisée. Ces dimensions ne peuvent donc pas être réduites davantage. La solution proposée consiste alors à enlever la métallisation de la puce à l'aplomb des boules de signal afin de réduire la capacité parallèle. Cependant, par défaut dans cette technologie, le Silicium est légèrement dopé. Pour ne pas que les performances soient trop dégradées, on applique donc à la zone sous les boules des entrées-sorties millimétriques les mêmes étapes de fabrication que pour l'intégration des plots de signal sans plan métallique inférieur. Ces étapes permettent de disposer d'un substrat non dopé sans coût supplémentaire. De plus, il a été proposé de minimiser l'espace dégagé sous les boules en gardant l'anneau de plots de puce à l'extrémité du Silicium et créer des ouvertures uniquement sous les boules des entrées des récepteurs et sorties des émetteurs (Figure 60).

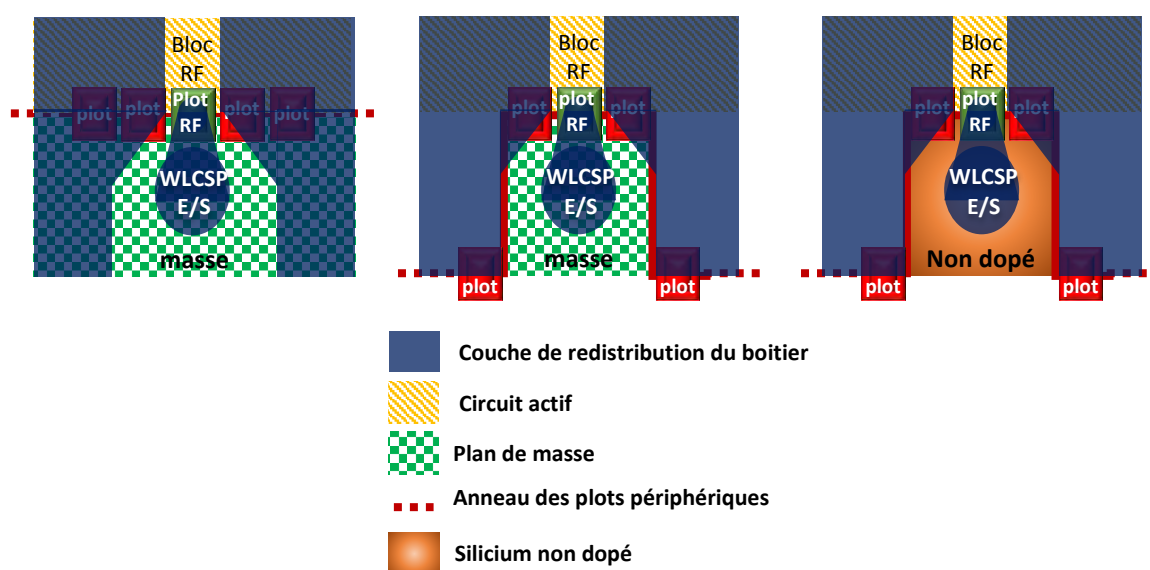


Figure 60 : étapes de conception de la puce prévue pour du WLCSP optimisé

Cette architecture permet de minimiser l'espace inutilisé au niveau du circuit intégré et de minimiser la capacité parallèle liée au niveau RDL, qui est le principal contributeur à la désadaptation du boîtier. Cette opération permet de réduire la capacité parallèle associée au RDL de 225 fF à environ 110 fF, ce qui devrait changer significativement le comportement de l'interface boîtier. L'inconvénient de cette approche est l'ajout de pertes supplémentaires au bilan des pertes de l'interface à cause de la résistivité du Silicium. Les pertes supplémentaires dû à cette résistivité sont estimées à environ 0.6 dB, ce qui reste très correct vu les avantages apportés par cette modification.

Le deuxième contributeur aux pertes de l'interface boîtier est le PCB. Des investigations basées sur l'analyse des champs EM au niveau du PCB ont montré leur sensibilité par rapport au positionnement des vias. La disposition des boules et les dimensions des vias contraignent leur positionnement. Néanmoins, on dispose d'un léger degré de liberté au ras du boîtier après les boules. Ce degré est exploité pour améliorer la transition sur PCB et minimiser ses pertes. Les investigations ont montré que la configuration de vias (disques verts) proposé dans la Figure 61 (b) donnait les meilleurs résultats possibles avec la contrainte du positionnement des boules (disques jaunes).

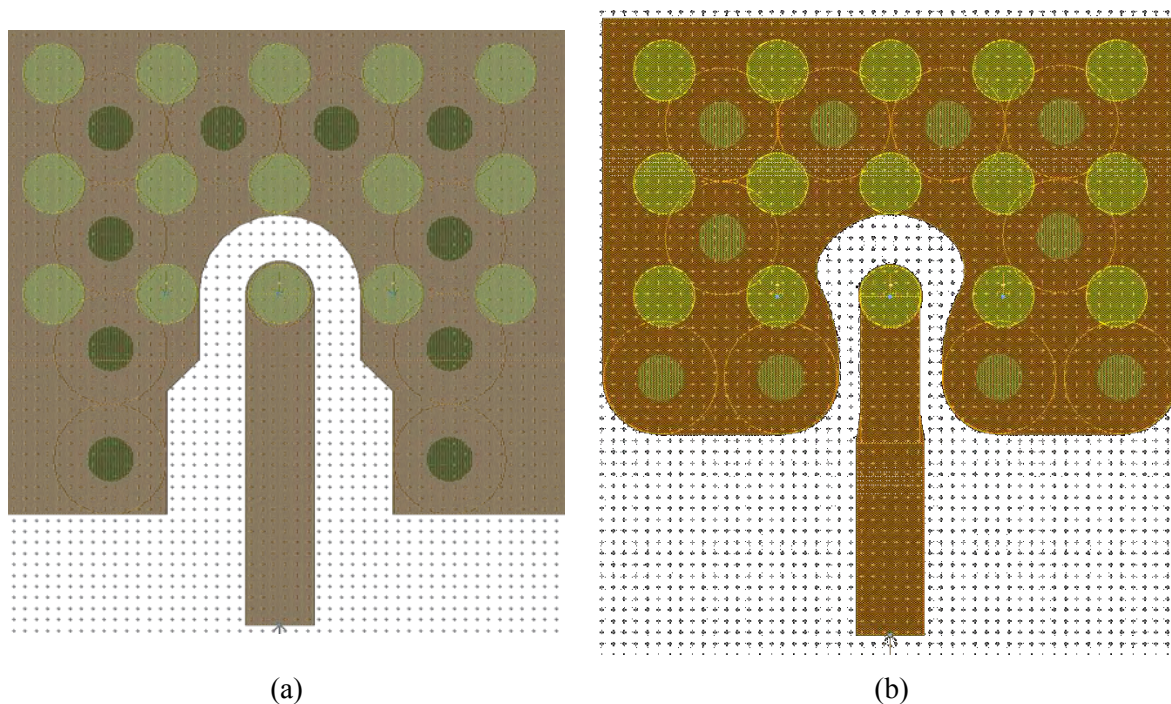


Figure 61 : dessin des PCB du boîtier WLCSP V1 (a) et du boîtier WLCSP V2 (b)

Ces deux modèles de PCB, contenant uniquement les métallisations du circuit imprimé, ont été simulés (Figure 62). On constate que les performances sont bien meilleures avec la nouvelle configuration proposée. Par ailleurs, on améliore la marge de sécurité par rapport à la résonance constatée sur le PCB.

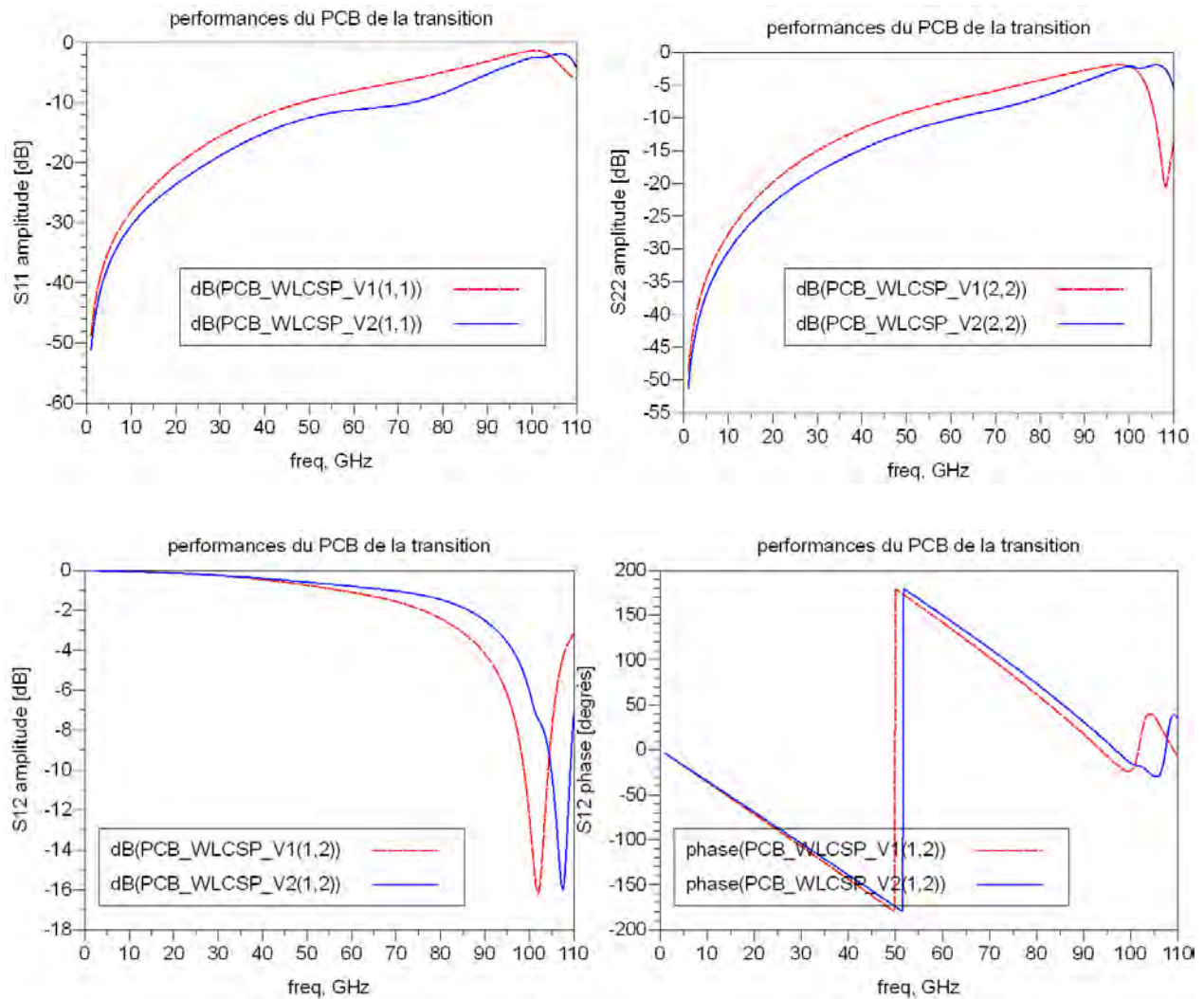


Figure 62 : comparaison des performances des PCB associés aux interfaces boîtier des versions 1 et 2 du boîtier WLCSP

On constate que l'adaptation, au-delà de 70 GHz, est bien meilleure avec l'architecture de PCB proposée. De plus, la résonance de la structure est décalée de plus de 5 GHz vers les fréquences supérieures.

Avec la réduction importante de la capacité parallèle du RDL et l'amélioration de la transition au niveau du PCB, il a été possible de transformer l'interface boîtier initiale, inutilisable en hautes fréquences, en une interface boîtier aux performances très correctes et qui avoisinent celles du RCP. La Figure 63 montre l'évolution de la transition après les deux changements majeurs que nous y avons apportés.

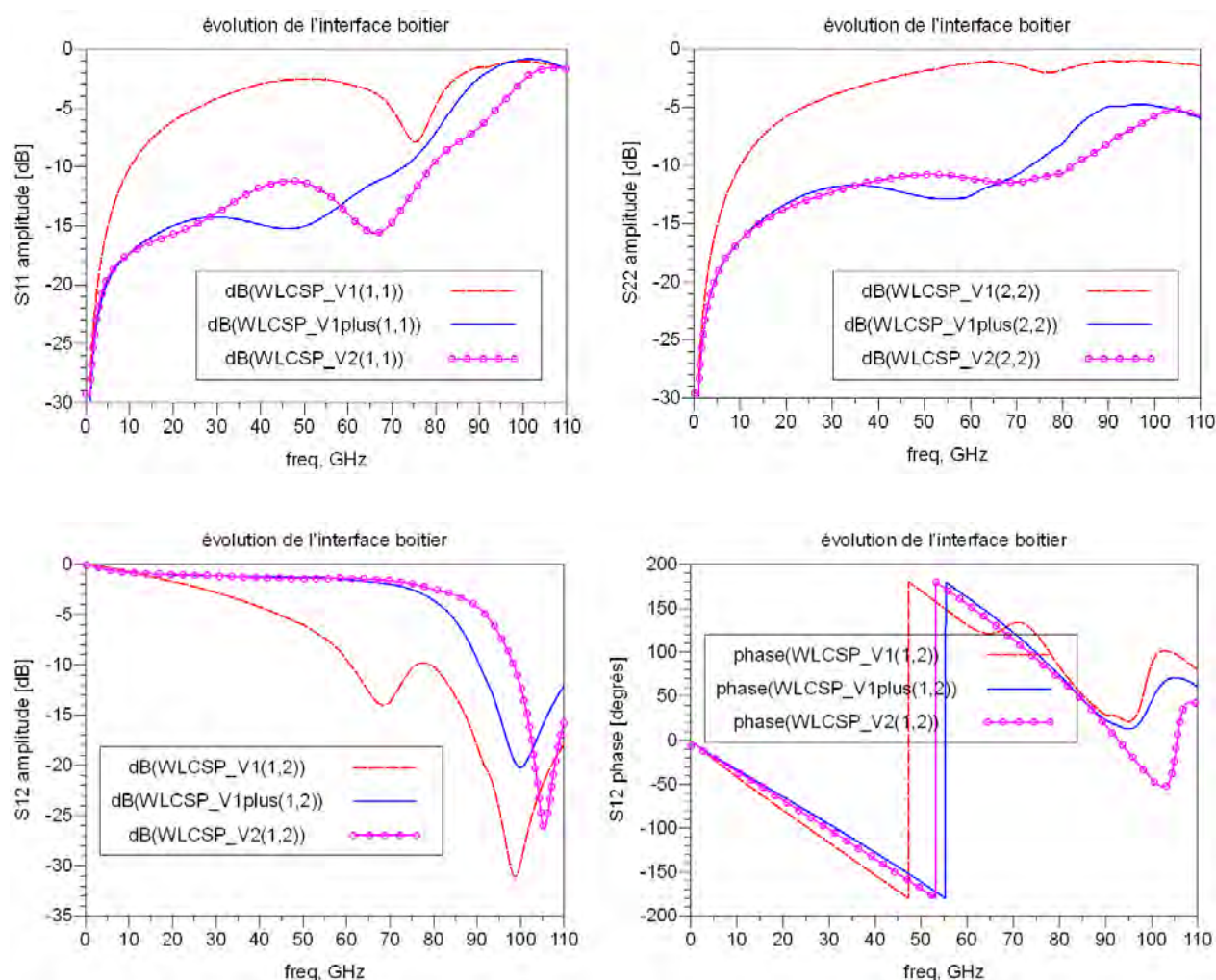


Figure 63 : évolutions du boîtier WLCSP (V1) avec la réduction de la capacité (V1plus) puis en rajoutant l'amélioration du PCB (V2)

Les pertes de l'interface boîtier ont été réduites de 10 dB à 2.5 dB, environ, dans la bande radar [76-81] GHz. Ces performances sont exceptionnelles pour ce type de boîtier et c'est la toute première fois qu'un boîtier de type « fan-in » WLP présente de telles performances à ces fréquences, à notre connaissance. L'adaptation est également très bonne, en particulier après les changements apportés au niveau du PCB. On constate que l'adaptation est meilleure que -10 dB dans la bande radar, quand l'interface est chargé par des terminaisons 50 Ω des deux côtés. Ces performances sont suffisantes pour concevoir un radar automobile fonctionnel. Ceci dit, elles restent à confirmer par des mesures.

3.4.3 Validation expérimentale basée sur une nouvelle méthodologie « Half-OSL »

Après avoir conçu un modèle de boîtier WLCSP présentant des performances très prometteuses grâce aux différentes innovations qui y ont été apportées, il s'agit maintenant de valider expérimentalement ce modèle et de confirmer les avantages qu'il apporte par rapport à la version initiale (V1).

3.4.3.1 Description de la méthodologie Half-OSL

La méthodologie OSL décrite précédemment a montré quelques limitations et notamment par rapport à la définition du plan de référence au niveau de la puce. Le plot de la puce était inclus ou pas dans la transition selon le standard, générant ainsi une erreur dans l'extraction. Afin de remédier à ce problème, une courte ligne microruban 50 Ω a été ajoutée après le plot. Le plan de référence du modèle a été défini à l'extrémité de cette ligne. Le plot de puce et cette ligne font maintenant partie de la transition et doivent être pris en compte lors de la conception du circuit et de la modélisation EM de l'interface (Figure 64).

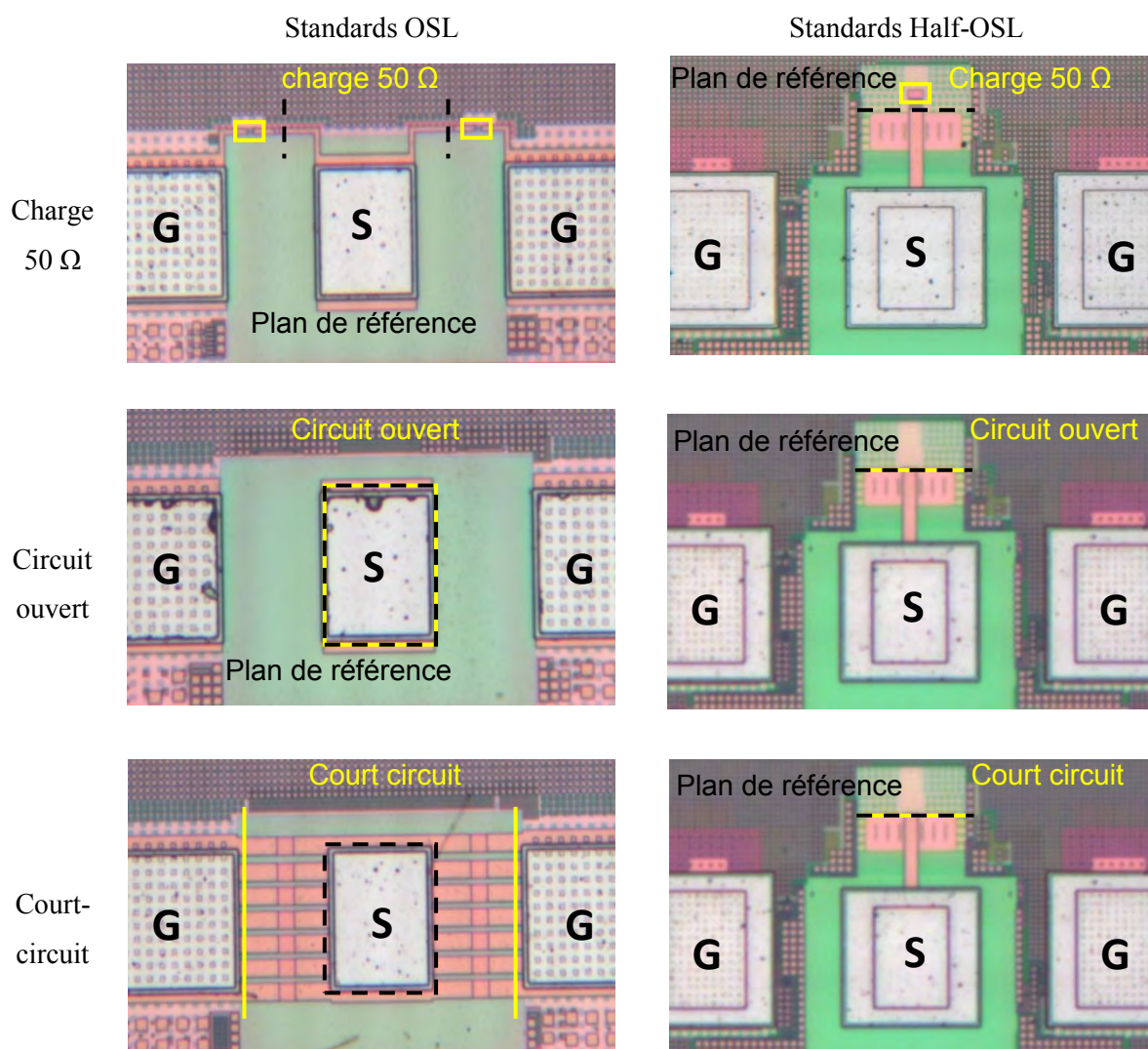


Figure 64 : photos des standards OSL et Half-OSL au niveau de la puce incluant les définitions du plan de référence

À partir de cette nouvelle définition du plan de référence, les standards OSL sur puce sont conçus. Il s'agit pour l'« Open » d'avoir un circuit ouvert au bout de la ligne, aucun élément n'est donc connecté à cette ligne. Pour le « Short », des vias connectent directement le bout de la ligne au plan de masse métallique inférieur de façon à limiter l'inductance à quelques pH. Pour le « Load », une résistance 50 Ω précise est connectée au bout de la ligne. La capacité et l'inductance associées à cette résistance sont négligeables, même à 110 GHz vu la taille de la résistance. Cette conception des standards OSL permet

l'application de l'algorithme développé dans le cadre de ces travaux de thèse. Cet algorithme est fondé sur l'hypothèse que non seulement la charge $50\ \Omega$ est parfaite mais également les standards « Open » ou « Short ». Pour ces deux derniers standards, il serait possible de considérer des imperfections mais elles ne pourraient pas être extraites de mesures directes à cause de l'extrême difficulté associée à la caractérisation. À partir des valeurs extraites des simulations EM et l'analyse des conséquences de faibles variations autour de ces valeurs, nous avons pu conclure que l'hypothèse d'un « Open » et d'un « Short » parfaits était valide pour les fréquences d'intérêt. L'application de cette hypothèse permet d'utiliser directement l'algorithme OSL présenté au plus haut mais en se basant uniquement sur les mesures au niveau du PCB et en considérant que les standards Open, Short sur puce présentent des caractéristiques idéales de circuit ouvert et de court-circuit respectivement, comme résumé sur la Figure 65.

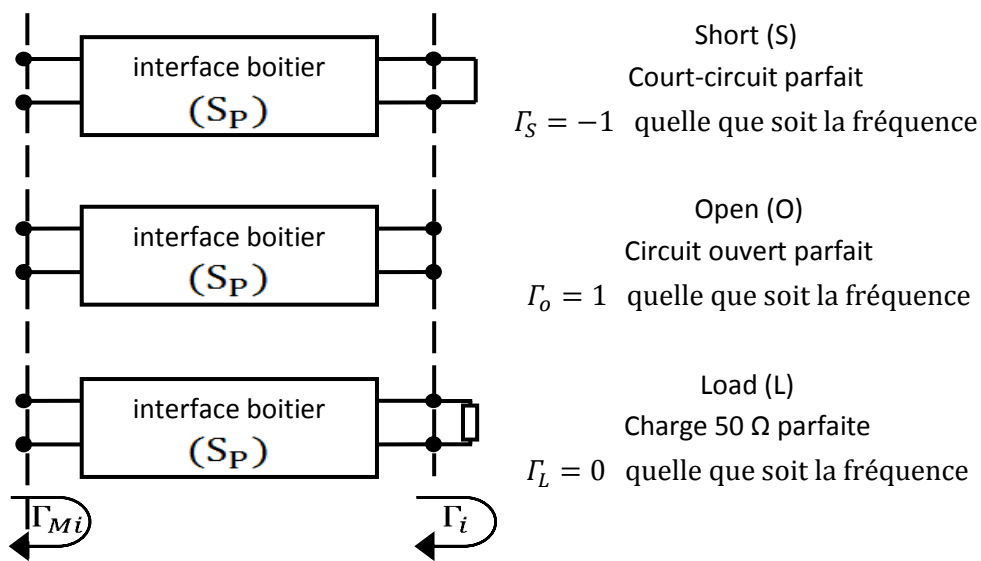


Figure 65 : schéma représentant l'extraction à base de l'algorithme Half-OSL proposé

La méthodologie développée ici permet d'économiser les mesures sur puce nue par rapport à la méthodologie OSL présentée plus haut, d'où l'appellation « Half-OSL ». Cette approche enlève l'ambiguïté liée au plan de référence au niveau de la puce et vise à réduire l'incertitude de la mesure. En effet, l'élimination des mesures sur puce nue élimine le risque de couplage entre les plots de puce et la pointe lors des mesures.

3.4.3.2 Validation expérimentale

Nous procédons maintenant à la validation expérimentale de ce nouveau boîtier développé en WLCSP en utilisant la méthodologie « Half-OSL ». Les performances résultant du modèle extrait à partir de cette méthodologie sont comparées aux résultats obtenus avec le modèle EM développé (Figure 66).

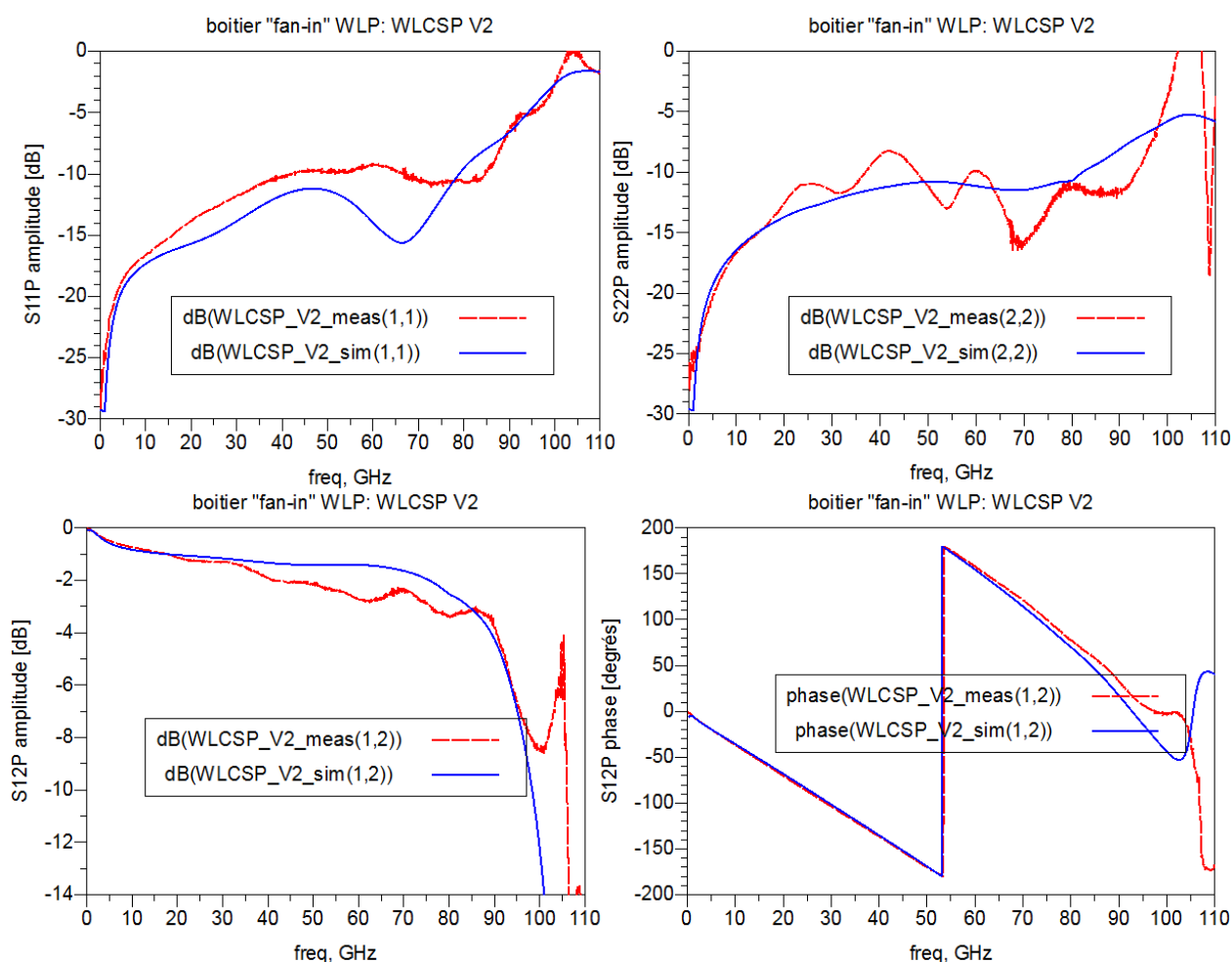


Figure 66 : simulations et mesures de la nouvelle interface boîtier conçue WLCSP V2

On relève une très bonne correspondance entre les simulations et l'extraction expérimentale basée sur la méthodologie « Half-OSL ». Cela prouve, d'une part, que le modèle EM est assez précis, même en incluant le silicium, et, d'autre part, que l'extraction du boîtier basée sur la méthode « Half-OSL » est plus précise que la méthodologie OSL classique. Ces travaux ont également démontré que les pistes de développement suivies pour la conception d'un boîtier utilisable en bande radar 77 GHz ont été couronnées de succès.

Ce travail a été considéré par la communauté des concepteurs des circuits intégrés et des architectes boîtier comme étant révolutionnaire vu le nombre de nouvelles opportunités qu'il ouvre, malgré ses pertes légèrement supérieures à la valeur visée (2 dB). Il permettra d'utiliser pour la première fois ce boîtier standard, considéré jusqu'à ce jour comme boîtier « basses fréquences », pour des applications aux fréquences millimétriques. De plus, aucune modification/évolution du processus de fabrication n'est à prévoir à aucun niveau de la chaîne : la fabrication de la puce, du boîtier et du PCB reste inchangée. L'économie réalisée en passant d'un boîtier « fan-out » à un boîtier « fan-in » sur une production à grande échelle peut s'élever à plusieurs millions d'euros.

3.5 Optimisation de la nouvelle interface boîtier WLCSP

Les résultats que nous venons de présenter montrent que la conception d'un boîtier « fan-in » WLP est finalement envisageable pour l'encapsulation du radar automobile dans la bande [76-81] GHz. Cependant, les pertes de la transition restent encore un peu élevées par rapport aux spécifications définies pour le boîtier. Nous avons donc poursuivi les investigations en explorant des voies d'optimisation qui permettent de réduire ces pertes, et qui permettent aux performances électriques du boîtier « fan-in » WLP de rivaliser avec celles des boîtiers de type « fan-out ». Le boîtier WLCSP sera optimisé à base de simulations EM et le modèle sera de nouveau comparé aux mesures. Par ailleurs, de nouvelles structures vont être spécifiquement développées pour l'application d'une méthodologie d'extraction basée sur le principe TRL. Ces structures vont être exploitées parallèlement aux structures « Half-OSL », et elles permettront la comparaison des deux extractions.

3.5.1 Améliorations proposées sur le boîtier WLCSP (version 3)

La conception de l'interface boîtier WLCSP est basée, au niveau de la puce, sur une réduction de la capacité parallèle entre le niveau RDL et la masse. Au niveau du PCB, la transition du microruban vers les boules du boîtier a été améliorée de sorte à minimiser la désadaptation. Ces modifications ont porté leurs fruits mais une optimisation supplémentaire s'impose pour réduire encore les pertes de la transition. L'analyse des pertes montre que le premier contributeur aux pertes est le Silicium à cause de sa résistivité assez faible ($10 \Omega \cdot \text{cm}$). Cette valeur a été extraite à partir des travaux réalisés pour la conception des plots (chapitre 2). Nous allons par conséquent nous focaliser sur la réduction des pertes liées au Silicium. La résistivité du Silicium non dopé ne peut être augmentée d'avantage sans augmentation significative du coût de production. En effet, il s'agirait alors d'utiliser des plaquettes de Silicium forte résistivité, impactant significativement le procédé de fabrication et la validité des modèles de tous les composants de la technologie BiCMOS utilisée. La réduction de la surface du niveau RDL en regard du Silicium est donc le seul levier dont on dispose encore pour réduire ces pertes. La surface de RDL délimitant les boules est déjà au minimum autorisé. Cependant, la consultation des manuels de conception du boîtier WLCSP a permis de s'apercevoir que cette technologie pose moins de contraintes que la technologie RCP au niveau du routage. Ce point sera donc exploité pour minimiser la surface de l'interconnexion du niveau RDL reliant le plot de la puce à la boule. De plus, les plans de masse sur le niveau RDL ont été rapprochés du signal pour mieux confiner les champs. La taille des plots de puce est également moins contrainte pour la technologie WLCSP et cela n'avait pas été exploité, jusqu'à présent. La Figure 67 présente un aperçu des changements apportés pour optimiser les performances de l'encapsulation WLCSP, la nouvelle génération étant nommée version 3 ou « V3 ».

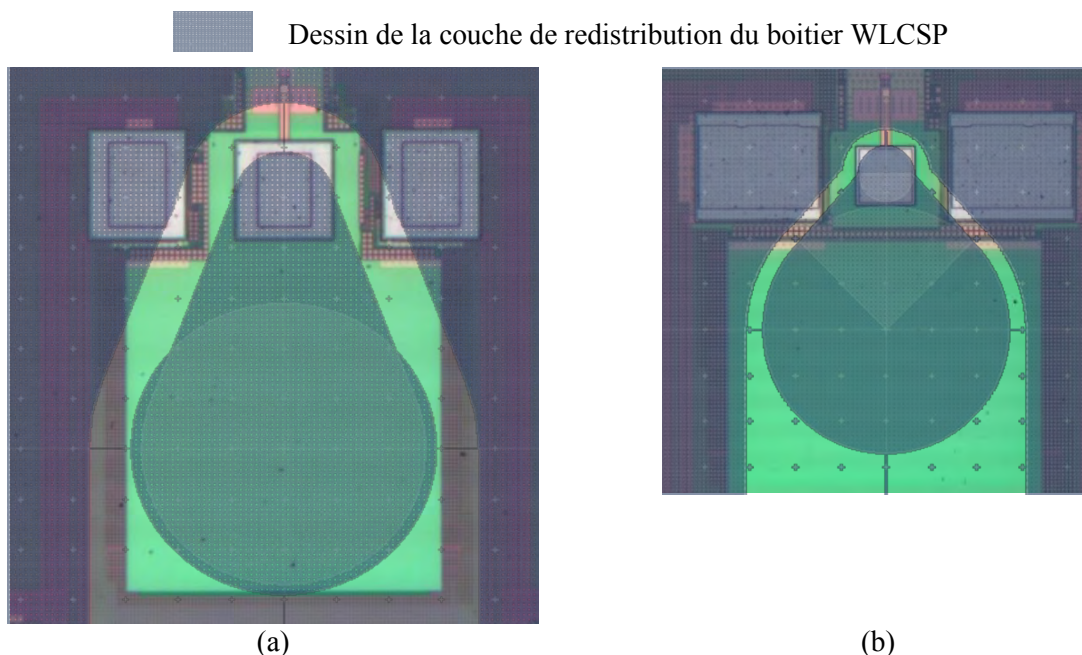


Figure 67 : photos à la même échelle des plots des puces conçues pour être encapsulées respectivement en boîtier WLCSP V2 (a) et en boîtier WLCSP V3 (b)

La nouvelle interface boîtier V3 obtenue permet d'améliorer la transition boîtier en réduisant ses pertes, comme le montre la Figure 68.

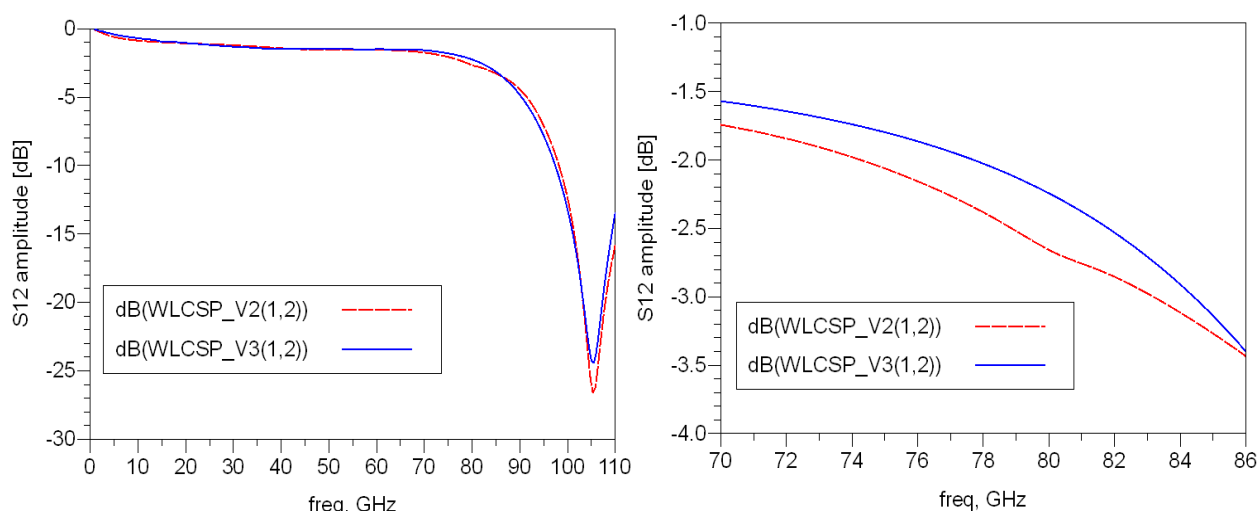


Figure 68 : comparaison des modèles de boîtier WLCSP V2 et V3

Grâce aux dernières optimisations, l'interface boîtier V3 présente des pertes réduites d'environ 0.5 dB supplémentaires dans la bande radar 77 GHz. Les performances de ce boîtier optimisé deviennent ainsi très comparables à celles d'un boîtier RCP. Le défi de concevoir un boîtier de type « fan-in » WLP utilisable pour l'encapsulation du radar automobile semble donc relevé. En dehors de cette application, les performances de ce boîtier étant très proches de celles des boîtiers de type « fan-out », ce boîtier « fan-in » WLP devient ainsi très attractif pour toutes les applications millimétriques.

Il reste maintenant à valider expérimentalement ces améliorations de l'interface WLCSP.

3.5.2 Comparaison de la méthodologie Half-OSL et d'une méthodologie basée sur le principe TRL

Une nouvelle méthodologie baptisée « Half-OSL » a été développée pour extraire le plus rigoureusement possible un modèle quadripôle du boîtier à partir de mesures en dipôle. Les résultats obtenus grâce à cet algorithme devraient être meilleurs que ceux résultants de l'algorithme OSL que nous avons développé et utilisé jusqu'à présent. Cependant, il nous a semblé important de pouvoir confronter les résultats issus de la méthodologie Half-OSL aux résultats d'une méthode de référence ayant fait ses preuves. Dans ce manuscrit, nous avons déjà abordé l'algorithme TRL en mentionnant son principal inconvénient par rapport à l'approche OSL qui est le positionnement du plan de référence. De plus, la réalisation de cette extraction nécessite de connecter ensemble deux interfaces boîtier et de disposer de structures spécifiques au niveau de la puce encapsulée. Ceci dit, la méthodologie TRL ne souffre pas des problèmes associés à la non-idéalité d'une charge 50 Ohm, n'en utilisant pas par construction, et ne demande pas de précision trop grande pour la caractérisation des standards « Thru », « Reflect » et « Line » utilisés dans l'extraction.

Des circuits radar encapsulés étaient utilisés lors des précédentes extractions. La configuration des accès millimétriques empêchait alors l'application d'une extraction TRL. Pour cette version finale de la transition, un nouveau boîtier a été spécifiquement développé pour faciliter l'extraction à partir du principe TRL. Cette configuration nous permettra de réaliser une extraction de référence pour l'évaluation d'une extraction Half-OSL menée sur le même boîtier.

Description de la méthodologie basée sur TRL

L'extraction des caractéristiques de l'interface boîtier, à l'aide de l'algorithme TRL, repose sur les équations développées par Engen en 1979 [3.14]. La Figure 69 reprend le principe de la méthodologie développée en lien avec l'algorithme TRL. L'utilisation de deux interfaces boîtiers identiques permet de disposer de deux modèles A et B, lors de chaque extraction. On peut constater que les modèles extraits tiennent compte de la moitié de la ligne microruban « Thru » en plus de l'interface boîtier. La contribution de cette ligne peut également être retirée grâce à l'algorithme TRL. En effet, le calibrage TRL génère, en plus des boîtes d'erreur A et B, la constante de propagation (γ) de la ligne microruban utilisée pour réaliser le retard. Cette valeur suffit pour retirer la contribution de la moitié de la ligne du « Thru » et disposer finalement d'un modèle ne contenant que l'interface boîtier.

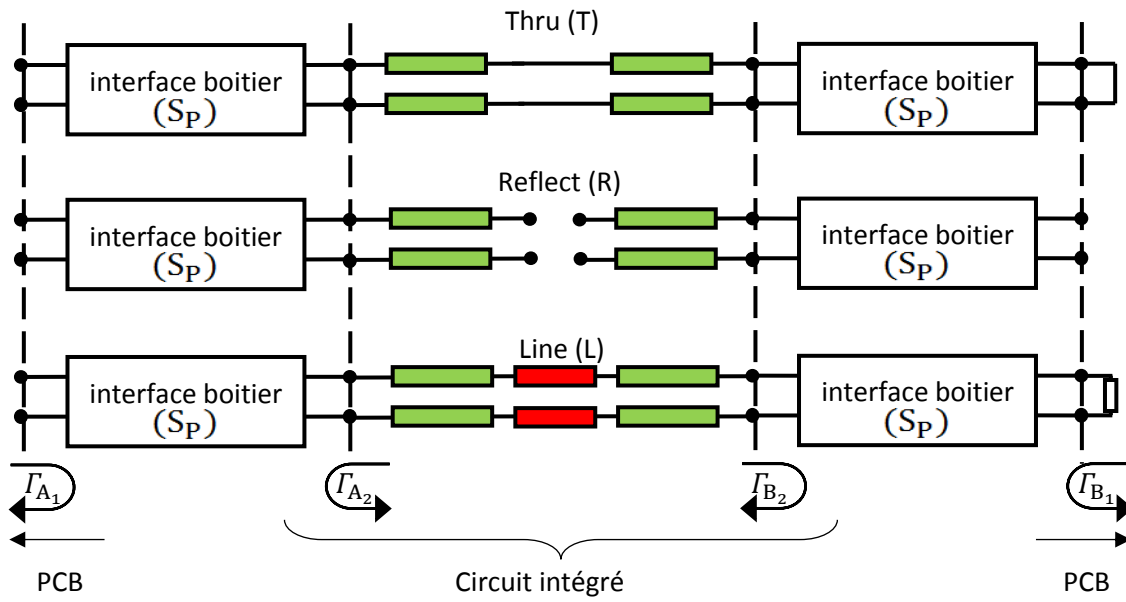


Figure 69 : principe de l'extraction TRL proposée

La Figure 70 reprend les dessins de masque et photos des standards TRL au niveau puce.

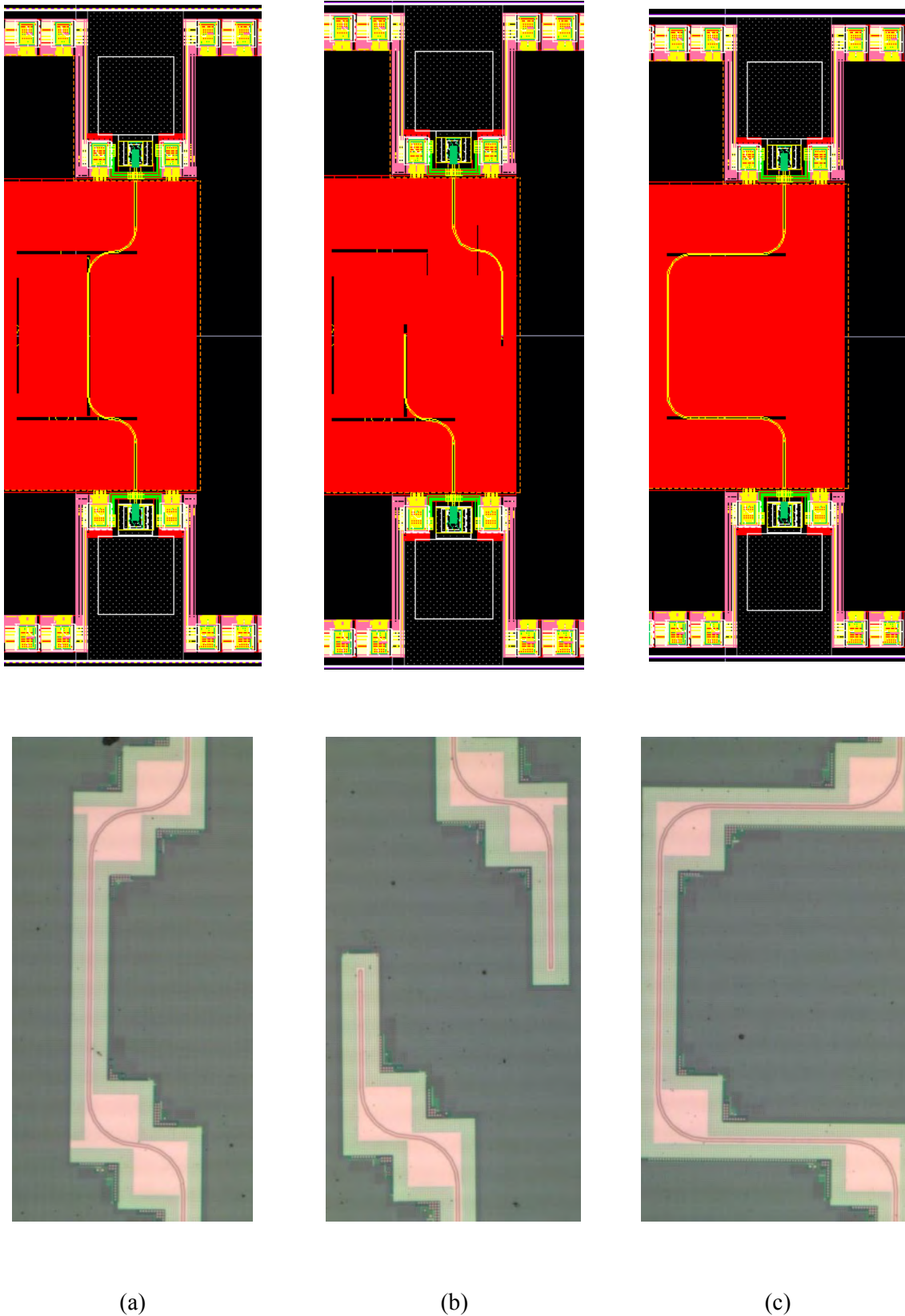


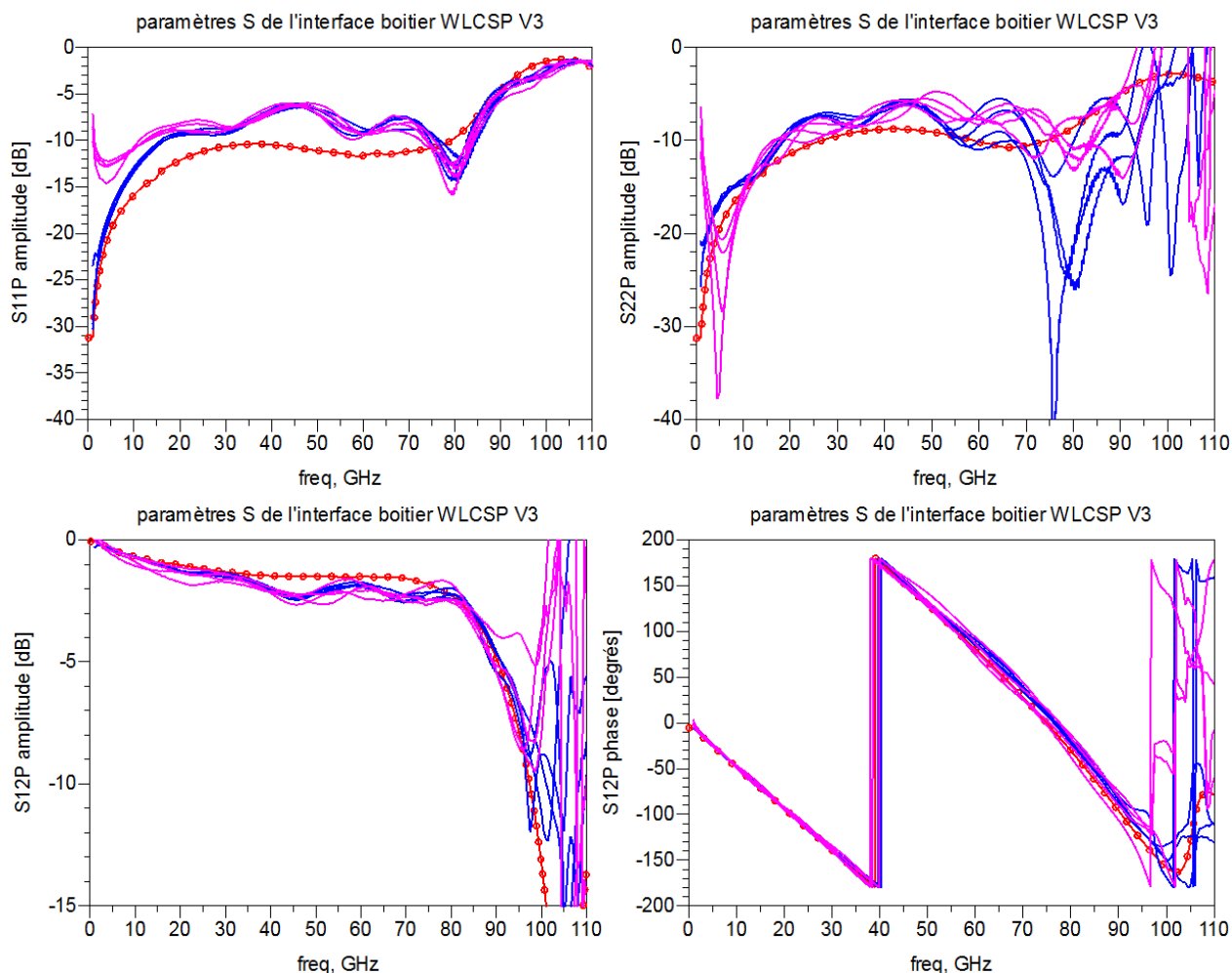
Figure 70 : les dessins de masque et photos des 3 standards TRL sur puce nue : Thru (a), Reflect (b) et Line (c)

Les interfaces de boîtier intervenant dans les deux extractions Half-OSL et TRL sont les mêmes. Cependant, ils utilisent des standards sur puce et des algorithmes très différents puisque pour la première méthode les standards de calibrage sont des dipôles alors que pour la seconde ils sont des quadripôles.

Cette différence fondamentale entre les deux approches rend la comparaison des résultats très intéressante.

3.5.3 Validation expérimentale basée sur les méthodologies Half-OSL et TRL

Les deux méthodologies basées sur Half-OSL et TRL ont été appliquées à la version finale de l'interface WLCSP version 3, et les extractions obtenues sont comparées aux résultats de simulations. Cette comparaison est très importante puisqu'elle validera en même temps l'approche innovante Half-OSL proposée, et les performances de l'interface boitier optimisée. Plusieurs pièces ont été mesurées pour avoir un aperçu sur la dispersion après l'application des méthodologies Half-OSL et TRL. Les résultats de ces mesures sont comparés aux simulations sur la Figure 71.



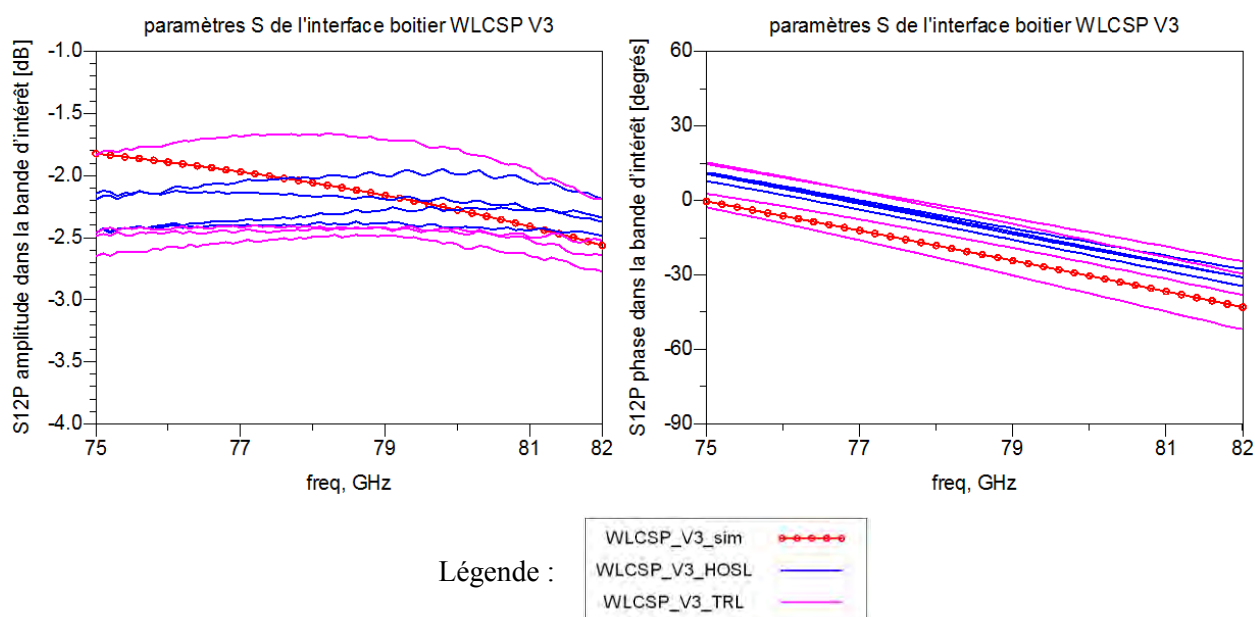


Figure 71 : comparaison des simulations et des mesures de l'interface boîtier WLCSP V3 par les méthodes Half-OSL et TRL

Intéressons-nous d'abord à la comparaison entre les deux extractions Half-OSL et TRL. Les deux méthodes donnent des paramètres S très similaires sauf aux fréquences inférieures à 10 GHz où la TRL est en dehors de sa région de validité compte tenu de la longueur de ligne trop faible. La dispersion des résultats est très similaire dans les deux cas, que ce soit en amplitude ou en phase. On constate que la dispersion sur les pertes d'insertion est d'environ ± 0.3 dB dans la bande d'intérêt, ce qui reste très raisonnable à ces fréquences. Ces observations confirment la validité des algorithmes et des hypothèses sur lesquelles repose l'approche Half-OSL.

La comparaison des simulations et des mesures démontre la bonne précision du modèle électromagnétique développé. En effet, les pertes d'insertion et le déphasage apportés par la transition sont très bien modélisés puisqu'ils sont alignés aux mesures. L'adaptation, que ce soit au niveau de la puce ou au niveau du PCB, est bien prise en compte. On voit de légères différences entre les coefficients de réflexion simulés et mesurés, mais compte tenu des valeurs absolues, autour de -12 dB, ces écarts n'ont pas d'impact sur la qualité du modèle. Par ailleurs, on constate que les pertes ont été effectivement réduites par rapport à la version antérieure du boîtier WLCSP. Elles se situent autour de 2.2 dB.

3.6 Conclusion

Le principal objectif derrière les travaux présentés dans ce chapitre était de proposer une interface boîtier de type « fan-in » WLP qui soit bien caractérisée et qui réponde aux spécifications du radar automobile 77 GHz, sans augmentation de coût. L'absence totale de travaux autour de cette problématique a poussé à adopter une approche incrémentale ayant comme point de départ le boîtier RCP, un boîtier de type « fan-out » WLP. L'étude de la transition du PCB vers la puce a commencé par la modélisation et la caractérisation de structures sur PCB. Cette étude a permis de vérifier la validité des propriétés données

pour les matériaux et de s'assurer de la rigueur de la modélisation électromagnétique. Une fois cette étape franchie, le boîtier RCP a été modélisé et caractérisé. La caractérisation de ce boîtier nous a éclairés sur les limites de la méthodologie développée pour l'extraction des caractéristiques des boîtiers. Par la suite, un nouveau boîtier « fan-in » a été conçu en incluant des innovations le rendant utilisable aux fréquences millimétriques. Le développement de nouvelles approches d'extraction à partir de la mesure des paramètres S a accompagné cette conception. Après quelques itérations, un boîtier « fan-in » WLCSP très performant et rigoureusement caractérisé a été conçu et réalisé. La Figure 72 revient sur les performances des principales évolutions de ce boîtier WLCSP et les compare aux performances du boîtier RCP.

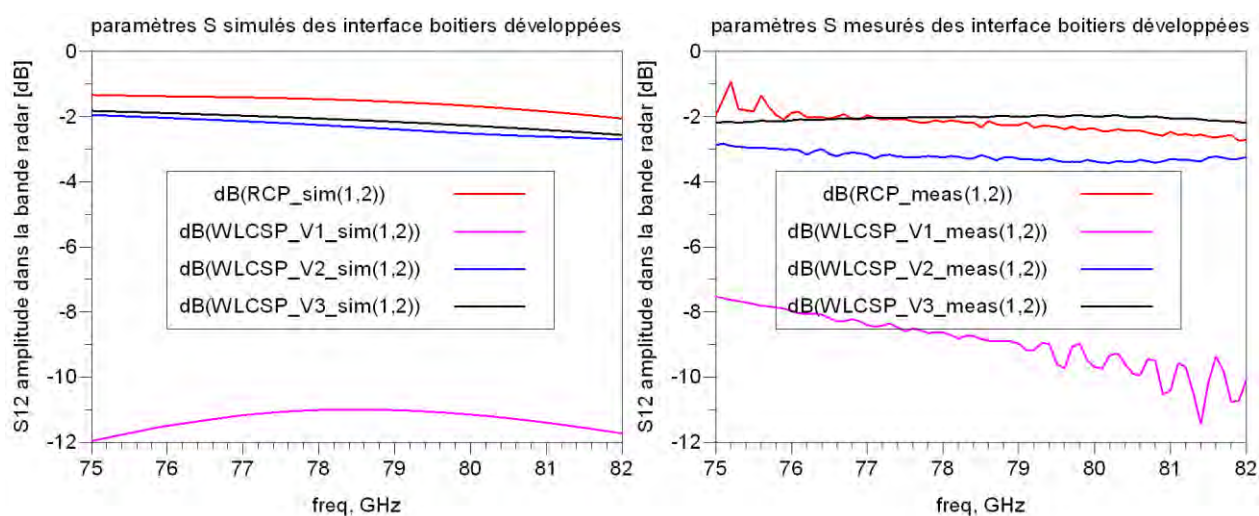


Figure 72 : évolutions des performances simulées et mesurées du boîtier WLCSP et comparaison au RCP

Ces graphes permettent de résumer l'évolution du boîtier « fan-out » en « fan-in », les contraintes associées à cette évolution ainsi que les performances obtenues après les différentes phases d'invention, d'innovation et d'optimisation. Ces courbes révèlent une amélioration des pertes mesurées de plus de 5 dB entre les boîtiers WLCSP V1 et V2 et une amélioration d'environ 1 dB entre les boîtiers WLCSP V2 et V3. Par ailleurs, on remarque que la dernière génération du boîtier WLCSP dispose de meilleures performances en pertes d'insertion qu'un boîtier RCP. La validation de toutes les étapes de conception par différentes approches de caractérisation innovantes donne une robustesse certaine à ces travaux.

4 Application : Conception de la chaîne de réception d'un radar automobile 77 GHz encapsulé

4.1 Introduction

Ce dernier chapitre est consacré à la conception d'une chaîne de réception respectant les spécifications du produit radar automobile 77 GHz, tout en considérant une encapsulation en boîtier WLCSP. La conception se base sur la dernière version développée pour le boîtier WLCSP telle que présentée au chapitre 3. Afin d'atteindre des performances à l'état de l'art, nous reprenons la conception de la chaîne de réception complète en nous penchant sur différentes topologies et en exploitant les travaux du chapitre 2 sur l'optimisation des performances des structures passives. Plus particulièrement, nous nous intéressons au mélangeur de fréquences et à l'amplificateur faible bruit (LNA pour Low Noise Amplifier), lorsque ce dernier est présent, étant donné qu'ils représentent les blocs les plus critiques de la chaîne. Les amplificateurs et les filtres en bande de base ne sont considérés que lors d'analyses du système. Leur contribution est prise en compte dans les performances globales de la chaîne, mais nous ne présentons pas les détails de leur conception. Enfin, la simulation électromagnétique au niveau circuit est au cœur de l'optimisation du mélangeur.

Une analyse des spécifications de la chaîne de réception et du mélangeur de fréquences est tout d'abord réalisée. Elle est ensuite suivie d'une étude des topologies de mélangeur pouvant répondre à ces spécifications, puis de la conception et l'implémentation de l'architecture sélectionnée. Enfin, le mélangeur de fréquences est intégré dans son boîtier et le composant dans son environnement WLCSP est caractérisé.

4.2 Spécifications et architecture de la chaîne de réception radar

Nous nous sommes intéressés, lors du chapitre 1, aux principaux paramètres intervenant dans les spécifications des blocs de l'émetteur-récepteur radar, et notamment de la chaîne de réception. Ces spécifications sont détaillées ici de sorte à disposer de toutes les informations nécessaires à la conception de cette chaîne au niveau du circuit intégré. Il s'agit principalement de réaliser une chaîne de réception aux performances, au moins aussi bonnes que le produit NXP décrit dans [1.18]. Ce produit constitue l'état de l'art des circuits radar automobile en boîtier en cours de production. Ce circuit de réception

radar développé et commercialisé par NXP est principalement composé de mélangeurs de fréquences, d'amplificateurs 77 GHz et d'amplificateurs en bande de base. Il est encapsulé en boîtier « fan-out » RCP dans une topologie « single ended ».

Des recherches bibliographiques ont montré une grande ressemblance entre les spécifications de ce produit NXP et le mélangeur du circuit intégré de l'article [1.34], discuté au chapitre 1. Par conséquent, cet article présente un excellent point de repère pour l'optimisation des performances du circuit. On remarque d'ailleurs que la chaîne de réception présentée utilise un balun LC en entrée du mélangeur pour réaliser la conversion « single ended » vers différentiel (Figure 73). Par ailleurs, le mélangeur de fréquences n'est pas précédé d'un LNA. La chaîne de distribution du signal de l'oscillateur local (OL) n'est pas étudiée ici. Nous nous basons donc sur les amplificateurs OL existants pour disposer du bon niveau de puissance. En effet, un des étages d'amplification OL pourra être doublé ou supprimé en fonction du niveau de signal OL requis par la nouvelle conception.

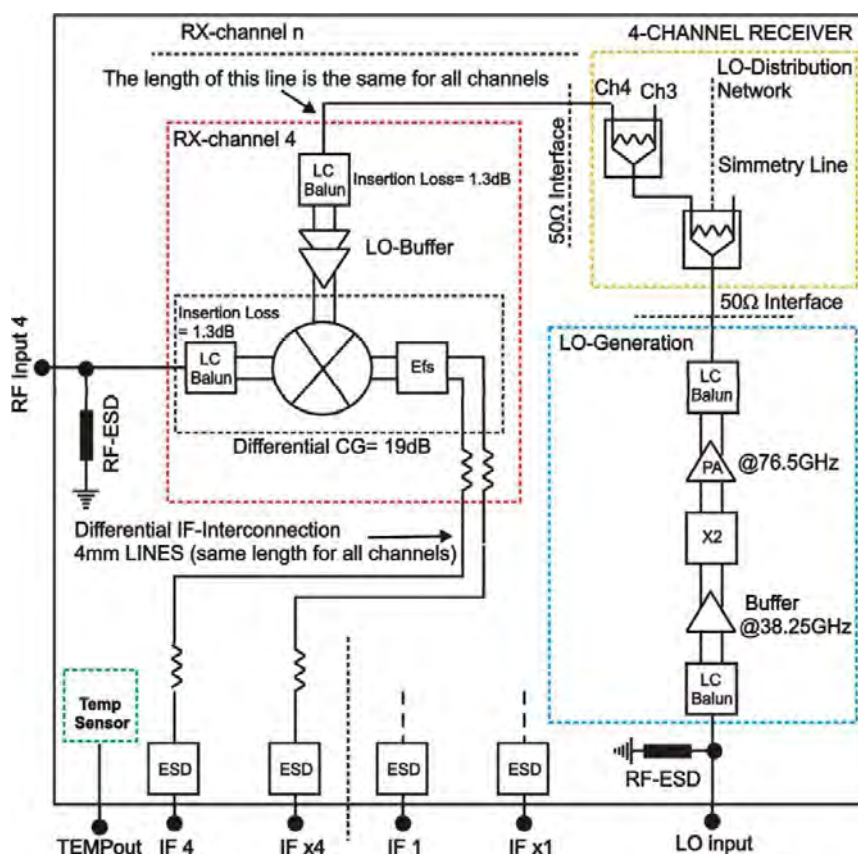


Figure 73 : Schéma de blocs de la chaîne de réception du radar monolithique NXP [1.34]

4.2.1.1 Spécifications en gain

La chaîne de réception doit produire un gain typique d'une valeur de 57 dB, avec une possibilité de programmation entre 26 dB et 57 dB. Le besoin de programmation est lié à la distance des cibles qui varie entre quelques dizaines de centimètres et quelques centaines de mètres. Le produit NXP utilise deux amplificateurs en bande de base programmables en tension (VGA pour Variable Gain Amplifier)

cascadés, avec des gains maximaux respectifs de 22 dB et de 16 dB. On en déduit que le mélangeur dispose d'un gain de 19 dB. En effet, dans une chaîne de réception radar, il n'y a pas de contrainte particulière sur le gain du mélangeur, étant donné que les VGA présentent un réglage du gain dans une dynamique importante et peuvent donc compenser une faible valeur de gain pour le mélangeur. Cependant, un compromis doit être réalisé car, d'une part, le gain du mélangeur doit être le plus grand possible pour minimiser la contribution en bruit des VGA, et d'autre part, ce gain est limité par la compression en puissance du mélangeur. La puissance en sortie ne doit pas être trop élevée pour garder un point de compression à 1 dB élevé. On considère donc pour le moment une spécification sur le gain de conversion du mélangeur à 19 dB et on vérifiera par la suite si celle-ci est compatible avec les spécifications de compression et de bruit. Les VGA sont en charge de détecter le niveau de puissance à leurs sorties et d'ajuster leurs gains respectifs en conséquence afin de présenter au convertisseur analogique-numérique un signal sans distorsion et d'amplitude suffisante.

4.2.1.2 Spécifications pour le point de compression

Revenons maintenant sur la spécification de point de compression en entrée de la chaîne de réception. Elle se situe autour de -3 dBm pour les radars embarqués dans l'automobile alors qu'elle est bien inférieure pour les systèmes de télécommunications mobiles. Cette différence a un impact non négligeable sur la conception. En effet, le niveau de compression en sortie du premier étage de la chaîne de réception est limité par la tension d'alimentation et la technologie du Silicium. Par conséquent, un niveau de compression élevé en entrée implique un faible gain dans ce premier étage [4.1].

Avec l'architecture du mélangeur présenté dans [1.34], qui sera détaillée plus loin, et une tension d'alimentation de 3.3 V, il est difficile d'obtenir une fréquence intermédiaire d'une tension pic supérieure à 1.3 V, comme le montre la Figure 74. En effet, le mode commun se situe autour de 2 V en conditions typiques en sortie du mélangeur, limitant ainsi la tension de sortie du signal en bande de base. Sur la base d'un calcul simplifié, cette tension pic de 1.3 V correspond à une puissance de 12 dBm. L'utilisation d'un mélangeur doublement équilibré augmente cette puissance de 6 dB puisque la sortie bande de base est alors différentielle. Dans cette configuration, il est donc possible d'avoir en théorie une puissance maximale à la sortie du mélangeur avoisinant 18 dBm. Avec le gain de conversion à 19 dB et une compression à 1 dB en entrée d'une valeur de -3 dBm, la puissance maximale en sortie du mélangeur est de 15 dBm. Grâce à ce calcul élémentaire, on peut déduire qu'on dispose d'une marge de 3 dB pour une conception avec cette technologie, mais d'autres paramètres intervenant dans la compression ont été négligés et les valeurs de tension collecteur-émetteur (V_{ce}) utilisées pour le calcul sont approximatives. Cependant, cette évaluation simple permet de vérifier la consistance du compromis spécifié entre le gain de conversion et la compression en entrée. On peut ainsi conclure que les performances de gain de conversion et de point de compression décrites par [1.34] semblent convenir pour ce compromis.

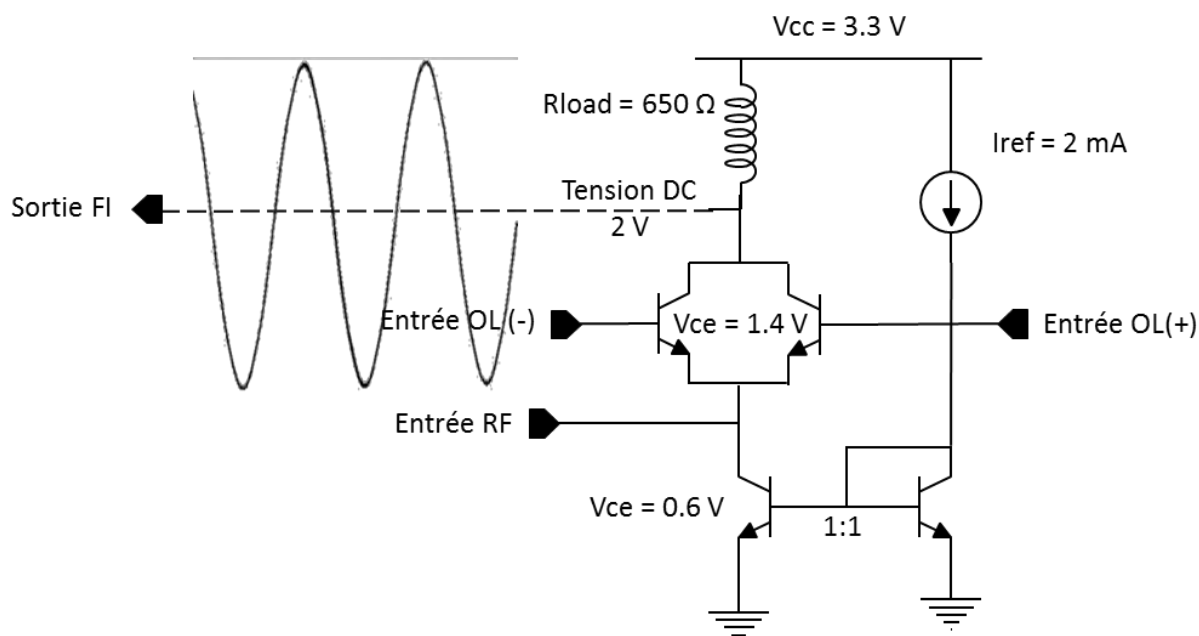


Figure 74 : schéma simplifié de la polarisation du mélangeur sans étage de gain

Enfin, on ne s'intéresse pas ici à la compression des VGA puisque cette spécification n'impacte pas la conception du mélangeur. Les VGA et les filtres passe-bas présentent un gain programmable et leur compression dépend donc des scénarii d'utilisation. Ce paramètre n'est donc pas pertinent pour cette étude.

4.2.1.3 Spécifications en bruit :

Le facteur de bruit (NF pour Noise Figure) de la chaîne de réception doit rester inférieur à 14 dB dans la gamme de températures allant de -40° à 125° C. Pour un radar automobile, il s'agit d'un des trois paramètres les plus critiques, les deux autres étant la puissance de sortie du PA et le bruit de phase de l'oscillateur local, comme décrit précédemment.

Il convient de noter que toutes les valeurs de NF qui seront données ici sont basées sur la convention « Bande Latérale Unique » (ou SSB pour single-sideband).

Le facteur de bruit total d'une chaîne est généralement calculé en utilisant la classique formule de Friis. Ce calcul n'est cependant possible que si on dispose du NF des différents étages. Il est très commun d'utiliser le NF aux fréquences millimétriques et RF pour caractériser les dispositifs en bruit. Or, pour les amplificateurs à basse fréquence, c'est plutôt la tension de bruit en entrée qui est utilisée pour déterminer le bruit ajouté par le bloc. Nous revenons donc à la définition du NF à partir des SNR en entrée et en sortie [4.2], la caractérisation en bruit de la chaîne complète étant plus naturelle à l'aide du formalisme du NF. Cette formulation se basera sur [4.2].

Par souci de simplification, les amplificateurs et les filtres bande de base sont considérés comme un seul VGA avec deux étages d'amplification et un gain total de 38 dB. De même, nous nous limitons au bruit

à 1 MHz en négligeant l'effet des filtres passe-bas à cette fréquence. La tension de bruit en entrée de ce VGA est nommée (e_{VGA}).

On considère le facteur de bruit linéaire F lié au NF par l'équation (23).

$$NF = 10 \log(F) \quad (23)$$

A la sortie du VGA, le facteur de bruit linéaire de la chaîne de réception (F_{RX}) peut s'écrire comme suit (24) :

$$F_{RX} = \frac{SNR_{in}}{SNR_{out}} = \frac{\frac{S_{in}}{N_{in}}}{\frac{S_{out}}{N_{out}}} = \frac{\frac{S_{in}}{N_{in}}}{\frac{S_{in} \times G_{mix} \times G_{VGA}}{N_{out}}} = \frac{N_{out}}{G_{mix} \times G_{VGA} \times N_{in}} \quad (24)$$

Avec N_{in} (W) la puissance de bruit totale en entrée de la chaîne, N_{out} (W) la puissance de bruit totale en sortie de la chaîne, G_{mix} le gain linéaire du mélangeur de fréquences et G_{VGA} le gain linéaire du VGA.

Par ailleurs, on sait que :

$$N_{in} = k_B \times T \times BW \quad (25)$$

$$N_{out} = [(N_{in} + N_{mix}) \times G_{mix} + N_{VGA}] \times G_{VGA} \quad (26)$$

Avec k_B la constante de Boltzmann ($k_B = 1.38 \times 10^{-23}$ J/K), T la température en Kelvin (K), BW (Hz) la bande passante d'intérêt, N_{mix} et N_{VGA} les puissances de bruit en entrée du mélangeur et du VGA, respectivement.

On obtient alors :

$$F_{RX} = 1 + \frac{N_{mix}}{N_{in}} + \frac{N_{VGA}}{N_{in} \times G_{mix}} \quad (27)$$

Le facteur de bruit du mélangeur de fréquences (F_{mix}) peut être rapidement obtenu à partir de (27) en considérant que le VGA est un élément sans pertes (c'est à dire $N_{VGA} = 0$ et $F_{mix} = F_{RX}$). Dans ces conditions, il est possible d'écrire :

$$F_{mix} = 1 + \frac{N_{mix}}{N_{in}} \quad (28)$$

Et donc :

$$F_{RX} = F_{mix} + \frac{N_{VGA}}{N_{in} \times G_{mix}} \quad (29)$$

Comme la puissance N_{VGA} peut être obtenue à partir de la tension de bruit en entrée du VGA (e_{VGA}) en utilisant :

$$N_{VGA} = \frac{(e_{VGA})^2}{Z_S} \times BW \quad (30)$$

L'équation permettant de calculer le facteur de bruit de la chaîne de réception (F_{RX}) à partir du bruit du mélangeur (F_{mix}) et de son gain (G_{mix}), ainsi que de la tension de bruit en entrée du VGA (e_{VGA}) est :

$$F_{RX} = F_{mix} + \frac{(e_{VGA})^2}{k_B \times T \times Z_S \times G_{mix}} \quad (31)$$

Cette équation est très utile pour rapidement identifier la contribution en bruit des amplificateurs en bande de base lorsque leur spécification est définie à partir de la tension de bruit en entrée.

Pour notre étude, on considère que l'étage VGA possède une tension de bruit en entrée maximale de 10 nV/ $\sqrt{\text{Hz}}$ à 1 MHz. Cette tension n'est pas l'état de l'art en matière de bruit mais présente un bon compromis entre la programmation des gains et des fréquences de coupure du VGA, la contrainte d'intégration et de surface, et enfin, la compression en entrée ainsi que le bruit. Des amplificateurs en bande de base avec de meilleures caractéristiques en bruit existent [4.3], mais pour lesquels on ne dispose pas d'information sur leur comportement en relation avec les autres spécifications du radar.

Le facteur de bruit de la chaîne doit rester inférieur à 14 dB. On utilise l'équation (31) pour connaître la spécification du NF SSB maximal du mélangeur de fréquences : avec une tension de bruit en entrée du VGA de 10 nV/ $\sqrt{\text{Hz}}$ et un gain de conversion du mélangeur de 19 dB, le facteur de bruit (NF SSB) du mélangeur de fréquences ne doit pas dépasser 12.8 dB. Il est donc primordial de ne pas dépasser cette valeur de NF afin de garantir la spécification de 14 dB. Pour prendre en compte l'effet de la température, on spécifie un NF à température ambiante et aux conditions typiques à 11 dB pour la fréquence intermédiaire de 1 MHz.

4.2.1.4 Architecture de la chaîne de réception

Pour convenir à l'application radar pour l'automobile considérée ici, la chaîne de réception à concevoir (Figure 75) doit disposer d'un gain maximal de 57 dB, partagé entre le mélangeur (19 dB fixe) et les amplificateurs programmables en tension (38 dB au maximum). Le mélangeur doit avoir un point de compression à 1 dB en entrée autour de -3 dBm. Le facteur de bruit de la chaîne ne doit pas dépasser 14 dB à 1 MHz. Pour y arriver, le NF du mélangeur de fréquences doit rester inférieur à 11 dB en conditions typiques et la tension de bruit en entrée des amplificateurs en bande de base doit rester inférieure à 10 nV/ $\sqrt{\text{Hz}}$ à 1 MHz. Toutes les spécifications relatives au mélangeur sont résumées sur le Tableau 4.

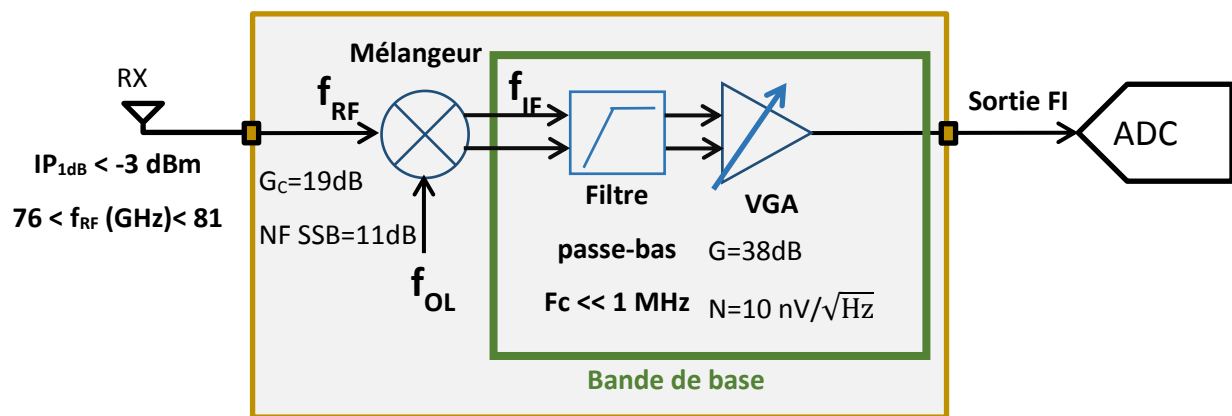


Figure 75 : diagramme de blocs d'une chaîne de réception à base d'un mélangeur actif

Paramètres	Bande de fréquences RF (GHz)	NF SSB (dB)	G _c (dB)	IP _{1dB} (dBm)	Tension d'alimentation (V)
Valeurs typiques	[76-81]	11	19	-3	3.3

Tableau 4 : spécifications du mélangeur de fréquences extraites des spécifications de la chaîne de réception

4.3 Mélangeur de fréquences à 77 GHz

Après avoir défini les principales spécifications du mélangeur, nous décrivons ici sa conception. Après l'étude des architectures et le choix de la topologie la mieux adaptée aux spécifications, nous présentons la conception du mélangeur puis la validation expérimentale de ses performances.

4.3.1 Étude des principales topologies de mélangeurs

Le panel des architectures des mélangeurs de fréquences est assez large mais la technologie monolithique utilisée est souvent la première à restreindre le choix, compte tenu des spécifications. On peut distinguer quelques grandes familles de mélangeurs sans être exhaustifs. La famille la plus connue est probablement celle des mélangeurs actifs doublement équilibrés [4.4], les cellules de Gilbert en particulier [4.5]. Ces mélangeurs sont réalisés en technologie bipolaire ou MOS. La deuxième famille de mélangeurs est celle des mélangeurs passifs MOS. Cette topologie est toujours accompagnée d'un LNA en amont afin de réduire la contribution du NF du mélangeur et compenser ses pertes de conversion [4.6]- [4.8]. D'autres topologies sont basées sur des diodes et notamment des diodes Schottky, connues pour avoir une fréquence de transition très élevée. Ces diodes sont généralement reportées sous forme de composants discrets et elles sont rajoutées dans de très rares cas aux technologies bipolaires et MOS. Des diodes Schottky très performantes ont été justement expérimentées en technologie BiCMOS 180 nm [4.9]. La technologie MOS 180 nm dont nous disposons apparaît trop limitée pour le développement de mélangeurs passifs, compte tenu de fréquences de transition des transistors de l'ordre de grandeur de la fréquence d'opération (Figure 76). D'autres travaux menés au LAAS sont revenus en détails sur ces

limitations et ont proposé des améliorations permettant d'opérer à des fréquences aux limites de la fréquence de transition, notamment en se basant sur des mélangeurs à sous-échantillonnage en technologie BiCMOS 130 nm [4.10]. Malheureusement, la fréquence de transition des MOS 180-nm est bien inférieure aux fréquences radar automobile 77-GHz. Enfin, il est possible de réaliser la conversion de fréquence à l'aide de mélangeurs en fonctionnement sous-harmoniques, basés sur un mélange de la fréquence RF avec le second harmonique de la fréquence OL [4.11], par exemple. L'avantage de cette solution est la réduction des contraintes sur le signal OL qui est à une fréquence moins élevée. Son plus grand inconvénient est le gain de conversion assez faible à cause de l'intervention de l'harmonique du signal de l'OL, non le fondamental, dans le mélange. Cela oblige à utiliser un LNA en amont, impactant directement la linéarité de la chaîne de réception. Des travaux réalisés en interne au sein de NXP ont confirmé cela.

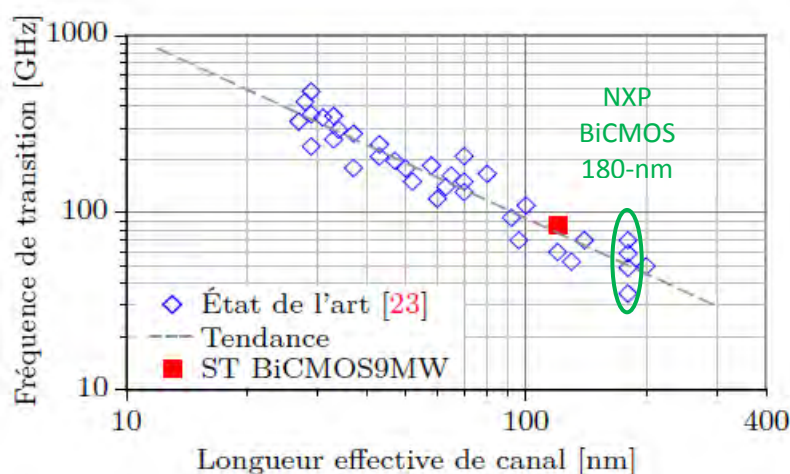


Figure 76 : Fréquences de transition en fonction de la longueur effective de canal des transistors NMOS [4.7]

4.3.2 Conception des mélangeurs retenus après la 1^{ère} étude

Après cette première étude, il s'avère que l'utilisation d'un mélangeur passif nécessite de passer à une technologie MOS avec un nœud plus avancé. Cette proposition sera investiguée en se basant sur une technologie BiCMOS 90-nm. On se penchera également sur la conception de mélangeurs à base de diodes Schottky afin de quantifier le NF qu'elles permettent d'atteindre. La dernière proposition est l'optimisation de la cellule de Gilbert, en incluant les changements proposés par Trotta [1.34] pour améliorer la linéarité : enlever l'étage de transconductance pour garder uniquement les transistors réalisant le mélange. On procédera à une étude plus détaillée de ces différentes propositions.

4.3.2.1 Mélangeur passif à base de diodes Schottky

Un mélangeur expérimental basé sur les diodes Schottky de [4.6] est développé. Il se base sur une architecture de mélangeur doublement équilibré et des transformateurs réalisant entre autres la fonction de conversion de « Single-Ended » vers différentiel (Figure 77).

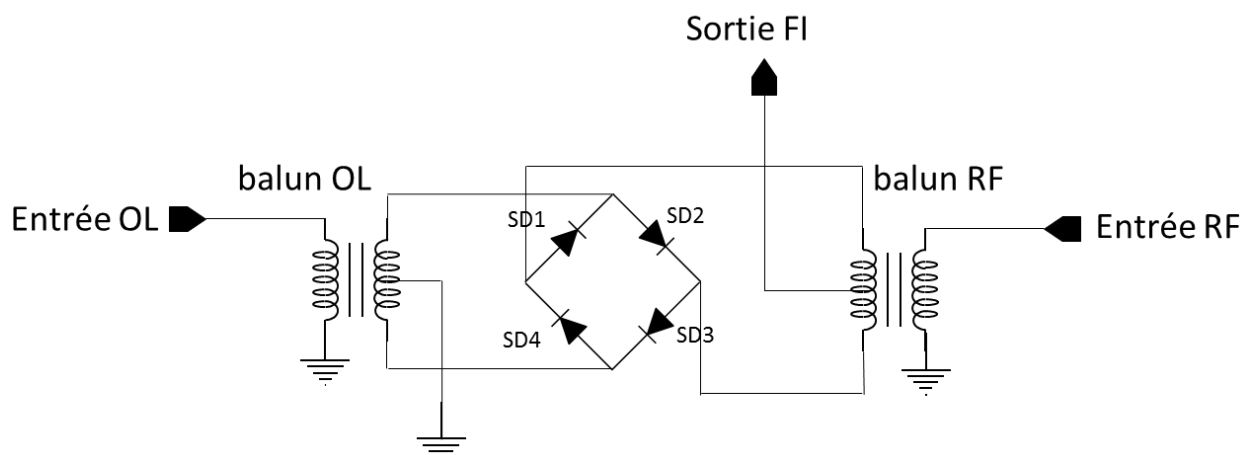


Figure 77 : schéma simplifié d'un mélangeur doublement équilibré à base de diodes Schottky et de baluns

Afin de réduire la contribution en bruit du mélangeur à base de diodes et d'améliorer le gain, un LNA est ajouté en amont du mélangeur (Figure 78). Le point de compression élevé du mélangeur permet d'utiliser un LNA tout en respectant la contrainte de linéarité de la chaîne de réception.

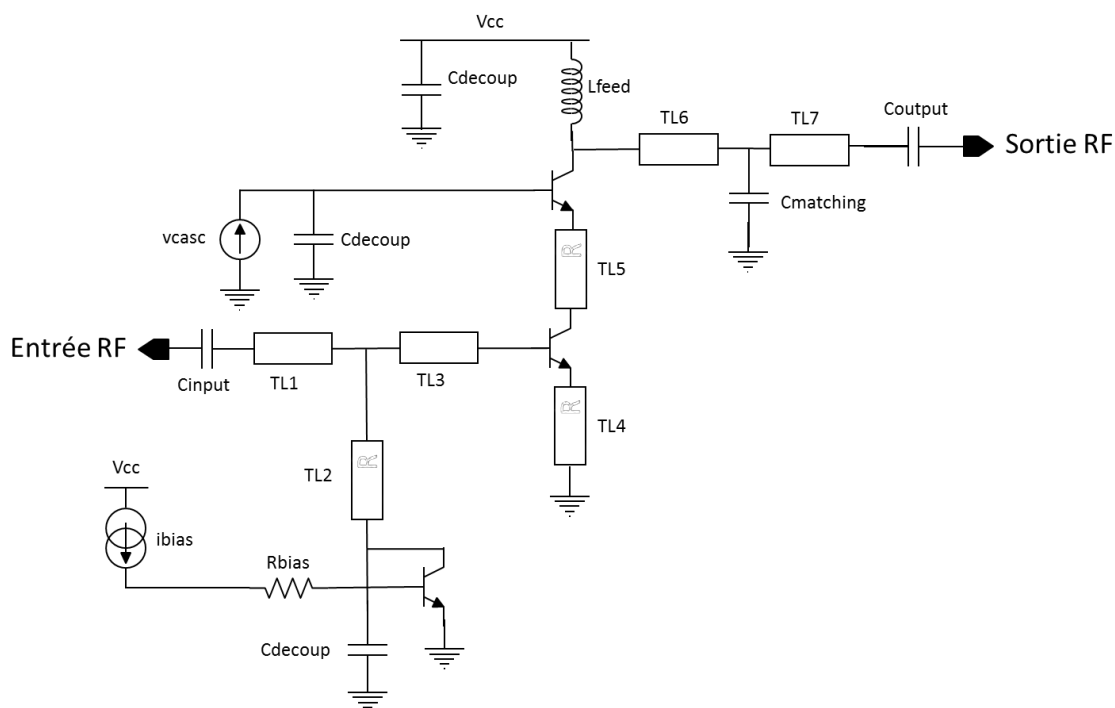


Figure 78 : LNA "Single-ended" en amont du mélangeur à diodes Schottky

Les deux blocs ont été conçus par la suite dans l'outil de simulation circuit ADS (Figure 79). Un transformateur, entre le LNA et le mélangeur, permet de polariser le LNA et de convertir le signal RF « single-ended » en signal différentiel. L'ensemble a été optimisé afin de réduire le bruit et maximiser le gain tout en respectant la contrainte sur la linéarité. Le dessin des masques a été réalisé à l'aide du logiciel Cadence Virtuoso (Figure 80).

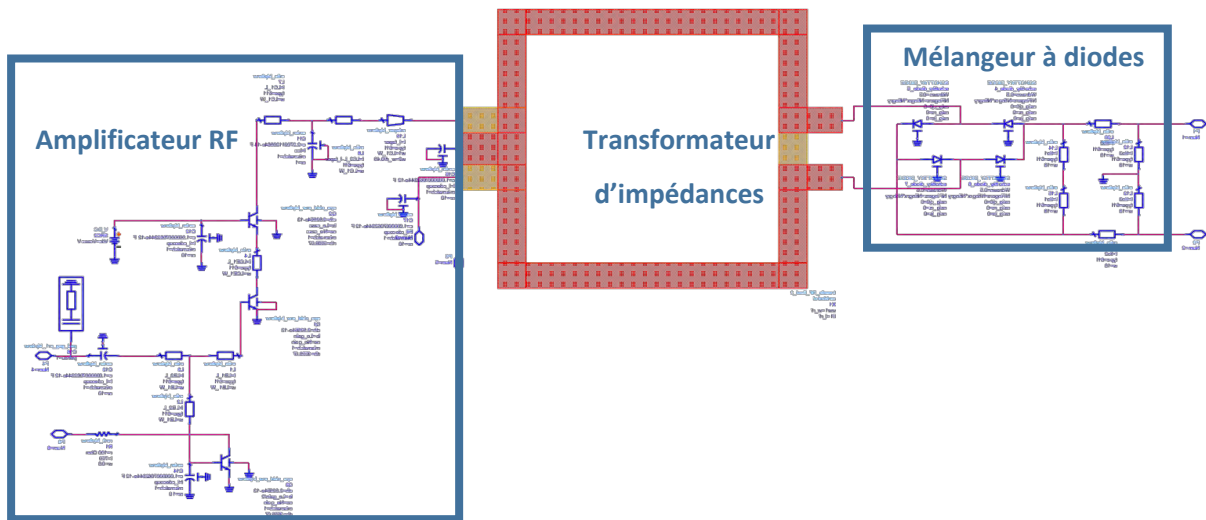


Figure 79 : schéma ADS du mélangeur à diodes Schottky avec le LNA

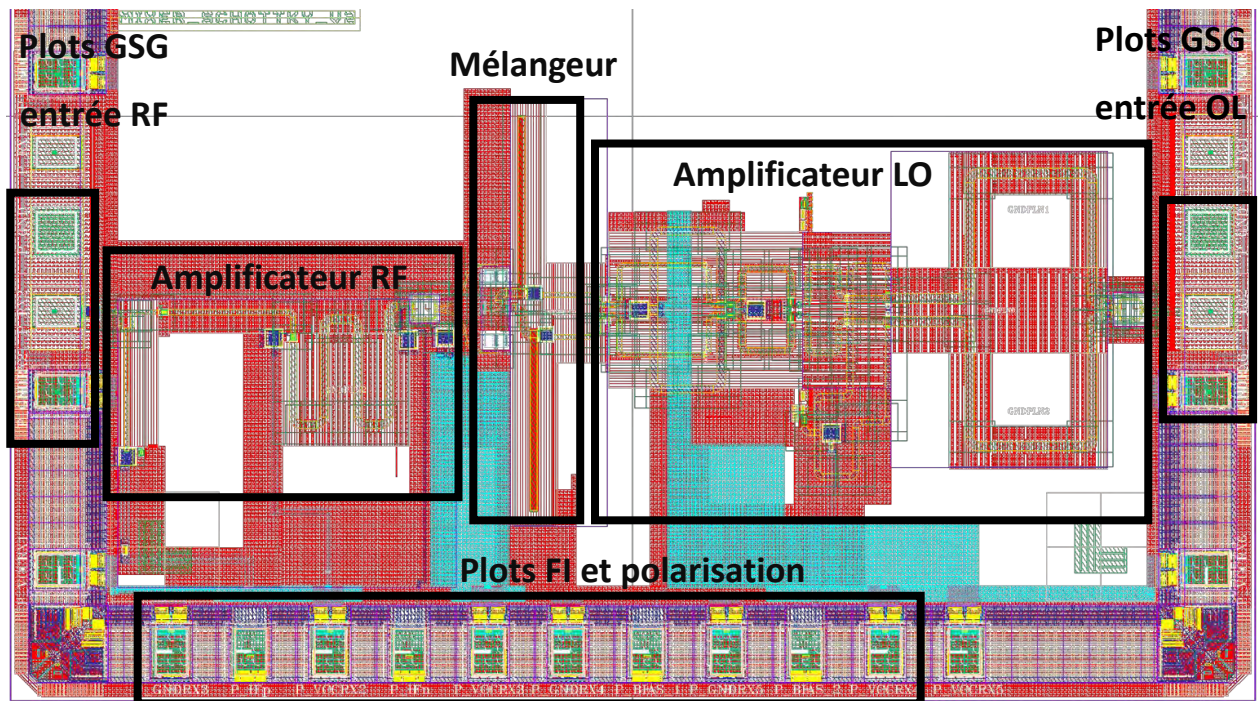


Figure 80 : dessin des masques du mélangeur à diodes Schottky avec le LNA

Comme reporté sur la Figure 81 (a), la chaîne de réception constituée du LNA et du mélangeur à diodes Schottky conduit à un gain de conversion de 3 dB (mixer_gain). Ce gain se décompose en 6 dB de pertes de conversion dans le mélangeur et 9 dB de gain dans le LNA (Gain_LNA sur la figure 81 (b)).

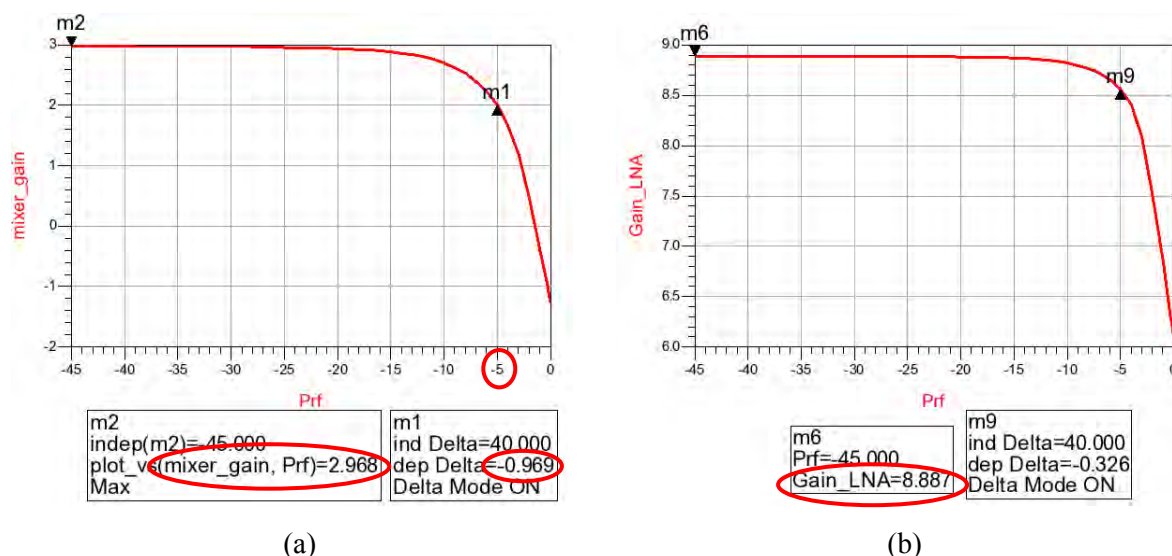


Figure 81 : résultats de gain et saturation du mélangeur avec le LNA (a) et du LNA seul (b)

Le facteur de bruit simulé pour l'ensemble {LNA, mélangeur} est de 9 dB à 1 MHz. Ces performances sont très bonnes mais le faible gain de conversion de la chaîne impose des contraintes importantes sur le bruit des blocs en bande de base. Le VGA a une tension de bruit en entrée de $10 \text{ nV}/\sqrt{\text{Hz}}$. Avec les performances décrites pour le mélangeur, l'ajout de ce VGA en aval remonte le bruit total de la chaîne à 24 dB. Cette dégradation très importante est liée au bruit très élevé du VGA, initialement prévu pour être placé devant un étage à 19 dB de gain. On remarque également que la sortie FI de ce mélangeur à diodes est « single-ended ». Il aurait été impossible de l'interconnecter directement au VGA sans ajouter un étage de conversion du signal FI en signal différentiel.

Afin de pallier ces problèmes, une étude d'architectures d'amplificateurs FI et de convertisseurs de « single-ended » vers différentiel a mené à des résultats innovants ayant fait l'objet d'un brevet [4.12]. La contribution d'un collègue de NXP, expert en conception en bande de base, a été précieuse pour la réalisation de ces travaux. Un amplificateur en bande de base très faible bruit, présentant des caractéristiques en bruit à l'état de l'art, avait été conçu et breveté par ce collègue. Il a été nommé « Boosted Follower » (BTF). Cet amplificateur, générant un gain fixe de 6 dB, est au cœur de l'architecture proposée. Deux étages ont été implémentés dans la chaîne, pour un gain de 12 dB, en aval d'un étage de conversion actif du signal en signal différentiel (Figure 82). Ce balun actif présente lui aussi un gain de 6 dB. L'ensemble a été optimisé en bruit. Par la suite, on appellera ce bloc LNA bande de base. La demande de brevet n'ayant pas encore été rendue publique, les détails de l'architecture ne peuvent pas être exposés d'avantage dans cette thèse.

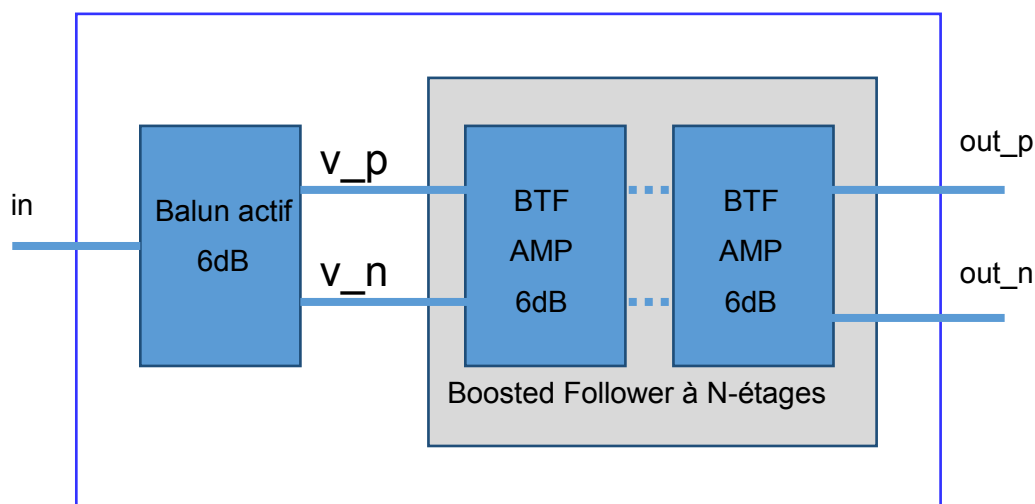


Figure 82 : schéma bloc du LNA en bande de base

Les simulations du bruit du LNA bande de base présenté confirment le très faible niveau de bruit généré par cet ensemble. En effet, la tension de bruit en entrée du système ne dépasse pas $1.4 \text{ nV}/\sqrt{\text{Hz}}$ à 1 MHz (Figure 83). La remontée de bruit à très faible fréquence résulte du bruit de scintillation, ou « Flicker Noise ». Le gain de ce bloc se situe autour de 18 dB.

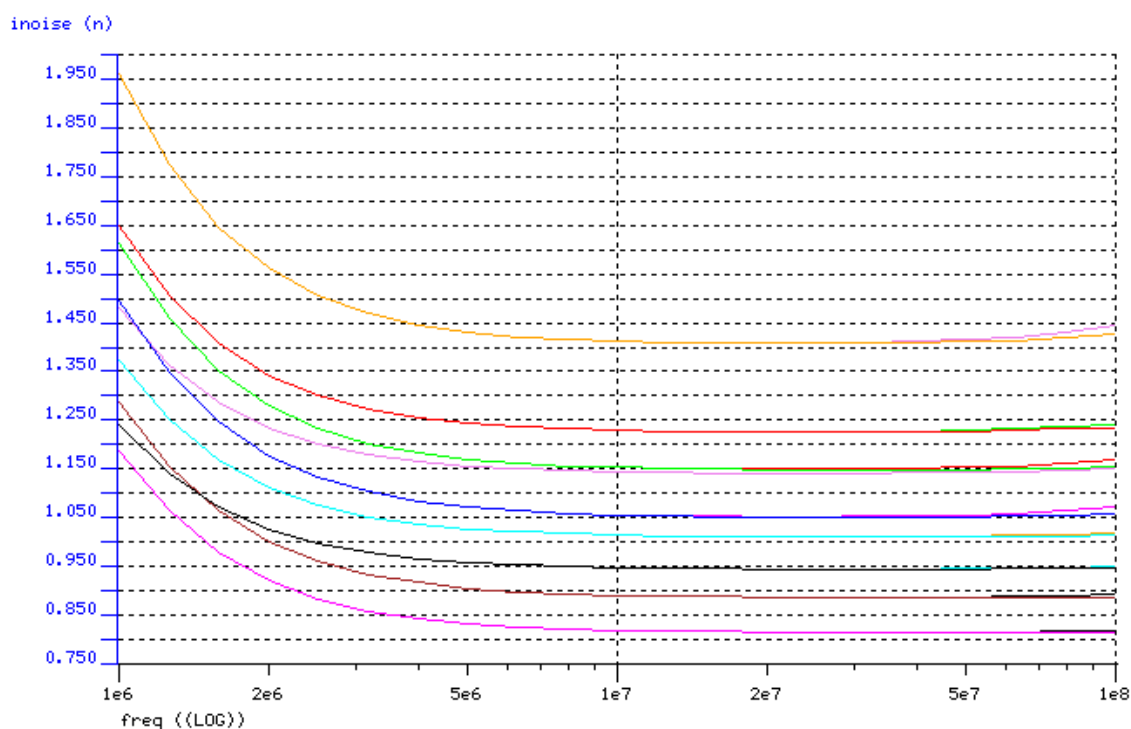


Figure 83 : tension de bruit en entrée du LNA en bande de base

On utilise la formule de calcul du bruit total de la chaîne (31) pour quantifier l'impact du LNA bande de base sur le NF de la chaîne de réception complète, lorsqu'il est placé entre le mélangeur et le VGA (Figure 84). La valeur du NF de la chaîne en sortie du LNA bande de base est de 11 dB. Malgré le très faible bruit de ce bloc, il rajoute 2 dB au bruit du mélangeur parce que le gain de l'ensemble LNA et mélangeur reste faible. Le gain de ce dernier ensemble ne peut pas être augmenté d'avantage à cause de

la contrainte de compression. En ajoutant le VGA à la chaîne de réception, le bruit total de la chaîne de réception est d'environ 12.2 dB.

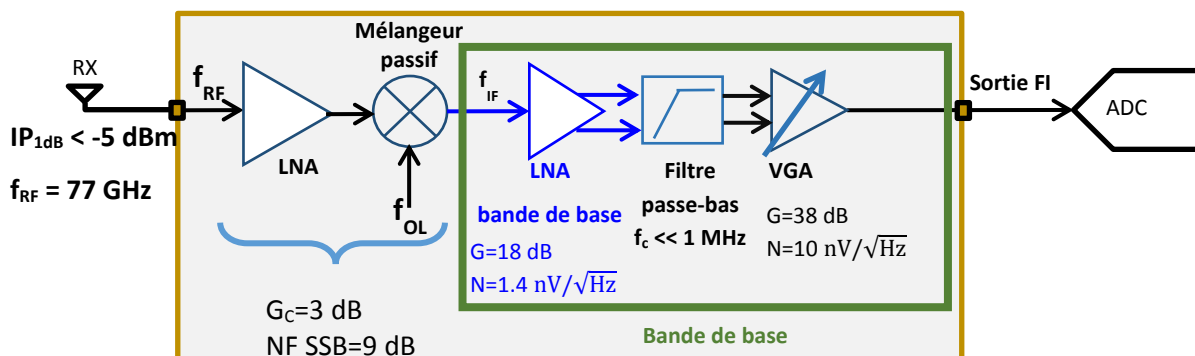


Figure 84 : Chaîne de réception incluant un LNA, un mélangeur passif, un amplificateur en bande de base faible bruit et un VGA

On peut conclure des performances simulées de ce mélangeur qu'il s'agit d'un bon candidat pour des applications radar automobile 77 GHz. Malheureusement, il n'apporte pas d'avantage majeur en termes de bruit par rapport au mélangeur actif présenté dans [1.34] et son aspect très innovant fait de lui un choix assez risqué pour une mise en production rapide. En effet, les diodes Schottky sont très rarement utilisées dans la technologie et on manque d'informations concernant leur robustesse en fabrication dans le cadre d'un mélangeur à 77 GHz. L'autre inconvénient concerne la linéarité du LNA bande de base qui est à ses limites avec une puissance en entrée à -5 dBm. Le filtre passe-bas devrait idéalement être placé en amont du LNA bande de base mais cela apporte des changements non négligeables au niveau de la conception de la chaîne. Ce choix d'architecture permet finalement d'atteindre les performances visées mais reste une solution risquée nécessitant potentiellement encore plusieurs itérations d'optimisation avant d'envisager la production en masse.

4.3.2.2 Mélangeur passif en BiCMOS 90-nm

Nous avons vu qu'il n'était pas possible de concevoir un mélangeur passif en MOS 180-nm à cause de sa fréquence de transition très inférieure à la fréquence d'opération radar 77 GHz. Le MOS 90-nm est à ses limites à cette fréquence mais il est susceptible de bien fonctionner en tant qu'interrupteur. Pour obtenir de meilleures performances qu'avec un mélangeur actif, un LNA est ajouté en amont du mélangeur et un amplificateur faible bruit en bande de base est ajouté en aval. Cette configuration est similaire à celle décrite pour le mélangeur à base de diodes Schottky.

On s'intéresse ici particulièrement au mélangeur. La sortie haute impédance de la fréquence intermédiaire du mélangeur passif à MOS permet d'obtenir un gain en tension élevé en sortie. Le mélangeur apporte donc un « gain » malgré sa structure passive. Cependant, la compression en entrée de ce mélangeur est limitée par les transistors et le LNA ne doit donc pas introduire un gain trop important. En effet, le gain du LNA ne doit pas dépasser 5 dB en conditions typiques. Pour cette raison,

on utilise uniquement une paire différentielle pour la conception du LNA contrairement aux structures classiques à montage cascade. Un amplificateur OL similaire à l'amplificateur RF est utilisé pour amplifier le signal et réduire les contraintes en puissance sur les générateurs externes d'OL. Plusieurs transformateurs sont introduits dans la structure pour réaliser la polarisation, le découplage DC et la conversion en mode différentiel (Figures 85 et 86).

Une méthodologie de co-conception des amplificateurs et du mélangeur a été développée. Elle est basée sur l'utilisation de transformateurs paramétrables en largeur et en diamètre afin d'optimiser les adaptations inter-étages. Cette approche se base sur une optimisation des performances au niveau globale en même temps qu'au niveau de chacune des fonctions.

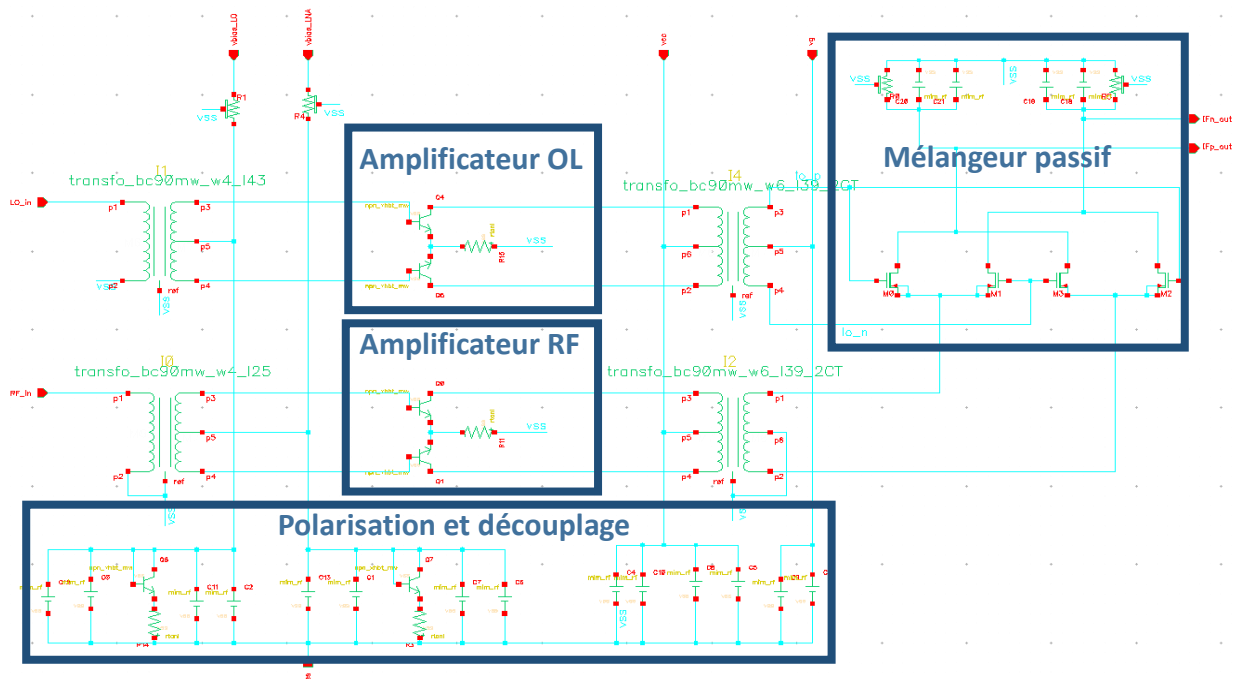


Figure 85 : schéma du mélangeur passif à MOS intégrant les amplificateurs des interfaces OL et RF

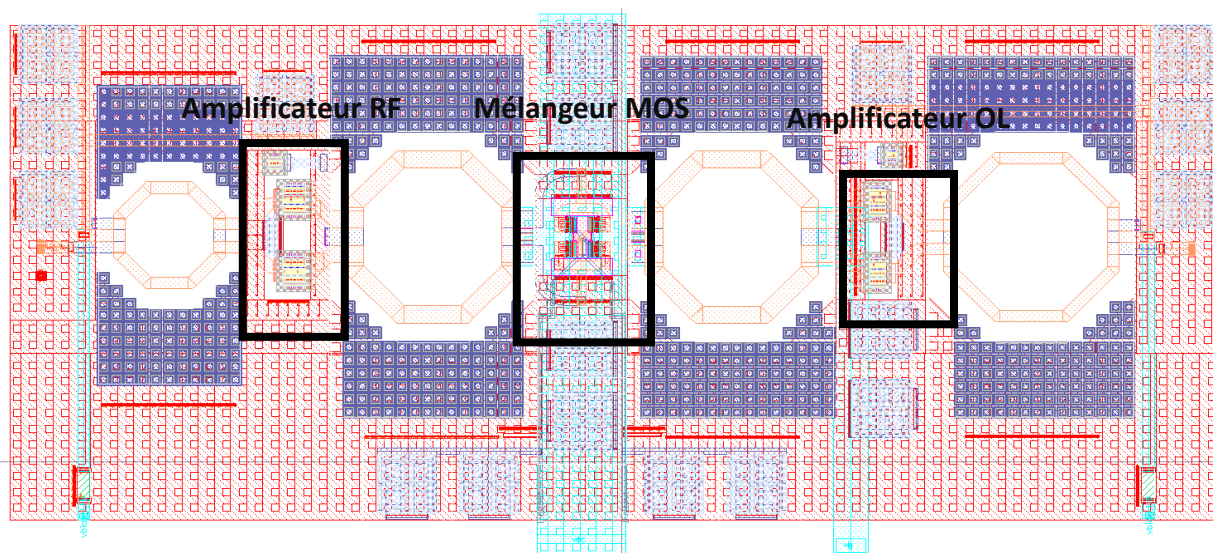


Figure 86 : dessin des masques du mélangeur passif à MOS intégrant les amplificateurs des interfaces OL et RF

Cette architecture a beaucoup de points de ressemblances avec un mélangeur actif, la paire différentielle en bipolaire (Figure 87 (a)) jouant la fonction d'étage de gain et les 4 transistors MOS (Figure 87 (b)) jouant la fonction de mélange. La conception d'une telle architecture est possible ici grâce à la faible longueur de grille des MOS (90 nm) et la présence d'une technologie bipolaire performante en gain.

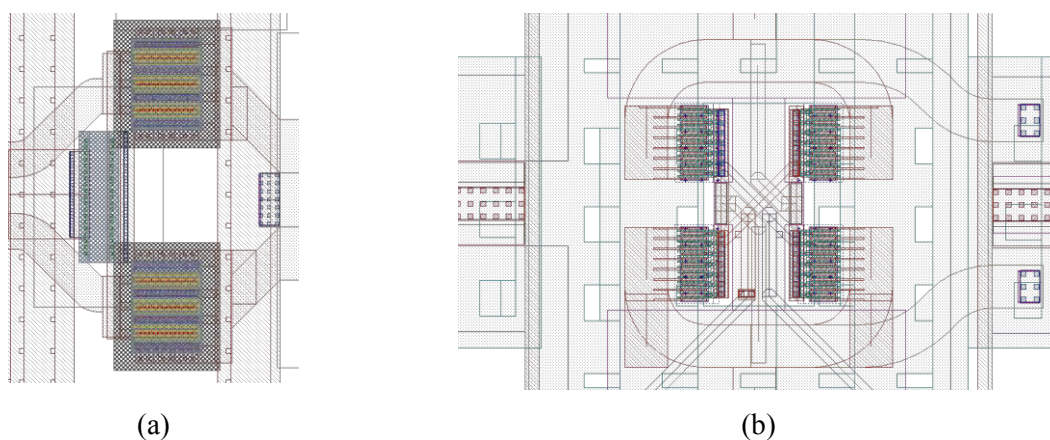
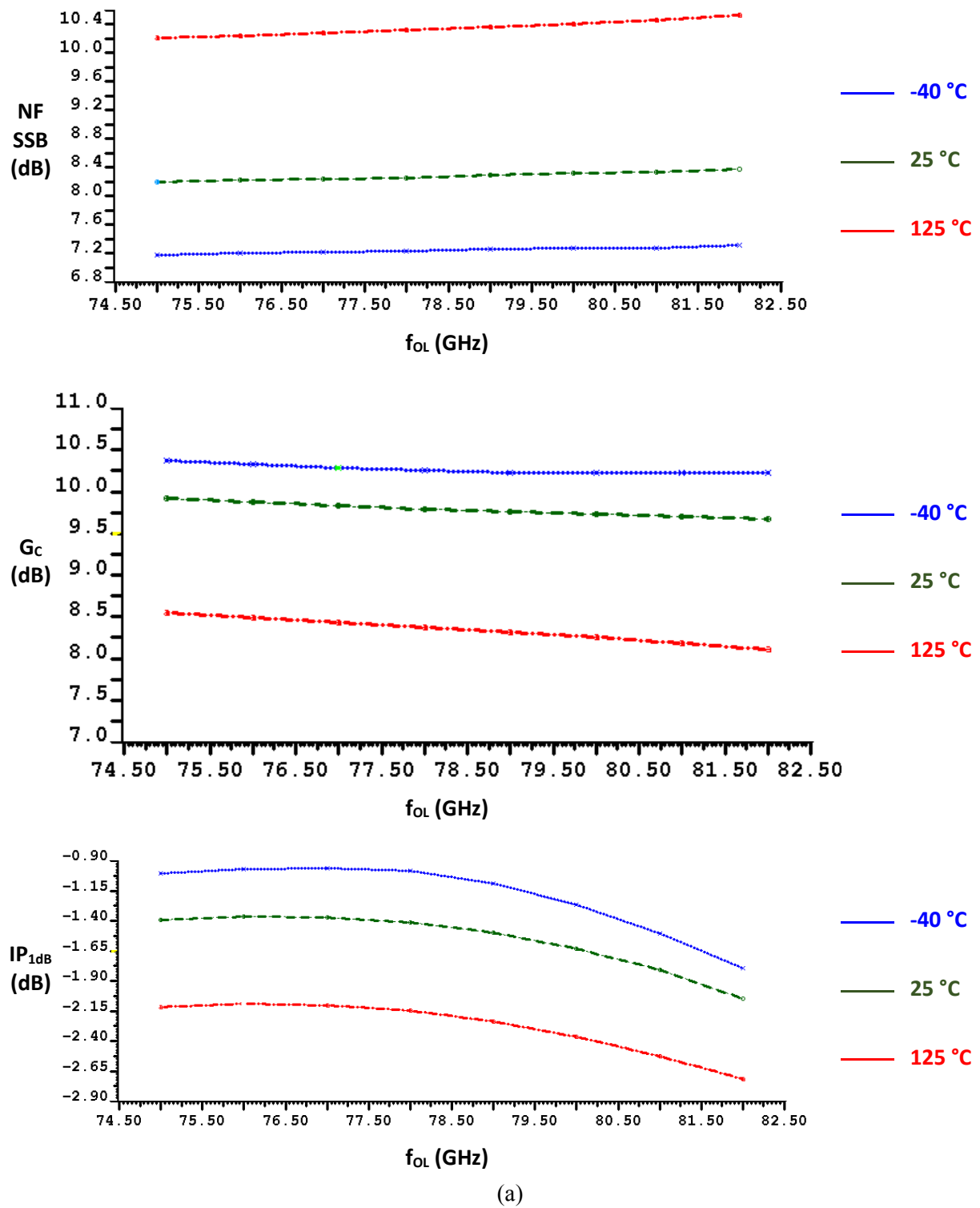


Figure 87 : dessins de masque des transistors MOS du mélangeur passif (a) et des transistors bipolaires de l'amplificateur RF (b)

Cette approche de conception a porté ses fruits puisqu'elle a permis l'obtention de performances à l'état de l'art pour un mélangeur passif à base MOS à 90 nm (Figure 88). Les simulations de bruit, de gain de conversion et de compression en entrée ont été réalisées pour toute la bande de fréquence radar [76-81] GHz et à 3 températures : -40 °C, 25 °C et 150 °C. Les résultats de bruit obtenus sont excellents puisque le facteur de bruit SSB ne dépasse pas 8.5 dB à 25° C avec un gain de conversion autour de 10 dB dans les mêmes conditions. Par ailleurs, un très bon niveau de puissance de compression est obtenu. La puissance de compression en entrée est supérieure à -3 dBm dans toutes les conditions de simulations.



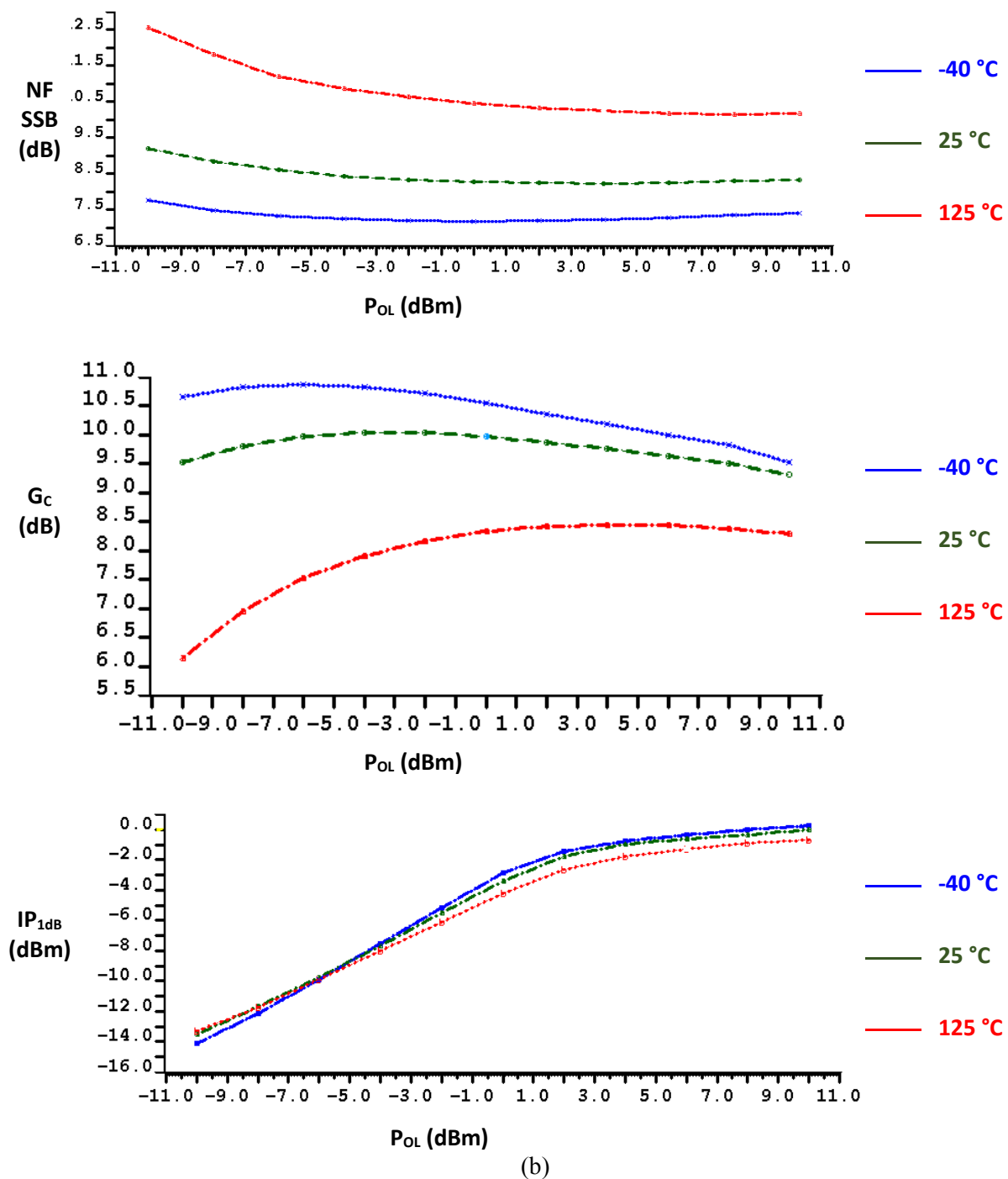


Figure 88 : résultats de simulation du facteur de bruit SSB (NF SSB), du gain de conversion (G_C) et puissance de compression à 1 dB en entrée (IP_{1dB}) du mélangeur MOS et les amplificateurs développés en fonction de f_{OL} à $P_{OL} = 2$ dBm (a) et en fonction de P_{OL} à $f_{OL} = 77$ GHz (b)

L'architecture choisie pour la conception de ce mélangeur a permis d'obtenir de très bons résultats en termes de gain et de bruit. Afin de prendre en compte l'effet des étages d'amplification en bande de base, on considère des conditions typiques où le gain de conversion et le bruit du mélangeur sont respectivement à 10 dB et à 8.2 dB. L'implémentation des deux étages de « Boosted Follower » rajoute 12 dB de gain et 1.5 nV/ $\sqrt{\text{Hz}}$ de tension de bruit en entrée. Le facteur de bruit (NF SSB) total de cette chaîne se situe ainsi autour de 9.1 dB et son gain est de l'ordre de 22 dB. Ces résultats sont excellents

et surpassent les performances du mélangeur à base de diodes Schottky, principalement grâce au gain en tension très élevé en sortie du mélangeur. Ce mélangeur est un excellent candidat pour les produits radar en technologie BiCMOS-90nm.

4.3.2.3 Mélangeur actif à base de cellule de Gilbert optimisée pour la linéarité

Nous étudions maintenant l'architecture du mélangeur actif basée sur une cellule de Gilbert, discuté lors de la phase de spécifications. Un mélangeur à base de cellule de Gilbert est un mélangeur doublement équilibré doté d'un étage de gain différentiel qui effectue une conversion tension-courant et apporte un gain (g_m) au courant de sortie. Cet étage est suivi de 4 transistors réalisant la fonction de mélange et une conversion courant-tension en sortie des transistors grâce aux résistances (R_L) de charge (Figure 89 (a)). Afin de minimiser la variation des performances avec la température et les dispersions du processus de fabrication, et d'améliorer la linéarité, [1.34] propose de supprimer les transistors de l'étage de gain (Figure 89 (b)). Cette modification permet également de réduire considérablement la consommation du mélangeur. Les transistors de mélange fonctionnent en tant qu'interrupteurs. Le courant total parcourant le bloc ne dépasse pas 5 mA, réduisant la consommation de 80%.

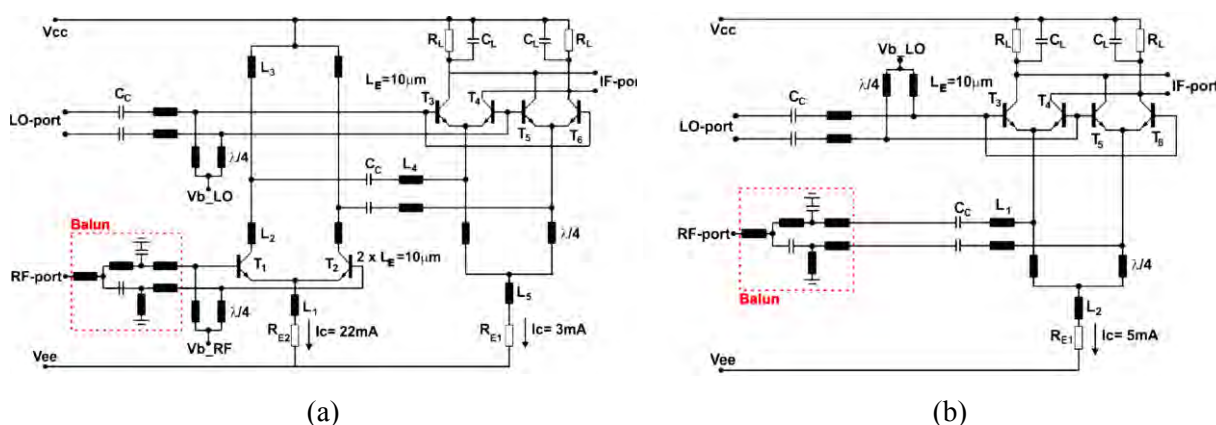


Figure 89 : topologie de la cellule de Gilbert classique (a) et topologie optimisée en linéarité (b) [1.34]

Les résultats de mesure de ces deux topologies montrent un meilleur IP3 pour la topologie sans étage de gain (Figure 90), comme attendu, et un gain de conversion autour de 19 dB dans les deux cas. Le seul inconvénient de cette topologie sans étage de gain est l'isolation de l'OL vers la RF qui est réduite.

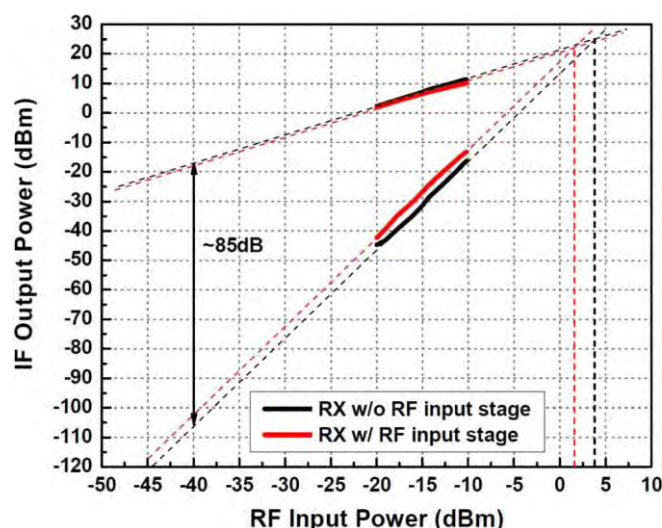


Figure 90 : résultats de mesures IP3 2-tons pour le mélangeur avec étage de gain (rouge) et sans (noir) [1.34]

L'approche suivie dans le cadre de cette thèse consiste à tirer profit du travail de modélisation EM réalisé au niveau circuit et exploiter la topologie proposée par [1.34].

En effet, on a évoqué au début de ce chapitre que cette architecture utilise un balun LC présentant 1.3 dB de pertes en entrée RF pour la réalisation de la conversion du signal « single-ended » en signal différentiel. De plus, des lignes quart d'onde ($\lambda/4$) et des capacités en série sont utilisées pour amener le signal continu de polarisation ou bloquer ce signal continu. Ces deux éléments introduisent des pertes supplémentaires de l'ordre de 1 dB. Lors du chapitre 2, un transformateur d'impédance a été conçu et caractérisé. Cet élément est en mesure d'effectuer ici les trois opérations simultanément. En effet, il réalise la conversion du signal en signal différentiel, alors que chaque enroulement permet d'acheminer le signal continu de polarisation et que d'un enroulement à l'autre ce signal continu est bloqué. De plus, comme nous l'avons montré, il introduit des pertes ne dépassant pas 1 dB entre l'accès « single-ended » et l'accès différentiel. Par conséquent, le remplacement du balun LC ainsi que des capacités et des lignes $\lambda/4$ uniquement par ce transformateur d'impédance permet, pour des conditions d'adaptation similaires, de réduire les pertes en entrée de 1.3 dB. Par ailleurs, on sait que toutes les pertes se produisant sur l'entrée RF sont directement traduites en NF (formule de Friis). L'intégration de ce transformateur peut ainsi réduire le NF de plus de 1 dB. Cependant, le défi de l'adaptation reste un point critique. Un balun LC permet de déterminer la transformation d'impédance avec quelques degrés de liberté alors que le transformateur réalise une transformation d'impédance fixe et il faut ajuster l'impédance, si besoin, en introduisant d'autres éléments passifs.

Un mélangeur à base de cellule de Gilbert sans étage de gain a été alors conçu en implémentant un transformateur sur l'entrée RF. Toute la conception a été revue et des modifications ont été apportées à l'adaptation sur les deux entrées RF et OL, à la charge en bande de base et aux conditions de polarisation. Une polarisation en courant au niveau des transistors bipolaires du mélangeur a notamment été rajoutée pour améliorer la stabilité du courant. Les simulations ont été réalisées à l'aide de Cadence Virtuoso (Figure 91). Un grand intérêt a été porté aux interconnexions entre les transistors réalisant les mélanges.

Une parfaite symétrie au niveau de ces accès est nécessaire pour disposer d'une bonne isolation entre les voies RF et OL (figure 92).

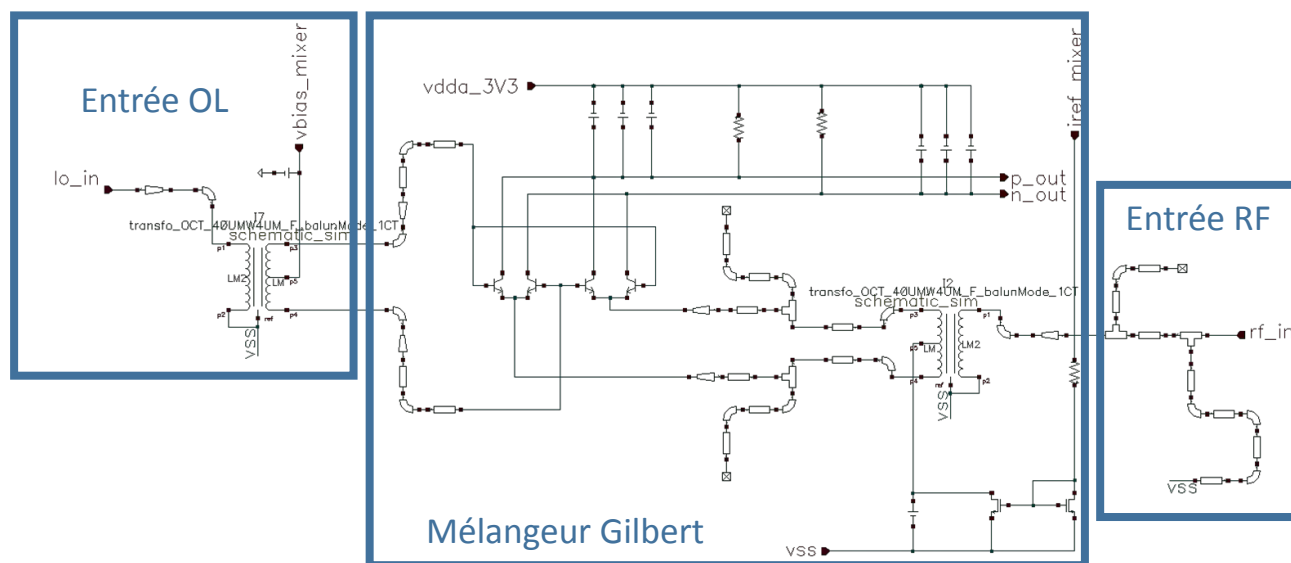


Figure 91 : schéma du mélangeur actif à forte linéarité

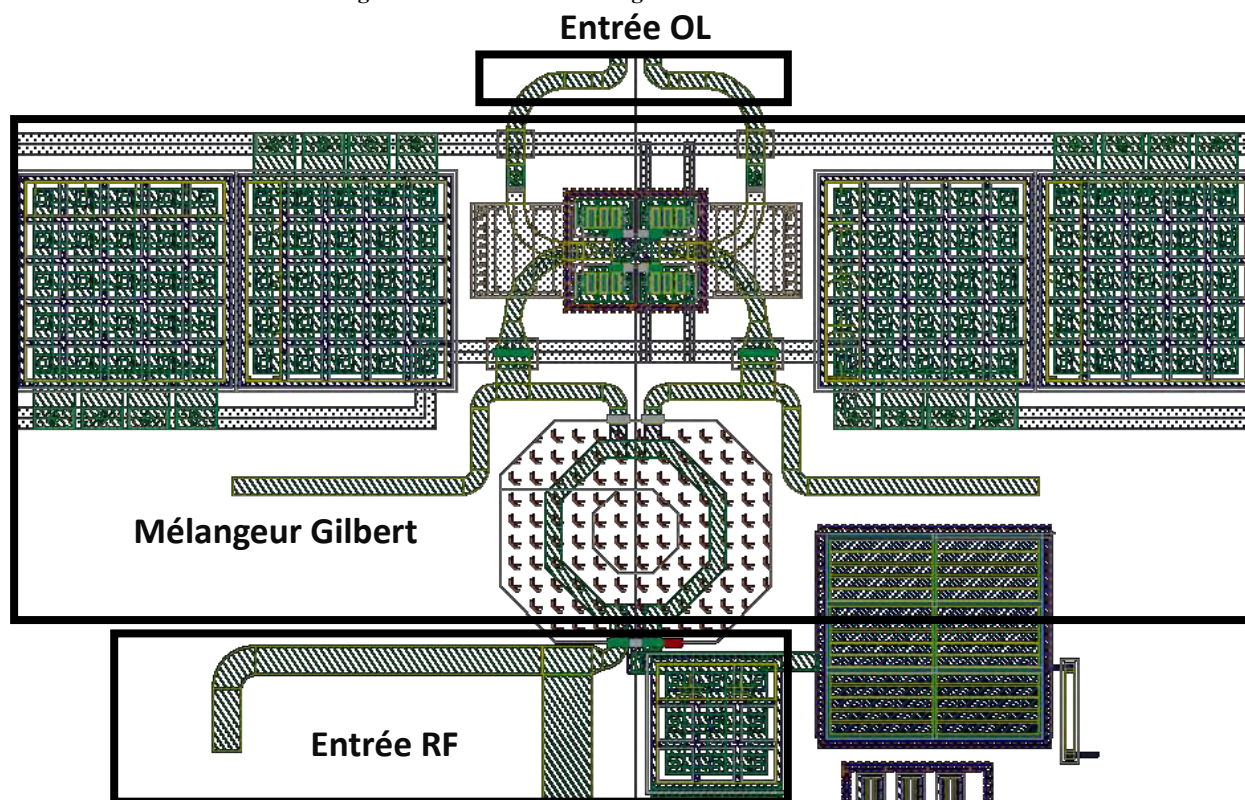
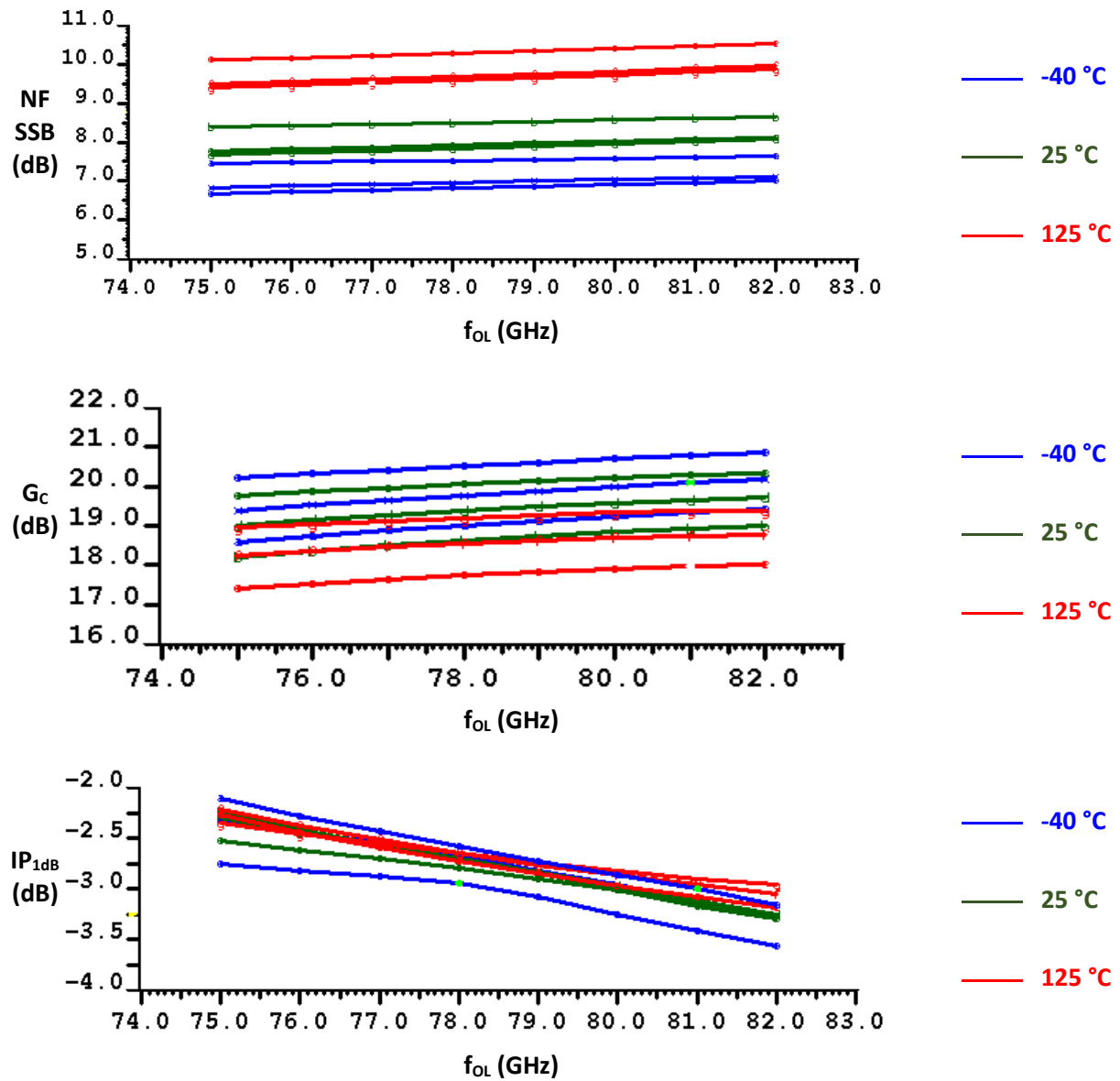


Figure 92 : dessin de masque du mélangeur actif à forte linéarité

Les simulations de ce mélangeur sont très prometteuses et confirment le grand intérêt, pressenti, de l'intégration d'un transformateur d'impédance à l'entrée de la chaîne de réception. Nous n'avons pas introduit ici d'amplificateur sur la voie OL car une chaîne d'amplification OL a déjà été conçue et caractérisée lors d'une précédente réalisation. Cette chaîne pourra être insérée sans difficultés dans la chaîne de réception du radar. Le mélangeur conçu a été simulé dans la bande [75-82] GHz et pour différents niveaux de puissance OL. Les courbes de la figure 93 sont obtenues pour les trois températures

-40° C, 25° C et 150° C, ainsi que pour les trois cas de processus de fabrication : typique, pire cas et meilleur cas.



(a)

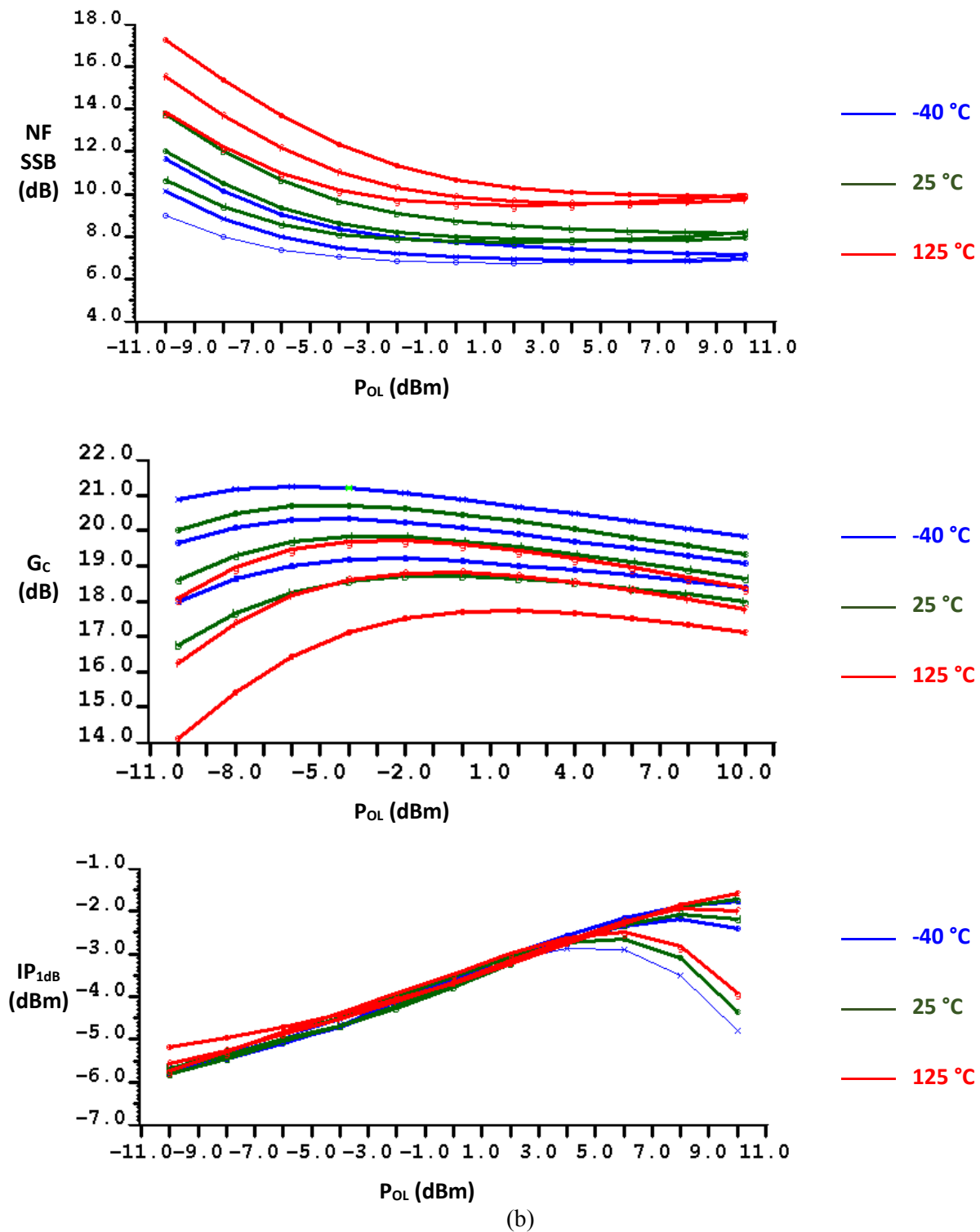


Figure 93 : résultats de simulation du facteur de bruit SSB (NF SSB), du gain de conversion (G_c) et puissance de compression à 1 dB en entrée (IP_{1dB}) du mélangeur à base de cellule de Gilbert, pour les trois cas de processus de fabrication : typique, pire cas et meilleur cas, en fonction de f_{OL} à $P_{OL} = 4$ dBm (a) et en fonction de P_{OL} à $f_{OL} = 77$ GHz (b)

Comme il peut être relevé sur la Figure 93, de très bonnes performances sont obtenues grâce à l'intégration d'un transformateur d'impédances à cette topologie. Un niveau de puissance assez élevé est requis au niveau de l'OL mais en réalité, c'est la tension OL qui réalise le mélange et l'optimisation

du dernier étage d'amplification OL permettra d'avoir la bonne tension à l'entrée du mélangeur avec une consommation en courant correcte. On remarque que le NF SSB est autour de 8 dB dans les conditions typiques, associé à un gain de conversion de 19 dB et une compression en entrée de l'ordre de -2.5 dBm. Ces performances dépassent les deux architectures à base de mélangeurs passifs présentées précédemment. Un des avantages de ce mélangeur par rapport aux précédents est que son gain de conversion est de 19 dB avec 8 dB de NF. L'ajout du VGA directement à sa sortie ne dégraderait le NF total de la chaîne que de 3 dB. La mise en boîtier, l'effet de la température et du processus de fabrication dégraderont le NF du mélangeur et le porteront à une valeur maximale avoisinant 13 dB. La contribution en bruit de la bande de base serait dans ce cas de l'ordre de 1 dB uniquement, portant le bruit total de la chaîne à 14 dB. Il n'est donc pas nécessaire d'introduire un étage en bande de base faible bruit à cette chaîne, en amont du VGA.

4.3.3 Choix du mélangeur et validation expérimentale sur puce nue

L'étude du mélange de fréquences aux fréquences millimétriques a mené à la conception des trois mélangeurs performants présentés précédemment. Nous nous intéressons maintenant aux résultats obtenus et au contexte technologique de chacune des architectures proposées. Cette étape doit nous permettre de sélectionner et de valider expérimentalement le mélangeur le plus approprié pour les contraintes de l'application radar automobile 77 GHz.

4.3.3.1 Critères de choix et décision

Après l'étude et la conception de trois architectures de mélangeurs, un seul de ces circuits sera considéré pour la conception de la chaîne de réception du radar automobile. Plusieurs critères sont pris en compte pour la prise de décision. En termes de bruit, on compare les blocs permettant d'avoir un gain de conversion d'environ 19 dB. Les amplificateurs rajoutés aux mélangeurs passifs sont considérés comme éléments du sous-bloc « mélangeur » (Tableau 5). Si les performances des trois mélangeurs sont finalement très comparables, il n'en est pas de même pour les niveaux de risque associés aux architectures et à la maturité des technologies Silicium utilisées.

mélangeur	NF SSB (dB)	G _C (dB)	IP _{1dB} (dBm)	Niveau de maturité
Spécifications niveau boîtier	11 (max)	19 (min)	-3 (min)	Intégrable à l'émetteur-récepteur d'un produit radar automobile en BiCMOS180
Passif Schottky	11	21	-5	Très faible : les diodes Schottky sont expérimentales et n'ont jamais été utilisées dans un mélangeur BiCMOS
Passif MOS	9	22	-1.5	Faible : les MOS de la BiCMOS90 ne sont pas encore rigoureusement modélisés au-delà de 60 GHz
Gilbert bipolaire	8	19	-2.5	Élevé : technologie BiCMOS180 et architecture déjà utilisés pour des circuits en production

Tableau 5 : comparaison des trois mélangeurs conçus à 77 GHz en conditions typiques

Concernant le mélangeur à diodes Schottky, nous ne disposons pas, à ce jour, d'un circuit fonctionnel à caractériser. Les performances de cette architecture, certes prometteuses, ont été jugées insuffisantes pour investir en temps et en jeux de masques dans cette architecture alors que nous disposons d'une architecture à base de mélangeur actif aux performances similaires et ayant fait ses preuves. Les mélangeurs passifs à MOS présentent des contraintes similaires puisqu'ils imposent de basculer tout l'émetteur-récepteur en technologie plus avancée, impactant significativement le cycle de développement du produit. Un autre point important dans le choix réside dans l'effort nécessaire au niveau de la bande de base. Les deux architectures à base de mélangeur passif nécessitent l'ajout de nouveaux blocs en bande de base. Cette modification remet en question l'architecture du système de réception en ce qui concerne la linéarité. Le positionnement des filtres passe-bas serait à revoir ainsi que les configurations de gain des différents étages. L'architecture intégrant le mélangeur actif ne nécessite aucune intervention au niveau des circuits en bande de base. Pour toutes ces raisons, il a été décidé de poursuivre le développement de la chaîne de réception en utilisant le mélangeur actif à base de cellule de Gilbert. L'architecture de ce dernier est moins innovante que les deux autres mais ce mélangeur présente le net avantage d'être proche d'une version déjà en production et le transformateur d'impédance que nous lui adjoignons est déjà caractérisé. Au final, c'est donc la solution la moins risquée et elle satisfait, en plus, les spécifications du produit. Néanmoins, les travaux sur les autres architectures sont loin d'être inutiles puisque l'équipe dispose grâce à ces travaux de nouveaux circuits performants, prêts à l'emploi et avec des dessins de masque déjà réalisés. Ces circuits seront exploités dès que les technologies sur lesquelles ils se basent auront gagné en maturité.

4.3.3.2 Validation expérimentale du mélangeur à cellule de Gilbert sur puce nue

Le mélangeur de Gilbert optimisé en linéarité et intégrant le transformateur d'impédance a été sélectionné parmi les différentes topologies étudiées pour réaliser le cœur de la conception de la chaîne de réception. Avant de l'intégrer à cette conception, il est judicieux de caractériser ce mélangeur au niveau plaquette pour s'assurer de ses bonnes performances et optimiser, si nécessaire, les conditions de polarisation.

Aucun épluchage n'est effectué au niveau des mesures sur puce. En effet, les plots de mesure avaient été caractérisés séparément dans le cadre des travaux présentés au chapitre 2. Les modèles de ces plots ont été intégrés à la conception. Ils sont par conséquent déjà pris en compte dans les simulations.

Nous nous intéressons en premier lieu à la réponse en fréquence et en température du mélangeur. Le banc de mesures à notre disposition n'étant pas capable de descendre aux basses températures, nous ne caractériserons ce mélangeur qu'aux températures de 25 °C et de 125 °C. Les mesures sont réalisées sur la bande de fréquence [75-82] GHz, étendant ainsi la plage fréquentielle des mesures de 1 GHz de part et d'autre de la bande radar pour l'automobile.

Le facteur de bruit est mesuré avec la méthode du facteur Y [4.13]. On rappelle que toutes les valeurs de NF sont en SSB. Cette méthode a fait ses preuves aux fréquences millimétriques et son instrumentation est régulièrement calibrée, notamment grâce au calibrage fréquent de l'ENR (pour Excess Noise Ratio) de la source de bruit. Un amplificateur en bande de base très faible bruit sur PCB, réalisant les trois fonctions de gain, de conversion du signal différentiel en signal « single-ended » et de transformation d'impédance vers 50 Ω est rajouté à la chaîne en sortie. Cette sortie en bande de base est finalement connectée à un analyseur de signal (figure 94).

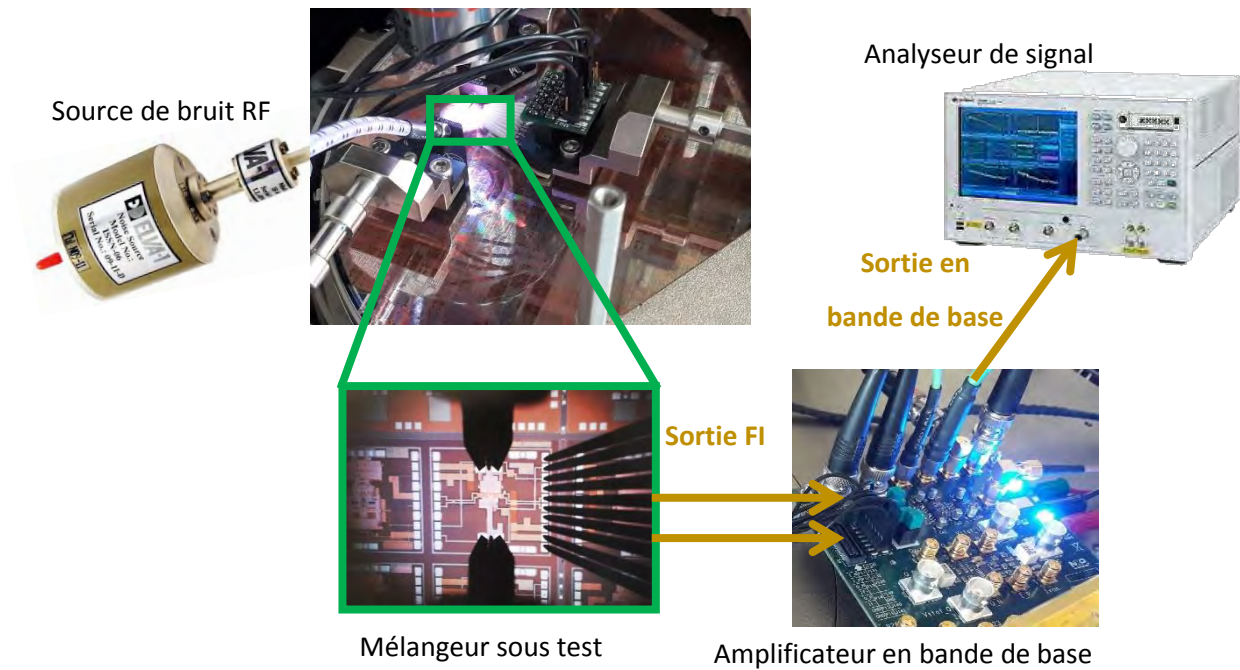


Figure 94 : banc de mesure du NF SSB et du G_c du mélangeur avec la méthode du facteur Y

Les résultats des mesures présentés sur la Figure 95 confirment les très bonnes performances obtenues en simulation. La valeur du NF est bien située entre 8.5 et 9 dB dans la bande pour une température de 25 °C et entre 10 et 11 dB à 125 °C. Ces résultats sont très proches des valeurs issues des simulations est les légers écarts constatés rentrent dans l'imprécision des mesures de NF à ces fréquences. Les résultats pour le gain de conversion (G_c) sont encore meilleurs avec des écarts encore plus faibles entre les mesures et les simulations. À 25 °C, les valeurs de G_c se situent entre 19 et 19.5 dB, et entre 18 et 19 dB à 125 °C.

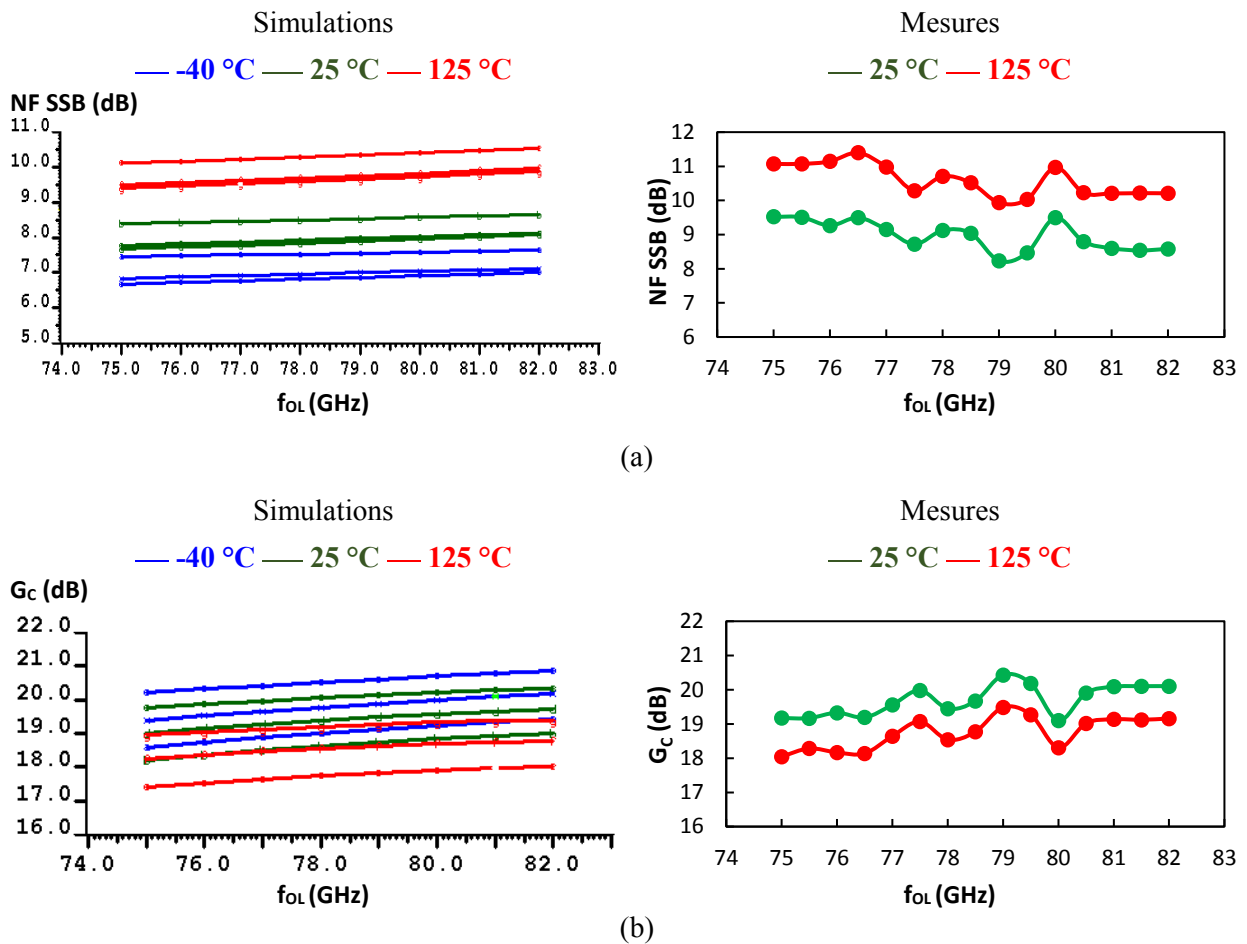


Figure 95 : comparaison des simulations et des mesures du facteur de bruit SSB (NF SSB) (a) et du gain de conversion (Gc) (b) du mélangeur actif en puce nue en fonction de la fréquence OL (f_{OL}) pour $f_{FI} = 1$ MHz pour différentes températures

Pour les mesures qui suivent, nous fixons la valeur de la fréquence OL à 76.5 GHz et la fréquence intermédiaire à 1 MHz, et nous étudions le comportement du circuit en fonction des conditions de polarisation. Nous appliquons des variations de courant et de tension de polarisation autour des valeurs nominales. Le courant de référence « $i_{bias_rf_0}$ » est de 800 μ A. Ce courant de référence est appliqué à un miroir de courant avec un rapport de 5. La tension de polarisation « v_{bias} » est également balayée autour de sa valeur nominale de 1.2 V. On peut se référer au schéma du mélangeur de la Figure 91 pour mieux visualiser ces entrées de polarisation.

Les mesures de NF en fonction du courant de polarisation montrent une faible dépendance entre la polarisation et le bruit (Figure 96). Des variations plus importantes sont observées en simulation. Cette information est très utile parce qu'elle montre que les contraintes sur le courant de référence sont relativement faibles, et que ce courant peut être réduit au besoin. Cependant, la valeur de ce courant est critique pour les valeurs de gain et de la puissance de compression, comme déjà démontré (Figure 74). Pour résumer, plus le courant de polarisation est faible, plus la valeur de la résistance de charge peut être augmentée pour augmenter le gain de conversion du mélangeur, sans que cela n'ait d'impact sur la puissance de compression.

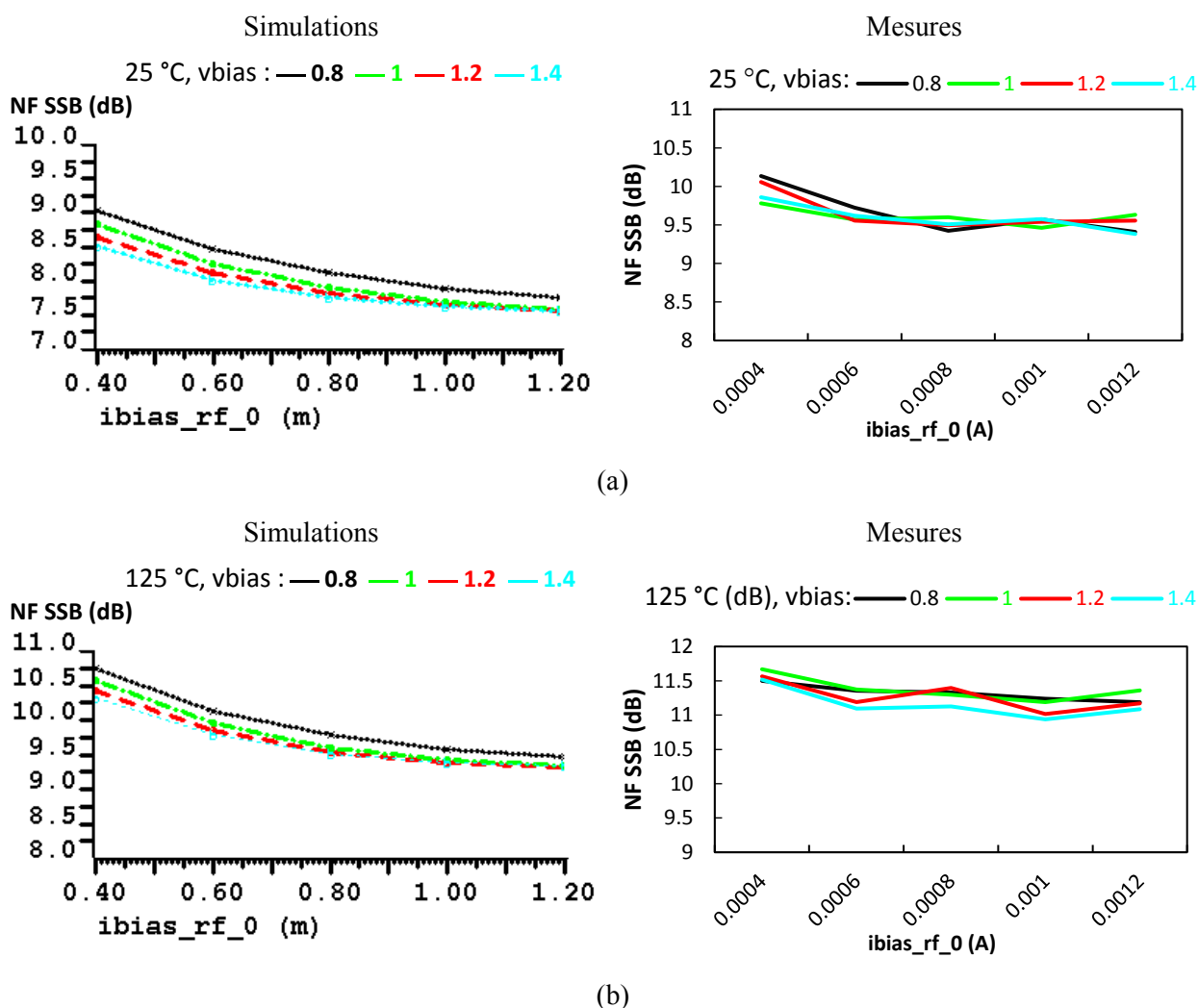


Figure 96 : comparaison des simulations et des mesures du facteur de bruit SSB (NF SSB) du mélangeur actif en puce nue en fonction du courant de polarisation de référence (i_{bias_rf}) et de la tension de polarisation (v_{bias}) pour $f_{OL} = 76.5$ GHz et $f_{RF} = 1$ MHz aux températures de 25 °C (a) et 125 °C (b)

On s'intéresse maintenant au gain de conversion mesuré directement à la sortie de la puce en bande de base, mesures réalisées à l'analyseur de spectre quand des signaux RF aux fréquences $f_{OL} = 76.5$ GHz et $f_{RF} = 76.501$ GHz sont appliqués aux entrées. Cette mesure permet également de déduire le point de compression. Les courbes de la Figure 97 montrent une très bonne correspondance entre les mesures et les simulations. Le gain de conversion est légèrement meilleur en mesures qu'en simulations mais l'impact de la polarisation est bien conforme. Les valeurs du point de compression en entrée du mélangeur IP_{1dB} montrent les mêmes tendances avec une compression inférieure de 1 dB en simulations par rapport aux mesures. Les variations de G_c et IP_{1dB} en fonction du courant sont consistantes entre les simulations et les mesures, mais on remarque que la dégradation du point de compression simulée est plus importante que celle mesurée pour des tensions de polarisation supérieures à 1.4 V. Cet écart est lié à la modélisation des transistors qui est à ses limites dans cette région de fonctionnement du mélangeur. Cependant, il ne pose aucun problème, puisqu'il est observé assez loin du point de fonctionnement nominal du mélangeur dans la chaîne de réception du radar.

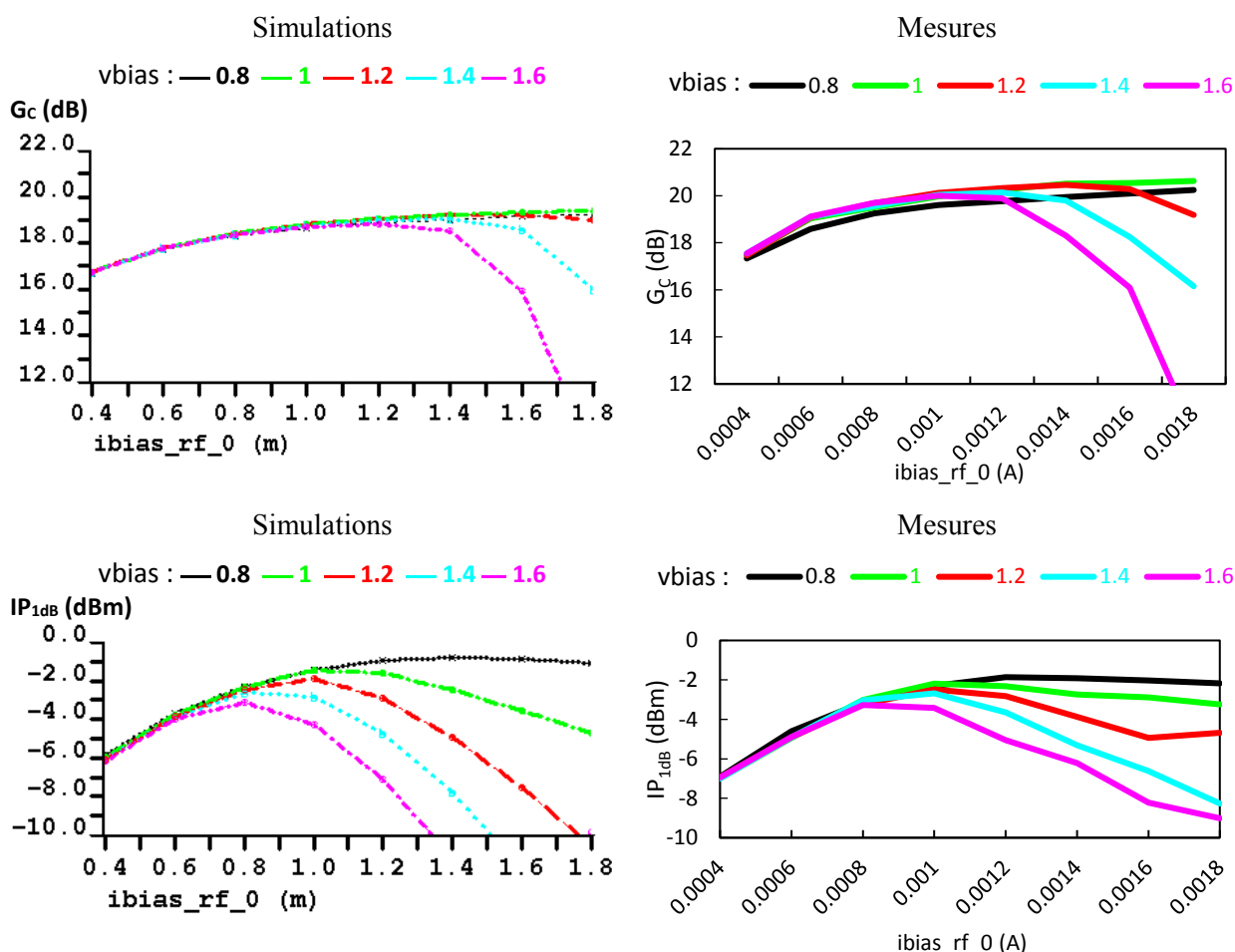


Figure 97 : comparaison des simulations et mesures de gain de conversion (G_c) et de compression en entrée (IP_{1dB}) en fonction du courant de polarisation de référence ($ibias_rf_0$) pour différentes tensions de polarisation ($vbias$)

La caractérisation du mélangeur confirme donc les très bonnes performances prédites lors de la phase de conception, notamment obtenues grâce à l'utilisation du transformateur d'impédance. Le travail de conception réalisé au niveau de la polarisation et de l'adaptation en bruit a permis d'obtenir des performances encore meilleures que ce qui avait été calculé (Tableau 6). Nous pouvons conclure que nous disposons d'un mélangeur aux performances à l'état de l'art, bien caractérisé et dont les simulations et les mesures sont très comparables. Tous ces éléments permettent d'envisager la conception de la chaîne de réception du radar avec une assurance élevée.

Mélangeur de Gilbert conçu	NF SSB (dB)	G_c (dB)	IP_{1dB} (dBm)
Spécifications niveau boîtier	11 (max)	19 (min)	-3 (min)
Simulations	8	19	-2.5
Mesures typiques	8.5-9	19.5	-3

Tableau 6 : résumé des performances simulées et mesurées du mélangeur conçu

4.4 Conception de la chaîne de réception en boîtier WLCSP

L'architecture de la chaîne de réception se base sur les exigences du produit radar, lui permettant de réaliser sa fonction principale : la détection de cibles. Les principaux éléments de cette chaîne, à savoir le boîtier et le mélangeur de fréquence, sont déjà caractérisés séparément. Le mélangeur conçu a un très bon niveau de bruit, autour de 9 dB à 25 °C, grâce aux innovations qu'il inclut. Il semble assez difficile de réaliser une meilleure performance sans revoir significativement la technologie utilisée. Pour le mélangeur en boîtier, la spécification impose un niveau de bruit ne dépassant pas 11 dB. En aucun cas, par conséquent, les pertes du boîtier ne doivent dépasser 2 dB. Pour cette raison, le défi d'amélioration des pertes de la transition était d'une criticité très élevée. Nous présentons ici l'intégration en boîtier du mélangeur conçu. Nous étudions ensuite l'impact des circuits en bande de base sur les performances globales de la chaîne de réception (composée de l'interface boîtier, du mélangeur et des blocs en bande de base).

4.4.1 Conception de la chaîne de réception

Aux fréquences millimétriques, le boîtier ne peut pas être directement et simplement inséré dans la chaîne. Comme nous l'avons déjà montré, la transformation d'impédance qu'il apporte impacte les performances du circuit et peut dégrader de façon importante le facteur de bruit. La prise en compte du boîtier demande par conséquent de revoir la conception du mélangeur, tout en essayant de minimiser les évolutions afin de garder un mélangeur le plus proche possible de la version caractérisée en puce nue.

La conception de la chaîne de réception est réalisée dans l'environnement Cadence Virtuoso. La méthodologie suivie pour l'introduction du modèle de boîtier à l'environnement de simulation consiste à importer les paramètres S de la dernière version caractérisée du boîtier sous forme de fichier « Touchstone ». Le symbole (figure 98) explicite le côté puce et le côté PCB afin d'éviter toute confusion par rapport à l'orientation du boîtier. Ce symbole, ainsi que le schéma et le fichier de données associés sont sous contrôle de version pour garder rigoureusement la traçabilité et permettre à l'équipe de développement d'utiliser la même structure, si besoin.

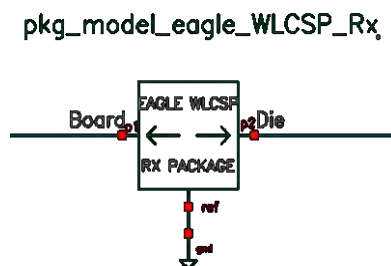


Figure 98 : symbole de la boîte des paramètres S du boîtier dans Cadence Virtuoso

Le banc de simulations est le même que celui utilisé pour la conception du mélangeur sur puce nue (cf. Figure 91) à l'exception des entrées RF et LO pour lesquelles les modèles de plot sont remplacés par les modèles de boîtier.

Comme schématisé sur la Figure 99, on procède, premièrement, à la réoptimisation de l'adaptation RF en entrée afin d'optimiser le facteur de bruit de l'application en boîtier. On réalise le même travail pour l'adaptation OL de manière à minimiser la puissance nécessaire au bon fonctionnement du mélangeur. Ces optimisations sont basées sur des modifications des longueurs et des largeurs des lignes entre le boîtier et les transistors de mélange. Après l'optimisation des adaptations des entrées, on procède à l'augmentation du gain de conversion qui est réduit d'environ 1.5 dB à cause des pertes introduites par le boîtier. L'amélioration de G_c passe par l'augmentation de la résistance de charge (R_L) qui augmente le gain en tension à la sortie. Pour garder alors la même tension émetteur-collecteur, et par conséquent le même point de compression en entrée, le courant de référence (I_{ref}) est légèrement réduit. La modification de R_L et I_{ref} n'impacte que très légèrement l'adaptation en entrée RF.



Figure 99 : étapes de réoptimisation du mélangeur en boîtier WLCSP

Les simulations en fin d'optimisation du mélangeur en boîtier montrent de très bonnes performances au niveau du PCB à l'entrée du boîtier (Figure 100).

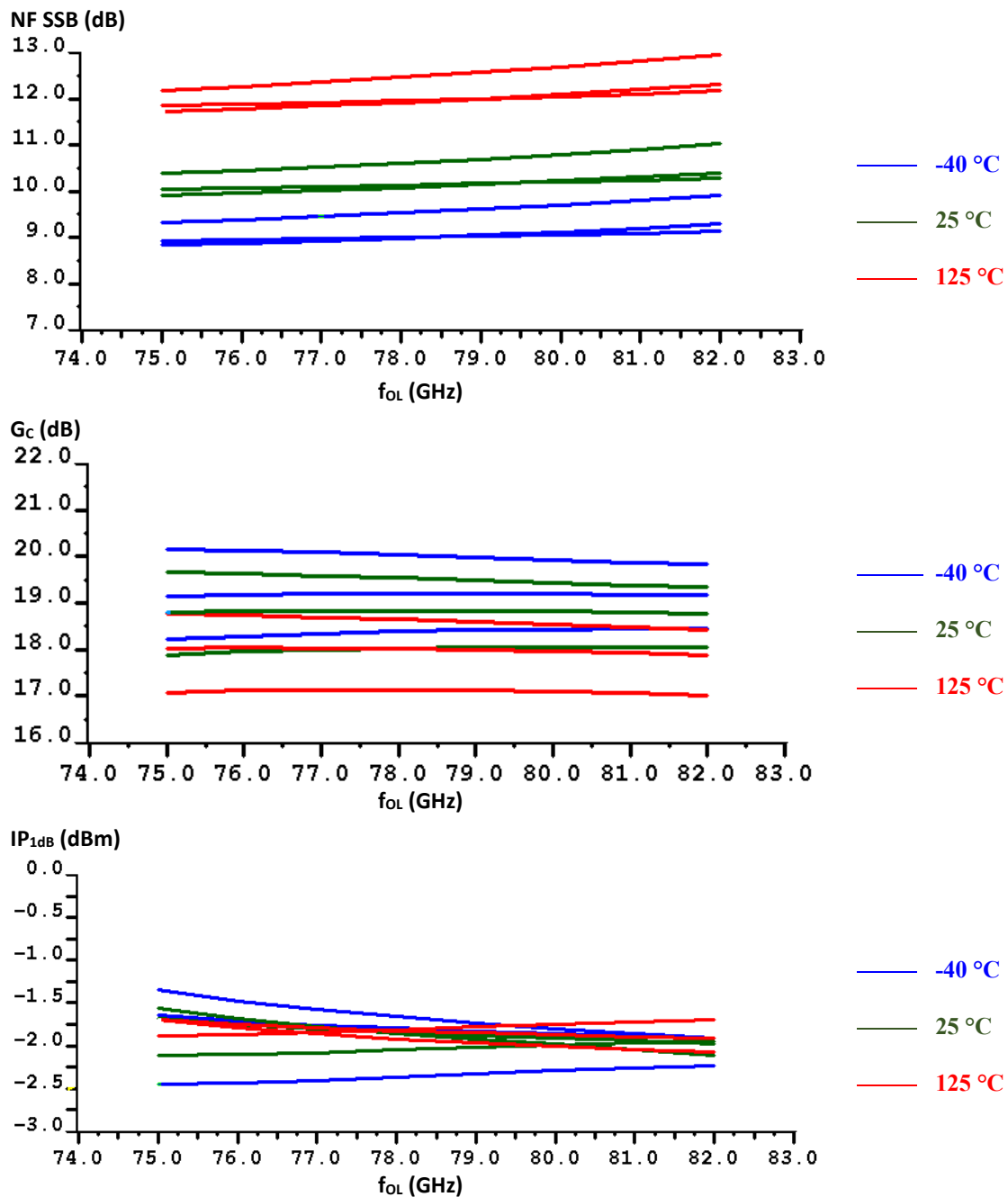


Figure 100 : résultats de simulations du facteur de bruit (NF SSB), du gain de conversion (G_C) et de la compression en entrée (IP_{1dB}) du mélangeur en boîtier en fonction de la fréquence OL (f_{OL}) pour f_{F1} = 1 MHz et pour différentes températures

Le facteur de bruit et le gain de conversion typiques sont respectivement autour de 10 et 19 dB à 77 GHz et 25 °C. La puissance de compression en entrée reste également supérieure à -3 dBm dans toutes les conditions simulées. Ces résultats dépassent les performances obtenues jusqu'à présent pour les produits radar en boîtier. L'utilisation pour la première fois d'un boîtier WLCSP pour l'encapsulation d'une puce en technologie BiCMOS à ces fréquences donne encore plus d'intérêt à ces résultats.

Procédons maintenant à l'étude de la chaîne complète, incluant les étages pour le traitement du signal en bande de base. Le schéma bloc de la chaîne de réception incluant le boîtier (Figure 101) énumère les contributeurs en bruit du récepteur radar. Nous pouvons étudier l'impact de la bande de base sur le NF

total en utilisant de nouveau l'équation (31). En considérant que les performances de l'interface boîtier associée au mélangeur sont un NF SSB de 10 dB et un G_C de 19 dB, et en considérant une tension de bruit en entrée du VGA de $10 \text{ nV}/\sqrt{\text{Hz}}$, le NF de la chaîne complète en conditions typiques est de 12 dB pour un gain de 57 dB.

Pour dimensionner correctement les autres paramètres intervenant dans l'équation radar, il est également nécessaire d'étudier le pire cas pour le facteur de bruit NF du mélangeur. Ce pire cas est basé sur des simulations de NF et de G_C pour une température de 150°C et au pire cas des dispersions du procédé de fabrication. Dans ces conditions et pour le mélangeur en boîtier WLCSP, on obtient, un NF de 12 dB et un G_C de 17 dB. Si l'on tient compte de plus du bruit donné pour le VGA dans ces mêmes conditions, on obtient un bruit total pour la chaîne de réception à 14 dB au pire cas. Ce résultat représente ici encore une très bonne performance pour le système radar en termes de sensibilité du récepteur.

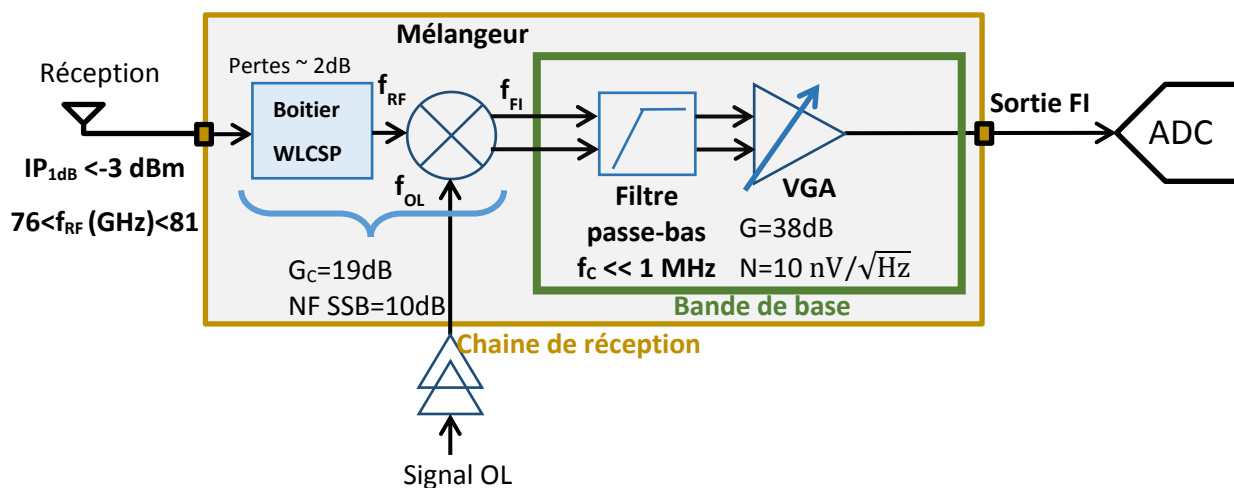


Figure 101 : schéma bloc de la chaîne de réception en boîtier

De plus, ces performances auront un impact très positif sur la portée du radar longue portée. En effet, si on compare les performances en bruit obtenues dans le pire cas, ici à partir d'un NF SSB de 14 dB, avec celles d'une chaîne de réception présentant un NF de 17 dB, on trouve que la conception proposée améliore d'environ 50 mètres la portée maximale du radar, soit une augmentation de portée d'environ 19%.

L'intégration des circuits en bande de base à la chaîne de réception complète ne sera pas étudiée plus en détails. Elle impliquerait des travaux dépassant le cadre de cette thèse. L'objectif que nous avons ici était de démontrer que le boîtier « Fan-In » WLCSP développé peut être utilisé pour encapsuler une chaîne de réception radar 77 GHz.

4.4.2 Validation expérimentale

La chaîne de réception a été fabriquée et son encapsulation réalisée. Nous avons donc procédé à sa caractérisation pour valider les valeurs données par les simulations et confirmer tout l'intérêt des travaux menés dans cette thèse.

On se focalise en premier sur le mélangeur en boîtier sans le reste de la chaîne qui traite les signaux en bande de base. L'objectif est de limiter les incertitudes liées au système et aux mesures, d'autant que la contribution du traitement en bande de base au bruit de la chaîne complète est limitée grâce au gain de conversion important du mélangeur. La validation expérimentale de la chaîne de réception complète, incluant la bande de base, est également réalisée mais pour une couverture de test plus limitée.

Sur la Figure 102, les mesures réalisées sur le mélangeur en boîtier montrent des valeurs du NF situées entre 10 dB et 11 dB comme attendu par la simulation (valeurs de NF comprises entre 10 et 10.5 dB) et un gain de conversion entre 18 dB et 19 dB (valeurs simulées autour de 19 dB). Les fluctuations observées sur les mesures dans la bande sont liées aux réflexions sur le chemin de test RF qu'il n'a pas été possible d'éplucher rigoureusement à cause de la nature scalaire du banc de mesures.

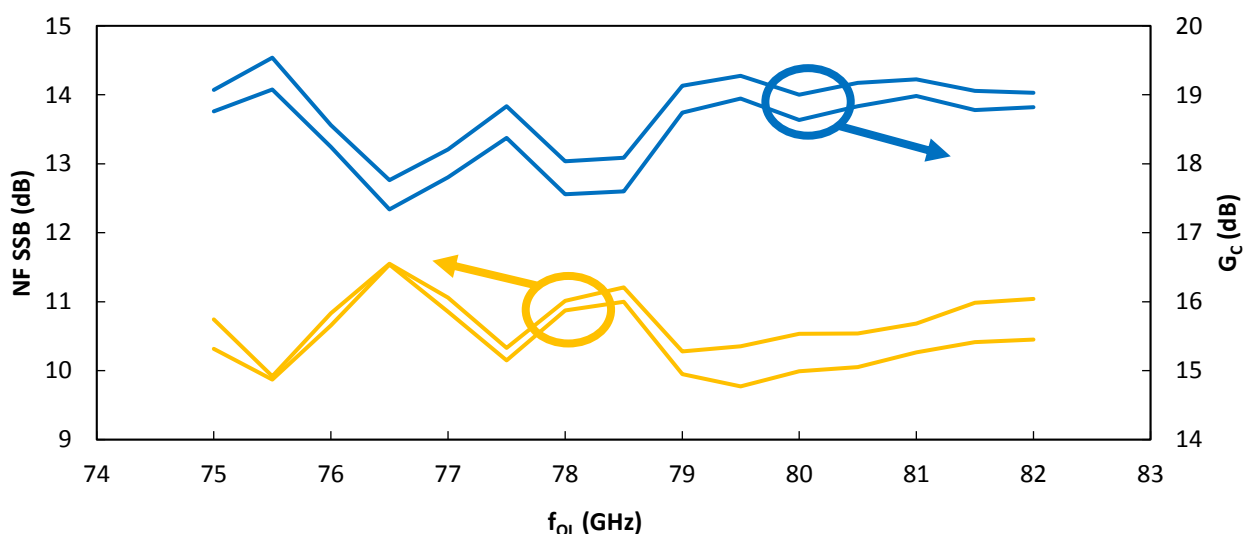


Figure 102 : facteur de bruit SSB (NF SSB) et gain de conversion (G_c) du mélangeur en boîtier en fonction de la fréquence OL (f_{OL}) pour $f_I = 1$ MHz et pour 2 pièces

Pour valider le bon fonctionnement du mélangeur dans les conditions de polarisation choisies, on fait varier le courant de référence (I_{ref}) autour de sa valeur nominale (Figure 103). On retrouve le même comportement tel qu'il a été observé dans les mesures sans boîtier, confirmant la faible dépendance entre le NF et le courant de référence.

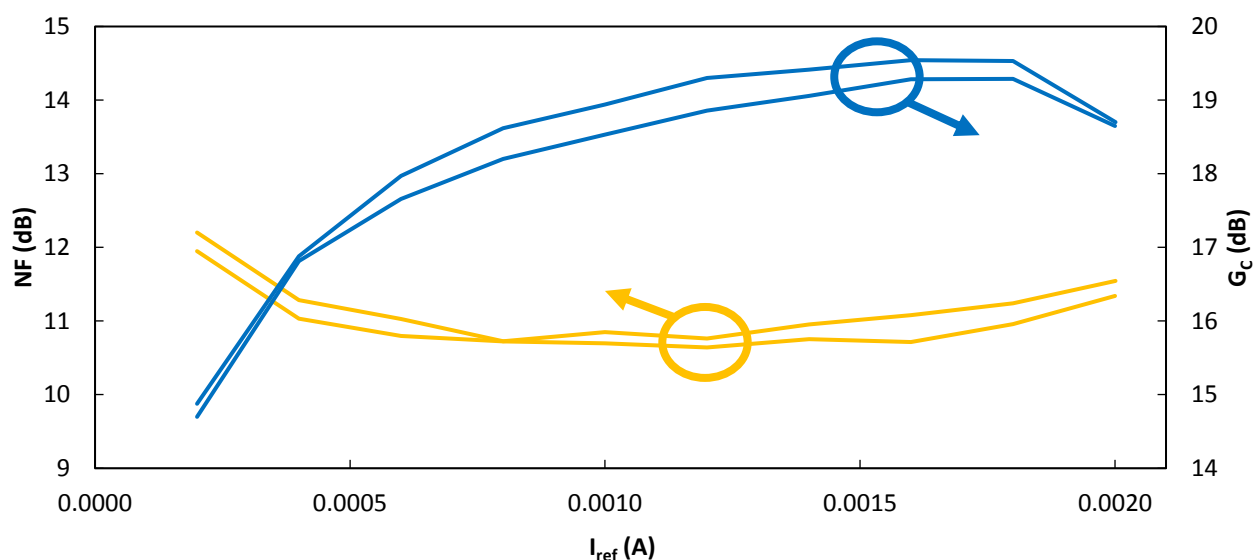


Figure 103 : facteur de bruit (NF SSB) et gain de conversion (G_c) du mélangeur en boîtier en fonction du courant de polarisation de référence (I_{ref}) à $f_{OL} = 77$ GHz pour $f_{FI} = 1$ MHz et pour 2 pièces

Nous faisons varier également la puissance OL afin de vérifier que le niveau appliqué permette bien l'obtention du NF optimal. On voit que le NF est à son plus faible niveau à partir de 6 dBm de puissance OL au niveau de l'accès PCB. Les autres mesures ont donc été réalisées dans des conditions correctes puisqu'ayant été réalisées avec une puissance du signal OL d'environ 8 dBm.

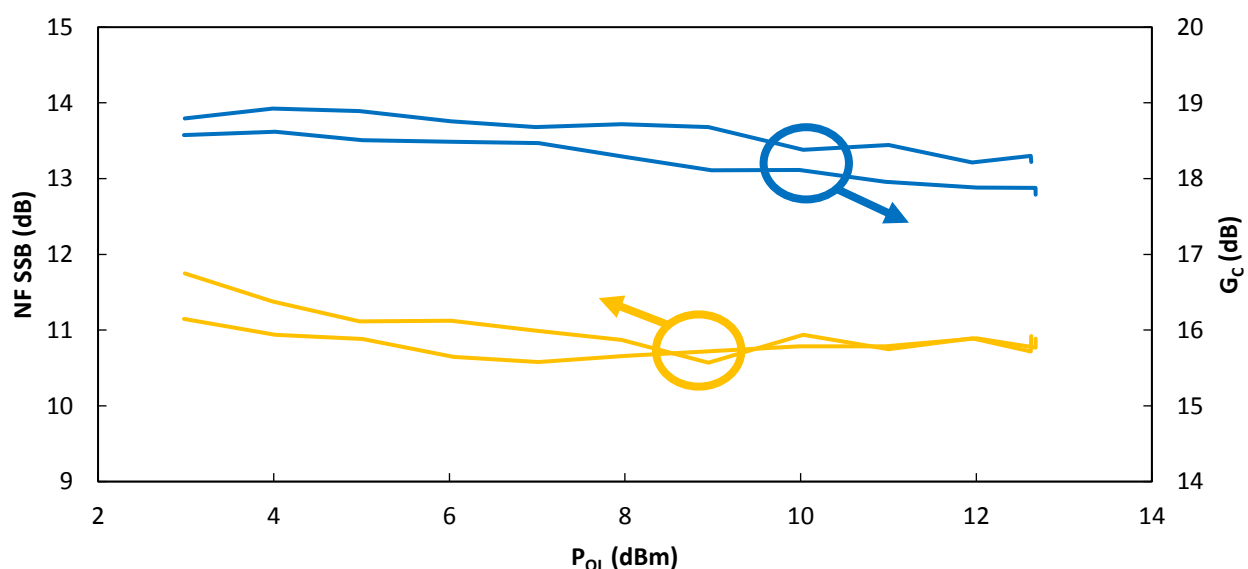


Figure 104: facteur de bruit (NF SSB) et gain de conversion (G_c) du mélangeur en boîtier en fonction de la puissance OL (P_{OL}) pour 2 pièces à $f_{OL} = 77$ GHz

Les performances du mélangeur en boîtier sont ainsi validées expérimentalement. Nous nous intéressons donc maintenant à la validation expérimentale de la chaîne complète, incluant les circuits de mise en

forme des signaux en bande de base. Les amplificateurs en bande de base produisent un gain de 32 dB et la fréquence de coupure à -3 dB du filtre passe-bas est à $f_{FI} = 300$ kHz. Nous présentons le facteur de bruit SSB et le gain de conversion de la chaîne. On remarque que le NF SSB et le G_C totaux de la chaîne sont alignés avec les calculs au niveau système à 1 MHz. En effet, on observe, un facteur de bruit compris entre 12 dB et 13 dB et un gain de conversion compris entre 49 dB et 51 dB sur la plage des fréquences du radar. Ce gain de conversion correspond à un gain de 32 dB produit par les amplificateurs en bande de base et un gain compris entre 17 dB et 19 dB produit par le mélangeur de fréquences. (Figure 105).

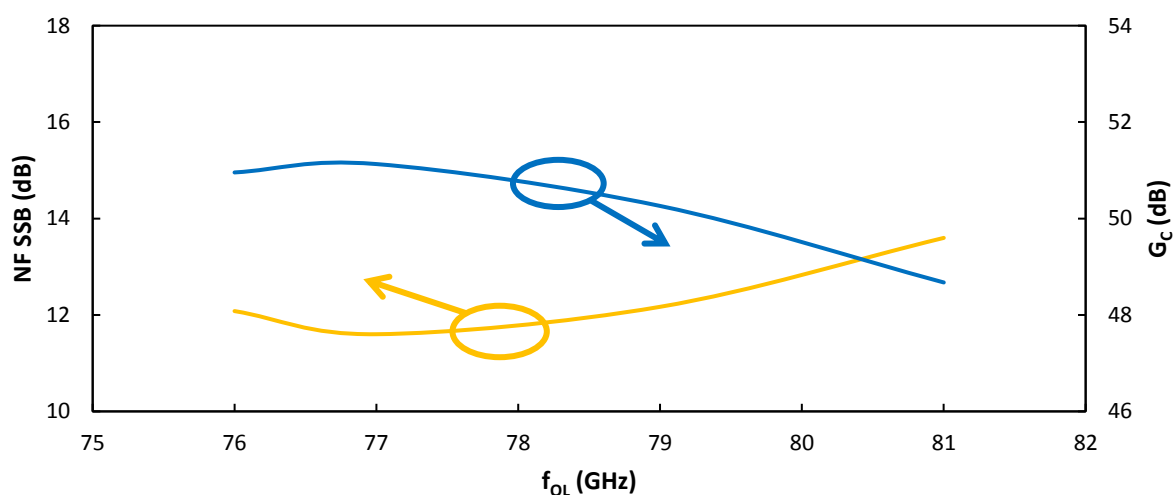


Figure 105 : facteur de bruit (NF SSB) de la chaîne de réception, incluant la bande de base, en boîtier en fonction de la fréquence OL (f_{OL}) pour $f_{FI} = 1$ MHz

La Figure 106 illustre le NF SSB et le G_C en fonction de la fréquence intermédiaire. On observe un bruit important en très faible fréquence, inférieure à 1 MHz, provenant de la bande de base. On remarque également l'effet du filtre passe-bas, dont la fréquence de coupure est à 300 kHz. Par ailleurs, on remarque que le NF s'améliore encore au-delà de 1 MHz et atteint la valeur de 11 dB à 79 GHz à partir de 2 MHz de fréquence intermédiaire alors qu'il est à 12 dB à 1 MHz.

Figure 106 : Facteur de bruit (NF SSB) et gain de conversion (G_C) de la chaîne de réception, incluant la bande de base, en boîtier en fonction de la fréquence intermédiaire (f_{FI}) pour différentes fréquences OL (f_{OL})

4.5 Conclusion

Ce dernier chapitre a présenté l'application des développements réalisés sur l'interface boîtier en « fan-In » WLCSP à un circuit intégré pour radar automobile dans la bande [76-81] GHz. Le mélangeur de fréquences a été choisi pour accomplir cette démonstration. Le choix de ce bloc nous a semblé le plus pertinent étant donné qu'il s'agit du premier élément dans la chaîne de réception et qu'il est donc le plus fortement en interaction avec le boîtier. De plus, pour une évaluation fine des avancées pouvant résulter de l'encapsulation conçue en boîtier « fan-In » WLCSP, nous avons réalisé une nouvelle étude système

suivie d'une étude approfondie des architectures de mélangeurs compatibles avec les spécifications radar, et qui pourront être adaptées aux interfaces du boîtier. Trois architectures très différentes de mélangeur doublement équilibré ont ainsi été analysées : deux topologies passives à diodes Schottky puis à transistors MOS froids, et une topologie active à base d'une cellule de Gilbert. Le travail du chapitre 2 autour de la modélisation EM au niveau du circuit a également été exploité dans l'optimisation de ces mélangeurs.

La comparaison des performances et l'analyse des contraintes de développement des mélangeurs conçus ont abouti au choix du mélangeur à base de cellule de Gilbert pour la suite du développement de la chaîne de réception en boîtier. Ce mélangeur utilise un transformateur d'impédance en entrée RF à la place de l'étage de gain classiquement implémenté dans la cellule de Gilbert. Il présente de très bonnes performances qui ont pu être validées par des mesures. De plus, il présente l'avantage de la grande maturité acquise grâce aux différents développements précédemment consacrés à des architectures similaires, menés au sein de l'équipe. Dans un processus d'optimisation globale, ce mélangeur a été optimisé par la suite placé dans son environnement constitué par le boîtier WLCSP. La caractérisation a permis de valider son très bon comportement en boîtier puisqu'on obtient un facteur de bruit entre 10.5 dB et 11 dB au niveau du PCB et un gain de conversion autour de 19 dB sur toute la bande RF [76-81] GHz. L'étude de la contribution des circuits de mise en forme du signal en bande de base a permis d'anticiper le comportement de la chaîne complète, confirmé ensuite par des mesures. Pour une fréquence RF de 79 GHz, la caractérisation en bruit de la chaîne complète, constituée du boîtier WLCSP, du mélangeur et des filtres et des amplificateurs en bande de base, conduit à une valeur du facteur de bruit total de la chaîne avoisinant 12 dB pour une valeur de la fréquence FI de 1 MHz.

Ces performances à l'état de l'art ont, avant tout, confirmé la possibilité de développer un produit radar pour l'automobile en boîtier WLCSP et ont démontré le grand intérêt de la modélisation EM, si elle est accompagnée de mesures régulières confirmant les performances qu'elle prédit.

Compte tenu des valeurs du facteur de bruit obtenues pour le mélangeur de réception encapsulé dans un boîtier « fan-In » WLCSP, le travail ayant fait l'objet de ce chapitre permet de prévoir une amélioration de 12% sur la portée du radar, par rapport aux produits radars actuellement en production. En remplaçant un mélangeur d'un radar du marché par ce nouveau mélangeur encapsulé dans un boîtier WLCSP, on peut espérer atteindre une portée de 225 mètres si ce radar avait initialement une portée de 200 mètres.

Conclusion générale

La généralisation des systèmes d'aide à la conduite à toutes les catégories de voitures et le développement de la voiture autonome constituent les principaux catalyseurs du développement de ces systèmes, le radar automobile 77 GHz en particulier. En effet, les équipementiers automobiles se doivent de proposer en permanence des modules radars plus performants et moins chers. Cette course effrénée ne s'arrête pas au niveau des équipementiers, elle remonte dans la chaîne de production jusqu'aux fournisseurs de circuits intégrés qui doivent fabriquer des émetteurs-récepteurs radars toujours plus performants et plus intégrés à chaque nouvelle génération. La rapidité avec laquelle évoluent les spécifications des circuits intégrés oblige les concepteurs à innover de plus en plus pour tirer le meilleur profit des évolutions des technologies semi-conducteurs, évolutions qui présentent un retard de plus en plus grand par rapport aux besoins du marché, surtout en considérant la contrainte de qualification des applications pour l'automobile. C'est dans ce contexte principalement axé sur l'innovation que cette thèse a été menée, tout en maintenant l'objectif que les résultats obtenus soient le plus directement possible applicables aux produits.

Le cœur des travaux menés durant cette thèse a consisté au développement d'une interface boîtier WLP pour l'encapsulation au niveau plaquette d'un émetteur-récepteur radar automobile à 77 GHz. Nous nous sommes intéressés dans un premier temps à la mise en place des outils de modélisation électromagnétique et à la modélisation de structures simples, avant de modéliser des boîtiers WLP « fan-out » et par la suite des boîtiers WLP « fan-in ». Ces derniers, plus faibles coût et meilleurs relativement à la dissipation thermique, n'avaient encore jamais été utilisés pour des applications en bandes millimétriques en technologie Silicium. La principale raison qui conduit à l'absence du boîtier WLP « fan-in » pour l'encapsulation à ces fréquences découle du très mauvais comportement électrique en hautes fréquences des solutions actuelles. Pour corriger cette situation, il est essentiel de développer une interface boîtier compatible avec les fréquences radar 77 GHz. Pour mener ce développement, l'encapsulation d'une chaîne de réception d'un radar pour l'automobile peut être une excellente application. En effet, cette chaîne est particulièrement sensible aux pertes des interconnexions de la puce avec le PCB. L'obtention d'une chaîne de réception encapsulée avec un bruit à l'état de l'art est donc en mesure de réaliser la démonstration que l'implémentation du boîtier WLP « fan-in » est possible à ces fréquences.

Ainsi, le premier chapitre apporte les éléments qui permettent de comprendre de façon approfondie les contraintes liées au module radar. Cette présentation permet, entre autres, de quantifier l'effet d'une dégradation des performances de l'émetteur-récepteur sur la fonctionnalité du module radar. Une partie importante du chapitre est donc consacrée à la définition du système radar, aux diverses fonctions qu'il intègre et à leur réalisation, avant de se focaliser plus particulièrement sur l'émetteur-récepteur radar.

La forte interaction entre le circuit et son environnement est également détaillée et l'accent est porté sur l'impact de l'interconnexion entre le PCB et le circuit sur les performances de l'émetteur-récepteur radar.

Dans le deuxième chapitre, nous présentons ensuite nos travaux axés sur la modélisation électromagnétique des interconnexions et autres structures passives implémentées au niveau du circuit intégré. Plusieurs objectifs ont orienté ce travail de modélisation. Il s'agit tout d'abord de se baser sur des composants simples qui permettent de confirmer que les caractéristiques utilisées pour la modélisation des matériaux sont correctes ou de les ré-ajuster dans certains cas. La validation de la qualité des résultats donnés par les simulateurs par rapport à la problématique de l'encapsulation aux fréquences millimétriques fait également partie des objectifs de ces travaux. La confrontation avec les résultats expérimentaux des performances simulées était un élément critique pour confirmer la précision des simulations. Des structures plus complexes, comme le transformateur d'impédance, ont été modélisées et validées par la suite. Ce transformateur présente des propriétés électriques très intéressantes mais reste peu utilisé dans l'industrie à cause de la nécessité de passer par des outils de simulation électromagnétiques et de valider expérimentalement ces modèles. La validation du modèle du transformateur d'impédances conçu a permis de disposer d'une structure aux propriétés très intéressantes et facilement utilisable pour de nombreuses applications comme les entrées et sorties des amplificateurs et les embranchements des déphaseurs.

La modélisation électromagnétique des structures au niveau du circuit intégré a permis d'entamer avec assurance la modélisation des boîtiers du radar automobile, axe d'étude décrit dans le troisième chapitre. La modélisation des interfaces boîtier a commencé avec l'étude d'un boîtier WLP de type « fan-out », déjà utilisé pour la génération actuelle du radar automobile 77 GHz. Ensuite, la modélisation du boîtier WLP « fan-in », nommé WLCSP (pour Wafer Level Chip Scale Package), a été menée. Plusieurs itérations, accompagnées de phases de validation expérimentale, ont permis d'aboutir à une version optimisée du boîtier en mesure de couvrir les applications aux fréquences millimétriques. Les performances obtenues rivalisent finalement avec celles des boîtiers « fan-out ». La validation expérimentale des boîtiers développés a également posé de nombreux problèmes de caractérisation et plusieurs méthodologies de caractérisation basées sur un algorithme OSL (pour « Open, Short, Load ») ont été proposées pour faire face aux difficultés et aux limitations rencontrées.

La confirmation ultime des bonnes performances du nouveau boîtier WLCSP développé a été envisagée au travers de son implémentation pour l'encapsulation d'un émetteur-récepteur radar. Par conséquent, une chaîne de réception radar a été développée et ce développement fait l'objet du quatrième chapitre. La conception du mélangeur de la chaîne a été au cœur de notre travail. Une étude menée sur l'architecture de la chaîne de réception a conduit à extraire trois topologies de mélangeurs pouvant convenir à l'application : mélangeurs passifs à diodes Schottky ou à transistors MOS froids (non polarisés entre drain et source) ou actif à base de cellule de Gilbert. La comparaison des performances

de ces mélangeurs après intégration a permis la sélection de l'architecture active à base de cellule de Gilbert. En effet, ce mélangeur présente les performances les plus compatibles avec l'application radar pour la technologie BiCMOS disponible, et la maturité qui lui est associée est suffisamment élevée pour permettre son implémentation rapide dans un produit radar. Ce mélangeur a été, par la suite, optimisé en incluant le modèle du boîtier WLCSP. La contribution des circuits de mise en forme des signaux en bande de base au bruit total de la chaîne a également été étudiée. Les performances mesurées pour la chaîne de réception en boîtier WLCSP sont à l'état de l'art. Elles confirment ainsi tous nos travaux de modélisation, de conception et de caractérisation. Ces performances démontrent aussi que le nouveau boîtier WLCSP développé est parfaitement compatible avec les applications aux fréquences millimétriques, notamment le radar 77 GHz pour l'automobile.

Les réalisations de cette thèse confirment que la co-optimisation du boîtier, du circuit intégré et du circuit imprimé permet d'accéder à de nouvelles opportunités d'innovation et d'amélioration des performances, inaccessibles autrement. Nous avons établi une méthodologie de modélisation des structures et nous l'avons exploitée pour la conception de boîtiers innovants et de mélangeurs à l'état de l'art. La même méthodologie peut être exploitée dans le futur pour investiguer, par exemple, la conception d'éléments passifs sur la couche de redistribution du boîtier WLP « Fan-In ». Les inductances des oscillateurs commandés en tension aux fréquences millimétriques peuvent en effet être réalisées au niveau de cette couche de redistribution du boîtier afin d'améliorer leur facteur de qualité. Cette implémentation devrait permettre d'améliorer le bruit de phase des émetteurs-récepteurs radars. Dans le chapitre 4, un transformateur d'impédances a été intégré à l'entrée de la chaîne de réception, directement après l'interface boîtier. La conception d'un transformateur d'impédances ou d'un balun LC en utilisant la couche de redistribution du boîtier et les métallisations du circuit intégré est susceptible d'améliorer les performances en bruit de la chaîne de réception. Pour aller plus loin, l'étude d'antennes planaires intégrées en boîtier WLCSP est une piste de miniaturisation intéressante pour les modules radar, qui permettrait aussi l'amélioration des performances en s'affranchissant d'une transition par le PCB pour l'interconnexion critique entre l'antenne et l'émetteur-récepteur. Si la taille des antennes radars 77 GHz s'avère être trop grande en comparaison aux dimensions du boîtier, les radars 122 GHz seront sûrement d'excellents candidats pour cette exploration.

Publications

Brevets

Charaf-Eddine Souria, Cristian Pavao Moreira, « BASEBAND AMPLIFIER CIRCUIT », EPO, brevet déposé le 1^{er} novembre 2016.

Charaf-Eddine Souria, Gilles Montoriol, Stéphane Thuries, “Semiconductor device package, electronic device and method of manufacturing electronic devices using wafer level chip scale package technology.”, déposé le 18 novembre 2015.

Conférences internationales

Charaf-Eddine Souria, Thierry Parra, Gilles Montoriol, Christophe Landez, “ Accurate Package Model extraction up to 110 GHz using one-port measurements. Application to a 77 GHz radar transceiver. ”, 2016 IEEE Radio and Wireless Symposium (RWS), 24-27 janvier 2016.

Conférences nationales

Charaf-Eddine Souria, Thierry Parra, Gilles Montoriol, Christophe Landez, « Nouvelle méthodologie pour l'extraction d'un modèle de boîtier jusqu'à 110 GHz. Application à l'encapsulation d'un radar 77 GHz. », 19èmes Journées Nationales Microondes (JNM 2015), Bordeaux, France, juin 2015.

Références bibliographiques

Bibliographie du chapitre 1

- [1.1] A. Mcauslin, “Safe Autonomous Systems: Centralized or distributed processing?”, October 27, 2016, www.blog.nxp.com
- [1.2] A. Bartsch, F. Fitzek, R. H. Rasshofer, “Pedestrian recognition using automotive radar sensors”, *Advances in Radio Science*, 2012
- [1.3] K. Kobayashi, T. Morita, H. Mukai, T. Kishigami, Y. Nakagawa, “79GHz-band Coded Pulse Compression Radar System Performance in Outdoor for Pedestrian Detection”, 2013 European Microwave Conference, 2013
- [1.4] D. Geronimo; A. M. Lopez; A. D. Sappa; T. Graf, “Survey of Pedestrian Detection for Advanced Driver Assistance Systems”, *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 2010, Vol. 32, Issue 7
- [1.5] “The European Table of Frequency Allocations And Applications In The Frequency Range 8.3 kHz To 3000 GHz (ECA TABLE)”, version approuvé en mai 2015.
- [1.6] <http://www.rfcafe.com/references/electrical/ew-radar-handbook/rf-atmospheric-absorption-ducting.htm>
- [1.7] C. Kärnfelt, A. Péden, A. Bazzi, G. E. Shhadé, M. Abbas, T. Chonavel, F. Bodereau, « 77 GHz ACC Radar Simulation Platform », 9th International Conference on Intelligent Transport Systems Telecommunications, (ITST), 2009.
- [1.8] RECOMMANDATION UIT-R P.525-2, « CALCUL DE LA PROPAGATION EN ESPACE LIBRE », (1978-1982-1994)
- [1.9] “ARS 408-21 Premium Long Range Radar Sensor 77 GHz Data Sheet”, http://www.conti-online.com/www/industrial_sensors_de_en/themes/ars_premium_en.html.
- [1.10] S. Trotta et al., “An RCP Packaged Transceiver Chipset for Automotive LRR and SRR Systems in SiGe BiCMOS Technology”, *IEEE Transactions on Microwave Theory and Techniques*, 2012, Vol. 60, Issue 3
- [1.11] https://fr.wikipedia.org/wiki/Radar_%C3%A0_ondes_entretenues#/media/File:Fmcw_prinziple.png
- [1.12] D. Kok; J. S. Fu, “Signal Processing For Automotive Radar”, *IEEE International Radar Conference*, 2005.
- [1.13] Q. N. Nguyen, H. Yoo, M. Y. Park, Y. Kim, F. Bien, “A 77 GHz Waveform Generator with MFSK”, *IEEE International Wireless Symposium (IWS)*, 2015.
- [1.14] “38-38.5 GHz RF 4-channel voltage controlled oscillator front-end for W-band radar applications Data Sheet”, rev. 2.0, août 2016, www.nxp.com/radar.
- [1.15] Z. Tong, R. Reuter, M. Fujimoto, “Fast chirp FMCW Radar in automotive applications”, 2015 IET International Radar Conference, 2015.
- [1.16] S. Pacheco, R. Reuter, S. Trotta, D. Salle, J. John, “SiGe technology and circuits for automotive radar applications”, 2011 IEEE 11th Topical Meeting on Silicon Monolithic Integrated Circuits in RF Systems (SiRF), 2011.

- [1.17] “76-77 GHz RF transmitter front-end for W-band radar applications Data Sheet”, rev. 2.0, août 2016, www.nxp.com/radar
- [1.18] “76-77 GHz RF receiver front-end for W-band radar applications Data Sheet”, rev. 2.0, sept. 2016, www.nxp.com/radar
- [1.19] “76 - 81GHz MMIC transceiver (4 RX / 3 TX) for automotive radar applications Data Brief”, février 2016, www.st.com
- [1.20] E. Hyun, “Parallel and Pipelined Hardware Implementation of Radar Signal Processing for an FMCW Multi-channel Radar”, ELEKTRONIKA IR ELEKTROTECHNIKA, Vol 21, No 2, 2015
- [1.21] D. Kok; J. S. Fu, “Signal Processing For Automotive Radar”, IEEE International Radar Conference, 2005.
- [1.22] L. Zhao; W. Liu; X. Wu; J. S. Fu, “A Novel Approach for CFAR Processors Design”, 2001 IEEE Radar Conference, 2001.
- [1.23] “MPC577xK-MCU - Fact Sheet”, <http://www.nxp.com>
- [1.24] “RO3000 Series Circuit Materials RO3003, RO3006, RO3010 and RO3035 High Frequency Laminates datasheet”, <https://www.rogerscorp.com>
- [1.25] “RT/duroid® 5870 /5880 High Frequency Laminates datasheet”, <https://www.rogerscorp.com>
- [1.26] J. Wang et al., “Design of a High Isolation 35/94 GHz Dual-Frequency Orthogonal Polarization Cassegrain Antenna”, IEEE Antennas and Wireless Propagation Letters, 2016.
- [1.27] J. Schoebel, P. Herrero, “Planar Antenna Technology for mm-Wave Automotive Radar, Sensing, and Communications”, Radar Technology, Guy Kouemou (Ed.), 2010.
- [1.28] Z. Chen; Z. Y. Ping, “24-GHz microstrip grid array antenna for automotive radars application”, 2015 IEEE 5th Asia-Pacific Conference on Synthetic Aperture Radar (APSAR), 2015.
- [1.29] S. Beer et al, “An Integrated 122-GHz Antenna Array With Wire Bond Compensation for SMT Radar Sensors”, IEEE Transactions on Antennas and Propagation, 2013, Vol. 61, Issue 12.
- [1.30] B. H. Ku, O. Inac, M. Chang, H. H. Yang, G. M. Rebeiz, « A High-Linearity 76–85-GHz 16-Element 8-Transmit/8-Receive Phased-Array Chip With High Isolation and Flip-Chip Packaging”, IEEE Transactions on Microwave Theory and Techniques, 2014.
- [1.31] S. Beer et al., “Design and Measurement of Matched Wire Bond and Flip Chip Interconnects for D-Band System-in-Package Applications”, 2011 IEEE MTT-S International Microwave Symposium, 2011.
- [1.32] B. H. Ku et al, “A 77–81-GHz 16-Element Phased-Array Receiver With Beam Scanning for Advanced Automotive Radars”, IEEE Transactions on Microwave Theory and Techniques, 2014.
- [1.33] X. Fan, “Wafer Level Packaging (WLP): Fan-in, Fan-out and Three-Dimensional Integration”, 2010 11th International Thermal, Mechanical & Multi-Physics Simulation, and Experiments in Microelectronics and Microsystems (EuroSimE), 2010.
- [1.34] L. Devlin, “The Future of mm-wave Packaging”, in Microwave Journal, vol. 57, no. 2, pp. 24-38, février, 2014.
- [1.35] S. Trotta et al, “An RCP Packaged Transceiver Chipset for Automotive LRR and SRR Systems in SiGe BiCMOS Technology”, IEEE Transactions on Microwave Theory and Techniques, 2012.
- [1.36] N. H. Huynh et al, “eWLB package for millimeter wave application”, 2015 European Microelectronics Packaging Conference (EMPC), 2015.

- [1.37] K. Tsukashima, M. Kubota, O. Baba, T. Kawasaki, A. Yonamine, T. Tokumitsu, Y. Hasegawa, "E-Band Receiver and Transmitter Modules With Simply Reflow-Soldered 3-D WLCSP MMIC's", 8th European Microwave Integrated Circuits Conference, 2013, Nuremberg, Allemagne.
- [1.38] S. Jin; J. Zhang; J. Lim; K. Qiu; R. Brooks; J. Fan, "Analytical Equivalent Circuit Modeling for Multiple Core Vias in a High-Speed Package", 2016 IEEE International Symposium on Electromagnetic Compatibility (EMC), 2016.
- [1.39] H. W. Kang et al., "Equivalent Circuit Model of the IC-Stripline Coupling to IC Package", 2014 IEEE 18th Workshop on Signal and Power Integrity (SPI), 2014.
- [1.40] T. Mandic, B. K. J. C. Nauwelaers, A. Baric, "Simple and Scalable Methodology for Equivalent Circuit Modeling of IC Packages", IEEE Transactions on Components, Packaging and Manufacturing Technology, Vol. 4, Issue 2, 2014.
- [1.41] B. Kwal, "La théorie des équations de Maxwell et le calcul des opérateurs matriciels", J. Phys. Radium, pp.445-448, 1934.
- [1.42] E. Bouty, « Équations fondamentales du magnétisme induit, d'après Maxwell », J. Phys. Theor. Appl., 1881, 10 (1), pp.284-294.
- [1.43] J. M. Dunn, "Where Did EM Simulation Tools Go?", IEEE microwave magazine, janvier 2014

Bibliographie du chapitre 2

- [2.1] R. Fournié, « Diélectriques: bases théoriques », Technique de l'ingénieur (D2 I), D213.
- [2.2] A. M. Niknejad, "Analysis, Simulation, and Applications of passive Devices on Conductive Substrate" PhD Dissertation, University of California at Berkeley, 2000.
- [2.3] S. Gevorgian, H. Jakobson, T. Lewin, E. Kollberg "Design Limitations for Passive Microwave Components in Silicon MMICs", 28th European Microwave Conference, Oct. 1998.
- [2.4] J. P. John, V. P. Trivedi, J. Kirchgessner, D. Morgan, I. To, P. Welch, "An Enhanced 180nm Millimeter-Wave SiGe BiCMOS Technology with f_T/f_{MAX} of 260/350GHz for Reduced Power Consumption Automotive Radar IC's", 2014 IEEE Bipolar/BiCMOS Circuits and Technology Meeting (BCTM), 2014.
- [2.5] V. Blaschke, "ADS Interoperability for RFIC Design with ADS2016.01" January 5, 2016.
- [2.6] T. Hirano, K. Okada, J. Hirokawa, M. Ando, "Accuracy Investigation of De-Embedding Techniques Based on Electromagnetic Simulation for On-Wafer RF Measurements", <http://dx.doi.org/10.5772/48431> [consulté le 17/10/2016]
- [2.7] H. Hasegawa, M. Furukawa, H. Yanai, "Properties of Microstrip Line on Si-SiO₂ System," IEEE Transactions on Microwave Theory and Techniques, vol. 19, no. 11, pp. 869–881, Nov. 1971.
- [2.8] D. D. Grieg, H. F. Engelmann, "TEM wave properties of microstrip transmission lines", Proceedings of the IRE, Volume: 40, Issue: 12, décembre 1952
- [2.9] Anne-Laure Franc, thèse sur « les lignes de propagation intégrées à fort facteur de qualité en technologie CMOS. Application à la synthèse de circuits passifs millimétriques. », Université Grenoble Alpes, 2011.
- [2.10] G. Avenier et al., "0,13 μ m SiGe BiCMOS technology fully dedicated to mm-wave applications", IEEE Journal of Solid State Circuits, 2009.
- [2.11] B. Leite, thèse sur « Conception et modélisation de transformateurs intégrés millimétriques en technologies CMOS et BiCMOS », 2011, université de Bordeaux.
- [2.12] O. Bushueva, C. Viallon, A. Ghannam, T. Parra, "On-Wafer Measurement Errors Due to Unwanted Radiations on High-Q Inductors", IEEE Transactions on Microwave Theory and Techniques, Vol. 64, Issue 9, 2016.

Bibliographie du chapitre 3

- [3.1] B. Wu, T. Mo “Printed circuit board electrical design for wafer-level packaging”, 12th International Conference on Electronic Packaging Technology and High Density Packaging (ICEPT-HDP), 2011
- [3.2] T. Rovensky, A. Pietrikova, I. Vehec, I. Kmec, “Measuring of dielectric properties by microstrip resonators in the GHz frequency”, 38th International Spring Seminar on Electronics Technology (ISSE), 2015.
- [3.3] S. O. Land, O. Tereshchenko, M. Ramdani, F. Leferink, R. Perdriau, “Printed circuit board permittivity measurement using waveguide and resonator rings” International Symposium on Electromagnetic Compatibility, EMC’14, 2014, Tokyo, Japan
- [3.4] A. M. Albishi, O. M. Ramahi, “Surface crack detection in metallic materials using sensitive microwave-based sensors”, 2016 IEEE 17th Annual Wireless and Microwave Technology Conference (WAMICON), 11-13 avril 2016.
- [3.5] T. Jilani, M. P. W. Wong, Z. M.A, Yen Cheong Lee, “Dielectric characterization of meat using enhanced coupled ring-resonator” 2014 IEEE Asia-Pacific Conference on Applied Electromagnetics (APACE), 2014.
- [3.6] K. Chang, L. H. Hsieh, “Microwave Ring Circuits and Related Structures”, John Wiley & Sons, 2004.
- [3.7] Aakashdeep, Saurav Kr. Basu, G. V. Ujjwal, Sakshi Kumari, V. R. Gupta, “Measurement of Effective Dielectric Constant : A Comparison”, 2011 IEEE Applied Electromagnetics Conference (AEMC).
- [3.8] E. Seler, M. Wojnowski, R. Weigel, A. Hagelauer « New Simulation Procedure for Accurate Package Modeling Considering Chip-Package Interaction », IEEE 22nd Conference on Electrical Performance of Electronic Packaging and Systems (EPEPS), 2013.
- [3.9] B.-H. Ku, O. Inac, M. Chang, and G. M. Rebeiz, “75-85 GHz flip-chip phased array RFIC with simultaneous 8-transmit and 8-receive paths for automotive radar applications”, 2013 IEEE Radio Frequency Integrated Circuits Symposium (RFIC), 2013, pp. 371–374.
- [3.10] Saman Jafarlou, Mohammad Fakharzadeh, Behzad Biglarbegan, Mihai Tazlauanu, “Characterization of flip-chip interconnect for mm-wave system in package applications”, 2016 IEEE International Symposium on Antennas and Propagation (APSURSI), juin 2016.
- [3.11] Z. Chen, Y. L. Wang, Y. Liu, and N. H. Zhu, “Two-port calibration of test fixtures with OSL method” 2002 3rd International Conference on Microwave and Millimeter Wave Technology, 2002. Proceedings. ICMMT 2002, 2002, pp. 138–141.
- [3.12] **C. E. Souria**, T. Parra, G. Montoriol, C. Landez, “ Accurate Package Model extraction up to 110 GHz using one-port measurements. Application to a 77 GHz radar transceiver. ”, 2016 IEEE Radio and Wireless Symposium (RWS), 2016.
- [3.13] **C. E. Souria**, G. Montoriol, S. Thuries, “Semiconductor device package, electronic device and method of manufacturing electronic devices using wafer level chip scale package technology.”, déposé le 18 novembre 2015.
- [3.14] G. F. Engen, C. A. Hoer, “Thru-Reflect-Line: An Improved Technique for Calibrating the Dual Six-Port Automatic Network Analyzer,” *IEEE Trans. on Microwave Theory and Techniques*, vol. 27, no. 12, pp. 987–993, décembre 1979.

Bibliographie du chapitre 4

- [4.1] K. H. Lin, T. H. Yang, C. W. Hsu, "Implementation of a Low-Power Folded-Cascode RF Front-End for LTE Receivers", 2014 International Symposium on Computer, Consumer and Control, 2014
- [4.2] J. Rogers, C. Plett, "Radio Frequency Integrated Circuit Design", Artech House Microwave Library, 2003
- [4.3] "AFE5401-Q1 Quad-Channel, Analog Front-End for Automotive Radar Baseband Receiver Data Sheet", www.ti.com, décembre 2013, révisé mai 2016.
- [4.4] S. K. Reynolds, J. D. Powell, "77 and 94-GHz Downconversion Mixers in SiGe BiCMOS", 2006 IEEE Asian Solid-State Circuits Conference (ASSCC 2006), 2006.
- [4.5] A. Mariano, T. Taris, B. Leite, C. Majek, Y. Deval, E. Kerhervé, J-B. Bégueret, D. Belot, "Low power and high gain double-balanced mixer dedicated to 77 GHz automotive radar applications", 2010 Proceedings of the ESSCIRC, 14-16 sept. 2010.
- [4.6] C. H. Li, C. N. Kuo, "16.9-mW 33.7-dB Gain mmWave Receiver Front-End in 65 nm CMOS", 2012 IEEE 12th Topical Meeting on Silicon Monolithic Integrated Circuits in RF Systems (SiRF), 2012.
- [4.7] C. Viallon, G. Meneghin, T. Parra, "NMOS Device Optimization for the Design of a W-band Double-Balanced Resistive Mixer", IEEE Microwave and Wireless Components Letters, Vol.24, N°9, pp.637-639, 2014.
- [4.8] G. Meneghin, C. Viallon, T. Parra, "A double balanced resistive down-conversion mixer integrated in BiCMOS SiGe technology for 79 GHz automotive radar", IEEE NEWCAS-TAISA'09, 2009, Toulouse, France.
- [4.9] V. P. Trivedi, J. P. John, K. H. To, W.M. Huang, "A Novel Integration of Si Schottky Diode for mmWave CMOS, Low-Power SoCs, and More", IEEE Electron Device Letters, 2011.
- [4.10] C. Viallon, A. Magnani, M. Borgarino, T. Parra « Mélangeur résistif à sous-échantillonnage dans la bande 46–65 GHz intégré en technologie BiCMOS 0,13 μm », 19èmes Journées Nationales Microondes (JNM 2015), 2015, Bordeaux, France
- [4.11] H. Y. Su; R. Hu, C. Y. Wu, "A 78 ~ 102 GHz Front-End Receiver in 90 nm CMOS Technology", IEEE Microwave and Wireless Components Letters, septembre 2011.
- [4.12] **C. E. Souria**, C. P. Moreira, « BASEBAND AMPLIFIER CIRCUIT », EPO, brevet déposé le 1^{er} novembre 2016.
- [4.13] « application note: Noise Figure Measurement Accuracy – The Y-Factor Method », Keysight Technologies, www.keysight.com, 3 août, 2014.

Title: « Conception of innovative packages for 77-GHz automotive radar. Application to the design of an optimized packaged radar receiver channel »

Charaf-Eddine Souria

Director: Thierry Parra, Professeur, Université Toulouse 3 – LAAS

Co-supervisor: Gilles Montoriol, NXP Semiconductors

Keywords: automotive radar, interconnection, millimeter-waves, WLP package, 77 GHz, mixer

The development of automotive radars, at the frequency band 76-77 GHz, has experienced a significant growth over the last decade. Ongoing developments have to cope with two main challenges. The first challenge is reducing the cost to equip more car categories with these radars. The second challenge is to improve radar performance in order to satisfy the increasing demands of the road safety authorities and to equip the autonomous car. The automotive radar transceiver is the masterpiece of the system. Therefore, significant pressure is exerted on the semiconductor suppliers to develop next generation radars with superior performances and at lower cost than previous generations. Improving the radar transceiver performances requires improving these four main parameters: Noise Figure (NF), Power Amplifier (PA) power, Phase Noise (PN) and heat dissipation. Lowering the cost can be achieved by reducing test time, chip and PCB sizes, and wafers and package costs. We propose, in this work, a reduction of package cost and PCB size and improvement of heat dissipation by using a FI-WLP. The Wafer Level Chip Scale Package (WLCSP), the best known FI-WLP, was chosen for this application. It is the first time, in Silicon semiconductors history, that a FI-WLP is used at such high frequencies.

The first chapter describes the radar system in general and its main components. It focuses on the contribution of the transceiver then the package to the radar performances. The second chapter provides a methodology for EM models validation based on the modeling and experimental validation of passive structures on-chip. Innovations, significantly improving the WLCSP electrical performances, are revealed in the third chapter. The characterization of WLP is, itself, a challenge and novel methodologies to perform it are proposed in the same chapter. Thereafter, a new WLCSP packaged mixer, where block core and RF input matching are co-optimized, is designed and presented in the fourth chapter. The obtained NF is at the state-of-the-art, whereas the very constraining FI-WLP is used. All WLCSP transition and mixer simulation results are validated through measurement. This characterization confirms the excellent performances expected from this novel package and circuit designs.

Titre : « Conception d’interfaces boîtiers innovantes pour le radar automobile 77-GHz. Application à la conception optimisée d’une chaîne de réception radar en boîtier »

Charaf-Eddine Souria

Directeur de thèse : Thierry Parra, Professeur, Université Toulouse 3 – LAAS

Co-encadrant de thèse : Gilles Montoriol, NXP Semiconductors

Mots-clés : radar automobile, interconnexion, fréquences millimétriques, boîtier WLP, 77 GHz, mélangeur de fréquences

Le développement des radars automobiles, à la bande de fréquences 76-77 GHz, a connu une croissance importante au cours de la dernière décennie. Les développements en cours doivent faire face à deux grands défis. Le premier défi est la réduction du coût pour équiper plus de catégories de voitures avec ces radars. Le deuxième défi est l’amélioration des performances du radar afin de satisfaire les demandes croissantes des autorités de sécurité routière et d’équiper la voiture autonome. L’émetteur-récepteur radar automobile constitue le cœur du système. Par conséquent, une pression importante est exercée sur les fournisseurs de semi-conducteurs pour développer des radars de nouvelle génération avec des performances supérieures et à un coût inférieur par rapport aux générations précédentes. Améliorer les performances de l’émetteur-récepteur passe par l’amélioration de ces quatre paramètres : le facteur de bruit, le niveau de puissance de l’émetteur, le bruit de phase et la dissipation thermique. La réduction de coût peut être obtenue en réduisant le temps de test, les tailles de la puce et du PCB et le coût du boîtier. Dans ce travail, nous proposons une réduction du coût du boîtier et de la taille du PCB, en plus de l’amélioration de la dissipation thermique grâce à une encapsulation intégrée au niveau plaquette (FI-WLP pour Fan-In Wafer Level Package). Le boîtier WLCSP (Wafer Level Chip Scale Package), le plus connu FI-WLP, a été choisi pour cette application. C’est la première fois dans l’histoire des semi-conducteurs que le FI-WLP est utilisé pour du Silicium à des fréquences aussi élevées.

Le premier chapitre décrit le système radar et ses principaux composants. Il met l’accent sur la contribution de l’émetteur-récepteur, puis le boîtier, sur les performances du radar. Le deuxième chapitre fournit une méthodologie pour la modélisation électromagnétique et la validation expérimentale de ces modèles, appliquée à des structures passives sur puce. Des innovations, améliorant significativement les performances électriques du boîtier WLCSP, sont révélées dans le troisième chapitre. La caractérisation du WLCSP est en soi un défi. De nouvelles méthodologies de caractérisation de ce boîtier sont alors proposées dans le même chapitre. Par la suite, un nouveau mélangeur encapsulé en WLCSP est conçu et présenté dans le quatrième chapitre. Le facteur de bruit obtenu est à l’état de l’art, malgré l’utilisation du très contraignant boîtier FI-WLP. Tous les résultats de simulation de la transition WLCSP et du mélangeur sont validés par des mesures. Cette caractérisation confirme les excellentes performances attendues du boîtier et du circuit conçus.