



Conception de puces multi-fonctions MMIC GaN en bande Ka

Boris Berthelot

► To cite this version:

Boris Berthelot. Conception de puces multi-fonctions MMIC GaN en bande Ka. Réseaux et télécommunications [cs.NI]. Université Paul Sabatier - Toulouse III; Université de Sherbrooke (Québec, Canada), 2019. Français. NNT : 2019TOU30255 . tel-02494037v2

HAL Id: tel-02494037

<https://laas.hal.science/tel-02494037v2>

Submitted on 22 Oct 2020

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)

Cotutelle internationale avec l'Université de Sherbrooke, Québec

Présentée et soutenue par :

BERTHELOT Boris

le lundi 25 novembre 2019

Titre :

Conception de puces multi-fonctions MMIC GaN en bande Ka

École doctorale et discipline ou spécialité :

ED GEET : Électromagnétisme et Systèmes Haute Fréquence

Unité de recherche :

LAAS-CNRS (UMR)

Directeur/trice(s) de Thèse :

TARTARIN Jean Guy et VIALON Christophe

MAHER Hassan et BOONE François

Jury :

Jean Guy Tartarin/ Hassan Maher: Directeurs de thèse

Christophe Viallon/ François Boone: Co-directeurs de thèse

Nathalie Deltime/ Christophe Gaquière: Rapporteurs

Dominique Langrez: Examineur

Thierry Parra: Examineur

Oana Lazar: Examineur

Dominique Drouin: Invité

Rémy Leblanc: Invité

Remerciements

Cette thèse CIFRE en cotutelle a été effectuée entre différents sites en France, dans l'entreprise OMMIC initiatrice du projet, au Laboratoire d'Analyse et d'Architecture des Systèmes du Centre National de la Recherche Scientifique (LAAS-CNRS) de Toulouse, dans le groupe Microondes Opto-microondes pour Systèmes de Télécommunications (MOST). Une partie de ces travaux a été réalisée au Québec, Canada, à l'Institut Interdisciplinaire d'Innovation Technologique (3IT) à Sherbrooke, au sein du groupe Microélectronique III-V. Je remercie M. Liviu Nicu, Directeur du LAAS-CNRS, M. Eric Tournier, Maître de Conférences, au LAAS-CNRS responsable du groupe MOST de m'avoir accueilli dans l'équipe. Je tiens aussi à remercier M. Richard Arès Directeur du 3IT et M. Hassan Maher, Professeur à l'Université de Sherbrooke, responsable du groupe Microélectronique III-V de m'avoir accueilli dans son équipe.

Mes sincères remerciements vont à Mme Nathalie Deltimple, maître de conférences à l'université de Bordeaux, et M. Christophe Gaquière, Professeur à l'université de Lille, pour avoir évalué ce travail en la qualité de rapporteur, ainsi qu'à Mme Oana Lazar, ingénieur TAS, M. Dominique Langrez, ingénieur TAS, M. Thierry Parra, professeur à l'université de Toulouse, Mr Dominique Drouin, Professeur à l'Université de Sherbrooke et M. Rémy Leblanc, ingénieur OMMIC pour avoir accepté de prendre part à ce jury en tant qu'examineurs.

Mes plus profonds remerciements à Jean-Guy Tartarin, Directeur de thèse et Professeur à l'Université Paul Sabatier de Toulouse III, qui m'a fait confiance avant, durant et après les travaux de thèse. Son soutien m'a permis de m'améliorer dans bien des domaines qu'ils soient scientifiques ou plus personnelles. Je ne vais pas m'étendre par écrit et préfère partager cette sympathie de vive voix.

Merci à Christophe Viallon, co-directeur de thèse pour son aide de tous les instants.

J'adresse mes remerciements à M. Hassan Maher, Directeur de thèse et Professeur à l'Université de Sherbrooke, et M. François Boone, Co-directeur de thèse, pour leur confiance et soutien.

Merci à toutes les personnes que j'ai pu rencontrer durant ces années de thèse que ce soit à Limeil Brévannes, Sherbrooke ou Toulouse.

Merci à Alexandre Rumeau pour sa disponibilité et son soutien technique. Merci à Samuel Charlot et Bernard Franc pour leur disponibilité.

Et enfin un merci à Audrey ma fidèle collègue de bureau et au bro Napoléon.

Introduction

Les antennes réseaux à commande en phase (ARCP) ont traditionnellement été utilisées dans le domaine des radars. Pour ce type d'application les bandes passantes nécessaires étaient relativement étroites. Plus récemment cette contrainte a changé puisque les ARCP se sont généralisées pour des applications plus large bande exploitées dans les réseaux de télécommunications. Des techniques telles que le *beamforming* (conformation de faisceau) ont été déployées dans le but d'améliorer l'efficacité de la transmission des données. Cette dernière s'appuie sur des ARCP pour conformer les faisceaux plus efficacement dans une direction donnée, augmentant ainsi la qualité des bilans de liaison. Une architecture classique d'ARCP est illustrée en Figure 1. La chaîne comporte différents éléments qui se classent en deux catégories.

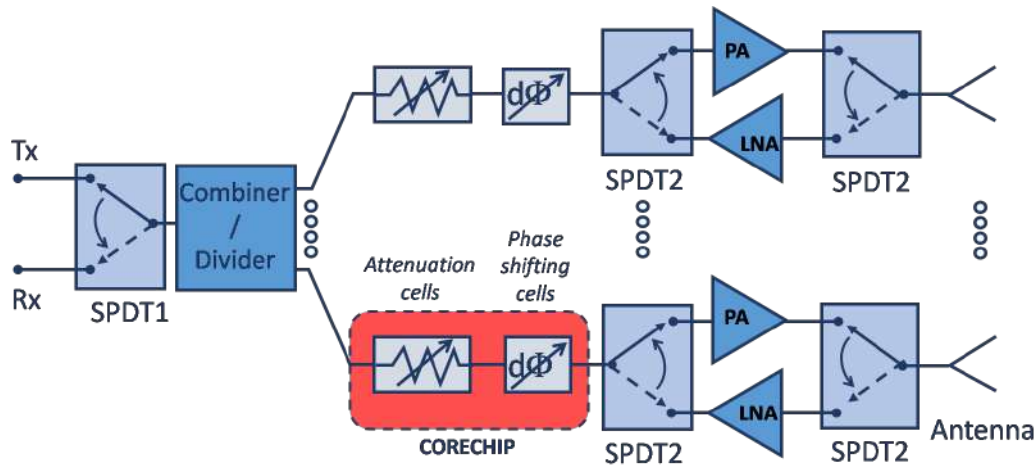


FIGURE 1 – Schéma d'une architecture d'émission-réception de type ARCP

La partie *front-end*, qui est constituée de tout ce qui traite de l'amplification (puissance en émission, et faible bruit en réception) et la partie *back-end* où le traitement du signal est réalisé. Cette dernière partie regroupe les fonctionnalités dites de contrôle (atténuation et déphasage) et est entourée en rouge sur la Figure 1. L'architecture présentée peut être utilisée à la fois en émission et en réception, l'orientation du signal vers l'une ou l'autre de ces options est garantie par des commutateurs ici notés SPDT pour *Single Pole Double Throw* ou «Commutateur une entrée deux sorties» en français. A faibles niveaux de puissances et selon la technologie utilisée, des limiteurs et circulateurs peuvent remplacer ces commutateurs.

La partie centrale du travail exposé dans ce manuscrit réside dans l'étude des fonctions de contrôle,

c'est-à-dire les puces MMIC regroupant ces fonctionnalités, et plus spécifiquement leur intégration en technologie GaN selon la filière développée par la société OMMIC. Ces fonctions sont souvent nommées *core-chips*, elles permettent un contrôle numérique de la phase et de l'atténuation des signaux transitant par ce module (il existe des versions avec un contrôle analogique des phases et atténuations). En contrôlant les paramètres de phase et d'atténuation selon une loi de commande prédéfinie sur chaque branche du système et en espaçant les différentes antennes d'une distance donnée (typiquement la demi-longueur d'onde quand l'intégration le permet), il est possible de former ou de recombinaison les signaux traités par chaque antenne. De plus, le contrôle électronique individuel des antennes autorise des balayages focalisés (faisceau directif) ou larges. La Figure 2 illustre le fonctionnement d'un tel réseau avec la recombinaison des ondes en un front uniforme.

Pour implémenter le *beamforming* des éléments déphaseurs et atténuateurs sont nécessaires. Ces différents éléments peuvent être réalisés tout aussi bien grâce à des circuits actifs que des circuits passifs et être à commande analogique ou numérique. La première permet une couverture angulaire et modulaire continue et globale définie par la tension de polarisation. Cette dernière permet aussi de corriger des erreurs dues aux dégradations de performances liées au vieillissement ou bien à des dérives relatives à la dispersion technologique lors de la fabrication. Cependant ces avantages se font au détriment d'une consommation statique importante et d'une sensibilité aux fluctuations ou dérives de la tension d'alimentation. Pour des applications satellites, le contrôle de la tension d'alimentation est rendu plus critique car les signaux de contrôle sont principalement numériques. La conversion de ces derniers en signaux analogiques nécessite donc l'implémentation de convertisseurs numérique/analogique qui pénalisent l'optimisation de l'espace occupé.

Une solution alternative consiste en l'utilisation de circuits passifs sous forme de cellules mises en cascade réalisant chacune des pas d'atténuation et de phase discrets. On parle alors de fonctions de contrôle numériques. Cette solution s'affranchit des problèmes dus aux tensions d'alimentation du fait du fonctionnement en commutation des transistors. Cette stabilité fonctionnelle est garantie au prix de pertes d'insertion plus élevées, caractéristiques inhérentes à l'utilisation de circuits passifs. Dans ce manuscrit, nous étudierons principalement les *core-chips* passifs même si les performances de certains actifs seront présentées à titre de comparaison. Il faut tout de même préciser que les puces «passives» présentées et issues de littérature présentent une certaine consommation. Ceci est dû au fait que les auteurs considèrent alors la consommation globale du module TxRx, et pas celle du *core-chips* à proprement parler.

L'intégration des cellules de phase et d'atténuation dans un *core-chip* complet nécessite un soin particulier. En effet, les cellules n'étant pas parfaites (erreur de phase/amplitude [1] et/ou désadaptations d'impédances [2]), les performances du système finales peuvent être dégradées. Des solutions algorithmiques pour identifier les agencements aux performances les plus favorables ont déjà été mis au point [3] mais celles-ci ne vont pas plus loin que l'obtention de l'ordre optimal. Identifier les cas les plus favorables, les cellules critiques et les potentielles retouches de cellules permettrait donc de gagner un temps précieux de conception et de repousser les performances des *core-chips*.

Ce manuscrit est séparé en 4 chapitres. Le premier dresse un état de l'art de la littérature concernant les *core-chip* allant des technologies aux topologies utilisées. Puis il présente notre méthodologie de conception de *core-chips* GaN en bande Ka. Le second chapitre présente les simulations des cellules de déphasage et d'atténuation que nous utiliserons pour réaliser les *core-chips* en mode *single-ended* ou différentiel. Le troisième chapitre traite de la problématique de la mise en commun des cellules et nous présentons notre algorithme d'ordonnancement des cellules qui nous permet d'identifier les agencements de cellules les plus favorables ou défavorables. A cela nous ajoutons l'évaluation de l'impact d'amélioration des cellules dans l'optique d'une ou de plusieurs retouches. Enfin nous terminons dans le chapitre 4 par les mesures des circuits réalisés, 5 au total. Nous avons réalisé 2 circuits *single-ended*

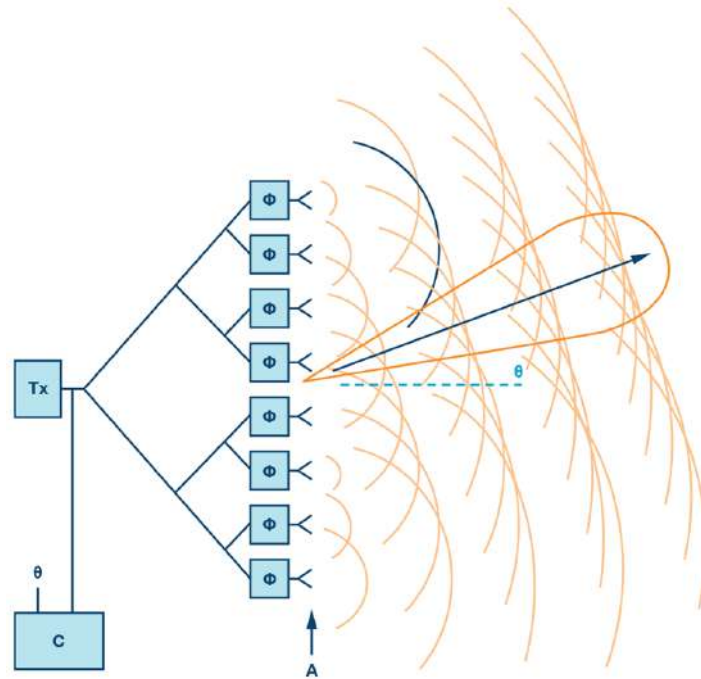


FIGURE 2 – Schéma explicatif de la recombinaison des ondes dans un ARCP, avec C le bloc de contrôle et θ l'angle visé et 3 circuits différentiels.

Bibliographie

- [1] Matthew A. Morton, Jonathan P. Comeau, John D. Cressler, Mark Mitchell, and John Papapolymerou. Sources of Phase Error and Design Considerations for Silicon-Based Monolithic High-Pass/Low-Pass Microwave Phase Shifters. *IEEE Transactions on Microwave Theory and Techniques*, 54(12) :4032–4040, dec 2006.
- [2] Inder J. Bahl. *Control Components Using Si, GaAs, and GaN Technologies*. Artech House, 2014.
- [3] Andrea Bentini, Mauro Ferrari, Walter Ciccognani, and Ernesto Limiti. A novel approach to minimize RMS errors in multifunctional chips. *International Journal of RF and Microwave Computer-Aided Engineering*, 22(3) :387–393, may 2012.

Table des matières

Remerciements	iii
Introduction	v
Bibliographie	vii
1 Les puces multi-fonctions	1
1 Les technologies utilisées	1
2 La place du GaN dans les chaînes d'émission-réception	2
3 Composition et critères d'évaluation des ARCP	3
4 Les déphaseurs	4
4.1 Les déphaseurs analogiques	4
4.2 Les déphaseurs numériques	5
5 Les atténuateurs	14
6 Les core-chips	16
7 La place du GaN dans les puces multifonctions	18
8 Méthodologie de conception	20
9 Modes de fonctionnement	25
Bibliographie	27
2 Conception des cellules de déphasage et d'atténuation	31
1 La société OMMIC	31
2 Présentation du kit de conception de la filière GaN D01GH	32
2.1 Les éléments passifs	32
3 Description simulation EM	34
4 Cahier des charges	36
4.1 Approche large bande	36
4.2 Approche mono-fréquence	37
5 Réalisation cellules individuelles single ended	37
5.1 Cellule 11,25° single ended	37
5.2 Résultats pour toutes les cellules single-ended	40
6 Conception des cellules différentielles	48
6.1 Cellule 180° différentielle	48
6.2 Résultats pour toutes les cellules différentielles	50
Bibliographie	58

3	Mise en commun et optimisation de l'ordre des cellules	59
1	Choix de l'ordre des cellules	59
2	Méthodologie de mise en commun des cellules développée	62
2.1	Méthodologie sans algorithme	63
2.2	Méthodologie avec algorithme	64
3	Application pour les cellules single ended	65
3.1	Sans utilisation de l'algorithme	65
3.2	En utilisant l'algorithme	71
3.3	Optimisation des cellules	78
4	Application pour les cellules différentielles	90
4.1	Sans utilisation de l'algorithme	90
4.2	En utilisant l'algorithme	90
	Bibliographie	99
4	Mesures des puces multifonctions	101
1	Environnement de mesures	101
2	Circuits single-ended	104
2.1	Première version	105
2.2	Deuxième version	105
3	Circuits différentiels	109
3.1	Première version	110
3.2	Deuxième version et troisième version	115
	Bibliographie	118
	Conclusion générale	119
	Résumé	121
	Liste des publications	123

Chapitre 1

Les puces multi-fonctions

1 Les technologies utilisées

La tendance allant toujours vers la miniaturisation des systèmes et vers une réduction des prix des modules électroniques RF, les concepteurs se retrouvent face à de nouvelles contraintes de conception. En effet il faut trouver des technologies et/ou des topologies répondant à la fois aux besoins qui s'expriment en termes de performances (efficacité, compacité, linéarité...) mais aussi en termes de coût pour certaines applications grand public. La technologie Silicium Germanium (SiGe) offre de grandes possibilités en termes de fonctionnalités (forte densité d'intégration par l'utilisation de technologies mixtes BiCMOS, maîtrise des erreurs de phase et d'amplitude) tout en gardant un coût bien inférieur aux technologies III-V (Arséniure et Nitrure de Gallium). Cependant, ses propriétés physiques l'empêchent de fournir des amplificateurs d'une efficacité suffisante pour la fabrication de module d'émission réception «tout SiGe». Ainsi il est nécessaire de panacher les technologies, en utilisant par exemple des *core-chips* SiGe associés à des modules de gestion de la puissance en GaAs ou plus récemment en GaN [1]. Ce point limite l'efficacité d'intégration initialement visée par la technologie BiCMOS. Pour des applications où le coût n'est pas un facteur limitant ferme (spatial et militaire), la technologie GaAs conserve quelques parts de marché. De par ses propriétés physiques, elle permet de réaliser à la fois les fonctions de contrôle mais aussi la gestion de puissance. Le Tableau. 1.1 recense les principales propriétés physiques des technologies majeures utilisées dans notre domaine d'étude (SiGe et GaAs).

Technologie	SiGe	GaAs	GaN
Bande interdite (eV)	1	1.4	3.4
Mobilité des électrons ($cm^2/V.s$)	2800	8500	2000
Vitesse de saturation des électrons ($10^7 cm/s$)	2.2	2	2.5
Champ de claquage ($10^6 V/cm$)	0.25	0.4	>5
Conductivité thermique ($W.cm^{-1}.K^{-1}$)	1.1	0.5	1.3

TABLE 1.1 – Propriétés physiques du SiGe, GaAs et GaN [2], [3]

A ce tableau est rajoutée la technologie GaN qui, à l'heure de l'initiation de nos travaux, n'était pas encore exploitée pour réaliser de telles fonction. Le GaN fait état de multiples avantages relativement aux autres technologies, en termes de bande interdite et de champ de claquage (associés à une meilleure dissipation thermique), en plus de son aptitude à fonctionner à hautes fréquences (mobilité et vitesse de saturation électroniques). La large bande interdite rend le matériau plus robuste à des perturbations

électromagnétiques externes. Le champ de claquage important implique directement des tensions et courants relatifs aux ondes RF admissibles importantes, ce qui permet l'obtention de niveaux de puissance en sortie bien supérieures à ceux obtenus par les technologies concurrentes SiGe et GaAs. La conductivité thermique élevée assure une bonne évacuation de la chaleur générée par le fonctionnement du composant, permettant ainsi une intégration plus dense que pour le GaAs.

2 La place du GaN dans les chaînes d'émission-réception

Dans ses travaux de thèse, Tyler Ross [4] met en évidence les intérêts d'une chaîne d'émission/réception exploitant pleinement la technologie GaN, par rapport à l'architecture conventionnelle induite par l'utilisation de technologies à faible bande interdite. Comme illustré en Figure 1.1a où les différences d'organisation de la chaîne selon la technologie utilisée apparaissent en bleu sur la Figure 1.1b, il est possible de ré-optimiser la chaîne d'émission et de réception.



FIGURE 1.1 – Différences d'agencement pour une chaîne Tx/Rx en technologie (a) non GaN (b) GaN [2]

En premier lieu, il est possible de mettre les déphaseurs au plus proche de l'antenne en raison de la meilleure tenue en puissance du GaN et de supprimer les limiteurs traditionnellement placés avant les amplificateurs faible bruit (LNA). Avec cette nouvelle disposition, le nombre de déphaseurs utilisés et donc les pertes induites sont grandement diminuées grâce à la réduction du nombre d'éléments sur la chaîne. La réduction des pertes améliore ainsi le facteur de bruit de la chaîne de réception et améliore l'efficacité de la chaîne de transmission. De plus, la capacité des systèmes passifs GaN à travailler avec des signaux à moyenne puissance, simplifie l'implémentation des amplificateurs de puissance (PA). En effet la contrainte sur le gain de ces derniers sera moins importante grâce au niveau de puissance d'entrée déjà élevé ce qui peut leur éviter de travailler à des forts niveaux de compression.

Deuxièmement, OMMIC fabriquant déjà des *T/R chips* (composants regroupant la partie gestion de la puissance, amplificateur de puissance et amplificateur faible bruit) en GaN, nous comprenons bien l'intérêt de réaliser la partie fonction de contrôle dans cette même technologie. D'une part cela permettrait une intégration monolithique comprenant toute la chaîne de réception sur une même puce. Et d'autre part, en utilisant les technologies classiques de *core-chips* SiGe, des incompatibilités de puissances apparaissent aux interfaces des éléments SiGe et GaN. Ces incompatibilités menant à l'ajout d'éléments limiteurs après les LNA et des pré-amplificateurs avant les PA, ces éléments dégradant la compacité du système.

Un des objectifs initiaux de la thèse envisageait l'intégration de blocs amplificateurs de compensation des pertes au sein de la partie fonction de contrôle du *core-chip* de façon à obtenir une dynamique de 6 bits de phase ($5,625^\circ$ de résolution), 5 bits d'atténuation (0,5dB de résolution) et une puissance de sortie de 10dBm. Dans la continuation de cet objectif, une intégration monolithique complète de la puce d'émission réception pourrait être envisagée. En effet, les éléments d'amplification en SiGe étant bien moins performants que les technologies III-V en terme de puissance, la majorité des puces actuelles

regroupent 2 technologies (SiGe-GaAs [5] ou SiGe-GaN [1]) dont la connectique est assurée par des transitions alumines encombrantes, et de fait peu compatibles avec une intégration maîtrisée telle que le nécessite l'exploitation de la bande Ka (du fait de la faible longueur d'onde qui contraint la dimension d'antenne, qui elle-même définit l'encombrement du module Tx/Rx). Ces dernières topologies de puces présentent donc deux inconvénients majeurs qui concernent la place occupée et les pertes associées [6]. Prouver la faisabilité d'une puce Tx/Rx « tout GaN » présentant des tailles concurrentielles, permettrait de s'affranchir de ces transitions qui dégradent les performances finales.

Après avoir dressé le contexte d'utilisation des modules *core-chips*, et plus spécifiquement leur réalisation en technologie MMIC GaN, nous allons maintenant nous intéresser plus précisément aux éléments qui les constituent. Ensuite nous présenterons les paramètres sur lesquels les performances des fonctions de contrôle sont évaluées et nous étudierons les différentes topologies de déphaseurs et d'atténuateurs puis nous conclurons par un comparatif des performances à l'état de l'art en bande Ka des atténuateurs et déphaseurs, mais aussi des *core-chips*.

3 Composition et critères d'évaluation des ARCP

Pour évaluer les performances des éléments constituant les *core-chips*, les caractéristiques suivantes sont étudiées :

-l'erreur de phase : c'est un élément primordial inhérent à la maîtrise de la fonction même de déphasage. Elle traduit la différence entre la phase désirée et la phase obtenue. Pour des déphaseurs multi-états, cette erreur correspond à l'écart relatif entre la phase de référence mesurée ou simulée et la phase mesurée ou simulée à l'état étudié (auquel on vient enlever la phase de consigne théorique pour obtenir l'erreur relative). Cette erreur peut s'exprimer en erreur absolue en degrés mais elle est généralement exprimée par l'erreur quadratique moyenne (*RMS error*) en degré l'expression générale d'une erreur *RMS* est donnée en Equation 1.1 . Ce type d'erreur est généralement employé puisqu'il permet de quantifier une erreur relative à de multiples états de phase. La notion de moyenne dans ce calcul d'erreur implique une mauvaise description d'éventuels états réfractaires qui seraient très éloignés de leur consigne. Cette erreur correspond aussi à la différence de phase entre les deux états des atténuateurs, qui devra elle aussi être minimisée.

$$RMS = \sqrt{\sigma_x^2 + \bar{x}^2} \quad (1.1)$$

Avec σ_x l'écart type de l'erreur et $\bar{x}$ la moyenne des erreurs.

-l'erreur d'atténuation : elle traduit la différence (ou le rapport) d'amplitude (en dB ou en valeur naturelle) entre les deux états de fonctionnement de chaque cellule. Elle peut également être engendrée par la fonction de déphasage, à l'instar de l'erreur de phase potentiellement générée par la fonction atténuateur. Selon le déphaseur étudié, elle sera calculée selon différents indicateurs. Pour les déphaseurs analogiques, elle correspond à l'atténuation/gain associé au changement de tension de polarisation. Pour les déphaseurs numériques, elle va correspondre à la différence de pertes d'insertion entre les deux états du déphaseur. C'est bien souvent ce paramètre qui vient complexifier le dessin du déphaseur puisqu'une erreur d'atténuation faible implique de garder des pertes d'insertion quasiment égales suivant les deux états commutés tout en garantissant une différence de phase contrôlée et plate sur la bande de fréquence. La diversité des trajets utilisés selon la consigne d'état peut ainsi avoir un impact sensible sur la différence d'atténuation d'un déphaseur. Pour les atténuateurs elle correspond à la différence entre l'atténuation obtenue et l'atténuation de consigne ; c'est l'objectif principal de conception

de chaque cellule. Cette erreur est prise en compte durant le calcul de l'erreur d'atténuation globale du *core-chip* (erreur d'atténuation des atténuateurs et erreur d'atténuation des déphaseurs).

-l'adaptation en impédance des accès (généralement exprimée en dB) : ce paramètre traduit le rapport d'onde réfléchi (relativement à l'onde incidente) quand le circuit est fermé sur une impédance de référence spécifique (Z_0). La valeur de $Z_0=50\Omega$ a été choisie par convention, et ce sera sur cette charge que le niveau d'adaptation d'impédance sera exprimé si rien n'est précisé dans les sections relatives à ce paramètre dans la suite du document. Ce paramètre est particulièrement critique lors de la mise en cascade de plusieurs éléments ; en effet comme il sera détaillé dans la suite du manuscrit, le comportement d'une cellule individuelle (de phase ou d'atténuation) dépend de ses conditions d'ouverture et de fermeture (impédances placées à son entrée/sortie). Ainsi, obtenir un bon niveau d'adaptation de chaque cellule constitutive du *core-chip* (par exemple $>20\text{dB}$ ou $<-20\text{dB}$ selon les conventions utilisées) permet de limiter l'influence des cellules les unes sur les autres.

-la bande de fréquence de fonctionnement : elle se définit par la largeur de bande de fréquence (en Hz) sur laquelle les 3 paramètres précédemment cités respectent le cahier des charges. Il est aisé d'imaginer que la réalisation simultanée des critères précédemment cités sera plus difficile à maîtriser conjointement sur une bande de fréquence importante. Ceci aura une grande incidence dans le cas de notre étude.

-la linéarité : La linéarité définit la plage de puissance sur laquelle le gain d'un déphaseur analogique et les pertes d'un déphaseur/atténuateur numérique sont constants. Elle est souvent définie par le point de compression à 1dB ou par le point d'interception d'ordre 3 qui définit le produit d'intermodulation d'ordre trois. Dans le cas de ce travail, ces paramètres sont des éléments importants puisque le GaN, par ses propriétés intrinsèques, offre une linéarité et des niveaux de puissances admissibles supérieurs à ceux de ses concurrents de type GaAs et SiGe. Cependant, l'étude ne validera cet aspect de linéarité que comme étant la conséquence d'une étude exclusive sur les quatre premiers points, c'est-à-dire qu'aucune optimisation ne peut être initialement engagée sur la linéarité a priori des fonctions premières visées.

4 Les déphaseurs

Le contrôle de la phase du signal est un élément primordial dans un système de communication à pointage électronique. En effet, c'est du contrôle de la phase de chaque élément du réseau que dépend la précision du pointage de faisceau.

Pour réaliser la fonction de déphasage, deux possibilités existent, les dispositifs de temporisation, ou lignes à retard (*Time delay*), et les déphaseurs. Nous parlons de temporisation quand le déphasage créé est proportionnel à la fréquence de fonctionnement (lors de l'utilisation d'une longueur électrique par exemple). Les temporisations sont principalement utilisées pour des applications où des délais supérieurs à 360° (une longueur d'onde) sont nécessaires. Par exemple dans le cas de radars (ou leurres radar), ces dernières permettent de faire du traitement de signal sur de nombreuses impulsions reçues (émises), en appliquant un délai à celles-ci. Dans le cas de sources de fréquence, les temporisations peuvent permettre de synchroniser des horloges de fréquences différentes. Dans les applications que nous visons, une couverture de 360° est suffisante et permet une largeur de bande de fréquence importante, ainsi notre choix se portera plutôt vers la deuxième catégorie d'éléments de contrôle de phase : les déphaseurs. Ceux-ci se divisent en deux grandes catégories : analogiques et numériques.

4.1 Les déphaseurs analogiques

Le principe général est de faire varier une charge (le plus souvent réactive) grâce à une tension permettant d'ajuster le déphasage souhaité, à l'instar des varactors. Grâce à ce principe ils peuvent

produire un déphasage continu nécessaire à certaines applications. Le déphasage réalisé est un déphasage absolu, c'est-à-dire qu'il est appliqué sans référence particulière en temps réel selon le contrôle de commande. Les déphaseurs analogiques sont à l'heure actuelle les plus utilisés, mais tendent à être remplacés par leurs homologues numériques. En termes de performances, ils sont très dépendants de la tension de polarisation comme évoqué précédemment, et donc très sensibles aux fluctuations affectant celle-ci. De plus, selon Pozar [7] les erreurs de phase et les pertes d'insertion générées par les déphaseurs analogiques dégradent le gain du réseau total.

4.2 Les déphaseurs numériques

L'appellation numérique vient du fait que ces déphaseurs ont un comportement à deux états distincts assimilables à une logique binaire (ON/OFF), cependant leur fonctionnement intrinsèque reste analogique. Cette fois ci, un déphasage relatif est produit puisque celui-ci est calculé par rapport à un état de référence. Le déphasage total résulte donc de la somme de déphasages d'un ensemble de cellules individuelles. Ainsi chaque combinaison binaire est associée à un déphasage résultant de la mise à 1 ou à 0 de la cellule pilotée par le bit. Un exemple de cette configuration est donné en Figure 1.2 pour une association de n cellules individuelles, et le calcul du déphasage total est précisé en Equation 1.2. Plusieurs topologies de circuits existent et peuvent être classées en quatre catégories : les lignes à retard, les lignes chargées, les déphaseurs à filtres et les topologies réfléchives. Ils existent d'autres types de déphaseurs qui sont moins utilisés dans l'industrie au vu de leurs performances moins intéressantes, ils ne seront pas traités dans cet état de l'art.

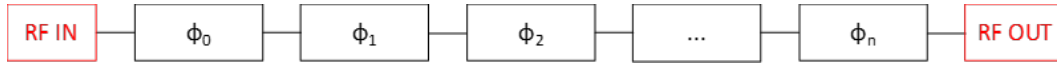


FIGURE 1.2 – Schéma d'un déphaseur numérique à n cellules (n bits)

$$\phi_{total} = b_0 \times \phi_0 + b_1 \times \phi_1 + b_2 \times \phi_2 + \dots + b_n \times \phi_n \quad (1.2)$$

Avec b_n la valeur du bit associé à la cellule n , ϕ_n le déphasage associé à la cellule n et ϕ_{total} le déphasage total.

La ligne à retard

L'intérêt principal de cette topologie réside dans sa simplicité de conception. En effet le déphasage est produit grâce à la différence de longueur électrique entre les lignes mises en jeu : elle peut être classée dans la catégorie des temporisations car le déphasage est proportionnel à la fréquence. Lorsque la cellule est désactivée le signal passe par l'état dit de référence (le plus souvent par le transistor série). C'est l'état de phase du transistor passant qui va donc déterminer la phase initiale à partir de laquelle le déphasage de la cellule est calculé.

Un compromis apparaît vite entre le déphasage du transistor et les pertes d'insertion (le moins de signal possible doit passer dans la ligne à retard quand le transistor est passant). La ligne à retard est ensuite dimensionnée pour produire un déphasage par rapport à l'état de référence. La Figure 1.3 illustre le fonctionnement de cette topologie, celle-ci est simple de conception mais ne permet pas d'obtenir de déphasages importants, principalement à cause de l'encombrement engendré par la ligne qui devient vite problématique avec l'augmentation du déphasage souhaité. De plus, la longueur électrique variant avec la fréquence, la différence de longueur électrique entre la ligne de référence et la ligne à

retard n'est pas constante, l'équation décrivant le fonctionnement de ce type de déphasage est donnée en Equation 1.3.

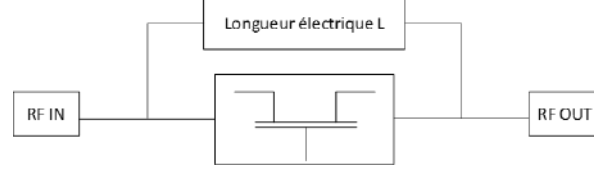


FIGURE 1.3 – Topologie de déphaseur à base d'une ligne à retard de longueur électrique L

Ceci empêche l'obtention d'un déphasage constant sur toute la bande de fréquence. Toutefois, lorsque des déphasages moins importants sont nécessaires cette topologie offre une facilité d'adaptation et de faibles pertes d'insertion assimilables au premier ordre à celles des commutateurs composant l'élément. La mise en parallèle d'un transistor en commutation et d'une ligne entraîne automatiquement une résonance due à la capacité C_{off} du transistor. Pour éviter que ce comportement se manifeste dans la bande de fréquence d'intérêt, le couple transistor/ligne est soigneusement choisi. Il n'y a pas de limitation physique au déphasage atteignable mais le plus souvent cette topologie est utilisée pour des déphasages inférieurs à 20° .

$$\Delta\phi = \frac{2 \times \pi \times f}{v_p} \times (l_{transistor} - l_{ligne}) \quad (1.3)$$

Avec $\Delta\phi$ le déphasage souhaité, f la fréquence du signal de porteuse, v_p la vitesse de phase et $l_{transistor}$ et l_{ligne} correspondant respectivement aux longueurs électriques du transistor et de la ligne.

Il est rare de trouver des circuits MMIC exploitant uniquement une cellule de déphasage et utilisant une ligne à retard. Et même dans le cas contraire, il est rare d'obtenir des renseignements sur la cellule individuelle, en raison de la simplicité de mise en œuvre d'un tel système. En effet cette topologie peu complexe ne nécessite pas d'aménagement circuit particulier. Il existe cependant des déphaseurs utilisant uniquement des topologies à base de ligne à retard [8] [9] [10]. Le plus souvent, l'élément de commutation utilisé est la diode PIN. Malgré sa facilité d'intégration et sa haute tenue en puissance, sa consommation statique lui fait perdre l'avantage face à un transistor en commutation. En Figure 1.4 est donné un exemple de ce type de topologie. Un grand nombre de diodes PIN est utilisé, celles-ci devant être pilotées par des courants typiques de 20mA, la consommation statique de l'ensemble est très importante. Les différences de longueur de ligne (haut et bas) sont bien visibles et sont responsables du déphasage des cellules successives (respectivement de 180° , 90° , 45° , 22.5° , 11.25° et 5.625° de la gauche vers la droite). L'espace occupé de 32.4mm^2 en fait une puce extrêmement encombrante à cette fréquence.

La Figure 1.5 illustre la faisabilité de circuit beaucoup plus compact grâce à des lignes à retard (1.544mm^2). Cependant cette réduction de taille s'explique par l'augmentation de la fréquence et donc de la réduction de la longueur d'onde du signal. Les applications visées par ce type de circuits sont plutôt les radars alors que le circuit de la Figure 1.4 semble plus approprié pour des applications de télécommunications de type 5G.

Une réalisation de lignes à retard exploitant des transistors comme élément de commutations est présentée en Figure 1.6. Dans le cas de ce circuit, les transistors utilisés sont dits non résonnants, cela se traduit par une isolation de l'élément de commutation bien moins dépendante de la capacité C_{off} du transistor utilisé. Du fait de leur grande taille et de leur consommation statique importante dans certains cas, ce type de topologie est rarement utilisée pour obtenir des performances à l'état de l'art.

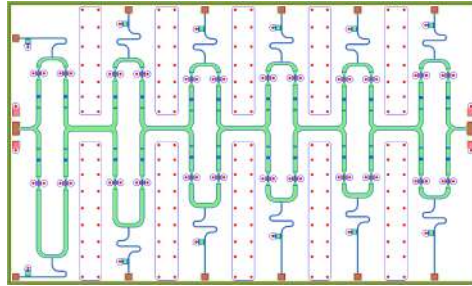


FIGURE 1.4 – Exemple de déphaseur 6 bits AlGaAs à base de ligne à retard fonctionnant de 27.5 à 29.5 GHz pour une taille de 7.34*4.41mm² [1]

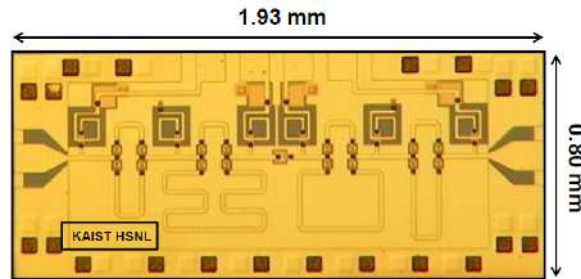


FIGURE 1.5 – Exemple de déphaseur 4 bits InGaAs fonctionnant à 81GHz [9]

La ligne chargée

Cette topologie consiste à utiliser une ou plusieurs charges réactives en shunt placées à un niveau stratégique du circuit comme illustré en Figure 1.7. Cette charge va, quand elle est connectée, influencer sur la phase du circuit tout en préservant l'amplitude du signal. Les pertes d'insertion de ce type de topologie sont encore une fois déterminées par les pertes des commutateurs. En effet, ceux-ci ne peuvent présenter ni charge purement réactive à l'état OFF ni de résistance nulle à l'état passant (R_{on}), ces commutateurs fixent donc les limites du circuit. Selon [11] après l'étude de la matrice chaîne du quadripôle formé par le circuit, trois paramètres sont à déterminer : l'impédance de la ligne, la valeur des susceptances des charges et la longueur électrique de la ligne entre les deux charges. Cette topologie est généralement utilisée pour des déphasages inférieurs à 45° et propose une meilleure platitude de phase. Cependant, il est souvent nécessaire que la distance entre les deux charges soit de l'ordre de $\lambda/4$ (2.9mm à 35GHz), et la taille devient donc un facteur préjudiciable au choix de cette topologie. Ces raisons font que cette topologie est rarement utilisée à la fréquence d'intérêt.

Ce type de topologie n'est que très rarement utilisée en bande Ka. La Figure 1.8 montre un exemple d'utilisation de cette topologie, le repère spatial en bas de la figure confirme bien la taille bien trop importante (6cm de longueur) de ce type de circuit. Malgré la fréquence bien en deçà de la fréquence ciblée, ces topologies ne sont pas citées dans le tableau regroupant les performances à l'état de l'art des différents types de déphaseurs.

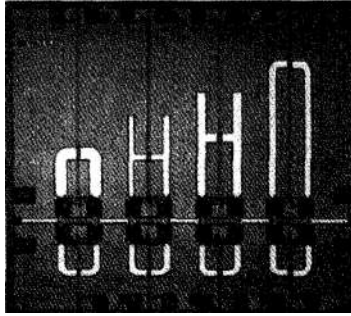


FIGURE 1.6 – Exemple de déphaseur 4 bits InGaAs HJFET fonctionnant entre 33 et 36 GHz pour une taille de $2.5 \times 2.2 \text{ mm}^2$ [10]

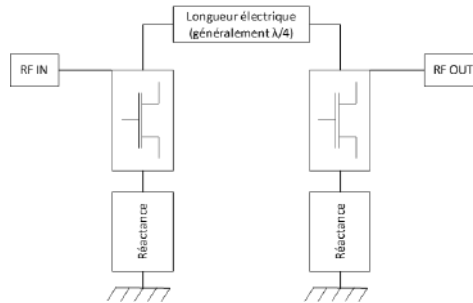


FIGURE 1.7 – Topologie de déphaseur à base de type ligne chargée

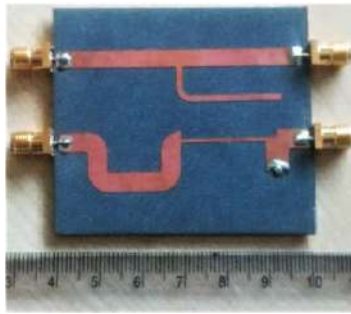


FIGURE 1.8 – Exemple de topologie à base de ligne chargée proposant un déphasage de 90° à 1.8GHz [12]

Topologies à filtre(s)

La topologie dite à filtre utilise les propriétés en phase des filtres passe-haut et/ou passe-bas. Ils proposent des déphasages constants sur de larges bandes de fréquences. Même si la topologie simple qui propose une commutation entre un chemin électrique de référence et un filtre ne peut pas dépasser un déphasage absolu de 90° , en commutant entre deux filtres proposant chacun un déphasage allant jusqu'à 90° et -90° il est possible d'obtenir un déphasage de 180° . Pour ce dernier déphasage, cette topologie est la plus adaptée puisqu'elle offre les meilleures performances en termes de platitude de phase et de largeur de bande. Ce qui en fait la topologie la plus utilisée dans les *core-chips*. Selon la

fréquence d'intérêt le choix entre filtre passe-bas et filtre passe-haut est fait. Sachant que pour les passe-haut, la capacité sur la ligne série peut être réalisée grâce à la capacité C_{off} du transistor. Pour garder un bon contrôle de la capacité série équivalente, la mise en parallèle du transistor et d'une capacité fixe est préférable. Il faut garder à l'esprit que le transistor détermine aussi les pertes d'insertion, avec le R_{on} présenté à l'état passant du transistor ; un compromis est donc à faire pour trouver la meilleure capacité en fonction des pertes d'insertion visées. La topologie en Figure 1.9 illustre un fonctionnement idéal, cependant quand le transistor série est passant, le filtre présente une impédance trop faible et dégrade les pertes d'insertion du système.

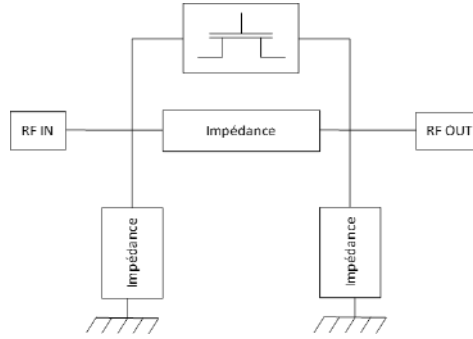
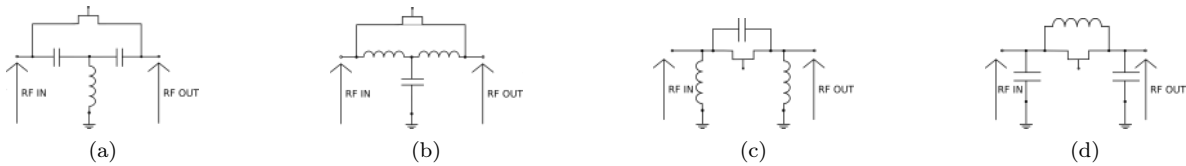


FIGURE 1.9 – Topologie de déphaseur à base de filtre commuté

Pour s'affranchir de cet effet il est nécessaire de fermer les branches shunt du filtre sur des circuits ouverts et non sur des masses. Néanmoins les masses sont nécessaires à l'état actif de la cellule, c'est pourquoi un second transistor est obligatoire. Ce dernier connecte les deux branches shunt du filtre, et assure une mise à la masse pour l'état actif de la cellule et une mise en haute impédance pour l'état bloqué de la cellule. Pour réaliser la mise à haute impédance on crée un « circuit bouchon » en mettant une inductance en parallèle au transistor. Ces masses et circuits ouverts n'étant pas parfaits, les performances du circuit sont dégradées.

Pour étudier ce type de topologies, il faut d'abord s'intéresser aux propriétés en phase des différents filtres. Les deux types de filtres les plus utilisés sont les topologies en T et les topologies en π . Ces filtres sont composés d'éléments réactifs (capacité et inductance) responsables du décalage en phase du signal qui les traverse. Selon la nature de l'élément placé sur la voie directe (inductif ou capacitif), le filtre sera qualifié de passe-bas (inductance) ou passe-haut (capacité). Les 4 différentes possibilités de filtres sont répertoriées en Figure 1.10a à 1.10d.

FIGURE 1.10 – Topologies à base de filtres commutés (a) en T passe-haut (b) en T passe-bas (c) en π passe-haut (d) en π passe-bas

Les formules permettant de trouver les valeurs des éléments du filtre sont données en Equation 1.4 et 1.5 pour les topologies en T et en Equation 1.6 et 1.7 pour les topologies en π .

$$Passe - haut : L = \frac{Z_0}{\omega \sin(\phi)} \quad ; \quad C = \frac{\sin(\phi)}{\omega Z_0 (1 - \cos(\phi))} \quad (1.4)$$

$$Passe - bas : L = Z_0 \frac{(1 - \cos(\phi))}{\omega \sin(\phi)} \quad ; \quad C = \frac{\sin(\phi)}{\omega Z_0} \quad (1.5)$$

$$Passe - haut : L = \frac{Z_0 \sin(\phi)}{\omega (1 - \cos(\phi))} \quad ; \quad C = \frac{1}{\omega Z_0 \sin(\phi)} \quad (1.6)$$

$$Passe - bas : L = Z_0 \frac{\sin(\phi)}{\omega} \quad ; \quad C = \frac{1 - \cos(\phi)}{\omega Z_0 \sin(\phi)} \quad (1.7)$$

Avec, ϕ le déphasage souhaité, ω la pulsation et Z_0 l'impédance sur laquelle le circuit sera adapté.

Une fois le déphasage ciblé obtenu, un élément de commutation est nécessaire pour choisir ou non de déphaser le signal. Dans ce type de topologie, le transistor en commutation est mis en parallèle au filtre, sa réactance doit donc être prise en compte lors de la simulation du filtre. Comme expliqué précédemment, la capacité C_{off} du transistor peut être utilisée pour un dispositif passe-haut. Cependant dans le cas d'un dispositif passe-bas, il faut veiller à ce que l'inductance série ne fasse pas résonner le transistor. De plus selon la fréquence de fonctionnement étudiée, certaines valeurs de composants seront impossibles à atteindre (spatialement ou juste trop importante). La topologie choisie devra donc prendre en considération les valeurs accessibles à la technologie GaN.

Il a été expliqué plus tôt que lorsque la cellule ne déphase pas, les branches parallèles des différents filtres doivent avoir l'extrémité non connectée au transistor série mise à la masse. Dans le cas contraire l'impédance présentée par le filtre viendrait perturber l'état de référence de la cellule. Un circuit bouchon est généralement utilisé pour réaliser cette fonction. La Figure 1.11 résume le fonctionnement de la cellule avec l'ajout du circuit bouchon.

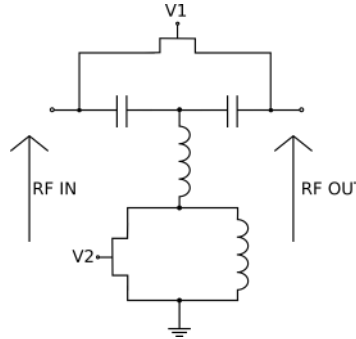


FIGURE 1.11 – Topologie en T passe-haut avec circuit bouchon

Les deux tensions de polarisation de grille V_1 et V_2 sont complémentaires et vont tour à tour rendre passant et bloqué les transistors concernés. Lorsque V_1 est égale à 0V le transistor en série est passant, le filtre en T est donc court-circuité par la résistance R_{on} équivalente dudit transistor. La tension V_2 est quant à elle inférieure à la tension de pincement (V_{off}) du transistor shunt, ce qui le met dans l'état bloqué. La mise en parallèle de la capacité C_{off} du transistor shunt et de l'inductance entraîne une résonance à la fréquence souhaitée. La haute impédance ainsi présentée permet de réduire l'influence du filtre sur l'impédance vue par le signal d'entrée. La cellule ne déphase pas et est dite à l'état de référence.

Lorsque V_1 est égale à V_{off} le transistor série est bloqué et la capacité C_{off} équivalente vient s'ajouter en parallèle à la branche série du filtre. V_2 est égale à 0V, la branche shunt du filtre est donc mise à la masse par la résistance R_{on} équivalente du transistor. La cellule est activée et se place dans la configuration de déphasage. Les Figures. 1.12a et 1.12b illustrent les deux états de la cellule par les schémas électriques équivalents. Le modèle en T choisi n'est qu'un exemple, le fonctionnement selon deux états du circuit est le même quelles que soient la configuration et la topologie choisies.

Théoriquement le circuit bouchon ne doit pas avoir d'influence sur l'état de phase de la cellule, en réalité il n'est pas possible d'obtenir un circuit ouvert et un court-circuit parfait. Pour cela, il est nécessaire de trouver un circuit bouchon permettant simultanément un respect de la consigne de phase et une bonne résonance, responsable d'une bonne adaptation et de pertes d'insertion réduites. Le choix de la tension V_{off} va fixer l'ondulation RF acceptable sans changement d'état du transistor. Elle est donc directement liée à la compression du circuit.

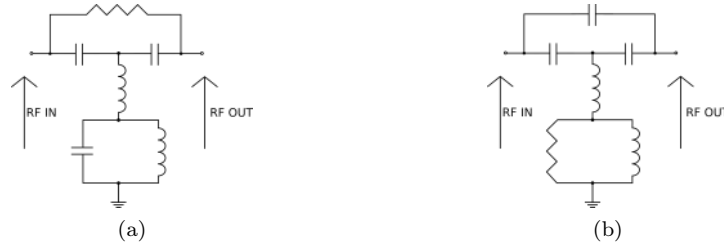


FIGURE 1.12 – Topologie en T passe-haut (a) à l'état de référence, le circuit bouchon résonne (b) le circuit bouchon se comporte comme un fil vers la masse

Cette topologie est la plus utilisée en bande Ka du fait de la platitude de phase sur la bande de fréquence et de la compacité des éléments lors d'intégrations monolithiques. Nous allons maintenant présenter quelques réalisations de déphaseurs utilisant ce type de topologies. La majorité des réalisations sont conçues en technologie SiGe et GaAs [13]-[14] et sont intégrées à des puces d'émission/réception (Tx/Rx), il est rare de trouver des déphaseurs seuls. La Figure 1.13a illustre un exemple de cellule réalisée en GaN. On peut voir qu'une grande partie de la surface du circuit est occupée par les différentes inductances (d'autant plus grandes en bande X).

Pour obtenir des déphasages supérieurs à 90° , il existe la possibilité de mettre en cascade deux déphaseurs dont la somme des variations en phase dépasse les 90° ou alors d'utiliser les propriétés en avance et retard de phase des filtres comme évoqué précédemment. En effet, en utilisant le retard en phase des filtres passe-bas et l'avance en phase des filtres passe-haut il est possible, en les disposant en parallèle, d'obtenir des déphasages supérieurs à 90° et ce avec une seule cellule. En commutant entre les deux voies, le déphasage obtenu sera la somme des valeurs absolues des déphasages respectifs de chaque filtre. L'élément assurant la commutation ne doit pas influencer sur le déphasage de chaque filtre, et doit garantir une isolation suffisante pour que les deux états de phase possibles du circuit soient indépendants. La topologie en π est schématisée en Figure 1.14; il y a été choisi deux topologies en T, ce qui n'est bien entendu pas imposé. De même pour la commutation, un simple transistor série est la solution la plus simple bien que des topologies de type Single Pole Double Throw (SPDT) bien plus élaborées puissent être préférées pour privilégier la linéarité, l'isolation ou d'autres paramètres. L'étude des SPDT est détaillée dans un chapitre entier de [16], nous ne nous intéresserons pas plus aux SPDT qui sont des commutateurs dont l'utilisation est plus orientée vers des applications de type puissance.

La seule obligation régissant cette topologie est d'avoir un déphasage positif sur un bras et négatif sur l'autre. L'utilisation de telles structures peut être envisagée pour notre application, cependant du fait de la mise en série de deux transistors (au moins) sur la voie passante les pertes d'insertion sont



FIGURE 1.13 – Exemple de cellule de déphasage (22.5°) en GaN (a) circuit réalisé en bande X de dimension $0,448 \times 0,678 \text{ mm}^2$ (b) modélisation électrique de la topologie en T passe-bas [15].

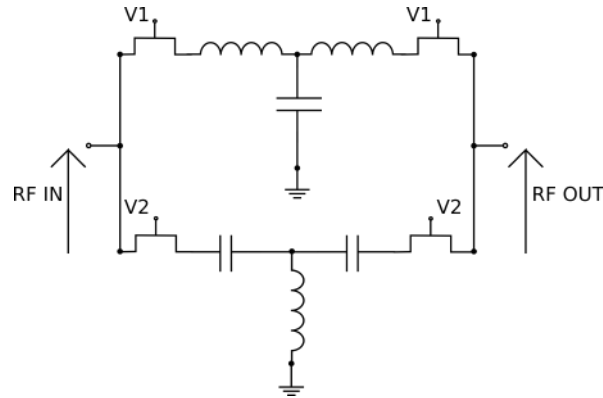


FIGURE 1.14 – Topologie filtre passe-haut/passe-bas en T

grandement augmentées. Néanmoins cette structure se base sur le déphasage relatif entre deux filtres, ainsi le déphasage du transistor (identique sur les deux voies) n'entre pas en compte dans le déphasage total de la cellule, en dehors d'éventuelles dispersions technologiques qui peuvent généralement être simulées si les modèles statistiques associés au kit de conception sont disponibles. La platitude de phase est donc améliorée. Un compromis apparaît donc entre platitude de phase et pertes d'insertion du circuit. Les équations utilisées pour trouver la valeur des éléments des filtres sont les mêmes que celles régissant les topologies passe-haut et passe-bas simples étudiées plus haut (équations 1.4 à 1.7 et Figure 1.10a à 1.10d).

Des variantes à cette topologie peuvent être utilisées comme l'illustrent par exemple les travaux de Zhou [13] et Zheng [17]. Dans ce type de configuration, le délai ou l'avance en phase est obtenue d'un côté par un filtre et sur l'autre branche par un élément de type ligne ou coupleur par exemple. La difficulté consiste à obtenir un niveau d'adaptation et de pertes d'insertion égal dans les deux états de fonctionnement du déphaseur.

Topologie "réflective"

La dernière topologie dite réflective, utilise les propriétés en phase des coupleurs comme l'illustre la Figure 1.15. A la sortie des ports transmis et couplé sont placées deux charges identiques calculées de telle façon à ce qu'elles présentent un coefficient de réflexion bien connu vis-à-vis de l'impédance du coupleur. Ces deux voies étant chacune déphasée de 90° par rapport à l'autre, nous pouvons exprimer la valeur du déphasage du signal de sortie par le port isolé (vers où les ondes couplées et transmises seront réfléchies). Ces déphaseurs permettent aussi d'atteindre des changements de phase allant jusqu'à 180° . Cependant, ils ont une bande de fonctionnement assez étroite, et occupent un espace plus large. De ce fait, les charges présentées à leurs sorties peuvent être variables et contrôlées en tension, de façon à pouvoir obtenir plusieurs déphasages différents avec la même structure. Dans ce cas les structures exploitent un fonctionnement analogique.

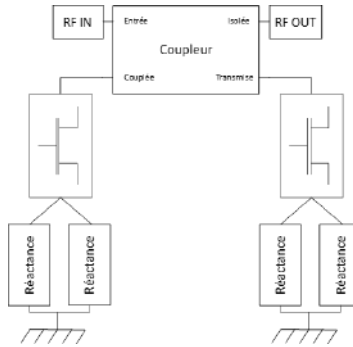


FIGURE 1.15 – Topologie de déphaseur dit "réflectif"

Bien que nous ayons vu précédemment que les topologies réflectives de déphaseurs (RTPS) sont plutôt utilisées dans des configurations de circuits analogiques, il est intéressant d'étudier certaines structures, et de les implémenter pour un fonctionnement assimilable à une application à commande numérique. Un élément reste commun à toutes ces structures, le coupleur. En effet les améliorations sur les topologies RTPS viennent pour l'essentiel d'une complexification des charges de sortie pour obtenir les performances souhaitées (déphasage, faible perte d'insertion, compacité...). De nombreuses publications [18]-[19] s'intéressent à l'optimisation des charges de sortie et des propriétés du coupleur pour obtenir des déphasages très importants (jusqu'à 360°), tout en gardant des pertes d'insertion inférieures à celles des topologies de types filtres commutés. En effet, commuter entre deux chemins différents, et devoir cascader plusieurs cellules pour augmenter la résolution du déphaseur dégrade fortement les pertes de celui-ci. Bien que les solution RTPS ne soient pas toujours aptes à des intégrations monolithiques, elles présentent des performances à l'état de l'art.

Autres topologies de déphaseurs

Il existe des types de déphaseurs plus exotiques, à base de microsystèmes électromécaniques (MEMS) [20]-[21], ou bien de ferrites [22]. Nous avons placé ces références à titre purement informatif et elles ne seront pas considérées dans l'état de l'art en raison de leurs performances insuffisantes, à l'heure actuelle, en bande Ka ; et également parce qu'elles ne sont pas implémentées dans le procédé technologique utilisé.

Comparatifs des performances de déphaseurs en bande Ka

Maintenant que nous avons présenté les différents types de topologies utilisables pour la conception des déphaseurs, nous pouvons comparer les performances de celles-ci selon la technologie choisie. Le Tableau. 1.2 présente les performances des déphaseurs micro-ondes à l'état de l'art. Dans ce tableau, nous avons jugé judicieux de ne pas inclure les performances des déphaseurs contenus dans les *core-chips*. En effet dans plusieurs travaux, les déphaseurs ne sont jamais implémentés individuellement mais ils sont plutôt intégrés au sein de *core-chips*. Il ne nous a pas semblé juste de les comparer puisque tous les aspects notamment de type consommation et, compacité sont faussés par l'intégration du déphaseur dans un système complet. Ces déphaseurs seront traités dans le comparatif sur les *core-chips*. Nous pouvons clairement voir que le CMOS présente de meilleures performances que les circuits III-V présentés, que ce soit sur les erreurs fonctionnelles ou sur la compacité. Les seuls points négatifs sont la bande de fréquence de fonctionnement (seulement 28GHz) et le point de compression qui n'est pas indiqué mais qui est certainement plus bas que celui des déphaseurs III-V.

Référence	[18]	[23]	[10]	[17]	[24]
Technologie	65nm CMOS	GaAs pHEMT	GaAs HJFET	InGaAs pHEMT	InGaAs PIN Diode
Topologie	RTPS	Filtres	Lignes à retard	Filtres et RTPS	Filtres
Fréquence (GHz)	28	32-37	33-36	31-40	26-30
Erreur RMS Phase	0.3°	3.5°	7°	4.7°	5.6°
Erreur RMS Amplitude	0.3 dB (non RMS)	0.4dB	0.9dB (@34.5 GHz)	0.6dB	0.7dB
Pertes d'insertion (dB)	8.05	7	15	8.8	7.8
Adaptation (dB)	6.7	7	10dB (@34.5 GHz)	9	9
Résolution Phase	11.25°	11.25°	22.5°	11.25°	11.25°
Compacité	0.16mm ²	1.41mm ²	5.5mm ²	3.315mm ²	2.3436mm ²
Consommation (W)	0	0	0	0	0°
P_{1db} (dBm)	Non spécifié	20	Non spécifié	16 (sortie)	21 (entrée)

TABLE 1.2 – Etat de l'art des déphaseurs en bande Ka

Maintenant que nous avons présenté les différentes topologies de déphaseurs nous pouvons nous intéresser aux seconds constituants des *core-chips* les atténuateurs.

5 Les atténuateurs

Selon l'application visée, les atténuateurs peuvent tenir plusieurs rôles :

- les systèmes de leurre : le principe est de faire croire aux récepteurs ennemis qu'un élément se trouve à une position, alors qu'il se trouve à une position décalée. Atténuer volontairement un signal peut permettre de réaliser cette fonction en remplaçant la signature électromagnétique réelle d'un objet (en amplitude ici, donc en faussant la distance perçue) par la signature factice générée.

- Le pointage électronique d'antenne : même si dans cette configuration l'objectif est d'obtenir un bilan de liaison au gain maximum, atténuer certains signaux alimentant un réseau d'antennes permet d'augmenter la directivité du faisceau en réduisant la puissance transmise aux lobes secondaires [7].

Dans cette partie nous étudierons les différents types d'atténuateurs micro-ondes existant :

Pour atténuer le signal, l'utilisation d'éléments résistifs (dissipatifs) est incontournable. Il existe plusieurs agencements de résistances qui répondent à la fonction recherchée. La topologie la plus simple est de placer une résistance en shunt sur la voie RF comme illustré en Figure 1.16. Cette résistance ne doit pas être du même ordre de grandeur que la résistance série de la ligne. En connectant ou déconnectant cette résistance, nous obtenons un atténuateur numérique. Plus la résistance shunt est grande plus l'atténuation est faible. Cependant, cette topologie ne permet pas d'obtenir des atténuations élevées puisque très vite les pertes d'insertion augmentent. Le gros avantage de cette technique, au-delà de sa

simplicité de mise en œuvre, est la très faible phase d'insertion générée. Dans les circuits de contrôle, les erreurs croisées (erreur de phase de l'atténuateur et erreur d'atténuation du déphaseur) sont très importantes. En effet, celles-ci rentrent en compte dans le calcul des erreurs RMS de la structure finale. Pour les atténuateurs, la phase d'insertion générée provient de la différence entre les deux chemins empruntés par le signal. L'utilisation de transistor en commutation introduit immanquablement des éléments réactifs qui perturbent la phase. Cet effet est d'autant plus important si le nombre de transistors sur les voies est différent, ou bien si ceux-ci sont de tailles différentes. La première topologie proposée en Figure 1.16 engendre peu de déphasage puisqu'il n'y a pas de transistor sur la voie directe et que le transistor shunt est vu comme un circuit quasi ouvert quand il est OFF. Cependant, en distribuant les résistances le long de la ligne directe (plusieurs transistors shunt en parallèle pour réduire la résistance équivalente shunt présentée), la gestion de cette phase devient plus compliquée ; c'est aussi pour cette raison que les niveaux d'atténuation importants sont inaccessibles avec cette topologie.

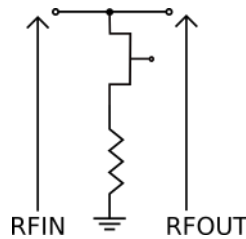


FIGURE 1.16 – Mise en série d'une résistance pour atténuer le signal incident

Lorsqu'une dynamique d'atténuation plus importante est nécessaire, deux topologies émergent : les réseaux résistifs en T et en π . Elles sont assez similaires et il est assez aisé de passer mathématiquement de l'une à l'autre, elles sont représentées en Figure 1.17. Ces topologies sont les plus utilisées car elles sont simples à mettre en œuvre et couvrent une très large gamme d'atténuation. Pour des applications numériques, la topologie en π est peut-être avantageuse, puisqu'une partie de la résistance série du réseau peut être obtenue grâce à la résistance R_{on} du transistor série. Comme explicité plus tôt, commuter entre une ligne simple et une cellule d'atténuation va engendrer une différence de phase ; c'est pour cela qu'intégrer la voie passante directement à l'atténuateur est une bonne alternative non seulement pour la réduction de la phase d'insertion mais aussi pour l'amélioration de la compacité du circuit. Une autre considération lors du choix entre π et T concerne les valeurs des résistances mises en jeu. En effet, pour des résistances à valeurs trop élevées, celles-ci peuvent présenter des comportements distribués, ce qui influe donc sur la phase des signaux.

FIGURE 1.17 – Topologie proposant des atténuation plus importantes : (a) réseau en π (b) réseau en T

Les formules théoriques régissant ces atténuateurs sont données en équations 1.8 et 1.9. Avec ce type de topologie, pour des niveaux d'atténuation importants, il est nécessaire de compenser la phase qui est différente entre les deux états.

$$R_{\text{seau en } \pi} : R_1 = \frac{Z_0}{2} \cdot \frac{10^{A/10} - 1}{10^{A/20}}; \quad R_2 = Z_0 \cdot \frac{10^{A/20} + 1}{10^{A/20} - 1} \quad (1.8)$$

$$R_{\text{seau en } T} : R_1 = Z_0 \cdot \frac{10^{A/20} - 1}{10^{A/20} + 1}; \quad R_2 = 2 \cdot Z_0 \cdot \frac{10^{A/20}}{10^{A/10} - 1} \quad (1.9)$$

Avec A l'atténuation souhaitée en dB ($A=10 \cdot \log(P_{in}/P_{out})$), et Z_0 l'impédance caractéristique sur laquelle adapter le réseau.

Obtenir une atténuation constante sur la bande de fréquence n'est pas la contrainte la plus critique. En effet, une résistance même non idéale présente peu de parasites réactifs qui pourraient faire varier son comportement en fréquence. La taille de cette dernière en revanche ne doit pas être trop importante pour éviter les problèmes d'effets distribués. La contrainte principale lors de la conception d'atténuateur est la phase d'insertion entre les deux états. En effet lors du calcul des valeurs de résistances, aucune phase n'est prise en compte. Néanmoins, du fait de la taille du circuit (différentes longueurs électriques parcourues par le signal selon l'état de la cellule), de l'ajout de transistors pour assurer la commutation, ou de lignes pour relier les éléments du circuit entre eux, des phases parasites viennent s'ajouter. Cette considération est d'autant plus sensible à hautes fréquences (bande Ka dans notre étude). Pour limiter ces effets, diverses solutions sont envisageables. La plus simple est de rajouter de la longueur électrique sur le chemin de référence. Cette solution fonctionne au prix de pertes d'insertion plus importantes. D'autres solutions ont été adoptées pour réduire cette phase d'insertion parasite [25]-[26]. Par exemple ajouter un élément réactif (plutôt capacitif dans la branche shunt du réseau), ou bien allonger les lignes d'accès (ajout d'une inductance) à la résistance en haut des circuits de la Figure 1.17 pour compenser l'écart de phase [27]. Cependant cette dernière solution augmente grandement les pertes d'insertion du circuit. Les performances d'un atténuateur sont évaluées sur les mêmes critères que ceux pour juger des performances des déphaseurs. Le poids de chaque performance étant cependant différent selon l'application voulue.

D'autres topologies d'atténuateurs, telle que le *Bridge T-Attenuator* peuvent être utilisées, mais sont plus exotiques et parfois difficilement conciliables avec les paramètres de notre étude. Pour le *Bridge T-Attenuator* par exemple, la nuance avec un simple *T-Attenuator* réside l'ajout d'une résistance parallèle sur la voie série. Cette légère nuance n'apporte pas grand-chose au circuit étant donné que nous plaçons déjà un transistor parallèle pour réaliser la commutation, ce qui revient quasiment à une topologie *Bridge-T*. Ces topologies sont présentées dans [28]. Le nombre de topologies d'atténuateurs étant plus restreint que ceux des déphaseurs, nous présentons directement dans le tableau 3 les atténuateurs à l'état de l'art. Comme pour les déphaseurs, ne figurent dans ce tableau que les atténuateurs « seuls » et pas ceux intégrés à un *core-chip*.

Une fois les deux éléments constitutifs des *core-chips* identifiés et présentés, nous pouvons nous intéresser à la mise en commun de ces éléments pour l'obtention de la fonction finalisée atténuation/déphasage numérique. Pour cela il faut regrouper toutes les cellules individuelles, avec les cellules de poids fort (dynamique) et de poids faible (sélectivité) de chaque fonction, et comme il est explicité dans la section suivante cela nécessite une méthodologie particulière.

6 Les core-chips

Dans cette section nous n'étudierons pas la mise en commun des cellules; en effet cet aspect nécessite une analyse plus complète qui sera détaillé dans le Chapitre 3 et sera suivie de la méthodologie d'intégration que nous avons développée. Dans cette partie nous étudierons les *core-chips* en tant que système, de leur fonctionnement général vers les réalisations existantes à l'état de l'art.

Référence	[29]	[23]	[30]	[26]
Technologie	GaAs pHEMT	GaAs pHEMT	GaAs pHEMT	SiGe BiCMOS
Fréquence (GHz)	DC-35	0.1-31	20-40	19-24
Erreur RMS Amplitude	0.2dB	0.9dB	1.2dB	0.5dB
Erreur RMS Phase	12°	NC	7.2°	4.1°
Pertes d'insertion (dB)	4.5	4.2	8.8	NC
Adaptation (dB)	12	9	15	19
Résolution Amplitude	1dB	0.5dB	0.5dB	0.5dB
Compacité	4.1mm ²	0.94mm ²	1.47mm ²	0.51mm ²
Consommation (W)	0	0	0	0
P_{1db} (dBm)	20 (entrée)	NC	24 (entrée)	NC

TABLE 1.3 – Etat de l'art des atténuateurs

Pour décrire grossièrement un *core-chip*, il est possible de le considérer comme un système avec une entrée et une sortie. Sur l'entrée sera envoyé un signal, et selon la commande appliquée par un mot binaire, ce signal sera transmis en sortie affecté d'une combinaison phase/amplitude. La précision atteinte sur la commande de phase et d'amplitude dépend de la résolution du *core-chip*. Cette dernière est associée à la taille du mot binaire d'entrée (et plus précisément aux bits de poids faibles). La Figure 1.18 illustre le fonctionnement du système. Pour évaluer les performances d'un *core-chip*, il suffit de regarder les performances des déphaseurs et atténuateurs ensembles. Ainsi le *core-chip* présentera une erreur globale de phase et une erreur globale d'amplitude. Attention cependant à bien préciser quelles erreurs sont prises en compte, car dans certaines documentations techniques l'erreur de phase du *core-chip* est évaluée uniquement sur les états du déphaseur et inversement l'erreur d'amplitude est calculée que sur les états d'atténuation.

Concernant la résolution, la dynamique globale du *core-chip* est morcelée en niveau de phase et d'atténuation respectant le plus souvent une loi binaire. Par exemple, un *core-chip* ayant une résolution en phase de 6 bits et en amplitude de 4 bits présentera 6 cellules de déphasage et 4 cellules d'atténuation. La couverture maximale pour ce type de système étant de 360°, la cellule au déphasage le plus important sera de 180° puis dans un ordre décroissant les cellule suivantes auront un déphasage divisé par 2 (90°, 45°, ..., 180/2ⁿ - 1 °). De même pour l'atténuation, excepté que le choix se fait plutôt selon la plus petite atténuation ('pas' d'atténuation, ou résolution). Ainsi la dynamique d'atténuation sera de $A_1 + A_2 + \dots + A_n$ (A_i en dB), avec $A_n = A_1 \cdot 2^{n-1}$ (avec $n > 1$) et A_1 l'atténuation minimale choisie. Ainsi pour le *core-chip* cité plus haut, 6 bits de phase donnent une résolution de 5.625° et une dynamique de 360°, tandis que 4 bits d'atténuation peuvent conduire à une dynamique de 4 dB en choisissant une résolution de 0.5dB, contre 16 dB avec 6 bits en conservant la même résolution.

Instinctivement, l'idée d'augmenter le nombre de cellules pour avoir une résolution maximale vient à l'esprit. Cependant, augmenter la résolution s'accompagne nécessairement d'une augmentation des pertes d'insertion du fait de la mise en commun d'un nombre plus important de cellules. Pour garantir une résolution significative, il faut que l'erreur soit inférieure à la moitié de la résolution. Ce critère devient donc très contraignant, spécialement en phase où par exemple, pour des résolutions de 5.625°, l'erreur ne doit pas dépasser 2.8°. La Figure 1.19 illustre cette contrainte de manière explicite avec des gabarits de résolution de phase et d'amplitude. Lors de l'examen d'un document technique ou lors de la conception d'un *core-chip*, il faut donc veiller à ce que l'erreur respecte bien cette contrainte de façon à pouvoir exploiter pleinement le nombre de cellules du *core-chip* réalisé, et optimiser ainsi les pertes globales du module conçu.

Les *core-chips* réalisent des fonctions de contrôle, fonctions intégrées dans un système plus global

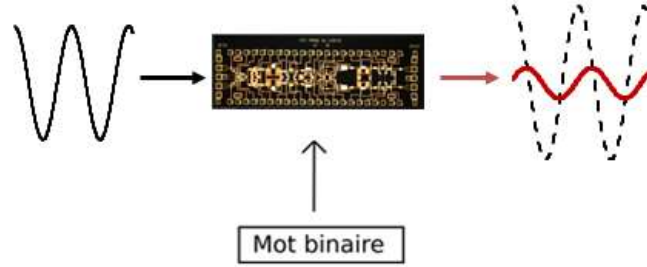


FIGURE 1.18 – Schéma de principe d'un core-chip n bits

d'émission réception in fine. Ainsi, dans la littérature, les performances des *core-chips* sont évaluées à travers la puce globale Tx/Rx dans laquelle ils sont intégrés. Malgré tout, les performances individuelles du *core-chip* sont importantes puisque, les autres éléments de la chaîne (amplificateur faible bruit, amplificateur forte puissance...) n'impactent pas les performances en phase et en amplitude.

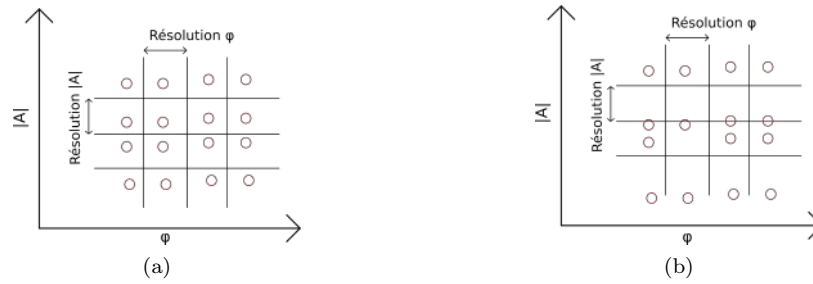


FIGURE 1.19 – Exemple de constellations d'états de phase/atténuation (cercle rouge) où (a) la résolution théorique visée (quadrillage) est cohérente avec la résolution finale synthétisée (O) (b) la résolution théorique est inférieure à la résolution finale synthétisée

En effet, lors de l'analyse d'une puce Tx/Rx, l'erreur RMS de phase et d'amplitude correspondent à celles du *core-chip*. Il est donc justifié de réaliser une revue bibliographique sur les puces Tx/Rx entières. Nous avons regroupé dans le Tableau. 1.4 les performances des *core-chips* à l'état de l'art. Il est intéressant de noter l'absence de la technologie GaN dans le Tableau. 1.4, et dans la littérature. Il n'existe pas à notre connaissance de *core-chip* réalisé en GaN. Il est assez facilement compréhensible qu'il ne puisse pas concurrencer le SiGe en terme d'intégration et de coût. Cependant, vis à vis du GaAs, rien ne l'empêche de se positionner comme une alternative concurrente puisqu'au-delà des fonctions de contrôle, le GaN présente de meilleures performances en puissance que son vis-à-vis le GaAs. Dans la partie suivante nous nous intéresserons plus particulièrement à l'utilisation du GaN pour réaliser des fonctions de contrôle en vue d'une intégration totale des éléments d'une puce Tx/Rx.

7 La place du GaN dans les puces multifonctions

En analysant le Tableau. 1.4, on s'aperçoit de l'absence du GaN pour des réalisations de *core-chip* à l'état de l'art. Comme explicité plus tôt, le GaN génère un réel engouement dans le domaine de micro-ondes, pour des applications actives de puissances (*front-end*). En effet, comme nous l'avons évoqué précédemment, des composants « discrets » tels que des déphaseurs ou des atténuateurs [2]

Référence	[14]	[13]	[31]	[32]
Technologie	0.13 μ m SiGe BiCMOS	GaAs pHEMT	GaAs pHEMT	0.13 μ m SiGe BiCMOS
Fréquence (GHz)	30-40	32-38	32-34	34-39
Erreur RMS Phase	4°	3.8°	3°	12°
Erreur RMS Amplitude	0.6dB	0.7dB	0.5dB	0.9dB
Gain (dB)	Rx :17 Tx :14	Rx :9 Tx :11.5	Rx :30 Tx :30	Rx :-1 Tx :2
Adaptation (dB)	9	9	NC	8
Résolution Phase	11.25°	11.25°	NC	11.25°
Résolution Amplitude	Pas d'atténuateur	0.5dB	NC	Pas d'atténuateur
Compacité	7mm ²	12mm ²	NC	4mm ²
Consommation (W)	Rx : 0.528 Tx : 1.587	0.15	Rx : 0.6 Tx : 2.5	Rx : 0.142 Tx :0.171
P_{1dB} (dBm)	Rx :-1 Tx :20.5 (sortie)	Rx :-3 Tx :0 (sortie)	NC	Rx :-18 Tx :4.7 (sortie)

TABLE 1.4 – Etat de l'art des core-chips en bande Ka

ont été réalisés individuellement. Cependant, il n'existe pas à notre connaissance de *core-chip* en GaN regroupant ces fonctions. Ce travail de thèse se positionne donc, à notre connaissance, comme un travail pionnier pour la réalisation de puce multi-fonctions en GaN.

Le GaN n'apparaît pas comme un bon candidat à la conception de *core-chips*. Certes en termes de puissance la comparaison plaide en sa faveur, néanmoins vis à vis de la capacité d'intégration il ne peut en aucun cas concurrencer le SiGe. D'une part le nombre de niveaux de métallisations est bien moindre, ce qui oblige à exploiter des surfaces plus conséquentes. De plus, la déclinaison horizontale des transistors à effet de champ (FET) induit des effets de déphasage non sensibles pour les technologies verticales HBT Silicium-Germanium. L'intérêt de l'utilisation du GaN ne peut pas être mis en avant en ce qui concerne les circuits de contrôle (déphasage et atténuation), mais il faut élargir à l'échelle système pour voir l'avantage d'une intégration globale comprenant les parties contrôle (*back end*) et puissance (*front-end*). Cette intégration monolithique permettrait de s'affranchir des transitions alumines nécessaires à l'intégration hybride (SiGe-GaN ou SiGe-GaAs) ce qui diminuerait la contrainte de gain sur la partie puissance. En effet dans les composants actuels d'émission réception, sont associées une partie *back-end* en SiGe et une partie *front-end* en GaN. De ce fait des incompatibilités de niveaux de puissance à l'interface de ces deux technologies apparaissent. En effet, d'une part la puissance en sortie des *core-chips* SiGe est insuffisante pour piloter les amplificateurs de puissance GaN dans leur zone de fonctionnement optimale. D'autre part, la puissance de saturation des *LNA* GaN (typiquement 15dBm) est trop importante pour être directement injectée en entrée des *core-chips* SiGe. Il est donc nécessaire d'ajouter respectivement des pré-amplificateurs et limiteurs pour s'assurer du bons interfacage des éléments. Une intégration monolithique en GaN permettrait de s'affranchir de ces modules supplémentaires coûteux en compacité. Il reste maintenant la question de la commande numérique des *core-chips*. En ce sens, des travaux tendent à développer des technologies GaN sur silicium compatibles CMOS, ce qui répondrait à cette contrainte d'intégration des circuits de contrôle entre autres. Mais d'ores et déjà, la réalisation de puces monolithiques en GaN est rendue possible par l'existence de transistors à enrichissement (*normally OFF*) qui permettent de réaliser la partie contrôle numérique des tensions de polarisation du *core-chip*. En effet, la présence d'un module entrée série sortie parallèle (*Serial Input Parallel Output*) est nécessaire pour réaliser la polarisation des grilles des transistors. Un mot binaire série est envoyé sur le module SIPO, et celui-ci le convertit en des trames de tensions parallèle grâce aux transistors à enrichissement disponibles en développement dans la technologie GaN D01GH d'OMMIC. Toutes les conditions nécessaires sont donc réunies pour la réalisation de *core-chips*.

Relativement aux performances tirées de l'état de l'art du Tableau. 1.4, il a été décidé par OMMIC de viser un cahier des charges ambitieux précisé dans le Tableau. 1.5 ci-dessous.

Si nous remarquons que ce cahier des charges correspond quasiment à chacune des meilleures

Technologie	GaN 0,1 μ m
Bande de fréquence (GHz)	30-40
Erreur RMS de phase	3°
Erreur RMS d'amplitude	0,5dB
Gain (dB)	Passive
Niveau d'adaptation (dB)	>10
Résolution en phase	5.625°
Résolution en amplitude	0,5dB
Compacité	NC
Consommation DC (W)	Passive
P1dB (dB)	>10

TABLE 1.5 – Cahier des charges OMMIC du Core-chip MMIC GaN bande Ka

performances du Tableau. 1.4, indépendamment de la technologie SiGe ou GaAs employée, il semble ambitieux de viser de tels objectifs de performances, et ce d'autant plus que la bande passante 30-40 GHz est très large. La tenue des performances pour chaque fréquence de la bande représente un réel challenge qui mérite une analyse fine de la méthode de conception à employer.

Face aux spécificités et problèmes de conception des cellules spécifiques aux technologies planaires (transistors à effet de champ), et donc au GaN, cette méthodologie de conception des circuits de contrôle a été définie et utilisée tout au long de ce travail. La technologie GaN étant jeune, les modèles de simulations (électriques et électromagnétiques) sont en cours de développement, et de fait évolutifs. Pour garantir une cohérence entre les étapes de dessin, et pour obtenir des circuits aux performances définies lors de la simulation, la méthodologie décrite dans la section suivante est utilisée.

8 Méthodologie de conception

Les topologies utilisées pour réaliser des fonctions de contrôle sont connues et n'ont pas radicalement évolué si l'on se réfère aux informations recueillies dans la littérature. Cependant, la méthodologie de conception est propre à chaque concepteur. La méthode choisie a été de partir de topologies idéales pour converger progressivement vers des topologies réelles. Nous initions l'étude avec des éléments localisés pour aller vers la complexité plus réaliste des modèles électriques puis électromagnétiques. Dans notre cas, la filière D01GH d'OMMIC a été utilisée ; c'est une filière GaN sur Si utilisant des transistors de 100 nm de longueur de grille. Cette filière étant très récente, les modèles des transistors ont évolué durant la période des travaux.

Comme il a été vu précédemment, pour réaliser des fonctions de contrôle, il faut disposer d'éléments réactifs et résistifs pour obtenir respectivement des déphasages ou des atténuations. Comme justifié précédemment, nous avons choisi des topologies à base de filtres commutés, de lignes à retard, de réseaux résistifs en π et de lignes chargées. Nous n'utilisons pas de coupleurs, ce qui nous permet de nous affranchir des problématiques électromagnétiques un peu plus délicates, inhérentes à ces topologies.

La méthodologie de conception que nous avons utilisée se décline selon les étapes suivantes :

-Utilisation d'éléments idéaux pour obtenir la fonction voulue, seul l'état de déphasage/atténuation est réalisé. En effet, aucun transistor n'est utilisé durant cette phase. C'est une approche bande étroite puisque la valeur des éléments localisés est définie par leur fréquence d'utilisation. La valeur choisie sert de point d'ancrage aux itérations suivantes, il faudra toujours contrôler que les valeurs des composants (simulations électriques ou électromagnétiques) restent cohérentes avec cette valeur initiale. Ceci est d'autant plus vrai que les modèles électriques idéaux de ces éléments ne varie pas avec la fréquence.

Cette notion est très importante car il est très facile de se perdre dans les itérations successives, et de compenser des erreurs par des valeurs de composants très loin des valeurs initiales. Cependant, malgré une performances globale correspondant au cahier des charges initial, obtenir un résultat sur des compensations dues à des valeurs parasites ou mêmes des compensations entre composants est un pari risqué, et ce d'autant plus si le circuit comporte un nombre important de briques de composants. En cas de circuit mesuré ne répondant pas au gabarit simulé, il est alors quasi impossible d'identifier l'origine du problème à corriger puisque la solution synthétisée peut être sensiblement éloignée des valeurs initiales calculées. En étant conscient de ces phénomènes, cette étape n'est en général pas très longue puisque les calculs permettant d'obtenir les valeurs des composants sont faciles d'accès.

-Intégration dans la même simulation des deux états du circuit (ON-OFF). Cela peut être réalisé par l'utilisation de SPDT idéaux ou l'utilisation de transistors du kit de conception (*design kit*, DK). Même si l'utilisation de SPDT idéaux renseignent sur la performance en commutation de la cellule, les informations apportées ne sont pas pertinentes puisque l'état passant est en général modélisé par une ligne. Il n'est pas réaliste de modéliser un transistor en commutation uniquement par une ligne. L'intégration de transistor apparaît donc incontournable, mais demande cependant plus d'efforts puisque ce dernier se comporte différemment selon son état contrairement au SPDT idéal. De plus suivant la fonction de contrôle souhaitée, le transistor pourra servir à la fois de commutateur entre les deux états de la cellule mais pourra aussi être intégré à la fonction elle-même (atténuation ou déphasage). Par exemple la capacité OFF pour un filtre passe-haut ou la résistance ON pour la résistance série d'un circuit en π résistif. Une manière efficace de vérifier que cette étape de conception vérifie toujours les conditions énoncées dans l'étape précédente passe par l'utilisation de l'abaque de Smith. En visualisant les impédances relatives aux différents étages nous pouvons nous assurer que les valeurs sont cohérentes avec les impédances initiales choisies, et ce pour toute la gamme de fréquence.

-Une fois les transistors dimensionnés, il faut remplacer les éléments localisés par les éléments réels du DK. Les performances vont se dégrader par l'adjonction d'éléments représentatifs des états du transistor ; Cependant, en raison d'un coefficient de qualité bien moindre du transistor modélisé par rapport aux éléments réels, la bande de fréquence de fonctionnement s'en trouve grandement améliorée (conjointement avec une réduction des pertes d'insertion). Ces éléments réels nous laissent ainsi une marge de manœuvre puisqu'il est possible de modifier leurs dimensions (longueur, largeur), ce qui permet de modifier les éléments réactifs de chaque composant (inductances parasites d'un condensateur ou d'une résistance, ou capacité parasite d'une inductance). Cette marge de manœuvre reste quand même limitée, car relative à la taille du composant utilisé. A la fin de cette étape, la fonction est décrite avec les transistors et nous avons une première perception de la bande de fréquence sur laquelle cette fonction est assurée. Pour gagner du temps, nous pouvons remplacer les modèles réels électriques des composants par leurs modèles électromagnétiques.

-La mise en commun de tous les composants nécessaires à la fonction est pour l'instant idéale. Durant cette étape, nous allons relier les éléments par des lignes réelles du DK. C'est généralement l'étape où les performances se dégradent le plus. Il peut être parfois nécessaires de redimensionner certains composants pour prendre en compte les effets des lignes, et réussir une intégration qui respecte à la fois les performances et les règles de dessin, en visant systématiquement la meilleure compacité. Les masses idéales peuvent aussi être remplacées par des via-holes. L'accès à ces derniers étant réalisé par des lignes, les inductances en série avec celles des via-holes doivent être considérées.

-Une fois toutes ces étapes réalisées, il est possible d'obtenir le layout du circuit, de façon automatique (autolayout) ou manuelle. La première solution est plus rapide mais demande une grande rigueur pour obtenir un layout fidèle au circuit électrique. La solution manuelle offre plus de souplesse et permet de gérer efficacement le routage et le positionnement de tous les éléments ; mais cette technique prend en général plus de temps et exige également beaucoup de rigueur pour respecter les dimensions

du circuit établies précédemment. La simulation EM de la cellule finale peut être obtenue et comparée avec la simulation électrique.

Durant ces étapes, il est nécessaire de limiter la baisse de performances au fur et à mesure des itérations. Pour cela il faut définir une marge d'erreur acceptable sur les performances entre chaque itération, respecter cette marge étant une condition nécessaire mais non suffisante pour passer à l'itération suivante. En effet, conserver des valeurs de composants cohérentes avec les valeurs initiales est primordiale pour s'assurer de rester dans le périmètre de la théorie sur laquelle nous avons basé la conception de nos cellules.

La méthodologie décrite ci-dessus permet le dessin de cellules individuelles qui composeront le *core-chip* final. Durant cette description, nous n'avons pas abordé les choix influençant les performances finales du circuit. Identifier les paramètres d'une cellule qui influent sur le fonctionnement global est nécessaire. La première condition concerne la condition de fermeture en entrée/sortie ; l'adaptation insuffisante d'une cellule peut placer le reste du circuit dans des conditions électriques différentes de celles selon lesquelles il a été conçu. Le changement de déphasage d'une cellule selon l'évolution de sa charge de sortie est illustré en Figure 1.20. Une règle empirique suggère une adaptation minimale de -18 dB par cellule pour garantir un fonctionnement quasi-indépendant de chacune des cellules. Cependant selon le pourcentage de bande passante visé, il peut devenir très difficile de maintenir un tel niveau d'adaptation, tout en gardant des performances acceptables en termes de fonctionnalités. Dans un souci de clarté et de simplification, dans la suite des explications nous qualifierons les erreurs d'amplitude et de phase comme étant les erreurs fonctionnelles. Il est donc nécessaire de différencier ces dernières des erreurs d'adaptation pour pouvoir identifier leurs impacts respectifs au niveau système. En effet, une fois les cellules mises en cascade, les erreurs fonctionnelles s'additionnent de façon aléatoire et il est impossible de prédire la marge d'erreur fonctionnelle finale.

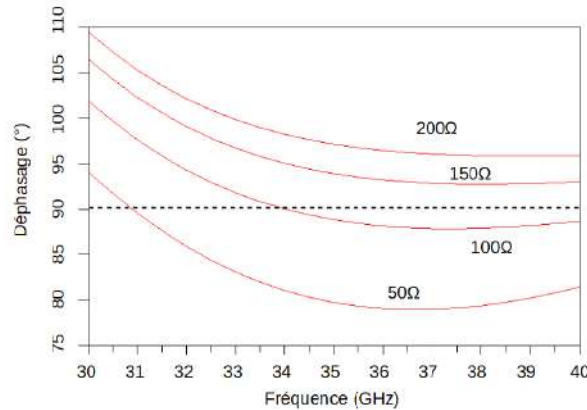


FIGURE 1.20 – Exemple de modification de performances en fonction de la variation de la charge de sortie pour la cellule 90° différentielle

Durant ces travaux plusieurs approches ont été utilisées :

-Approche large bande : Le cahier des charges définit une bande de fréquence allant de 30 GHz à 40 GHz. Cependant cette approche se révèle compliquée une fois la mise en commun des cellules effectuée, car étant donné que les erreurs fonctionnelles de chaque cellule vont se sommer de manière aléatoire, il sera compliqué d'expliquer le profil d'erreur du circuit complet. Cette difficulté vient de la largeur de bande, qui impose des performances et des profils d'erreurs fonctionnelles variables selon les différentes fréquences et qui ne permet aucune certitude sur l'origine des erreurs obtenues au

final. Au-delà des erreurs fonctionnelles, les différents niveaux d'adaptation sur la bande de fréquence, rendent encore plus compliqué l'interprétation des performances (impact avéré ou pas des réflexions entrée/sortie, et ce généralisé sur les n cellules du *core-chip*). Cette approche a été utilisée durant le design des *core-chips* différentiels pour répondre au premier cahier des charges. Malgré l'obtention de résultats probants, l'interprétation des performances reste compliquée.

-Approche mono-fréquence : La seconde façon d'opérer est de viser un fonctionnement à une fréquence centrale sur laquelle les erreurs seront minimisées ; de cette façon, nous pouvons nous affranchir des erreurs fonctionnelles à cette fréquence, et isoler ainsi plus aisément l'influence de l'adaptation. Le choix des profils d'erreurs fonctionnelles de chaque cellule conditionne la largeur de bande du système. Idéalement un profil d'erreur plat est souhaité, cependant les profils les plus courants sont des fonctions affines ou des paraboles autour de la fréquence centrale. Le circuit étant optimisé à la fréquence de 35GHz, les profils d'erreur les plus courants sont des paraboles. Le minimum de la parabole correspondant à l'erreur nulle souhaitée et l'ouverture définissant la largeur de bande du système (les erreurs allant croissantes plus la distance au minimum augmente). La Figure 1.21 illustre cette problématique de platitude d'erreur.

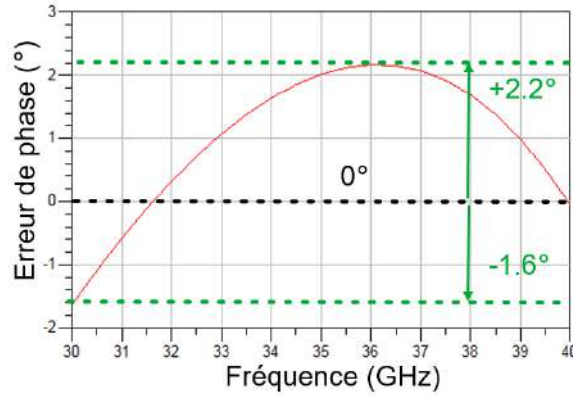


FIGURE 1.21 – Exemple de profil d'erreur pour une cellule de déphasage, en noir le profil d'erreur plat, en vert les erreurs maximales et en rouge l'erreur de phase

Avec la seconde approche il est possible d'optimiser efficacement la mise en commun des cellules, grâce à un programme de dénombrement dans l'environnement Matlab. Les concepts utilisés sont proches de ceux explicités dans [33], cependant l'algorithme que nous avons écrit ne nous permet de réaliser l'optimisation qu'à une seule fréquence. La méthode proposée vise à simuler l'ensemble des agencements possibles, et de rechercher la combinaison optimale de cellules (en considérant 2 états de commande pour chaque cellule). Le principe de cette méthode est schématisé en Figure 1.22.

Une fois la première étape de conception achevée (synthèse de chaque cellule individuelle), la matrice S est transformée en une matrice cascade ABCD. Le logiciel Matlab est ensuite utilisé pour dénombrer et calculer les caractéristiques de tous les agencements possibles de cellules de façon systématique. Pour chaque agencement, toutes les combinaisons sont également caractérisées selon les deux états possibles de chaque cellule (activation et désactivation). Trois indicateurs sont utilisés pour évaluer les performances obtenues ; l'erreur de phase, l'erreur d'amplitude, et le module de l'erreur vectorielle (EVM), ces marqueurs sont illustrés en Figure 1.22, et l'équation 1.10 indique l'expression mathématique liant les différents paramètres. Le critère EVM correspond à la distance qui sépare un état vectoriel réel donné relativement à l'état commandé. Une moyenne des valeurs pire-cas (erreurs maximales) et écart types sont extraites des calculs effectués sur l'ensemble des combinaisons d'agen-

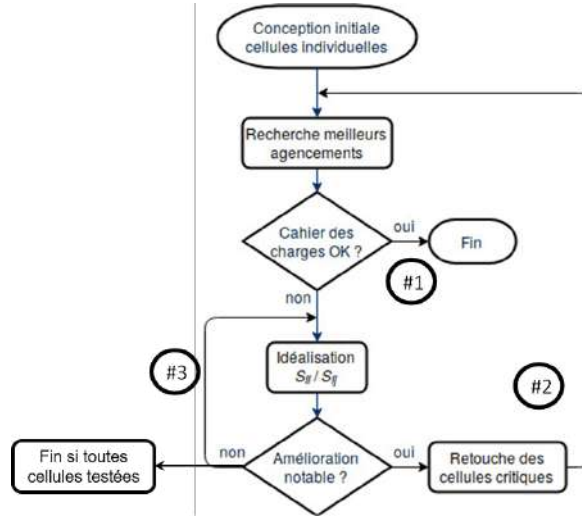
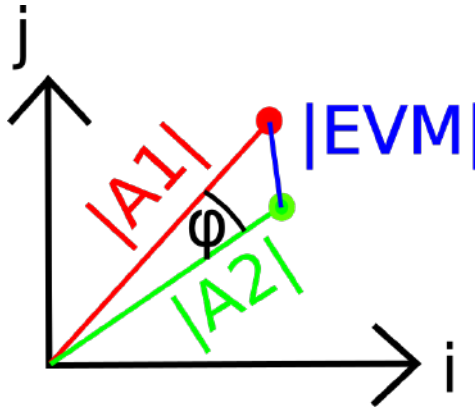


FIGURE 1.22 – Description de la méthodologie d’optimisation proposée

cement des n cellules constituant le *core-chip*. La Figure 1.23 permet d’apprécier la variation de ces trois critères selon l’ensemble des solutions simulées dans le cas d’un déphaseur 6 bits (représentées par valeurs décroissantes de l’erreur moyenne). Les courbes ayant une tendance similaire pour les valeurs maximales et minimales d’erreurs, l’étude menée sur le seul critère de l’erreur moyenne est acceptable. De plus, nous pouvons remarquer trois zones principales : les pires cas et les meilleurs cas respectivement représentés en début et fin de graphique qui montrent une grande variabilité (relative à environ 50 combinaisons chacune), tandis qu’une zone intermédiaire semble peu sensible aux combinaisons associées (plus de 600 combinaisons).

FIGURE 1.23 – Marqueurs d’erreur utilisés, avec ϕ l’erreur de phase, A_1 et A_2 les amplitudes respectives du vecteur référence et du vecteur réel, et en bleu l’amplitude du vecteur d’erreur (EVM)

$$EVM = \sqrt{A_1^2 + A_2^2 - A_1 \cdot A_2 \cdot \cos(\phi)} \quad (1.10)$$

Grâce aux marqueurs présentés plus haut il est possible de trouver des agencements de cellules optimaux. De même, il est possible d’identifier des combinaisons à éviter, qui présentent les performances les plus dégradées. Des variations très importantes peuvent être obtenues entre les meilleures et pires

combinaisons de cellules (50% sur la Figure 1.24). Notre outil d'analyse combinatoire nous permettra d'éviter les agencements les plus défavorables ; ceux-ci peuvent être analysés statistiquement, et peuvent servir à évaluer l'impact individuel des cellules, voire d'une combinaison de cellules.

Obtenir des variations d'erreurs importantes sur l'ensemble des agencements renseigne quant à l'apparition probable de cellules « critiques ». Ces cellules peuvent être critiques de par leurs erreurs initiales fonctionnelles (changement topologique) ou du fait qu'elles soient associées à d'autres cellules qui les mettent dans un plan de charge défavorable (amélioration des critères de fermeture ou modification de l'ordre opérationnel des cellules). Cette première analyse est importante puisqu'elle permet dans un premier lieu de savoir si cette méthodologie d'optimisation va apporter un gain de performance globale.

Dans l'affirmative, les agencements optimaux de cellules sont identifiés et les positionnements ou groupes de cellules dégradant les performances sont aussi renseignés. Ceci correspond au retour #1 de la Figure 1.22 pour lequel seul notre outil combinatoire est exploité afin d'ordonner les cellules.

Dans le cas contraire où la méthodologie n'apparaît pas suffisante au premier tri, il est nécessaire de vérifier si certaines améliorations (module ou phase des S_{ii} et/ou S_{ij}) de caractéristiques de cellules individuelles ne sont pas susceptibles de profiter à l'ensemble du système. Cette nouvelle itération est traduite par les retours #2 et #3 sur la Figure 1.22. Si l'amélioration est suffisante, cela aboutira à une nouvelle optimisation (#2), tandis qu'un travail itératif devra être mené sur les cellules critiques (adaptation et/ou variations de phase/amplitude #3) jusqu'à l'obtention de résultats aboutissant à la condition de sortie #2.

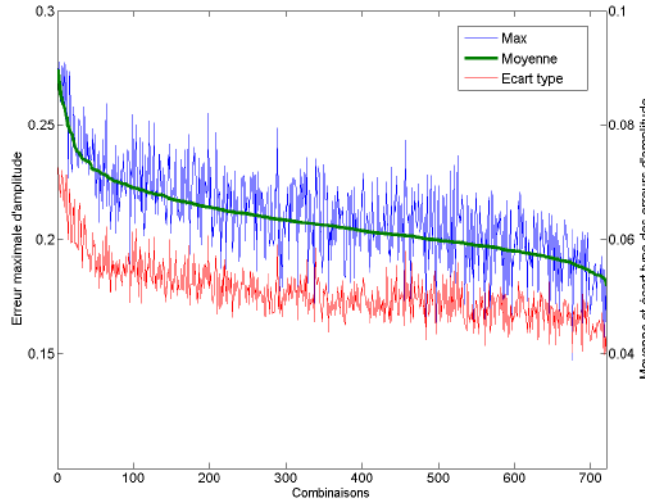


FIGURE 1.24 – Corrélation entre les indicateurs : moyenne (bleu), maximum (rouge) et écart type (vert) avec un classement selon l'erreur moyenne décroissante.

9 Modes de fonctionnement

Les différentes conceptions de ce manuscrit ont été développées selon deux modes de fonctionnement différents, le mode *single ended* et le mode différentiel. Si l'utilisation du mode *single ended* est la plus intuitive car c'est dans ce mode de fonctionnement qu'opère la majorité des *core-chips*, il reste

intéressant de se questionner quant à la pertinence de l'utilisation du mode différentiel. Le choix a premièrement été motivé par une problématique initiale d'OMMIC qui visait à réduire le rayonnement de couplage des cellules les unes par rapport aux autres et à l'intérêt du transfert des structures différentielles utilisées en SiGe vers le GaN. En confinant le champ électromagnétique grâce au mode différentiel nous pensions nous affranchir de ce couplage parasite. De plus, sur le papier, le mode différentiel opère sans l'utilisation de *via-holes* éléments difficiles à modéliser fidèlement, en plus de leur encombrement non négligeable sur les trajets RF.

En réalité cet avantage est à nuancer puisque pour garantir la bonne polarisation statique des transistors, l'insertion de masses statiques est obligatoire. Le mode différentiel présente un second avantage, vis à vis des niveaux de puissance admissible, en effet les puissances transitant sur chaque bras étant 3dB en dessous de leur équivalent *single ended*, la linéarité des circuits différentiels est donc augmentée en théorie de 3dB. Pour profiter de cet avantage, la symétrie des signaux (en amplitude) sur chaque bras doit être garantie systématiquement, ce qui peut devenir une contrainte lors du dessin. En effet, la symétrie en amplitude implique une symétrie géométrique de dessin. Cependant cela reste une condition nécessaire mais non suffisante à l'obtention de la symétrie électrique. De ce fait, tous les éléments du circuit doivent être symétrisés, chaque transition métallique entraînant de la conversion de mode différentiel vers du mode commun. Cette aspect de symétrie est très important puisque, comme explicité plus tôt, en mode différentiel, les *via-holes* ne sont plus nécessaires en RF puisque des masses virtuelles apparaissent à l'interface des deux branches. Cependant il faut respecter une symétrie vis-à-vis de ces masses pour être sûr qu'elles soient formées à l'endroit souhaité, garantissant une bonne mise à la masse dynamique.

En ce qui concerne les polarisations statiques, la symétrie peut parfois être difficile à respecter, selon la géométrie du transistor (nombre de doigts et agencement de l'accès de grille). La polarisation doit aussi être dédoublée étant donné que le nombre de transistors utilisé est multiplié par deux, la consommation est donc logiquement doublée par rapport à des topologies *single ended*. Ce dernier point ne revêt pas d'importance en mode commutation pour lequel la consommation est réduite aux fuites et aux états de transition ON/OFF ou OFF/ON. Le réseau de polarisation doit être symétrique par rapport à lui-même, mais aussi symétrisé par rapport aux pistes RF pour ne pas perturber le fonctionnement différentiel du circuit. La configuration différentielle occupe donc un espace plus important et offre moins de liberté de design. D'un autre côté elle offre une meilleure linéarité et une immunité aux bruits de type mode commun. Le Tableau. 1.6 résume les avantages et inconvénients de ces deux approches. Ces derniers seront confrontés au retour d'expérience à l'issue des travaux présentés dans ce manuscrit.

	Avantages	Inconvénients
Single-Ended	-la majorité des éléments présente des interfaces single-ended -facilité de dessin du layout	-d'avantage de via-holes -gestion du mode commun pas toujours maîtrisée
Différentiel	-cellules 180° très simplifiée -théoriquement pas de via-holes -confine mieux le champ EM -puissance injectable 3dB supérieure	-forte contrainte sur le dessin du layout -dédoublage de tous les éléments (transistors, polarisation,...)

TABLE 1.6 – Avantages et inconvénients des deux approches utilisées

Durant ces travaux de thèse nous allons donc explorer les possibilités de conceptions selon 2 axes principaux, l'approche (*single-ended* ou différentiel) et l'optimisation de l'ordre des cellules (algorithme et méthode). La Figure 1.25 présente les layouts des circuits réalisés durant cette thèse pour illustrer la stratégie de conception visée durant cette étude. Il aurait été tout aussi intéressant de développer une comparaison sur la version *single-ended* centrée 35 GHz et une version large bande 30-40 GHz afin

d'apprécier les performances selon les deux stratégies de conception.

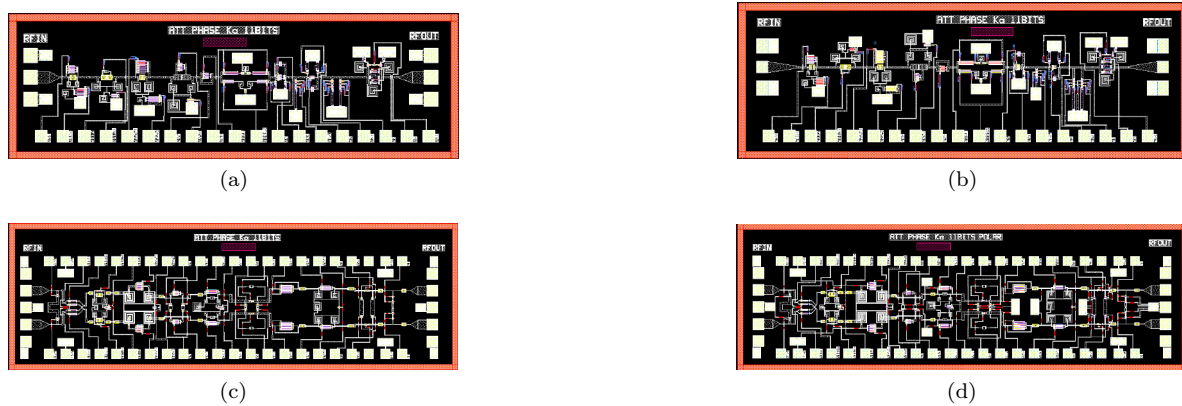


FIGURE 1.25 – Circuits réalisés selon l'objectif et l'axe de conception choisis ; Single-ended (a) sans optimisation (b) avec optimisation ; Différentiels (c) sans optimisation (d) avec optimisation

Durant ce chapitre, nous avons présenté tous les éléments nécessaires à la compréhension des *core-chips*, des éléments les constituant à leur mode de fonctionnement. Nous avons ensuite présenté notre méthodologie de conception de *core-chip* en technologie GaN sur la bande Ka. Durant les différentes conceptions nous avons fait deux choix topologiques différents, les circuits *single-ended* et les circuits différentiels. Ces derniers seront étudiés avec ou sans notre méthodologie d'optimisation de mise en commun des cellules et cela dans le but de déterminer les topologies respectant notre cahier des charges. Le chapitre suivant traite de la première étape de conception, les simulations des cellules.

Bibliographie

- [1] Martin Oppermann and Ralf Rieger. Multifunctional MMICs – Key Enabler for Future AESA Panel Arrays. In *2018 IMAPS Nordic Conference on Microelectronics Packaging (NordPac)*, pages 77–80. IEEE, jun 2018.
- [2] Egor Alekseev, Shawn S.H. Hsu, Dimitris Pavlidis, Tadayoshi Tsuchiya, and Michio Kihara. Broad-band AlGaIn/GaN HEMT MMIC Attenuators with High Dynamic Range. In *30th European Microwave Conference, 2000*, pages 1–4. IEEE, oct 2000.
- [3] Jean-Guy Tartarin. LA TECHNOLOGIE GAN ET SES APPLICATIONS POUR L'ELECTRONIQUE ROBUSTE, HAUTE FREQUENCE ET DE PUISSANCE. mar 2008.
- [4] Tyler Neil Ross. Gallium Nitride Phase Shifters. 2015.
- [5] Kuan Bao, Jun Zhou, Liangui Wang, Anfeng Sun, Qiang Zhang, and Ya Shen. A 29–30GHz 64-element Active Phased array for 5G Application. In *2018 IEEE/MTT-S International Microwave Symposium - IMS*, pages 492–495. IEEE, jun 2018.
- [6] Ralf Rieger, Andreas Klaaben, Patrick Schuh, and Martin Oppermann. A Full-Array-Grid-Compatible Wideband Tx/Rx Multipack Using Multifunctional Chips on GaN and SiGe. In *2018 48th European Microwave Conference (EuMC)*, pages 1453–1456. IEEE, sep 2018.
- [7] D.M. Pozar and B. Kaufman. Design considerations for low sidelobe microstrip arrays. *IEEE Transactions on Antennas and Propagation*, 38(8) :1176–1185, 1990.

- [8] Daniel Kramer. Ka-Band P-I-N Diode Based Digital Phase Shifter. In *2018 48th European Microwave Conference (EuMC)*, pages 1285–1288. IEEE, sep 2018.
- [9] Jung Gil Yang, Jooseok Lee, and Kyoungsoon Yang. A W-band InGaAs PIN-MMIC digital phase-shifter using a switched transmission-line structure. In *2012 International Conference on Indium Phosphide and Related Materials*, pages 99–101. IEEE, aug 2012.
- [10] K. Maruhashi, H. Mizutani, and K. Ohata. A Ka-band 4-bit monolithic phase shifter using unresonated FET switches. In *1998 IEEE MTT-S International Microwave Symposium Digest (Cat. No.98CH36192)*, volume 1, pages 51–54. IEEE.
- [11] H.A. Atwater. Circuit Design of the Loaded-Line Phase Shifter. *IEEE Transactions on Microwave Theory and Techniques*, 33(7) :626–634, jul 1985.
- [12] Anuradha Ahlawat, Rahul Gupta, and Mohammad S. Hashmi. Wideband phase shifter with stub loaded transmission line. In *2017 IEEE Asia Pacific Microwave Conference (APMC)*, pages 283–286. IEEE, nov 2017.
- [13] Min Zhou, Jiongiong Mo, and Zhiyu Wang. A Ka-band low power consumption MMIC core chip for T/R modules. *AEU - International Journal of Electronics and Communications*, 91 :37–43, jul 2018.
- [14] Chao Liu, Qiang Li, Yihu Li, Xiao Dong Deng, Hailin Tang, Ruitao Wang, Haitao Liu, and Yong Zhong Xiong. A Ka-Band Single-Chip SiGe BiCMOS Phased-Array Transmit/Receive Front-End. *IEEE Transactions on Microwave Theory and Techniques*, 64(11) :3667–3677, nov 2016.
- [15] Tyler N. Ross, Gabriel Cormier, Khelifa Hettak, and Jim S. Wight. High-power X-band GaN switched-filter phase shifter. In *2014 IEEE MTT-S International Microwave Symposium (IMS2014)*, pages 1–4. IEEE, jun 2014.
- [16] Duixian. Liu. *Advanced millimeter-wave technologies : antennas, packaging and circuits*. J. Wiley & Sons, 2009.
- [17] Qin Zheng, Zhiyu Wang, Kangrui Wang, Gang Wang, Hui Xu, Liping Wang, Wei Chen, Min Zhou, Zhengliang Huang, and Faxin Yu. Design and Performance of a Wideband Ka-Band 5-b MMIC Phase Shifter. *IEEE Microwave and Wireless Components Letters*, 27(5) :482–484, may 2017.
- [18] Robin Garg and Arun S. Natarajan. A 28-GHz Low-Power Phased-Array Receiver Front-End With 360° RTPS Phase Shift Range. *IEEE Transactions on Microwave Theory and Techniques*, 65(11) :4703–4714, nov 2017.
- [19] Maryam Tabesh, Amin Arbabian, and Ali Niknejad. 60GHz low-loss compact phase shifters using a transformer-based hybrid in 65nm CMOS. In *2011 IEEE Custom Integrated Circuits Conference (CICC)*, pages 1–4. IEEE, sep 2011.
- [20] Anas J. Abumunshar, Niru K. Nahar, Daniel Hyman, and Kubilay Sertel. 18–40GHz low-profile phased array with integrated MEMS phase shifters. In *2017 11th European Conference on Antennas and Propagation (EUCAP)*, pages 2800–2801. IEEE, mar 2017.
- [21] Saeed Akbari, Mostafa Amirpour, E.Abbaspour Sani, M. N. Azarmanesh, and sina Akbari. A ka-band phase shifter based on a new four-state MEMS switch. In *2017 Iranian Conference on Electrical Engineering (ICEE)*, pages 541–545. IEEE, may 2017.
- [22] Sharif Iqbal Mitu Sheikh, A. A. P. Gibson, M. Basorrah, G. Alhulwah, K. Alanizi, M. Alfarsi, and J. Zafar. Analog/Digital Ferrite Phase Shifter for Phased Array Antennas. *IEEE Antennas and Wireless Propagation Letters*, 9 :319–321, 2010.

- [23] TGL2223 - Qorvo.
- [24] Jung Gil Yang and Kyounghoon Yang. Ka-Band 5-Bit MMIC Phase Shifter Using InGaAs PIN Switching Diodes. *IEEE Microwave and Wireless Components Letters*, 21(3) :151–153, mar 2011.
- [25] Na Chen. A millimeter-wave 6-bit GaAs monolithic digital attenuator with low insertion phase shift. In *2013 International Workshop on Microwave and Millimeter Wave Circuits and System Technology*, pages 440–443. IEEE, oct 2013.
- [26] Ye Yuan, Shan-Xiang Mu, and Yong-Xin Guo. 6-bit Step Attenuators for Phased-Array System With Temperature Compensation Technique. *IEEE Microwave and Wireless Components Letters*, 28(8) :690–692, aug 2018.
- [27] Juseok Bae and Cam Nguyen. A Novel Concurrent 22–29/57–64-GHz Dual-Band CMOS Step Attenuator With Low Phase Variations. *IEEE Transactions on Microwave Theory and Techniques*, 64(6) :1867–1875, jun 2016.
- [28] Inder J. Bahl. *Control Components Using Si, GaAs, and GaN Technologies*. Artech House, 2014.
- [29] CHT3029-99F – DC-35GHz 4 Bit Digital Attenuator.
- [30] MACOM - Product Detail - MAAD-011021-DIE.
- [31] Remy Leblanc, Noelia Santos Ibeas, Ahmed Gasmi, and Joel Moron. Ka Band Chip-Set for Electronically Steerable Antennas. In *2014 IEEE Compound Semiconductor Integrated Circuit Symposium (CSICS)*, pages 1–4. IEEE, oct 2014.
- [32] Dong-Woo Kang, Jeong-Geun Kim, Byung-Wook Min, and G.M. Rebeiz. Single and Four-Element \$Ka\$-Band Transmit/Receive Phased-Array Silicon RFICs With 5-bit Amplitude and Phase Control. *IEEE Transactions on Microwave Theory and Techniques*, 57(12) :3534–3543, dec 2009.
- [33] Andrea Bentini, Mauro Ferrari, Walter Ciccognani, and Ernesto Limiti. A novel approach to minimize RMS errors in multifunctional chips. *International Journal of RF and Microwave Computer-Aided Engineering*, 22(3) :387–393, may 2012.

Chapitre 2

Conception des cellules de déphasage et d'atténuation

Dans ce chapitre, nous présenterons en premier lieu la société OMMIC [1] basée à Paris et qui a motivé le sujet de cette thèse avec les partenaires institutionnels du LAAS-CNRS de Toulouse et du LN2 Sherbrooke, puis nous nous intéresserons à l'environnement de simulation GaN (D01GH) utilisé tout au long des différentes conceptions. Ensuite nous définirons les cahiers des charges choisis pour la réalisation des différents *core-chips*. Nous finirons enfin par présenter les performances des cellules individuelles simples (*single-ended*) et différentielles qui seront par la suite intégrées dans les systèmes *core-chip* complets.

1 La société OMMIC

La société OMMIC est une fonderie spécialisée dans les matériaux III-V, spécifiquement le GaAs, le GaN et l'InP. Son savoir-faire s'étend de la croissance par épitaxie jusqu'à la fabrication de circuits MMIC. Ce qui lui permet de fournir des produits à l'état de l'art pour des applications de niches et spécifiques à chaque client. L'évolution des standards de télécommunications, avec notamment l'arrivée très prochaine de la 5G, motive les développements effectués en direction de la filière technologique GaN, tant au niveau du procédé de fabrication que des fonctions électroniques conçues à partir de ce procédé. Les domaines d'applications visés vont de l'amplification (HPA et LNA) au traitement du signal (mélangeurs et convertisseurs de fréquence) en passant par des modules d'émission réception hautement intégrés. Les bandes de fréquence couvertes sont larges (du kHz jusqu'au-delà de 150 GHz) et autorisent des débits allant jusqu'à 80Gb/s. Dans le chapitre précédent nous avons montré l'intérêt du GaN pour la réalisation de composants micro-ondes nécessaires aux fonctions de contrôle. Dans cette optique la société OMMIC souhaite renouveler son portfolio en intégrant progressivement des composants en technologie GaN et ce pour à plus long terme pouvoir substituer complètement les éléments GaAs. La bande de fréquence identifiée pour ce travail de thèse (et une bande possible pour la 5G) est la bande Ka et plus spécifiquement le segment 30-40 GHz.

2 Présentation du kit de conception de la filière GaN D01GH

Le kit de conception (Design Kit, ou DK) D01GH fourni par OMMIC s'intègre dans le logiciel de CAO ADS de Keysight et le logiciel AWR de National Instruments. L'ensemble du travail de simulation, électrique et électromagnétique (EM), présenté dans ce manuscrit a été effectué dans l'environnement ADS. Les composants mis à disposition du concepteur sont disponibles sous forme de modèles électriques et EM. Nous ne présenterons pas en détail tous les composants mais nous nous attacherons à décrire brièvement les transistors et les éléments passifs qui ont été utilisées dans les circuits conçus.

Les transistors HEMT de la technologie GaN D01GH sont construits sur un substrat AlGaN de $2\mu\text{m}$ d'épaisseur positionné au-dessus d'un substrat de silicium aminci de $100\mu\text{m}$ d'épaisseur. La vue en coupe de la technologie, illustrée en Figure 2.1, fait apparaître deux niveaux métalliques nommés IN et MET1 (en jaune). Les niveaux OH, LI et MD sont utilisées pour la réalisation de composants passifs (résistances et capacités). Comme pour les filières technologiques GaAs, la face arrière du substrat est métallisée et supporte la connexion de masse. Cette masse peut être acheminée vers la face supérieure par l'intermédiaire du via VH.

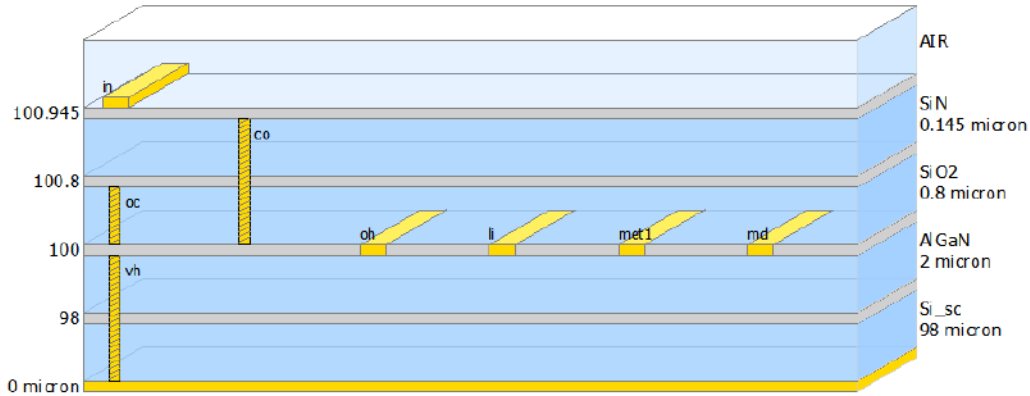


FIGURE 2.1 – Empilement des couches de la filière GaN D01GH

2.1 Les éléments passifs

Résistances

Il existe deux types de résistances, la première est réalisée à partir d'une couche de Nickel Chrome (couche MD) tandis que la seconde exploite la résistivité de la couche GaN (couche LI). La résistance en Nickel est très précise mais produit de faibles valeurs. Il existe deux degrés de liberté sur ces composants, la longueur et la largeur. Les tailles limites sont une largeur de $5\mu\text{m}$ et une longueur de $6\mu\text{m}$, pour la limite haute en taille ce sont les effets distribués qui apparaissent qui la fixe. Il existe un modèle électrique et un modèle EM pour cet élément. Le dessin des masques de cette résistance est présenté en Figure 2.2a.

La résistance à couche GaN permet d'obtenir des valeurs beaucoup plus élevées mais au détriment de la précision. Elles sont donc utilisées là où la valeur de la résistance n'a pas d'effet néfaste sur la fonction première de la cellule, par exemple en tant que résistance de grille. Des modèles électriques et EM sont disponibles pour leur utilisation, le dessin des masques associé est présenté en Figure 2.2b



FIGURE 2.2 – Dessins des masques des résistances (a) NiCr (b) GaN

Capacités

Deux types de capacité sont disponibles dans le DK. La première utilise la mise en regard des couches IN et MET1 pour réaliser l'effet capacitif. Ces deux niveaux métalliques sont séparés par deux couches d'isolant, du SiN et du SiO₂. La seconde capacité utilise les mêmes niveaux métalliques pour réaliser les armatures. Cependant, la couche SiO₂ est court circuitée en utilisant le via OC reliant le niveau MET1 à la couche SiN, ce qui permet d'obtenir des valeurs de capacité plus importantes grâce à la diminution de distance entre les deux armatures. Les degrés de liberté pour le dessin de ces capacités sont les largeurs d'accès, et la largeur des armatures. Les dessins des masques de ces deux condensateurs sont présentés en Figure 2.3.

FIGURE 2.3 – Dessins des masques de capacités (a) SiO₂ (b) MIM

Pour les capacités MIM que nous avons majoritairement utilisées durant la conception, il existe des différences de performances fréquentielles entre les modèles électriques et les modèles EM. La valeur initiale pour les deux modèles est la même cependant le modèle électrique exploite l'élément TFC d'ADS, ce dernier est moins précis que la simulation EM. La Figure 2.4 illustre les différences de valeurs suivant ces deux modèles. A 35GHz, cet écart important peut justifier les différences entre les simulations électriques et EM du *core-chip* complet.

Inductances

Les inductances disponibles sont de type enroulement carré avec une entrée IN et une sortie en MET1. Les degrés de liberté permettent le choix de la largeur des lignes IN ($>5\mu\text{m}$), l'écartement entre les lignes ($>5\mu\text{m}$) et la largeur de la piste MET1 de sortie qui va déterminer la taille de l'interconnexion CO entre IN et MET1. Il est important de vérifier la valeur réelle de cette inductance par rapport à la valeur demandée pour être sûr de se trouver dans la zone de fonctionnement inductive. Ainsi les simulations électriques devront toujours être confirmées par des simulations EM puisque le modèle électrique n'est utilisable que jusqu'à 25 GHz. De plus nous nous sommes rendu compte que pour certaines valeurs basses ($<0.2\text{nH}$) et hautes ($>1\text{nH}$) les résonnances n'étaient pas bien prises en compte dans les modèles électriques que nous avons utilisés durant les premières conceptions ou qu'il pouvait y avoir d'importantes variations de valeur de l'inductance à la fréquence d'intérêt. Ce phénomène est illustré en Figure 2.5.

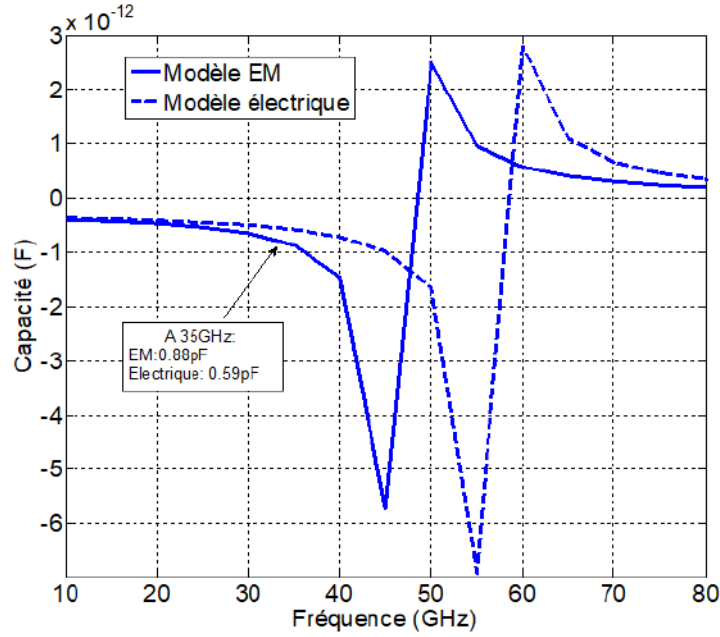


FIGURE 2.4 – Comportements fréquentiels des modèles électrique et EM d'une capacité MIM de valeur nominale 0.38pF

3 Description simulation EM

Pour les simulations EM nous avons utilisé le module Momentum d'ADS. Ce dernier permet de résoudre les équations de Maxwell dans des structures planaires intégrées dans un substrat diélectrique multicouches. Pour cela le logiciel utilise la méthode des moments pour réaliser une discrétisation numérique des équations [2]. Il existe deux modes de résolution des équations, le mode Radio-Fréquence (RF) qui utilise une approximation quasi statique et le mode Micro-Ondes (*Microwave*) qui utilise les équations pleines-ondes (*full-waves*).

La différence entre les deux modes de simulation vient de la définition des inductances et capacités $L_{i,j}$ et $C_{i,j}$. En mode RF, ces valeurs ne sont pas dépendantes de la fréquence ce qui permet de ne calculer qu'une seule matrice d'interaction pour toute la plage de fréquence. En revanche ce n'est plus vrai en mode micro-ondes où il faut calculer la matrice à chaque fréquence. C'est de cette différence que vient la durée de calcul bien plus importante pour le mode micro-ondes. Le mode RF est adapté pour des simulations de circuits de petites tailles ($< \lambda/2$) mais peine à modéliser fidèlement des circuits de tailles plus importantes du fait de l'apparition d'effets distribués. De plus, ce mode ne prend pas en compte les radiations émises par le circuit. A l'échelle d'une cellule, ce n'est pas la dimension qui est préjudiciable au mode RF mais plutôt la prise en compte du rayonnement qui n'est pas assurée. Or ce dernier peut renseigner sur certains défauts de la cellule c'est pourquoi nous avons utilisé le mode micro-ondes tout au long de nos simulations.

Pour envoyer l'excitation électrique en entrée et/ou sortie du circuit et collecter les réponses, il est nécessaire d'utiliser des ports de connexions. Nous n'entrerons pas dans les détails du choix des ports, puisqu'une étude menée au sein d'OMMIC a conclu que l'utilisation de sources d'excitation (Ports) TML *zero-length* (de longueur nulle) est à privilégier. Ce qui veut dire que ce sont des ports calibrés

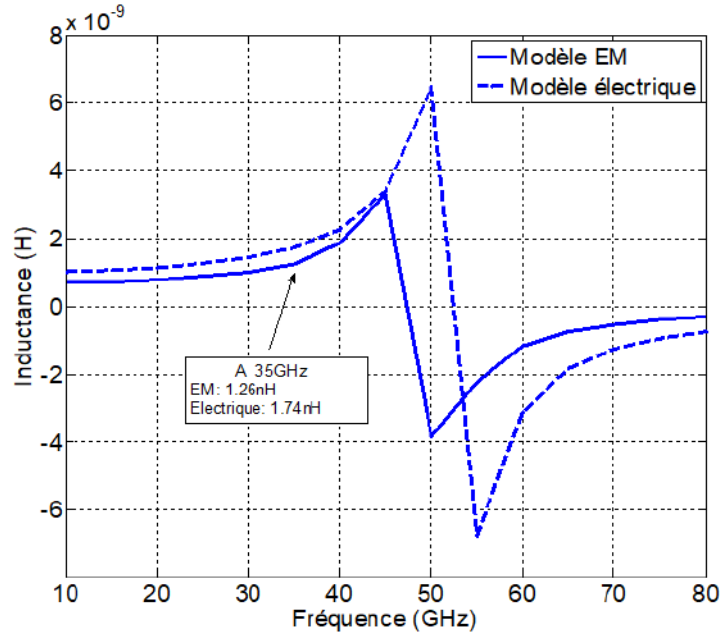


FIGURE 2.5 – Comportements fréquentiels des modèles électrique et EM d’une inductance de valeur nominale 0.99nH

mais pour lesquels le calibrage n’ajoute pas de longueur physique de ligne (d’où le *zero-length*). Ces ports suppriment uniquement la capacité de bord de l’interface entre le port et la structure testée. Pour le maillage nous l’effectuons à la fréquence de simulation la plus haute (en général 40 GHz) en utilisant 50 cellules par longueur d’onde.

Concernant la fiabilité des différentes simulations (EM et électriques), nous supposons que l’EM traduit toujours plus fidèlement le comportement des différents éléments du circuits. Cependant le nombre de ports utilisés va influencer les résultats de simulation. Par exemple, pour le cas de la simulation des cellules individuelles, celles-ci seront modélisées par leurs ports d’entrée et sortie mais aussi par les ports d’accès aux éléments actifs de chaque cellule, les transistors.

Plus tard dans les explications, nous parlerons de simulations électriques des cellules individuelles. Il faut donc comprendre dans ce cas-là une mise en commun idéale de modèles EM de chaque cellule. Ainsi, le nombre de ports EM utilisées sera bien plus important que dans une mise en commun des cellules directement sur le dessin des masques du système complet puisque les ports entrée sortie de chaque cellule auront disparus. De part ce nombre de ports plus élevé, les différences de résultats de simulations entre les simulations dites électriques et EM peuvent devenir importantes. Face à ces différences de résultats, nous avons choisi de faire confiance aux résultats de simulations EM pour la validation finale du circuit. Ainsi, certaines différences entre les simulations perçues dans la suite de cette étude peuvent être expliquées et ne tirent pas la sonnette d’alarme sur un éventuel problème de modélisation.

4 Cahier des charges

Comme nous l'avons explicité précédemment nous avons utilisé deux approches durant la conception des différents *core-chips*, une orientée vers une largeur de bande conséquente et l'autre focalisée sur une seule fréquence. Ce choix implique l'élaboration de cahiers de charges différents selon l'approche utilisée. Dans un premier temps il est nécessaire d'établir un cahier des charges pour les cellules individuelles puis un cahier des charges à l'échelle système (performances du *core-chip*).

4.1 Approche large bande

Pour cette approche nous devons garantir les performances sur la bande de fréquence 30 à 40 GHz, pour cela nous avons commencé par observer les performances atteignables avec des topologies idéales. Ces valeurs serviront ensuite d'objectifs pour la réalisation des cellules individuelles finales. Dans le flot de conception d'OMMIC, l'étape correspondante se nomme la *Preliminary Design Review* (PDR). Durant celle-ci nous établissons les topologies utilisées, les performances et la robustesse de ces dernières face à des variations de composants utilisés. A la fin de cette étape, un tableau récapitulatif établissant les performances individuelles à atteindre pour chaque cellule est donné. Durant toutes les étapes charnières qui suivent, une matrice de conformité des performances est établie, elle permet de vérifier si les performances obtenues à l'étape concernée vérifient le cahier des charges initial. Les deux réalisations effectuées en suivant cette méthode sont les deux *core-chips* différentiels 11bits, 30-40 GHz.

Pour donner suite à la PDR, les exigences en performances sont données dans le cahier des charges sous la forme du Tableau 1. Ces contraintes sont exigées sur toute la bande 30-40GHz. Le niveau d'adaptation a été fixé en comparant le faisable à l'idéal, en effet, plus tôt, nous avons spécifié qu'un niveau de -18 à -20 dB garantissait une bonne indépendance de chaque cellule vis à vis des problématiques dues aux effets de charge. Cependant, exiger un tel niveau d'adaptation tout en garantissant des erreurs fonctionnelles aussi basses sur toute la bande de fréquence est impensable. Le seuil de 15dB apparaît donc comme un compromis réaliste permettant de limiter les effets de charges tout en gardant une marge de liberté de conception plus importante à utiliser pour satisfaire les exigences fonctionnelles. Durant ce *design* nous n'avons pas quantifié l'erreur de phase acceptable pour les atténuateurs. Effectivement, malgré un effort de réduction de celle-ci nous n'avons pas formalisé d'objectifs précis la concernant. Il faut cependant la conserver en dessous de $5,625^\circ$ pour ne pas perturber l'effet du bit de poids faible en phase. De la même façon, l'erreur d'amplitude des atténuateurs a été fixée arbitrairement pour ne pas dégrader la résolution en amplitude du système complet (0.5dB).

Cellule	Pertes d'insertion (dB)	Adaptation (dB)	Erreur de phase ($^\circ$)	Erreur d'atténuation (dB)
$5,625^\circ$	0,2	-15	5	x
$11,25^\circ$	0,5	-15	5	x
$22,5^\circ$	0,6	-15	5	x
45°	1	-15	5	x
90°	1,5	-15	5	x
180°	2	-15	5	x
0,5dB	0,2	-15	x	0,07
1dB	0,4	-15	x	0,07
2dB	2	-15	x	0,07
4dB	2,5	-15	x	0,07
8dB	2,5	-15	x	0,07

TABLE 2.1 – Cahier des charges préliminaires core-chip différentiel 30-40 GHz

4.2 Approche mono-fréquence

Cette seconde approche vient pallier certains défauts de la méthode précédente. En effet, il est difficile, lors de la mise en commun des cellules, de déterminer l'origine des erreurs obtenues. Comme nous l'avons dit il existe les erreurs fonctionnelles (phase et amplitude) dues à des erreurs déjà existantes lors de la conception de cellules individuelles et il existe les erreurs dues à la désadaptation. Imaginons utiliser l'approche large bande, une fois les cellules rassemblées, nous obtenons une certaine quantité d'erreurs.

Dans une démarche d'amélioration des performances, nous recherchons donc la source de ces erreurs. Mais comment peut-on décorrélérer les sources d'erreurs ? Nous n'avons pas de fonction « coût » qui exprime l'erreur finale selon l'adaptation par exemple. Ainsi il est quasiment impossible de trouver, de façon systématique, les curseurs sur lesquels jouer pour corriger ces erreurs. Une solution que nous avons trouvée est de s'affranchir des erreurs fonctionnelles, cependant ce n'est possible que pour une bande de fréquence très réduite, voir pour une seule fréquence. De cette manière, à cette fréquence, les erreurs nulles fonctionnelles s'additionnent pour une somme d'erreurs logiquement nulle. Il en découle que l'existence d'erreurs lors de la mise en commun renseigne sur l'existence d'une autre source d'erreurs : le niveau d'adaptation. Il faut quand même préciser un élément, même si les erreurs fonctionnelles ciblées sont nulles, les erreurs fonctionnelles croisées (erreur de phase des atténuateurs et erreurs d'atténuation des déphaseurs) ne sont pas toujours nulles simultanément aux erreurs fonctionnelles. Nous en reparlerons par la suite quand nous aborderons la méthodologie d'optimisation. Nous n'avons donc pas établi de cahier des charges précis pour cette approche le seul objectif clairement identifié était d'obtenir une erreur fonctionnelle nulle (au moins pour l'erreur direct, une erreur croisée nulle étant souhaitable mais facultative).

5 Réalisation cellules individuelles single ended

Dans cette partie nous allons présenter un exemple de réalisation de cellule *single-ended* du *core-chip*. En effet nous ne présenterons pas cette méthodologie pour toutes les cellules car ce serait redondant. Cependant nous l'effectuerons pour la cellule 11,25° *single-ended* pour illustrer les étapes de la méthodologie précédemment explicitée.

5.1 Cellule 11,25° single ended

Pour réaliser cette cellule, nous avons choisi une topologie basée sur un filtre passe-haut commuté. Cette topologie présente une bonne platitude de déphasage sur la bande de fréquence visée et cela en conservant des pertes d'insertion acceptables. Nous commençons tout d'abord par calculer la valeur des éléments composant le filtre passe haut, L_1 et C_1 illustrés en Figure 2.6. Cependant, puisque le transistor de la voie de référence est intégré dans la valeur de la capacité du filtre passe haut, il est nécessaire de modifier la valeur de C_1 en conséquence. De plus le transistor va rajouter une résistante R_{on} quand il sera passant, ce qui rend la topologie différente d'un simple filtre passe haut. Ainsi il paraît plus cohérent d'intégrer les transistors dans la première étape de la méthodologie de conception de sorte que les deux états de ce dernier soient pris en compte lors de la simulation. Le schéma électrique qui servira de référence sur le simulateur ADS est illustré en Figure 2.6. Des résistances GaN de 4k Ω sont placées sur les grilles des transistors pour éviter au signal RF de se coupler à la commande statique de la grille du transistor, ce qui aurait pour effet de moduler la polarisation et donc l'état de l'impédance présentée par le canal. Ce phénomène n'est pas souhaitable pour des raisons de linéarité. Le modèle électrique du transistor proposé dans le DK se présente sous la forme d'une impédance

paramétrique fonction des tensions grille-source (V_{gs}) et drain source (V_{ds}). Ces deux tensions sont-elles-même déclarées sous la forme de variables externes.

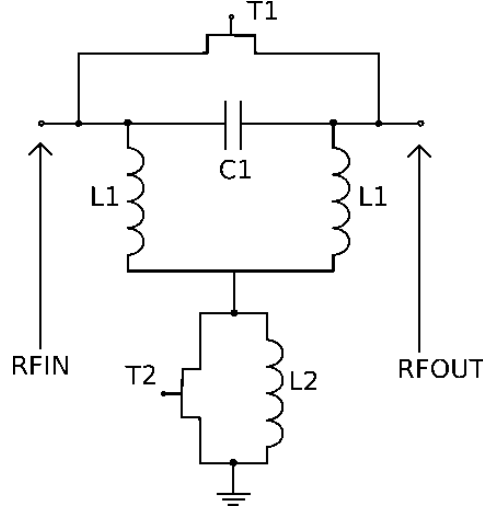


FIGURE 2.6 – Topologie utilisant un filtre passe haut commuté utilisé lors de la conception de la cellule 11.25°

Lors des simulations nous avons jugé les performances de la cellule au regard de l'erreur de phase, des pertes d'insertion, de l'erreur d'atténuation et de l'adaptation. Ces 4 différents paramètres apparaîtront à chaque étape de la méthodologie de conception. En calculant les valeurs des éléments nécessaires du filtre passe haut, nous obtenons des inductances L_1 de 2,31nH et une capacité C_1 de 0,47pF. Contrairement à la valeur de la capacité qui est assez proche de la valeur calculée, les inductances présentent une valeur plus de 3 fois inférieure à celle calculée. Cela est explicable par le circuit bouchon qui influence les valeurs des inductances. Cette influence est bénéfique puisqu'il est très difficile d'obtenir une inductance de 2,3nH à 35GHz sans effets parasites. Les résultats de simulation ADS sont présentés en Figure 2.7. Nous pouvons voir que sur chaque courbe, le comportement passant/bloqué de la cellule est visible puisque 2 courbes sont présentes. Nous n'avons pas précisé quel état était la référence et l'autre l'état déphasé puisque c'est un déphasage relatif entre deux états. Il est important de noter les variations importantes existantes entre les simulations idéales et les simulations EM. Ces différences demanderont donc une correction pour obtenir des performances plus proches des simulations idéales.

Les étapes de conception sont telles que suit :

- définition de la topologie finale (filtre passe haut et circuit bouchon) avec les éléments idéaux du simulateur. Nous choisissons de faire figurer le circuit bouchon dans la première itération pour que tous les éléments soient pris en compte dans cette simulation, malgré le fait qu'elle soit dite « idéale ». Les courbes ne sont pas présentées en Figure 2.7.

- Utilisation des éléments du DK D01GH pour remplacer les passifs d'ADS (capacités et inductances). Cette étape a pour but de valider la valeur des passifs qui seront utilisés dans la topologie réelle finale.

- Connexion des différents éléments entre eux à l'aide de lignes. Dans cette étape nous avons choisi de commencer à dessiner le dessin des masques final de la cellule réelle pour estimer les longueurs de ligne nécessaires à connecter tous les éléments entre eux. Une fois ces lignes obtenues nous pouvons réaliser la simulation électrique comportant les éléments du DK et les lignes. C'est normalement la

simulation la plus proche de la simulation EM finale.

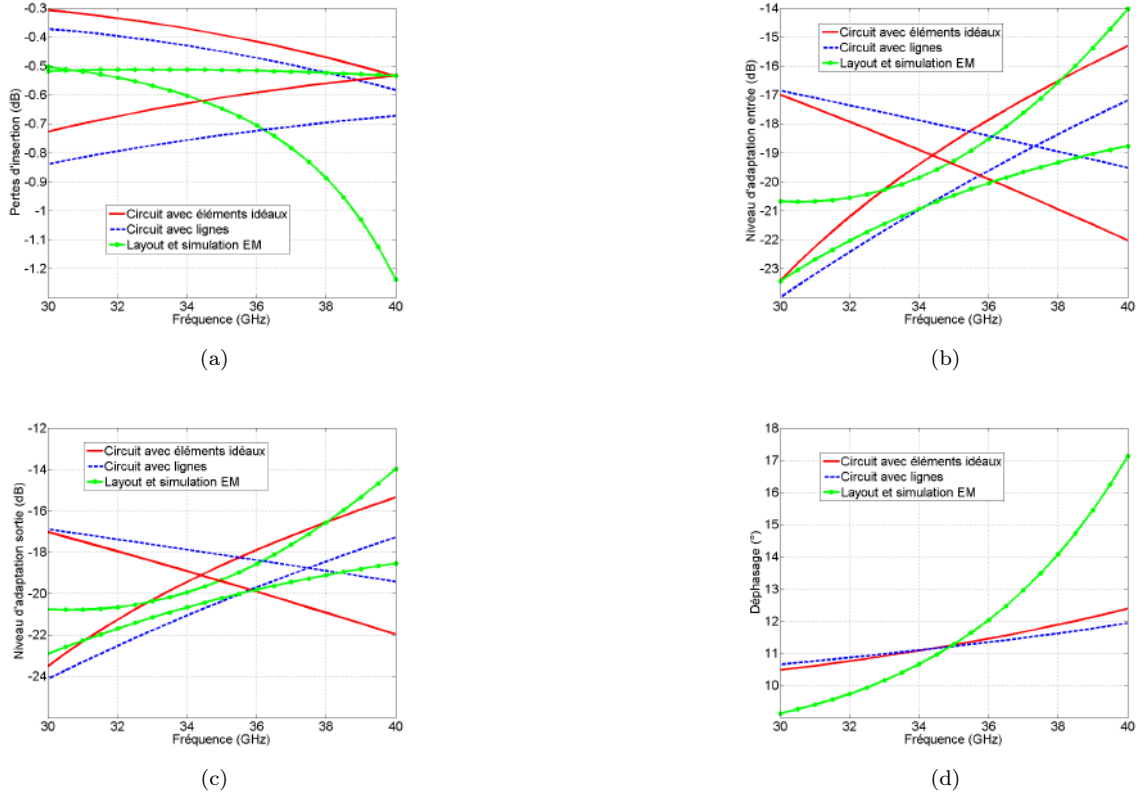


FIGURE 2.7 – Performances suivant les différentes itérations lors de la conception de la cellule 11.25° (a) pertes d'insertion (b) déphasage (c) adaptation d'entrée (d) adaptation de sortie

Durant les différentes itérations, les valeurs des éléments réactifs et la taille des transistors ont évoluées. Ces changements sont reportés dans le Tableau. 2.2. La simulation EM a été réalisée sur le dessin des masques présenté en Figure 2.8 dans les conditions de simulation explicitées plus haut. Un soin particulier a été accordé à la symétrisation de la partie filtre du circuit pour éviter des couplages parasites dégradant le déphasage.

De plus un *via-hole* de taille conséquente (100 μ m) placé au plus proche du transistor a été utilisé pour réduire l'inductance associée à cet élément. Le chemin potentiel de couplage entre les inductances L1 a été dimensionné de façon à profiter de l'inductance mutuelle sans perturber le passage des champs magnétiques. Les différents éléments nécessaires à la fonctionnalité du circuit ont été orientés de façon à limiter au maximum l'utilisation de lignes de transmission pour assurer la connectique.

Cette même procédure a été suivie pour la simulation des 11 cellules du *core-chip*. Les résultats individuels de chaque cellule seront présentés par la suite. Nous allons maintenant nous intéresser à la conception de cellules individuelles pour le mode *single-ended*.

Itérations	L1 (nH)	C1 (pF)	L2 (nH)	T1 (n * W μ m)	T2 (n*W μ m)
#1	0,76	0,5	1,5	1* 15	1*30
#2	0,72	1,5	0,634	1*15	1*39
#3	0,5	0,4	0,52	1*15	1*39
#4	0,52	0,98	0,545	1*25	1*70

TABLE 2.2 – Evolution de la valeur des composants de la cellule 11,25°

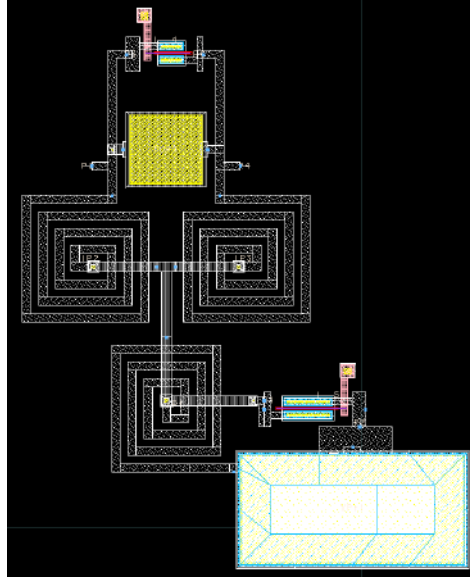


FIGURE 2.8 – Dessin des masques utilisé lors de la simulation EM de la cellule 11,25°

5.2 Résultats pour toutes les cellules single-ended

En suivant cette méthodologie de conception, nous avons réalisé les cellules individuelles du *core-chip single-ended*. Dans cette partie, nous présenterons les résultats de simulation et les subtilités de conception inhérentes à chaque cellule puis nous terminerons par présenter un tableau récapitulatif des performances à 35GHz. Nous rappelons que les objectifs en *single-ended* étaient d'obtenir une erreur fonctionnelle nulle (au premier ordre tout du moins, erreur de phase pour les déphaseurs et d'atténuation pour les atténuateurs).

Cellule 5.625°

Pour la réalisation de cette cellule, nous utilisons une ligne à retard. Cette topologie ne présente pas de difficulté particulière et n'a pas besoin de masse RF. Cependant, selon la connexion avec d'autres cellules il peut être nécessaire d'ajouter une masse DC pour garantir les polarisations du transistor. Nous reviendrons plus en détails dans la partie sur la mise en commun des cellules puisque c'est un problème récurrent de conception. Les performances de la cellule sont données en Figure 2.9.

Les performances sont très bonnes et en accord avec les objectifs fixés à 35GHz et cela pour les deux types d'erreurs fonctionnelles.

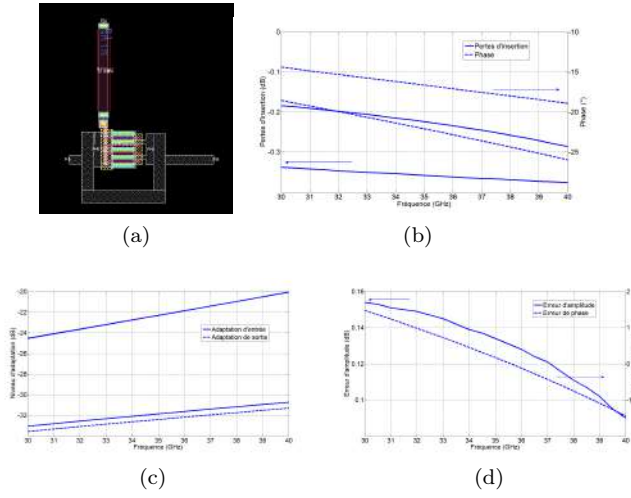


FIGURE 2.9 – Cellule 5.625° single-ended (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

Cellule 11.25°

Cette cellule utilise une topologie à filtre passe haut commuté avec deux capacités MIM en série et deux inductances spirale pour réaliser le filtre (Figure 2.6). Le circuit bouchon est obtenu par la mise en parallèle d'un transistor et d'une inductance spirale, le dessin des masques a déjà été donné en 2.8. Les performances de la cellule sont illustrées en 2.10.

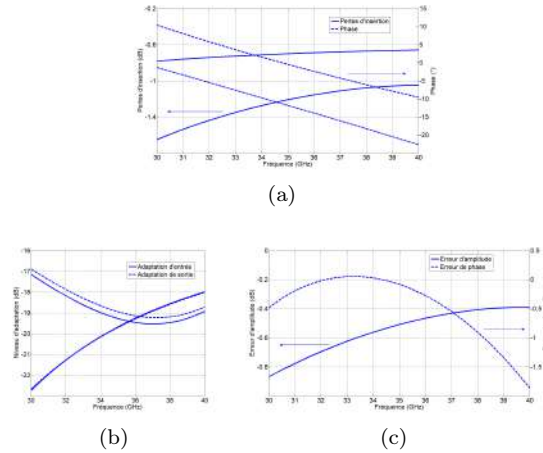


FIGURE 2.10 – Cellule 11.25° single-ended (a) pertes d'insertion et phase (b) niveaux d'adaptation et (c) erreurs fonctionnelles

Les performances sont bonnes, même si l'erreur d'amplitude aurait pu être réduite.

Cellule 22.5°

La topologie utilisée est exactement la même que la précédente, seules les valeurs changent. Les résultats de simulation sont affichés en Figure 2.11.

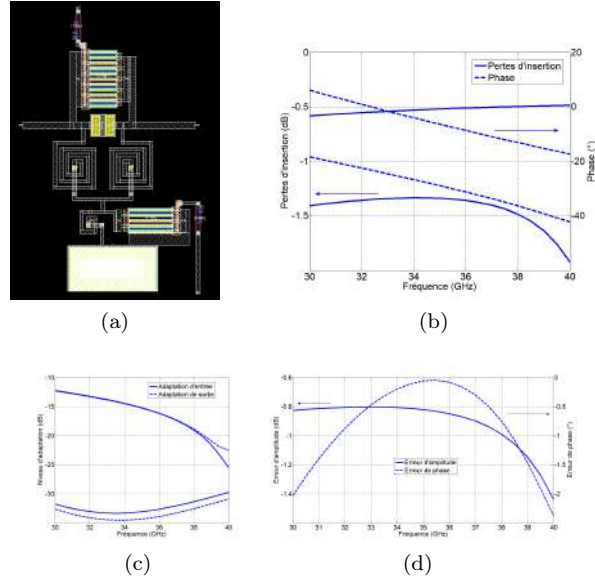


FIGURE 2.11 – Cellule 22.5° single-ended (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

Le niveau d'adaptation est à la limite du seuil des -15dB à 35GHz , aux fréquences plus basses, l'adaptation pourra devenir un problème lors de la mise en commun.

Cellule 45°

Pour cette cellule, nous avons utilisé la même topologie que les précédentes, en revanche nous avons utilisé des capacités SiO_2 pour la réalisation du filtre passe haut. Les performances de la cellule sont présentées sur la Figure 2.12.

Cette cellule présente de très bonnes performances en termes d'erreurs fonctionnelles à 35GHz . De plus l'adaptation est sous le seuil fixé par l'objectif et cela sur toute la bande.

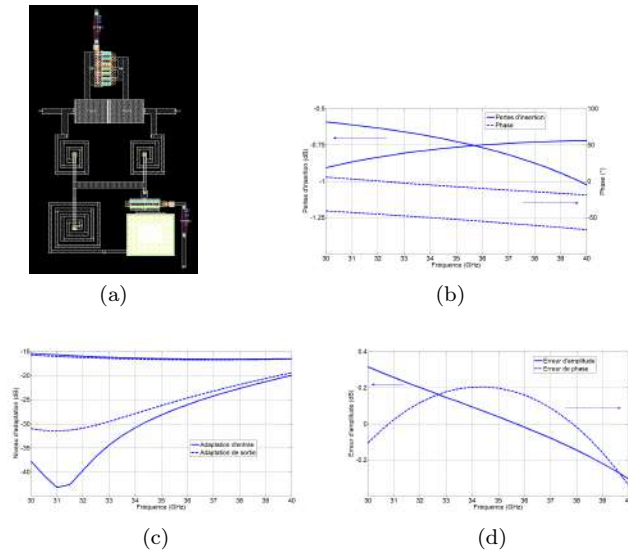


FIGURE 2.12 – Cellule 45° single-ended (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

Cellule 90°

La cellule a été réalisé avec la même topologie que les précédentes, les capacités utilisées pour la réalisation du filtre passe haut sont de type MIM. Les performances sont données en Figure 2.13.

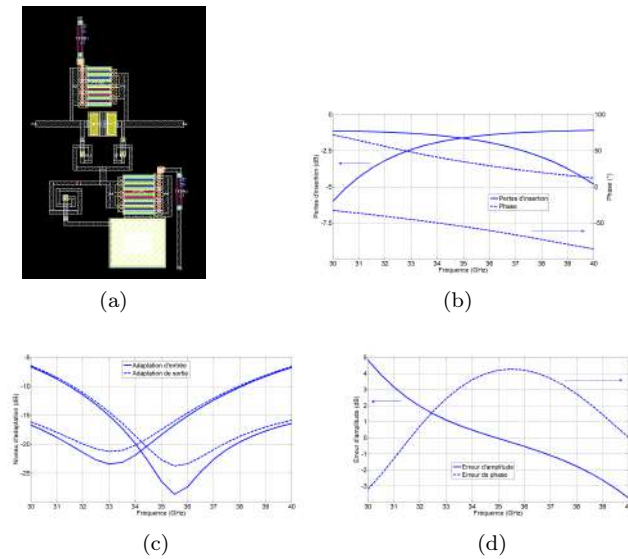


FIGURE 2.13 – Cellule 90° single-ended (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

Nous pouvons voir que les performances de cette cellule sont très centrées sur 35GHz avec une large variation sur la plage de fréquence. En revanche à la fréquence d'intérêt, les erreurs fonctionnelles et l'adaptation sont très bonnes.

Cellule 180°

Pour réaliser cette cellule, nous avons utilisé une topologie à base de filtres passe haut et passe bas. En utilisant la différence de phase de ces derniers, nous pouvons dépasser la limite théorique de déphasage de 90°. Nous avons utilisé des topologies en T pour réaliser les filtres car elles étaient spatialement plus faciles à réaliser. Cette topologie repose donc sur la mise en parallèle de deux filtres, ce qui augmente les pertes d'insertion de la structure. Le choix des transistors utilisés pour réaliser la fonction de commutation va fixer les performances en termes de pertes d'insertion et de niveau d'adaptation. Les performances sont illustrées en Figure 2.14.

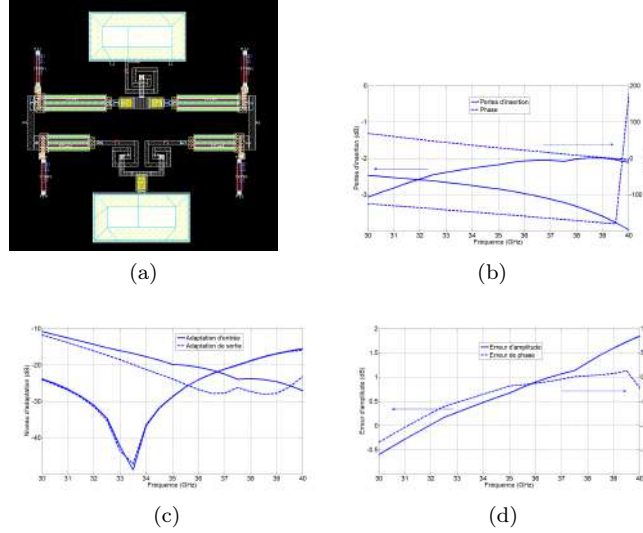


FIGURE 2.14 – Cellule 180° single-ended (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

Nous pouvons voir une inversion de la phase vers les hautes-fréquences mais ce n'est pas gênant pour l'interprétation des résultats. De plus, ce n'est pas dans notre fréquence d'intérêt. Les erreurs fonctionnelles méritent peut-être une amélioration, spécialement l'erreur de phase qui n'est pas nulle.

Cellule 0.5dB

Pour la réalisation de cette cellule nous avons utilisé la topologie à résistance shunt. Pour cela, la résistance utilisée est une résistance Nickel car nous avons besoin d'un contrôle relativement précis de la valeur de résistance pour maîtriser le niveau d'atténuation. Nous avons « empilé » trois transistors pour réaliser la commutation et augmenter le niveau d'isolation. Les performances de la cellule sont illustrées en Figure 2.15.

La simplicité de la structure assure de très bonnes performances, sur tous les paramètres observés.

Cellule 1dB

La topologie utilisée met en série deux résistances shunt pour réaliser la fonction désirée. La distance entre les deux résistances va fixer le niveau d'adaptation et la phase d'insertion. De plus, plus le niveau d'atténuation demandé est important plus les pertes d'insertion seront élevées. Les performances de la cellule sont données en Figure 2.16.

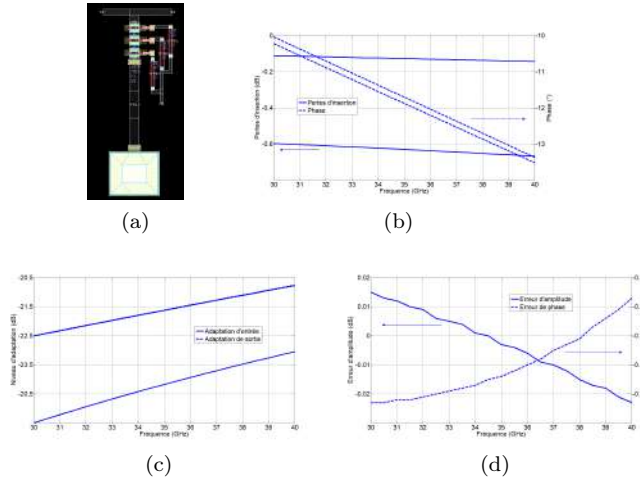


FIGURE 2.15 – Cellule 0.5dB single-ended (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

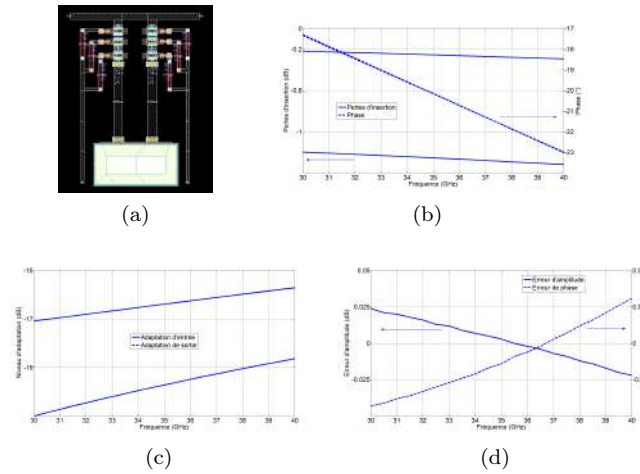


FIGURE 2.16 – Cellule 1dB single-ended (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

Comme nous pouvons le voir, malgré de très bonnes performances en termes d'erreurs fonctionnelles, le niveau d'adaptation se dégrade significativement. C'est pour cela que cette topologie n'est pas utilisée pour les autres atténuateurs. Nous ne pouvons pas utiliser de topologie à base de réseaux en π résistifs pour cette cellule car le niveau de perte d'insertion de ceux-ci est trop important. En effet, cette valeur d'atténuation exige de faire un compromis entre adaptation et pertes d'insertion.

Cellule 2dB

Pour cette cellule, un réseau résistif en π est utilisé pour réaliser l'atténuation. Pour la commutation, nous intégrons les transistors sur les deux voies du réseau. Nous mettons un transistor parallèle sur la voie série et des transistors série sur les voies shunt. Même si cette topologie présente des pertes d'insertion plus élevées, du fait de la résistance sur la voie série, elle permet d'atteindre des niveaux

d'atténuation importants. Les performances sont regroupées en Figure 2.17.

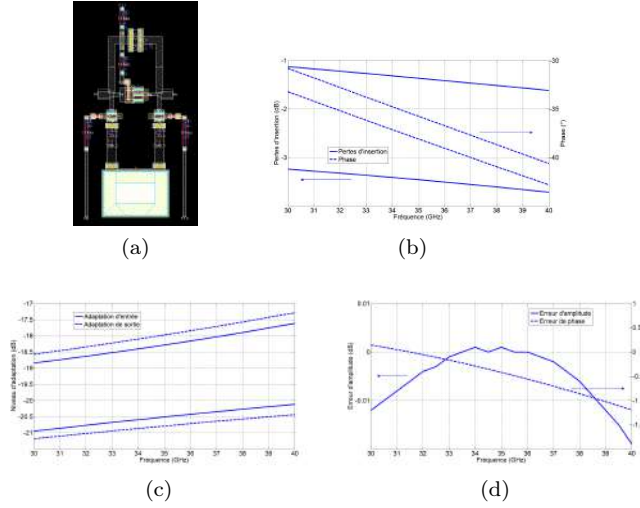


FIGURE 2.17 – Cellule 2dB single-ended (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

Comme nous l'avons prévu, les pertes d'insertion sont relativement élevées. Cependant les erreurs fonctionnelles et le niveau d'adaptation respectent les objectifs fixés à 35GHz. Une problématique inhérente à cette topologie est la gestion de la phase d'insertion générée entre les deux états. Pour compenser la différence de longueur électrique entre les deux états nous rajoutons des longueurs de lignes sur l'état en avance de phase. Mais cette manœuvre rajoute des pertes d'insertions.

Cellule 4dB

Nous utilisons la même topologie que la cellule précédente, néanmoins plus le niveau d'atténuation demandé est important plus la phase d'insertion générée augmente. La différence réside donc dans la longueur de ligne ajoutée pour pouvoir compenser cette phase. Les performances de la cellule sont illustrées en Figure 2.18.

Nous pouvons voir que par rapport à la cellule précédente, les pertes d'insertion sont moins élevées mais en contrepartie le niveau d'adaptation est légèrement dégradé. Cela met bien en évidence le compromis à réaliser pour ce genre de topologie, puisque le niveau de pertes d'insertion est directement lié au niveau d'atténuation. Les performances affichées par la cellule respectent le cahier des charges.

Cellule 8dB

La dernière cellule est basée sur le même principe que les précédentes, cependant nous utilisons deux résistances parallèles sur la voie série du réseau en π . De plus nous utilisons des inductances en parallèle avec chaque transistor pour améliorer l'isolation de ceux-ci à l'état bloqué. Cela dégrade un peu la compacité de la cellule, sans toutefois occuper beaucoup plus d'espace que les cellules précédentes. Les performances de la cellule sont présentées en Figure 2.19.

Nous pouvons voir que malgré l'ajout des inductances (0.68nH en parallèle du transistor de la voie de référence et 0.641nH sur les transistors shunt) le niveau d'adaptation est insuffisant par rapport à l'objectif. Les pertes d'insertions sont relativement élevées principalement à cause du niveau d'atté-

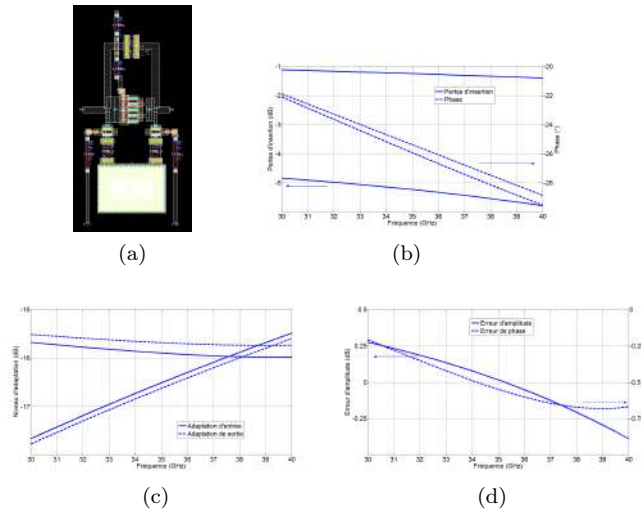


FIGURE 2.18 – Cellule 4dB single-ended (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

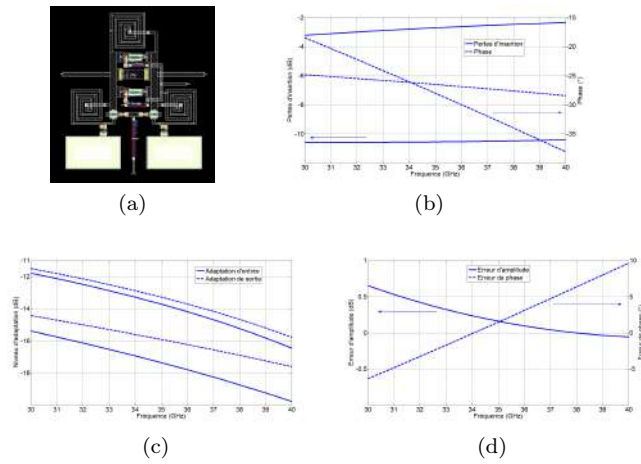


FIGURE 2.19 – Cellule 8dB single-ended (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

uation requis. Durant la conception de cette cellule, nous avons dû faire un compromis pour obtenir des erreurs fonctionnelles minimales tout en gardant une adaptation moyenne.

Pour conclure cette partie, le Tableau. 2.3 résume les performances de chaque cellule à 35GHz. Ces erreurs « individuelles » seront la limitation de performances du circuit lors de la mise en commun des cellules (dans l'hypothèse où il n'y aurait pas de compensation d'erreurs inter-cellules). Nous pouvons voir que les erreurs fonctionnelles ne sont pas nulles comme escomptées. Nous ne pouvons donc pas logiquement nous attendre à des erreurs système nulles une fois la mise en commun effectuée.

Cellule	Pertes d'insertion (dB)	Adaptation (dB)	Erreur de phase (°)	Erreur d'atténuation (dB)
5,625°	0,35	-22	0,15	0,135
11,25°	1,3	-19	0,05	0,5
22,5°	1,3	-15	0,02	0,82
45°	0,75	-16	0	0,02
90°	1,75	-17	0,2	0
180°	2,8	-20	1,5	0,7
0,5dB	0,15	-21,5	0,173	0
1dB	0,25	-16,5	0	0
2dB	1,3	-18,4	0,4	0
4dB	1,2	-15,8	0,55	0
8dB	2,6	-13,4	1	0,15

TABLE 2.3 – Résumé des performances des cellules du core-chip single ended à 35GHz

6 Conception des cellules différentielles

Comme pour la partie précédente nous ne détaillerons pas les étapes de simulation pour toutes les cellules. Cependant nous étudierons une cellule en particulier, la cellule de déphasage de 180°, qui nous paraît la plus intéressante ici puisqu'elle tire profit de la nature même du mode différentiel.

6.1 Cellule 180° différentielle

Le déphasage de 180° entre les deux voies conductrices de ce mode est exploité pour réaliser la différence de phase de 180° souhaitée. Ainsi, la topologie plus complexe utilisée en *single-ended* est réduite à une commutation entre voies en mode différentiel, ce qui a pour résultat d'améliorer les performances.

La fonction réalisée est un *quad-switch* (une commutation entre deux entrées et deux sorties possibles) représentée sur la Figure 2.20. On alterne entre deux états, l'état de référence ou le signal transite sur les transistors directs (polarisation V_1), il n'y a pas de commutation de voie et l'état de déphasage où l'on passe par les transistors croisés (polarisation V_2) et les voies sont commutées. Les tensions V_1 et V_2 sont complémentaires et assurent donc que les signaux empruntent uniquement le chemin de référence ou le chemin commuté. Lors de la commutation de voies, la phase du signal dit positif (nous prendrons un déphasage de 0° par convention) voie sa phase commuter pour la phase du signal dit négatif (déphasage de 180°). Ainsi nous obtenons bien une différence relative de 180° entre les deux états.

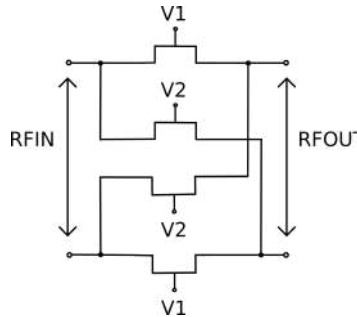


FIGURE 2.20 – Schéma électrique de la cellule 180° différentielle

Cette topologie est simple à réaliser puisqu'il n'y a qu'un seul élément influençant les performances de la cellule : le transistor. En premier lieu on cherche à réaliser un *quad-switch* présentant un niveau de pertes minimum. Dans cette situation, ce ne sont pas les résistances passantes R_{on} des transistors qui vont orienter notre choix mais plutôt les capacités C_{off} . En effet, dans cette configuration le niveau d'isolation d'un transistor est plus important que les pertes d'insertions qui lui sont associées, c'est à dire que la quantité de signal « perdu » par un transistor bloqué mal isolé est plus importante que celle d'un transistor passant aux pertes importantes. Ce raisonnement est valable pour les transistors de ce DK et peut différer selon la nature du ou des paramètres des transistors en commutation utilisés (pertes, isolation, puissance...). Un transistor mono-doigt est donc choisi, avec une longueur de grille moyenne (50 à 60 μm). Effectivement nous n'optons pas pour le transistor le plus petit (meilleure isolation) mais plutôt pour un mono doigt avec un R_{on} moyen. Les étapes suivantes consistent à rajouter les lignes assurant les connexions.

La criticité de la topologie réside dans le croisement des deux niveaux métalliques. A l'image d'un mélangeur, nous recherchons une symétrie électrique quasi parfaite. Ce chevauchement de niveaux métalliques n'a pas pu être simulé électriquement même s'il existe un modèle disponible dans le DK. Les performances de ce circuit de la simulation idéale à celle du dessin des masques complet sont reportées en Figure 2.21.

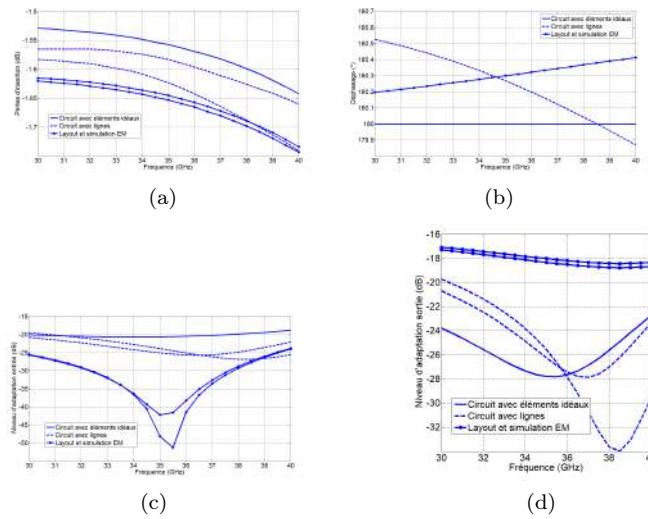


FIGURE 2.21 – Performances suivant les différentes itérations lors de la conception de la cellule 180° (a) pertes d'insertion (b) déphasage (c) adaptation d'entrée (d) adaptation de sortie

Classiquement ce type de cellules réalisées en *single-ended* présente des pertes d'insertions de 2,5 à 3dB. Ici nous pouvons donc bien voir le gain vis à vis des pertes d'insertion. Concernant le déphasage il est théoriquement toujours égal à 180° et toute variation par rapport à ce déphasage théorique vient du croisement entre les niveaux métalliques qu'il faut optimiser. Il en va de même pour l'adaptation qui est fixée par la taille des transistors utilisés mais aussi par la configuration des lignes en sortie. En effet nous assistons à une grosse différence entre l'adaptation d'entrée et de sortie. Nous pouvons l'expliquer en examinant la Figure 2.22 où les lignes d'accès d'entrée et de sortie sont de taille différente, influençant ainsi l'adaptation d'impédance de manière différente.

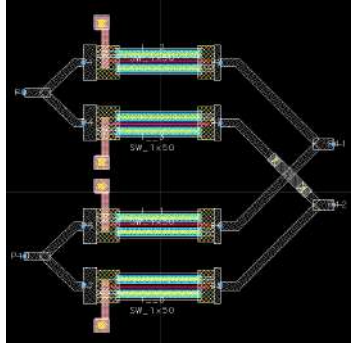


FIGURE 2.22 – Dessin des masques utilisé lors de la simulation EM de la cellule 180°

6.2 Résultats pour toutes les cellules différentielles

Nous avons appliqué notre méthodologie de conception lors des simulations de toutes les cellules différentielles individuelles. Nous allons maintenant présenter les résultats de simulation. Toutes les cellules ont été dessinées en essayant de garder le même écart entre les accès pour pouvoir les interconnecter. En conséquence, certaines lignes assurant la connexion inter-cellules ont été ajoutées et viennent dégrader la fonctionnalité de la cellule. Nous avons bien entendu cherché à réduire au maximum les longueurs de lignes utilisées. Nous n'avons pas jugé nécessaire de présenter chaque dessin des masques individuel de cellule puisque celles-ci seront présentées dans le dessin des masques final du *core-chip*.

Cellule 5.625°

Pour cette cellule nous avons dédoublé la topologie de la ligne à retard sur chaque voie différentielle. Cette cellule ne présente pas de difficulté et fonctionne exactement comme une cellule *single-ended*. Il faut uniquement veiller à ne pas trop rapprocher les deux voies, d'une part pour le couplage entre les deux voies et d'une autre part pour respecter la distance entre les deux voies différentielles d'entrée et de sortie qui doit être compatible avec les entrées/sortie des autres cellules pour éviter tout ajout de lignes supplémentaires. Les performances de cette cellule sont illustrées en Figure 2.23. Malgré un déphasage évoluant avec la fréquence, les performances sont suffisantes et respectent le cahier des charges.

Cellule 11.25°

Cette cellule est réalisée grâce à un filtre passe haut commuté. La topologie est similaire à celle présentée dans la partie *single-ended*. La différence vient du circuit bouchon utilisé, en effet nous n'utilisons qu'un seul circuit résonnant avec de part et d'autre les topologies à filtre. En mode différentiel une masse virtuelle vient se former quand le circuit bouchon ne résonne pas. Dans le cas où il résonne, un circuit ouvert isole les deux voies différentielles. Il n'y a donc pas de *via-holes* qui séparent les deux branches (qui ferment les accès shunt du filtre) contrairement au *single-ended*. Les condensateurs du filtre passe haut sont deux capacités MIM en série, pour symétriser le circuit. Les performances de la cellule sont reportées en Figure 2.25. A l'exception d'une adaptation de sortie de tout juste -15dB à 30GHz et des pertes d'insertion un peu élevées, la cellule respecte les exigences du cahier des charges. Des capacités ont été placées en entrée et sortie pour améliorer l'adaptation et limiter l'effet des lignes d'accès.

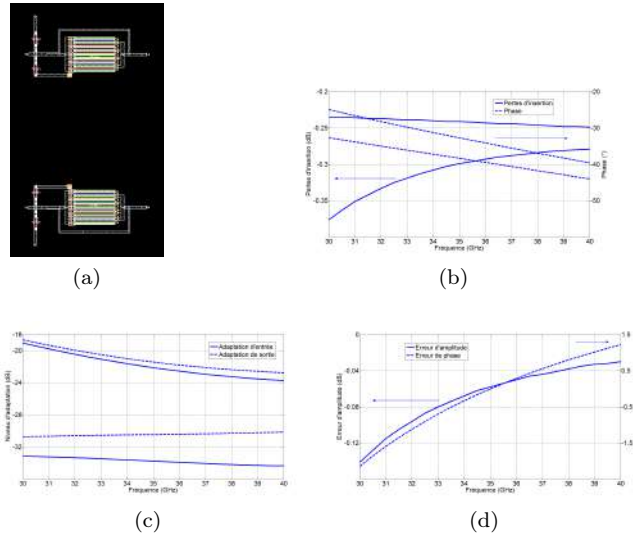


FIGURE 2.23 – Cellule 5.625° différentielle (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

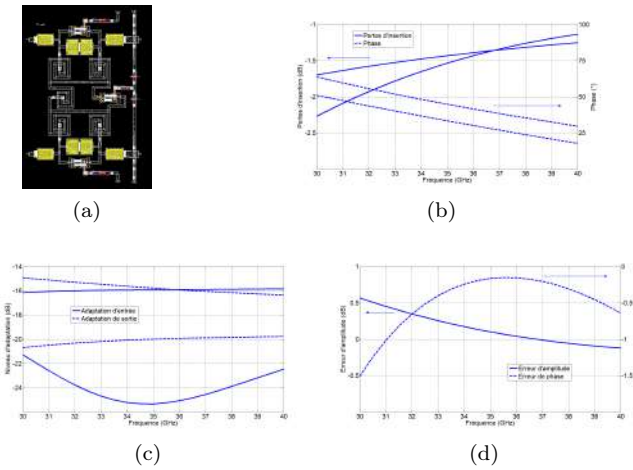


FIGURE 2.24 – Cellule 11.25° différentielle (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

Cellule 22.5°

Même topologie que la cellule précédente, cependant nous n'utilisons pas de capacité en parallèle avec le transistor pour réaliser le filtre passe haut. De plus le transistor du circuit bouchon ne nécessite pas une valeur d'inductance très importante, nous avons donc pu utiliser un petit bout de ligne pour le faire résonner. L'adaptation a aussi été améliorée grâce à des capacités en entrée et en sortie. Les performances de la cellule sont données en Figure 2.26. L'adaptation est insuffisante pour un des deux états, excepté cela les performances sont satisfaisantes.

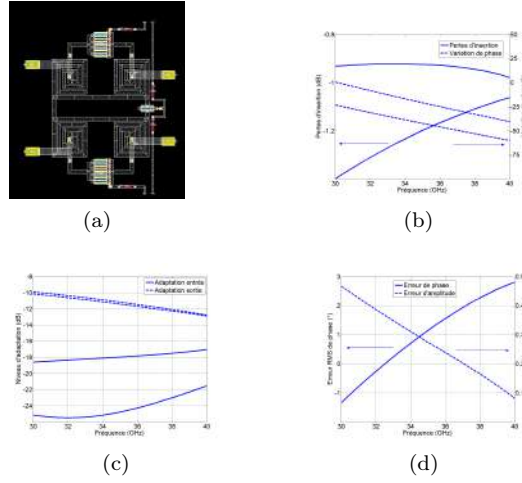


FIGURE 2.25 – Cellule 22.5° différentielle (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

Cellule 45°

Cette cellule utilise aussi la topologie à base de filtres commutés avec une adaptation entrée sortie améliorée par des condensateurs. La topologie est similaire à celle de la cellule 22.5° cependant cette fois ci une inductance spirale est nécessaire pour faire résonner le transistor central. A l'exception de l'adaptation trop faible pour un état, toutes les autres performances sont alignées sur le cahier des charges, elles sont représentées en Figure 2.26.

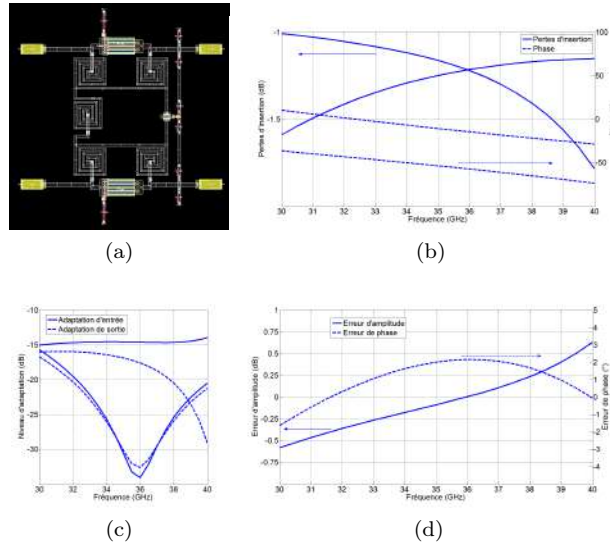


FIGURE 2.26 – Cellule 45° différentielle (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

Cellule 90°

Pour cette cellule, même topologie utilisée mais cette fois ci les capacités des filtres passe haut sont des condensateurs SiO_2 pour un contrôle plus précis de la faible capacité. Nous avons rencontré de nombreux problèmes lors de la conception de cette cellule, cela explique les piètres performances affichées en Figure 2.27.

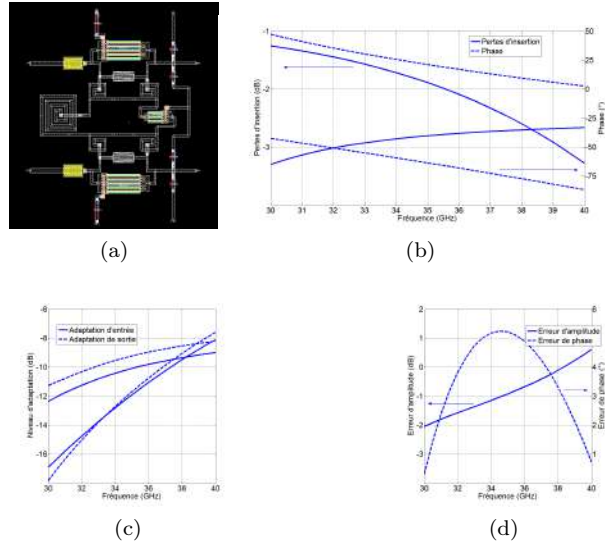


FIGURE 2.27 – Cellule 90° différentielle (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

Nous avons éprouvé des difficultés à concilier erreurs fonctionnelles et adaptation. Même si cette difficulté a été rencontrée avec les cellules précédentes, plus le déphasage est important plus la difficulté augmente. Ces difficultés apparaissent sous la forme d'erreurs fonctionnelles importantes ($>2\text{dB}$ et $>5^\circ$). Malgré tout, nous n'avons pas obtenu un niveau d'adaptation suffisant. Dans la majorité des cas, ces deux paramètres sont difficiles à obtenir simultanément. En effet, obtenir une erreur fonctionnelle basse implique souvent une dégradation de l'adaptation et vice versa. Pour cette cellule, la bande de fréquence importante a été un élément bloquant qui ne nous a pas permis de respecter le cahier des charges.

Cellule 180°

Comme nous l'avons expliqué plus tôt (Figure 2.22) cette cellule utilise le principe du *quad-switch* pour commuter entre les deux voies différentielles. Les résultats présentés en Figure 2.28 sont légèrement différents de ceux présentés en Figure 2.21 car comme nous l'avons expliqué précédemment, nous avons modifié les accès de la cellule pour la rendre compatible avec les autres cellules. Comme nous pouvons nous y attendre, l'ajout de lignes ajoute des pertes et augmente l'erreur de phase malgré cela, les performances restent dans les exigences du cahier des charges.

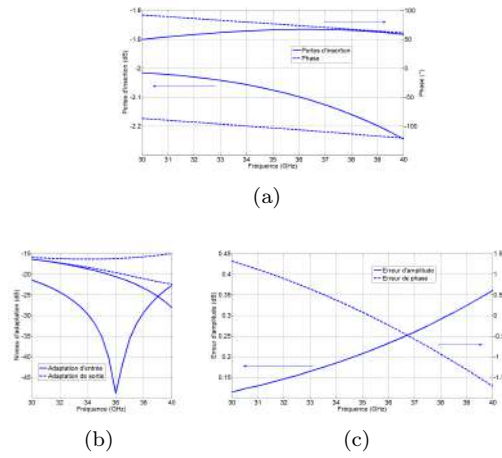


FIGURE 2.28 – Cellule 180° différentielle (a) pertes d'insertion et phase (b) niveaux d'adaptation et (c) erreurs fonctionnelles

Cellule 0.5dB

Pour cette cellule d'atténuation nous avons utilisé la topologie de résistance en shunt puisque l'atténuation nécessaire n'est pas très importante. Cette topologie ne présente pas de subtilité particulière car la phase d'insertion générée est faible et ne nécessite pas d'aménagement topologique précis. Les performances de la cellule sont reportées en Figure 2.29.

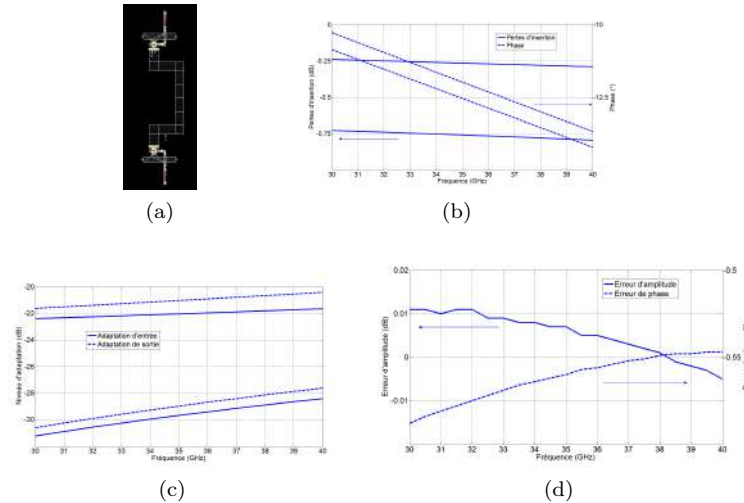


FIGURE 2.29 – Cellule 0.5dB différentielle (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

Cellule 1dB

Cette cellule est le résultat de la mise en série de deux topologies de la cellule précédente. La subtilité vient de la gestion de la phase d'insertion générée à cause de la distance entre les deux branches shunt du circuit. Les performances sont illustrées en Figure 2.30.

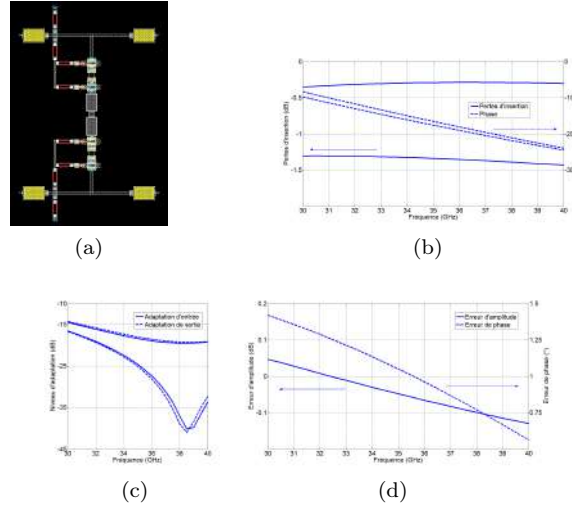


FIGURE 2.30 – Cellule 1dB différentielle (a) dessin des masques de la cellule et performances en (b) pertes d’insertion et phase (c) niveaux d’adaptation et (d) erreurs fonctionnelles

Comme nous pouvons le voir, l’erreur de phase augmente par rapport à la cellule précédente. Cette tendance va s’accroître plus l’atténuation demandée devient importante. La phase générée dépend du chemin emprunté par le signal, ce chemin différant de manière croissante avec l’augmentation de l’atténuation demandée. Sur cette cellule il fallait concilier entre l’adaptation et la phase d’insertion qui étaient déterminées par la distance entre les deux branches shunt du circuit. Cette distance est limitée par les dimensions des transistors (les polarisations spécifiquement), qui sont tous pilotés avec la même tension simultanément. La distance maximale est limitée par l’adaptation qui va se dégrader avec l’augmentation de la distance entre les branches shunt.

Cellule 2dB

Au-delà d’une atténuation de 1dB, la topologie à base de résistance shunt ne permet plus de concilier phase d’insertion et adaptation. C’est pourquoi nous avons utilisé une topologie avec réseau en π résistif. Cette topologie présente plus de pertes d’insertion mais permet une atténuation plate sur la bande tout en conservant une bonne adaptation. Les performances de la cellule sont illustrées en Figure 2.31. L’erreur de phase a été réduite en augmentant la longueur électrique de la voie de référence pour pouvoir « équilibrer » les distances parcourues par le signal entre les deux états. Nous obtenons ainsi des performances respectant le cahier des charges.

Cellule 4dB

Cette cellule est basée sur la même topologie que la précédente, la différence réside dans la compensation de phase d’insertion en rallongeant la longueur de la voie de référence. Les performances sont répertoriées en Figure 2.32. Malgré une erreur de phase réduite au moyen de la méthode explicitée plus tôt, le niveau d’adaptation est insuffisant vis-à-vis des exigences. Cependant c’est le meilleur compromis que nous avons trouvé pour respecter les exigences en termes d’erreurs fonctionnelles.

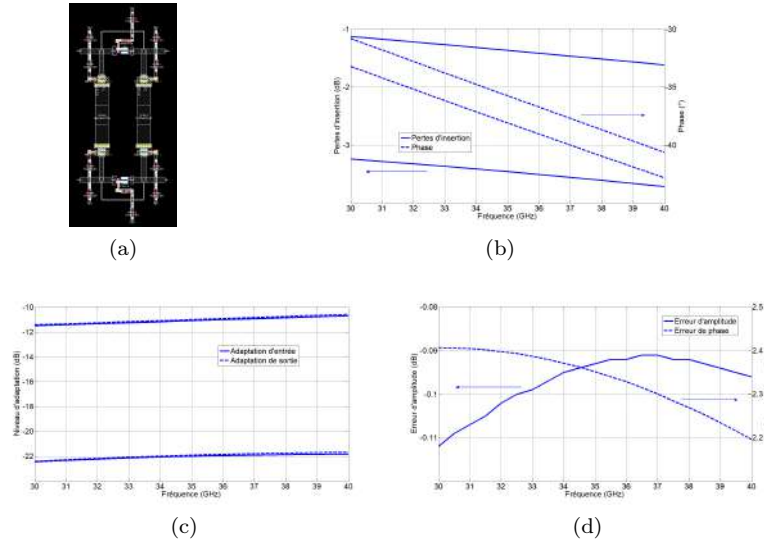


FIGURE 2.31 – Cellule 2dB différentielle (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

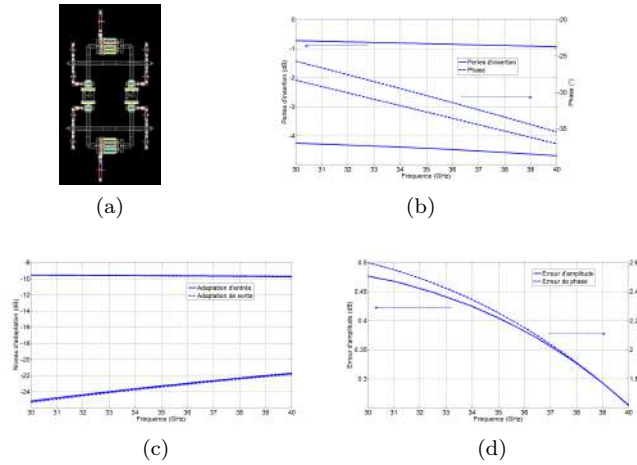


FIGURE 2.32 – Cellule 4dB différentielle (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

Cellule 8dB

Nous avons utilisé la même topologie que la cellule précédente et pour équilibrer la phase d'insertion nous avons ajouter plus de longueur de ligne autour de la résistance. Les performances de la cellule sont données en Figure 2.33. Les performances sont globalement correctes même si le niveau d'adaptation pour un des deux états est insuffisant.

Pour conclure ce chapitre, les performances des cellules différentielles sont reportées dans le Tableau 2.4. Nous pouvons nous rendre compte que les objectifs fixés initialement étaient ambitieux. En effet nous n'avons pas réussi à vérifier les objectifs en termes de niveaux d'adaptation et de pertes d'insertion principalement. En réduisant la bande, nous pourrions certainement nous rapprocher des

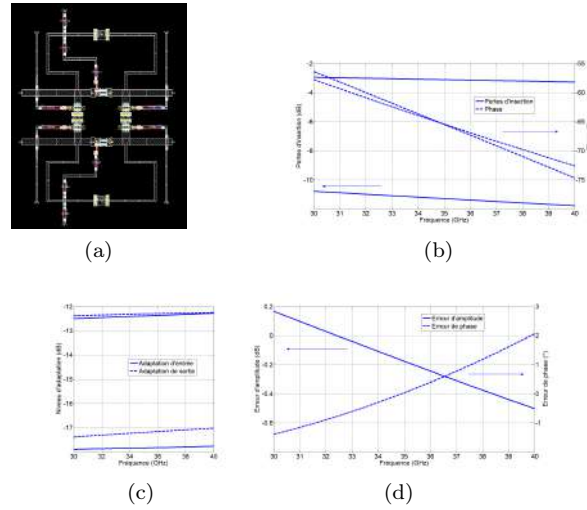


FIGURE 2.33 – Cellule 8dB différentielle (a) dessin des masques de la cellule et performances en (b) pertes d'insertion et phase (c) niveaux d'adaptation et (d) erreurs fonctionnelles

objectifs du cahier des charges initial. Le *core-chip* étant un système il faut maintenant regrouper ces cellules individuelles en un ensemble tout en cherchant à garder un fonctionnement optimal. En pratique des effets de charges entre les cellules influencent les performances du système de façon incontrôlée. En s'apercevant que l'étape de mise en commun des cellules n'était pas si anodine, nous avons développé un outil numérique permettant de trouver l'ordre optimal de cellule. Nous allons présenter ce processus dans le chapitre suivant qui traite des problématiques lors de la mise en commun des cellules.

Cellule	Pertes d'insertion (dB)	Adaptation (dB)	Erreur de phase (°)	Erreur d'atténuation (dB)
5,625°	0,37	-18	1,8	0,14
11,25°	2,25	-15	1,5	0,55
22,5°	1,4	-10	2,8	0,47
45°	1,75	-14	2,1	0,55
90°	3,25	-7,8	5,2	2
180°	2,24	-15	1,7	0,36
0,5dB	0,27	-20,2	0,58	0
1dB	0,35	-15	0,6	0,12
2dB	1,6	-10,5	2,4	0,11
4dB	1	-9,8	2,6	0,48
8dB	3,2	-12,2	2	0,5

TABLE 2.4 – Résumé des performances des cellules du core-chip différentiel à 35GHz

Les Tableaux 2.5 et 2.6 regroupent les meilleures performances respectives des cellules *single-ended* à 35 GHz et différentielles (sur la bande 30 à 40 GHz). Nous pouvons voir que les cellules différentielles présentent globalement des meilleures performances que les *single-ended* sauf si l'on observe les erreurs de phase qui ont l'air mieux optimisé pour les cellules Tableau 2.4 .

Cellule	Pertes d'insertion (dB)	Adaptation (dB)	Erreur de phase (°)	Erreur d'atténuation (dB)
5,625°	0,35	-22	0,15	0,135
11,25°	1,3	-19	0,05	0,5
22,5°	1,3	-15	0,02	0,82
45°	0,75	-16	0	0,02
90°	1,75	-17	0,2	0
180°	2,8	-20	1,5	0,7
0,5dB	0,15	-21,5	0,173	0
1dB	0,25	-16,5	0	0
2dB	1,3	-18,4	0,4	0
4dB	1,2	-15,8	0,55	0
8dB	2,6	-13,4	1	0,15

TABLE 2.5 – Performances des cellules single-ended à 35 GHz

Cellule	Pertes d'insertion (dB)	Adaptation (dB)	Erreur de phase (°)	Erreur d'atténuation (dB)
5,625°	0,22	-35	0	0,03
11,25°	1,1	-25	0,15	0
22,5°	0,95	-25	0,1	0
45°	1	-38	0	0
90°	1	-18	0,4	0
180°	1,8	-49	0	0,11
0,5dB	0,25	-33	0,55	0
1dB	0,3	-42	0,6	0
2dB	1,1	-23	2,2	0,09
4dB	0,8	-25	1,6	0,25
8dB	3,2	-17,9	0	0

TABLE 2.6 – Meilleures performances des cellules différentielles sur la bande 30-40 GHz

Bibliographie

- [1] Ommic.
- [2] Theory of Operation for Momentum - ADS 2011 - Keysight Knowledge Center.

Chapitre 3

Mise en commun et optimisation de l'ordre des cellules

Ce chapitre aborde la mise en commun des cellules individuelles étudiées dans le chapitre 2. Comme nous l'avons expliqué dans les chapitres précédents, cette étape cruciale consiste à regrouper les cellules tout en essayant de les faire fonctionner au plus proche de leurs performances optimales. Typiquement, une intégration idéale des cellules devrait aboutir à un système présentant des erreurs fonctionnelles système égales (ou au moins fortement ressemblantes) aux erreurs fonctionnelles individuelles. Or il est connu que des phénomènes de charge (désadaptation inter-étages) vont venir entraver cette intégration idéale par le biais des flux d'ondes incidentes et réfléchies (modification des performances liées aux altérations de charges vues par chacune de ces cellules). Durant ce chapitre nous allons présenter les modalités de mise en commun des cellules pour former les *core-chips single-ended* et différentiel. Notre approche nous a permis d'appréhender quatre types de situations différentes ; deux pour chaque mode de fonctionnement. En effet dans les deux cas, notre étude s'est divisée en deux axes aboutissant chacun à la fabrication d'un circuit MMIC. Les deux axes se différencient par l'utilisation ou non de notre algorithme d'optimisation d'ordre des cellules évoqué au chapitre 1. Avant de présenter en détails ces quatre conceptions, nous allons analyser les différences qui distinguent l'assemblage des cellules avec ou sans algorithme. Ensuite, nous analyserons les différentes situations, tout d'abord au stade des simulations.

1 Choix de l'ordre des cellules

Pour trouver l'ordre optimal des cellules les concepteurs utilisent différentes méthodes et stratégies, généralement personnelles et basées sur le retour d'expérience. En ne considérant que le critère d'adaptation, nous pouvons supposer intuitivement que placer les cellules les moins bien adaptées entres celles présentant la meilleure adaptation est une bonne solution. En effet, dans cette logique les cellules les moins bien adaptées verront leurs plans de charges moins dégradés par la bonne adaptation des cellules les entourant. Cependant comme expliqué en [1] et également dans le chapitre 1, le niveau d'adaptation ne peut pas être considéré comme le seul paramètre influençant la mise en commun des cellules. En effet, les cellules produisant un déphasage/atténuation important présentent généralement plus de pertes d'insertion et un niveau d'adaptation moindre. Cependant, l'impact négatif d'une adaptation dégradée sera en partie atténué par l'effet des pertes d'insertion dans le trajet réfléchi. Cet effet reste

toutefois marginal pour des cellules « optimisées » en pertes d'insertion (et erreurs de phase). Néanmoins, une combinaison de cellules désadaptées, mais conjuguées l'une sur l'autre pourrait produire une amélioration de leurs performances. Il est ainsi difficile d'émettre un jugement d'agencement de cellules a priori, en nous orientant vers un ordre particulier uniquement selon les critères d'adaptation.

Comme l'indique la Figure 3.1, l'ajout de phase au signal va faire déplacer le coefficient de réflexion sur un cercle à TOS constant. En conséquence nous pouvons voir des variations allant jusqu'à 6dB sur le coefficient de réflexion d'entrée se traduisant par des variations de plus de 0.1dB sur la transmission directe. Ce cas a été envisagé pour des cellules se trouvant dans le même état de fonctionnement (les deux en position référence). Nous pouvons donc imaginer les variations lorsque la totalité des cellules sont cascadiées et qu'elles se trouvent dans des états de fonctionnement différents.

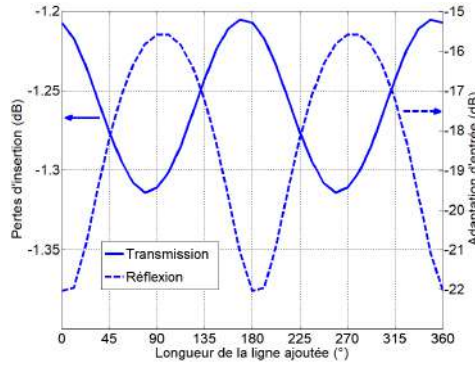


FIGURE 3.1 – Impact de l'ajout d'une ligne idéale de longueur variable (0 à 360°) sur les coefficients de transmission et de réflexion d'entrée entre deux cellules de déphasage

En raison des considérations précédentes, la prise en compte des critères d'erreurs de phase, d'amplitude et de désadaptation, et ce sur une plage de fréquence 30-40 GHz complexifie le choix d'un agencement optimal en design traditionnel (i.e. sans support quantitatif de l'amélioration ou de la dégradation des performances du système). En effet, pour un *core-chip* 11 bits, le nombre très élevé de combinaisons d'agencements possibles ($11! = 39.916.800$) est rédhibitoire. Nous proposerons dans ce travail un algorithme qui nous permettra d'émuler un nombre probant de combinaisons, en distinguant toutefois les modules déphaseurs ($6! = 720$ agencements possibles) et atténuateur ($5! = 120$ agencements).

Pour pallier cette difficulté de choix d'agencement, nous avons donc développé un algorithme basé sur l'utilisation d'optimiseurs intégrés aux logiciels de CAO. Bien que performante, cette solution est extrêmement couteuse en temps et va aussi dépendre de la méthode d'optimisation utilisée. Comme évoqué précédemment, le nombre d'agencements possibles étant de $n!$, avec n le nombre de cellules individuelles, les quantités de calculs exigées vont très vite augmenter. De plus, le risque de tomber sur un optimum local est extrêmement fort face à un aussi grand nombre d'agencements possibles et selon l'ordre de cellules de départ choisi.

Les cellules étant des éléments à deux états, il faut aussi envisager des différences d'adaptation au sein d'une même cellule. Concrètement, lors de la mise en cascade de deux cellules, il faut prendre en compte 8 états différents comme illustré dans le Tableau 3.1. Le nombre de cas à traiter passe donc de $n!$ à $n! \cdot 2n$, ce qui augmente grandement le temps de simulation pour un optimiseur.

Certains concepteurs choisissent un ordre qu'ils jugent optimal en explorant quelques agencements de façon non exhaustive. De cette façon ils gagnent du temps mais ils ne sont pas certains d'obtenir le meilleur agencement possible comme nous le verrons par la suite de ce travail : en admettant que les combinaisons offrent 10% de cas très favorables, 10% de cas défavorables et 80% de résultats

Position 1	Position 2
C1 ON	C2 ON
C1 ON	C2 OFF
C1 OFF	C2 ON
C1 OFF	C2 OFF
C2 ON	C1 ON
C2 ON	C1 OFF
C2 OFF	C1 OFF
C2 OFF	C1 ON

TABLE 3.1 – Tableau schématisant tous les états possibles lors de la mise en commun de deux cellules C1 et C2, ON/OFF spécifie l'activation ou non de la cellule

intermédiaires, il faudrait générer entre 50 et 100 simulations pour distinguer statistiquement 5 à 10 combinaisons favorables qui pourraient être interprétées (analyse croisée des ordres des cellules pour éviter les anomalies identifiées précédemment). Néanmoins cette « chance » d'organisation reposera plus sur l'expérience du concepteur que sur des faits statistiques. La Figure 3.2 illustre un exemple de différences de constellations d'état de phase/atténuation pour un agencement de cellules entraînant des performances qualifiées de typiques (agencement permettant des performances satisfaisantes) face au cas où l'agencement de cellules engendre les plus mauvaises performances possibles. Nous pouvons voir les différences en termes d'alignement des niveaux de phase et d'atténuation, ainsi que les zones non couvertes qui sont bien plus nombreuses pour l'agencement pire cas. Les zones non couvertes présentes pour la constellation pire cas sont sensiblement préjudiciables aux performances du système ; en effet l'absence de couverture de certaines zones va grandement limiter les constellations d'état possibles, et ce car la résolution du système opérationnel est directement liée à la taille de la zone non couverte la plus importante (atténuation et phase). Cela signifie également que les bits de poids faibles sont globalement inefficaces et pénalisent de fait la conception finale (performances et taille). C'est pourquoi nous préférons tirer pleinement profit des performances des cellules individuelles en s'assurant d'être dans le meilleur agencement possible et d'ainsi pouvoir augmenter le nombre d'états utilisés exploitables.

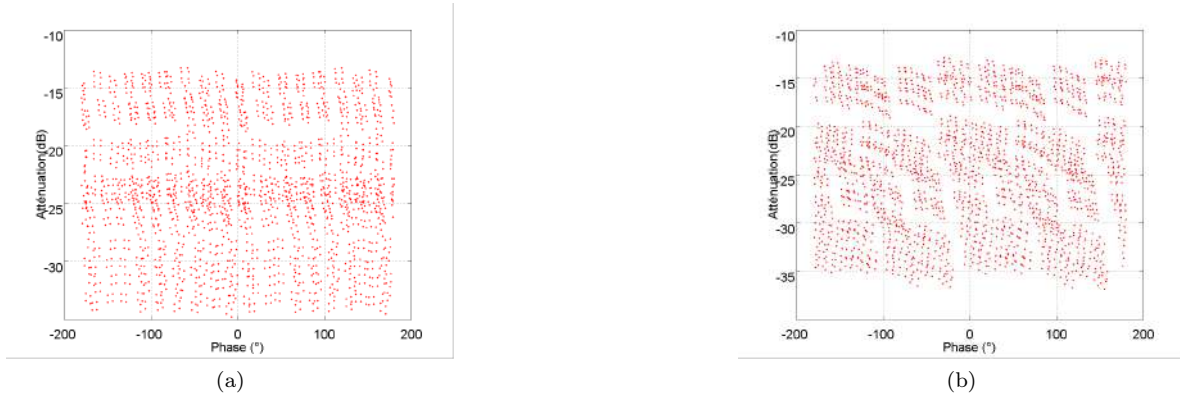


FIGURE 3.2 – Exemple d'impact de l'agencement de cellules aux performances (a) typiques et (b) pires cas à 35 GHz sur la constellation phase/amplitude (simulations électriques)

Une alternative serait de concevoir les cellules non pas selon un plan d'adaptation fixe de 50Ω mais en fonction des cellules adjacentes. De cette façon il ne serait plus nécessaire de chercher l'ordre optimal

de cellule. Cependant cette solution ajoute une nouvelle contrainte de conception puisque les cellules seraient conçues de manière itérative. Ainsi nous pourrions difficilement analyser leurs performances de façon individuelle et seule la simulation système pourrait valider les choix de conception.

2 Méthodologie de mise en commun des cellules développée

Notre méthodologie d'intégration a été développée à la suite d'une difficulté à remplir les exigences du premier cahier des charges large bande (*core-chip* différentiel). En effet, nous nous sommes rendu compte a posteriori de notre première conception du fort compromis généré par les critères retenus de bande de fréquence, d'erreurs fonctionnelles et d'adaptation, et ce au niveau des résultats de simulation du système *core-chip*. Dans le but de nous affranchir au maximum de ce compromis, et d'obtenir des erreurs fonctionnelles minimales, nous avons décidé de réaliser un *core-chip single-ended* aux performances centrées à 35GHz (erreurs fonctionnelles minimales –idéalement nulles– à cette fréquence). Cependant, même avec des performances centrées à 35GHz nous avons retrouvé des erreurs fonctionnelles sans pouvoir les attribuer à une source particulière (erreurs croisées évoquées en chapitre 1).

Nous avons donc décidé d'implémenter un algorithme d'ordonnement nous permettant d'être sûr de travailler avec l'agencement de cellules optimal. Nous avons donc pu utiliser l'algorithme selon les deux approches différentes (large bande 30-40 GHz et conception centrée 35 GHz), même si les résultats escomptés ne sont pas les mêmes du fait des objectifs de performances différents. En effet, en cherchant à appliquer celui-ci à des fréquences différentes (33-36 GHz ou 32-35 GHz par exemple) nous avons cherché à faire émerger des ordres de cellules qui privilégiaient non pas une seule fréquence mais plutôt la bande entière, quitte à relaxer les exigences sur les performances. Cette technique est donc quelque peu différente selon si les cellules individuelles ont été optimisées à 35 GHz ou pour 30-40 GHz.

Pour ce dernier cas, la marge de manœuvre amenée par l'algorithme d'organisation des cellules est peu importante, puisque s'il n'existe pas d'agencement permettant un gain sur la bande de 30 à 40 GHz, l'optimisation s'arrête. En revanche dans le cas où les cellules sont optimisées à 35 GHz, nous gardons la liberté de l'élargissement de la bande en fixant la dégradation de performances acceptables garantissant un fonctionnement plus large bande. Finalement, ce sera donc l'agencement présentant le meilleur gain de performances sur la bande par rapport aux performances à 35 GHz qui sera choisi. De plus grâce à notre méthode de retouche des cellules décrite dans le chapitre 1, il est possible (seulement pour les cellules centrées à 35 GHz) d'effectuer des modifications de cellules individuelles de façon à améliorer l'erreur fonctionnelle système ou bien d'élargir la bande de fonctionnement. Il est important de garder à l'esprit les différentes étapes de la mise en commun des cellules, celles-ci sont d'ailleurs assez similaires aux étapes présentées dans la méthodologie de conception du chapitre 1. Pour simuler les cellules dans une approche système, nous procédons tel que suit :

1- chaque cellule est représentée par une boîte de paramètre S. Ces paramètres S (*single ended* ou différentiels) sont issus de la dernière itération de la méthodologie de conception, à savoir la simulation EM globale de la cellule individuelle.

2- dans un premier temps les cellules sont regroupées et reliées par un simple fil dans la fenêtre schematic d'ADS. C'est une mise en commun électrique idéale.

3- Une fois qu'un ordre satisfaisant est obtenu lors de l'itération précédente, nous commençons la mise en commun directement sur le dessin des masques des masques comportant la totalité des cellules. Cette étape a pour vocation de trouver si certaines longueurs de lignes sont nécessaires à assurer la connexion critique/optimales entre cellules. Si c'est le cas, les lignes seront rajoutées sur la simulation électrique d'ADS pour anticiper l'ajout des lignes dans le dessin des masques des masques final.

4- la dernière étape consiste en la simulation EM du dessin des masques des masques final du *core-chip*. En comparant avec les simulations de l'étape précédente nous pouvons voir s'il existe des couplages inter-cellules non pris en compte dans la simulation électrique sur le *schematic* ADS (étape 2). Cependant ce constat est à relativiser puisque les ports de connexion jouent eux-mêmes un rôle non négligeable dans les résultats de simulation (idéalisations des conditions de fermeture des cellules). En effet nous avons testé l'effet des ports sur les simulations. Pour cela nous avons comparé les paramètres S de transmission et réflexion entre 30 et 40 GHz d'une ligne de 100 μm de longueur et de 25 μm de largeur avec ceux de deux lignes de 50 μm et de 25 μm de largeur, connectées par des ports. La Figure 3.3 illustre les résultats de tests. Même si la différence de niveau d'amplitude est négligeable, le décalage en phase (jusqu'à 1.4° à 40 GHz) est important. Nous pouvons donc imaginer l'impact de l'ajout de lignes et même l'effet de la connexion des cellules par des ports (simulation électrique) à la fréquence d'intérêt (la périodicité spatiale étant de l'ordre de la centaine de microns).



FIGURE 3.3 – Simulation pour apprécier l'effet des ports de connexion sur les résultats de simulation

2.1 Méthodologie sans algorithme

Lors de la conception de nos deux designs sans utilisation de l'algorithme, nous avons adopté la méthode qui suit. Nous sommes partis de l'ordre binaire des cellules puis nous avons effectué des permutations de cellules jusqu'à obtenir un résultat satisfaisant sur les erreurs fonctionnelles. Après réflexion et rétro simulation par algorithme, cette méthode ne semble pas la plus performante pour identifier un agencement qui optimise les performances (idéalement qui répond au cahier des charges). D'une part car elle a une forte tendance à tendre vers un optimum local fixé par le choix de l'ordre de départ. Et d'autre part car le nombre d'itérations est limité puisqu'il est réalisé manuellement. C'est pourquoi, en suivant cette méthode, nous ne pouvons jamais être certains d'obtenir le meilleur agencement de cellules possibles. Ce sont ces limitations qui nous ont poussé à développer l'algorithme d'ordonnancement des cellules. Dans le cas où l'utilisateur ne veut pas utiliser d'algorithme, nous avons pensé à une approche moins arbitraire qui nous permet d'atteindre plus vite les agencements présentant de bonnes performances (même si l'obtention du meilleur ordre n'est pas assuré). Nous allons donc expliciter cette seconde méthode, tout en précisant que nous ne l'avons pas mise en œuvre durant les deux différentes conceptions sans algorithme.

Pour initier la recherche non exhaustive (sans algorithme) du meilleur agencement de cellules, il faut choisir un point de départ. Il est vrai que ce point de départ peut avoir une influence sur le point d'arrivée dans une organisation itérative, puisque nous ne balayons pas toutes les possibilités contrairement au cas par utilisation de l'algorithme : l'optimisation est « directionnelle ». Cependant en choisissant plusieurs points de départ différents nous nous assurons d'une variété de « tendances » explorées ce qui nous permet de réduire la causalité de l'organisation des cellules.

Nous avons donc choisi 3 principes d'organisation différents : le premier selon l'ordre binaire des cellules (avec deux configurations possibles, les déphaseurs ou les atténuateurs en première position), un second ordre selon l'adaptation les cellules les moins bien adaptées qui sont entourées des mieux adaptées, et enfin un troisième ordre selon les pertes où les cellules présentant le plus de pertes sont placées aux extrémités de la chaîne (entrée et sortie). Ces différents choix nous permettent aussi de vérifier les hypothèses avancées dans le paragraphe précédent. Pour ce faire nous avons déployé ces trois points de départ pour chacun des différents designs et nous avons comparé les résultats de simulations à ceux que nous avons trouvé durant le design des circuits fabriqués. La Figure 3.4 schématise notre approche pour évaluer l'apport de l'utilisation des 3 différents points de départ pour réaliser l'organisation manuelle des cellules en comparaison à celle que nous avons utilisée pour le circuit fabriqué.

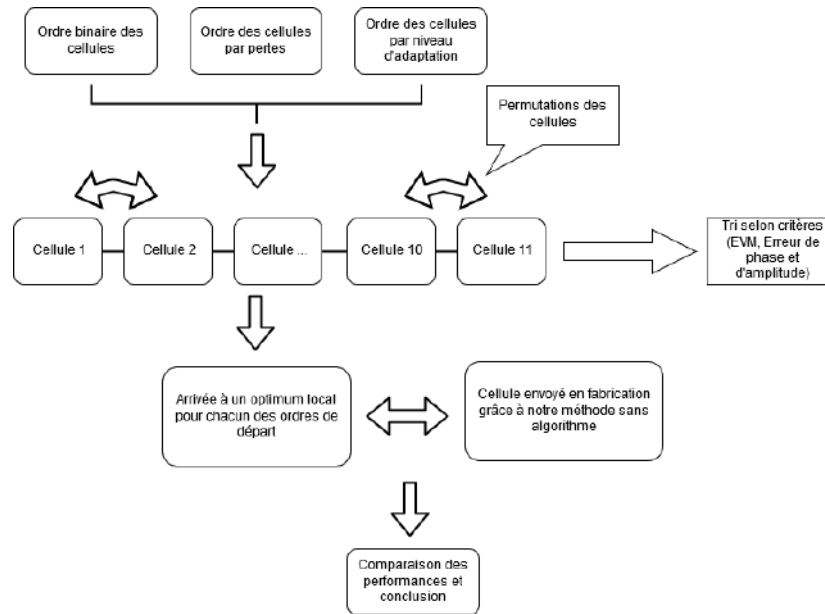


FIGURE 3.4 – Schéma de l'utilisation de notre méthode d'organisation manuelle des cellules

2.2 Méthodologie avec algorithme

Comme nous l'avons expliqué dans la partie précédente, le développement de l'algorithme (présenté au Chapitre 1) répond aux limitations de la méthodologie de choix non exhaustifs de l'ordre des cellules. En effet pour des *core-chips* à haute résolution (grand nombre de bits de contrôle) le nombre d'agencements possibles augmente très rapidement (en fonction factorielle). La probabilité de tomber sur l'ordre optimal de cellules « par hasard » est donc grandement réduite, ce qui a donc tendance à limiter le potentiel des performances atteignables. Pour s'assurer de trouver l'ordre optimal dans un intervalle de temps restreint, nous avons pensé à cet algorithme schématisé en Figure 3.5.

Au-delà de trouver l'ordre optimal, c'est aussi le temps gagné qui représente un gros avantage de cet outil. C'est aussi un outil d'analyse qui nous permet de jauger ce qui peut être amélioré au niveau système, tout comme il permet de pointer la criticité de certaines cellules ou agencement de cellules qui doivent donc être retravaillées en amont ! En effet, pouvoir très rapidement trouver le meilleur

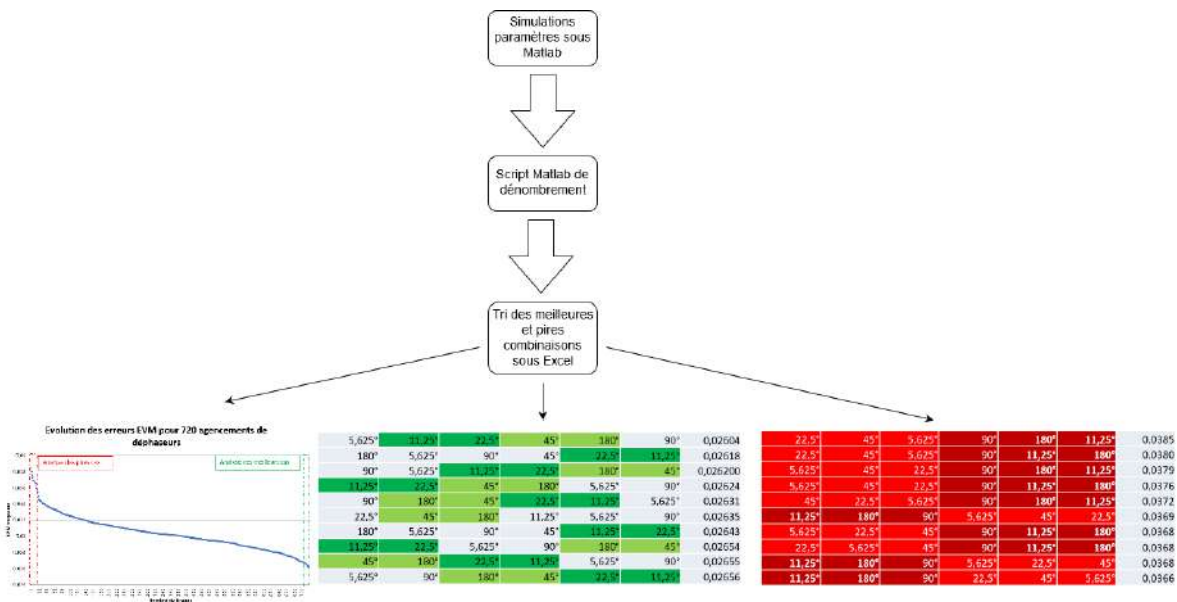


FIGURE 3.5 – Schématisation des étapes pour obtenir les meilleurs et pires agencements de cellules

ordre nous permet d'identifier les différents facteurs qui influencent la performance globale en jouant sur les performances individuelles des cellules : ainsi, l'idéalisation des paramètres d'adaptation ou de performances en transmission pour évaluer l'impact de la cellule réelle permet de repousser les performances initialement obtenues en corrigeant les facteurs limitants ou en améliorant les éléments bénéfiques aux cellules. La Figure 3.6 illustre la retouche de la cellule 180°. Nous pouvons y voir une diminution de l'erreur d'amplitude à 35 GHz.

3 Application pour les cellules single ended

3.1 Sans utilisation de l'algorithme

Nous allons tout d'abord présenter l'agencement que nous avons réalisé et envoyé en conception. Puis nous nous intéresserons aux 3 agencements de base : l'ordre binaire des cellules, l'ordre par pertes et l'ordre par adaptation. Nous allons donc présenter pour chacun de ces agencements initiaux les performances puis les améliorations possibles en déplaçant les cellules les unes par rapport aux autres. Les performances seront évaluées sur les erreurs RMS fonctionnelles et sur la constellation phase/atténuation. Bien que difficilement interprétables pour de légères variations de performances, celles-ci nous renseignent sur les pertes d'insertion, les zones non couvertes et les zones de recouvrement qui comme nous avons pu le voir plus tôt sont les éléments limitant un tel système.

Ordre choisi pour le premier design single-ended

Après avoir commencé par l'ordre binaire et effectué quelques permutations, nous nous sommes arrêtés sur l'ordre présenté en Figure 3.7. Cet ordre nous semblait présenter les meilleures performances, qui sont illustrées en Figure 3.8. Cette appréciation est personnelle au concepteur puisque nous aurions

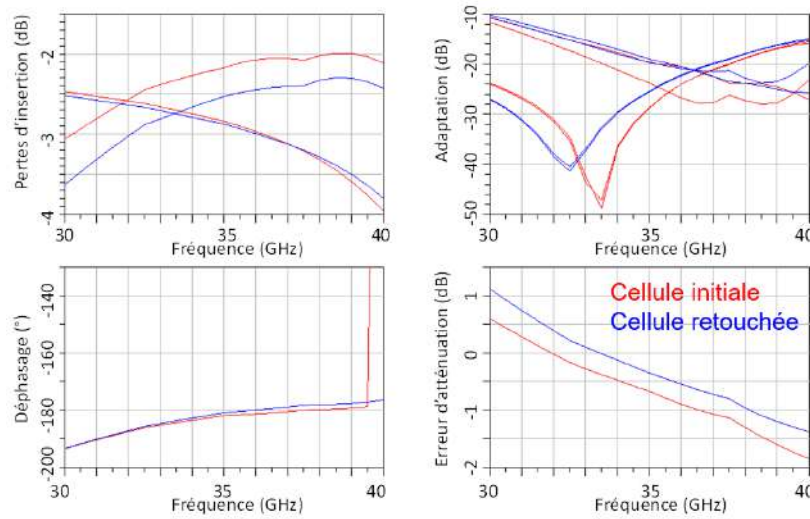


FIGURE 3.6 – Retouche de la cellule 180° single ended avec en rouge les performances initiales et en bleu les performances finales

pu continuer à chercher à améliorer les performances. Dans notre cas, en orientant notre recherche initiale en commençant par un ordre binaire, nous avons vraisemblablement atteint un optimum local avec cet ordre de cellules (à en juger par la solution plus efficace obtenue par l'exploitation de l'algorithme comme nous le verrons plus tard). Nous constatons aussi que les atténuateurs et déphaseurs ne sont pas inter-organisés, à l'instar des optimisations volontairement réalisées par algorithme pour les atténuateurs d'une part, les déphaseurs d'autre part. Cette observation peut laisser penser à une limitation de notre choix d'ordre, puisque soit nous ne tirons profit que de la moitié des éléments (il y a moins de combinaisons possibles en séparant les cellules de la sorte) soit c'est un hasard de performances qui aboutit sur un optimum local (ou global si nous sommes chanceux). Nous avons fait figurer les performances des cellules mise en commun de façon idéale et en conditions de dessin des masques réelle. De cette façon nous pouvons apprécier les variations de performances entre ces itérations de *design*.

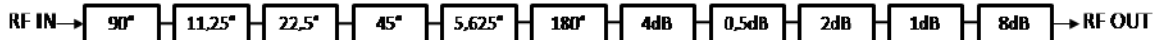


FIGURE 3.7 – Ordre final des cellules pour le core-chip single ended à 35GHz

L'analyse des figures respectives Figure 3.8a-Figure 3.8c et Figure 3.8b-Figure 3.8d montre une cohérence entre erreurs RMS et état de constellation. Nous pouvons voir que les performances sont globalement meilleures pour les simulations électriques que pour les simulations EM. Ce phénomène s'explique principalement par les couplages ayant lieu entre les cellules dans les simulations EM. En essayant de réduire les distances de lignes utilisées il se peut que des couplages indésirables viennent perturber le fonctionnement. Ces couplages n'étant pas prises en compte en simulation électrique. Une autre explication vient pourrait venir des ports de simulations d'entrée et de sortie des cellules. En connectant directement les cellules les unes aux autres sur le dessin des masques dans la simulation EM, les paramètres ne sont plus évalués aux entrées et sorties de chaque cellule mais à l'entrée et la sortie du système complet. En simulation électrique cependant, il existe toujours l'interface entre

chaque cellule, l'effet des ports de connexion est donc pris en compte.

A 35GHz, nous obtenons respectivement des erreurs RMS de phase de 3.9° et 5° et d'amplitude de 0.77dB et 2.07dB respectivement pour les simulations électriques et EM. Ces erreurs sont relativement bien traduites sur la constellation puisque la couverture est meilleure pour la simulation électrique. En effet malgré de légères disparités au niveau de l'amplitude, il y a globalement beaucoup moins de zones non couvertes et de recouvrement en simulation électrique qu'en simulation EM. Même si la dynamique en amplitude est similaire, la simulation EM présente 1 à 2dB de plus de pertes d'insertion. Concernant la largeur de bande, les performances en simulation électrique sont globalement meilleures que celles présentées en simulation EM.

Lors de l'analyse des constellations, les zones non couvertes et de recouvrement n'ont pas le même critère pénalisant. En effet dans le cas d'une zone de recouvrement, il existe plusieurs points correspondant à un seul point de consigne. Cette redondance, bien que pénalisante, peut être compensée électroniquement en changeant la loi de commande. En revanche, pour une zone non couverte, il est impossible d'adresser les zones « blanches » puisque même grâce à des opérations de translation de points, il y aura toujours une zone laissée déserte. Ainsi, même si ces deux problématiques vont toutes les deux impacter la résolution du système, de part ses redondances, le problème des zones de recouvrement peut être corrigé et ainsi réduire son influence sur la résolution. Généralement, les zones de recouvrement plus importantes sont accompagnées de zones non couvertes également importantes en raison d'un glissement des lois de certains bits. Tant que ces erreurs cumulées restent inférieures à la moitié du bit de poids faible (atténuation ou phase), aucun impact n'est ressenti sur le résultat final. Tandis que lorsque les erreurs cumulées excèdent cette demi-valeur générée par le bit de poids faible, il y a perte d'information. C'est le cas en simulation EM de la Figure 8.d autour d'une atténuation de 21 dB et également autour de la phase 70° notamment. Nous pouvons voir sur les Figure 8.e et Figure 8.f que les constellations reflètent bien les valeurs des erreurs RMS fonctionnelles à ces fréquences. En effet les constellations ne sont plus du tout homogènes et présentent de nombreuses zones non couvertes de tailles importante et des décalages de niveaux d'amplitude conséquents.

La corrélation entre erreurs RMS fonctionnelles est visible en comparant les deux constellations de bord de bande. Une erreur RMS de phase importante (à 30 GHz pour la simulation EM par exemple) se traduira par des zones non couvertes plus importantes que celles apparaissant à 40GHz. Nous pouvons faire le même constat pour l'erreur RMS d'amplitude, qui cette fois ci fixera la hauteur des ondulations en amplitude. Ce constat peut être fait sur tous les courbes suivantes, ce qui confirme la pertinence du critère d'erreurs RMS fonctionnelles.

En premier lieu, il est important de noter que les deux types de simulations donnent des résultats disparates ; il est difficile de donner une confiance forte à la conception dans son ensemble. Même si nous aurions tendance à privilégier la simulation EM qui tient compte notamment des phénomènes de couplage, nous savons que différents logiciels EM aboutissent à une diversité sensible des résultats sur des cas d'analyse non simplistes : il est fort probable que ce type de simulation ne soit pas aussi fiable qu'escompté.

Finalement, nous nous retrouvons quand même avec une erreur non nulle à 35 GHz et ce malgré le fait que les cellules soient optimisées, et conçues pour présenter une erreur nulle à cette fréquence. Nous attribuons ces erreurs au non-respect du cahier des charges concernant les erreurs fonctionnelles croisées qui sont venues se sommer aléatoirement pour rendre les erreurs fonctionnelles système non nulles à 35GHz. Cette limitation a donc nourri notre réflexion concernant d'une part une meilleure optimisation de l'ordre des cellules de façon manuelle et d'autre part pour l'élaboration d'un algorithme permettant de trouver le meilleur agencement possible et d'envisager d'éventuelles retouches de cellules pour corriger certaines erreurs fonctionnelles réfractaires. Dans la suite, nous présentons les différentes méthodes d'optimisation manuelle de l'ordre des cellules qui nous apparaissent mieux adaptées pour

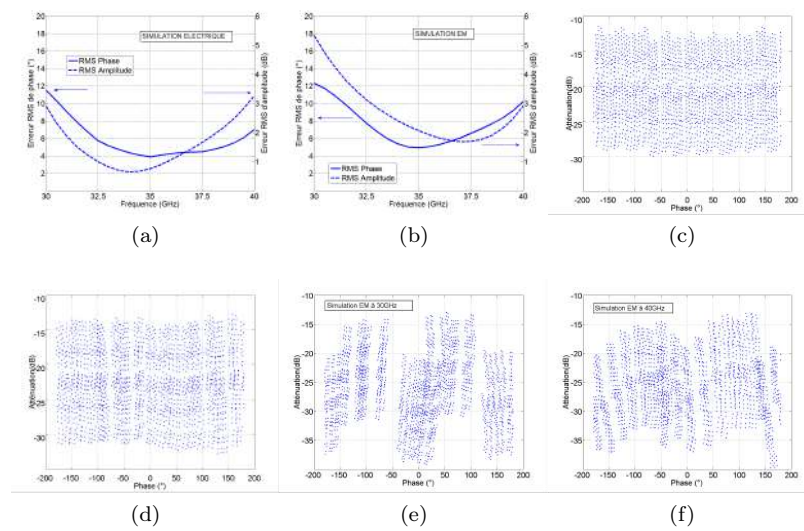


FIGURE 3.8 – Simulations des core-chips single-ended lors de la mise en commun des cellules en erreurs RMS fonctionnelles et en constellation phase amplitude. Les cas (a) et (c) correspondent aux simulations électriques et les cas (b) et (d) correspondent aux simulations EM à 35GHz. Les cas (e) et (f) correspondent aux constellations EM en bords de bande, 30 et 40GHz.

obtenir un ordre de cellules plus performant.

Ordre initial par classement binaire

Nous appelons classement binaire, l'ordre qui positionne les cellules par valeurs croissantes d'atténuation/déphasage (bit de poids le plus faible en premier et bit de poids le plus fort en dernier, comme illustré en Figure 3.9). Nous avons trouvé que l'ordre où les déphaseurs précèdent les atténuateurs, présentait de meilleures performances, c'est pour cela que nous n'avons fait figurer que les performances de cet ordre. Celles-ci sont reportées en Figure 3.10, nous pouvons donc voir qu'à 35 GHz les erreurs RMS de phase et d'atténuation sont respectivement de 4° et 1 dB.

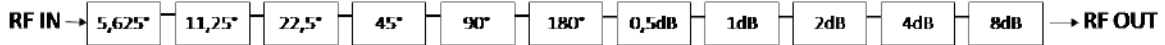


FIGURE 3.9 – Ordre des cellules par classement binaire

Relativement à la simulation de la constellation de la Figure 3.8d, la couverture ne présente pas cette fois-ci de zone non couverte significative et la configuration de cet agencement serait donc à prioriser par rapport à celui de la Figure 3.7. Concernant la constellation, nous avons toujours moins de -13 dB de pertes d'insertion, quelques désalignements d'amplitude et un léger recouvrement à -19 dB non gênant (mais qui explique une légère réduction de la dynamique d'atténuation par rapport à la Figure 3.8d). Comparé à l'ordre de la Figure 3.8b, la légère différence d'erreur RMS d'amplitude (0.3dB) se traduit peut-être par le recouvrement plus prononcé pour cet agencement. En effet en observant la partie basse de la constellation de la Figure 3.10, nous pouvons voir que le niveau maximal de pertes est de -28 dB contre -32.5 dB pour l'agencement précédent. Cette différence semble provenir d'une translation verticale vers le haut de tout le pan inférieur de la constellation de la Figure 3.10.

Etant donné qu'il n'y a qu'une seule superposition d'états et cela à un seul niveau d'amplitude, nous pouvons supposer que c'est l'œuvre d'une cellule désactivée pour la première moitié des états

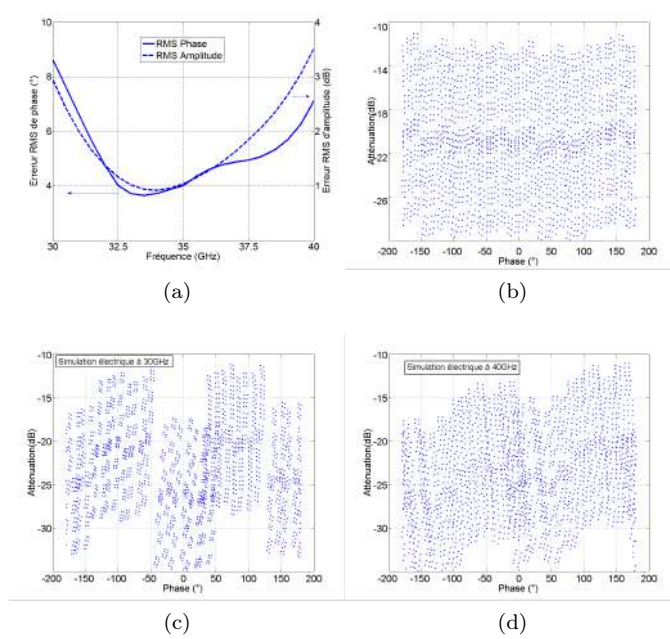


FIGURE 3.10 – Simulations électriques du core-chip dont les cellules sont classées par ordre binaire en termes de (a) erreurs fonctionnelles RMS sur la bande 30-40 GHz (b) constellation phase/amplitude à 35 GHz. Les cas (c) et (d) correspondent aux constellations en bords de bandes, 30 et 40GHz.

puis activée pour la moitié restante (typiquement la cellule 8dB). Du fait de la différence de pertes d'insertion lors de l'activation de cette cellule, la continuité des états d'atténuation n'est pas assurée. Ceux-ci sont au contraire décalés vers le haut, ce qui cause ce recouvrement.

Ordre initial de classement par pertes

Pour ce classement-ci, nous avons choisi de placer les cellules présentant le plus de pertes aux extrémités de la chaîne. De cette façon les ondes réfléchies sont plus atténuées ce qui améliore l'adaptation globale du système. Nous nous sommes donc référés aux performances présentées au chapitre précédent pour déterminer l'ordre respectant nos exigences. Nous avons donc placé les cellules comme illustré en Figure 3.11.

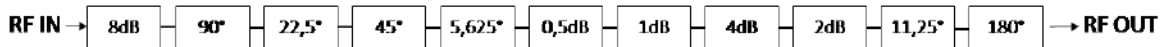


FIGURE 3.11 – Ordre de cellules placées par pertes

Pour trouver cet agencement nous avons essayé d'équilibrer les pertes de part et d'autre de la chaîne. Ce n'est pas le seul agencement possible et nous aurions pu en choisir un différent. Cette configuration de cellule nous permet d'obtenir les performances données en Figure 3.12.

Pour cette conception, nous obtenons une erreur RMS de phase de 2.1° et une erreur RMS d'amplitude de 0.87 dB à 35 GHz en simulation EM (contre 5° et 2dB en technique d'organisation implémentée dans le circuit, et contre 4° et 1dB pour la première technique d'organisation par classement binaire). Sur le simple critère d'erreurs RMS, cette conception semble plus avantageuse à 35 GHz. Cependant,

en analysant la constellation, nous pouvons voir que les pertes d'insertion sont plus importantes puisqu'elles atteignent -14 dB dans certains cas.

Cette constellation présente quelques désalignements en amplitude et une zone non couverte suivie d'une zone de recouvrement toujours aux alentours des -19 dB. Nous pouvons donc voir que malgré une diminution par deux de l'erreur de phase RMS à 35 GHz, nous n'arrivons pas à améliorer la constellation de manière significative. En effet, visuellement, la constellation de la Figure 3.10 paraît mieux couverte (régulière en pas de phase et d'amplitude) que celle de la Figure 3.12 (zones non couvertes à 19dB et -27 dB). Le recouvrement étant différent de celui de l'agencement précédent, nous pouvons constater que la sommation des pertes d'insertions des cellules diffère selon l'ordre des cellules.

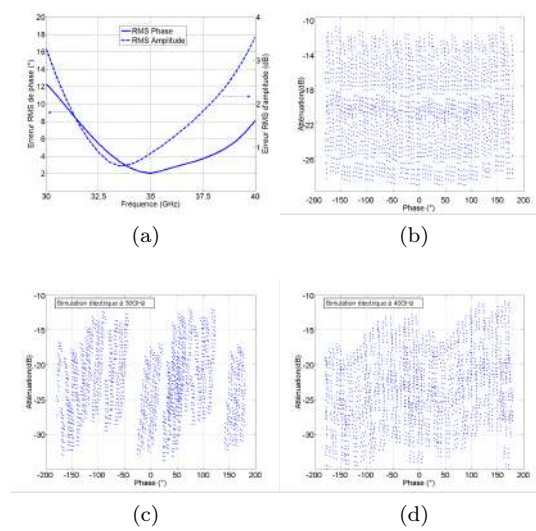


FIGURE 3.12 – Simulations électriques du core-chip dont les cellules sont classées par ordre de pertes d'insertion en termes de (a) erreurs fonctionnelles RMS sur la bande 30-40 GHz (b) constellation phase/amplitude à 35 GHz. Les cas (c) et (d) correspondent aux constellations en bords de bande, 30 et 40GHz.

Ordre initial de classement par niveau d'adaptation

Nous avons choisi ici de placer les cellules présentant les meilleurs niveaux d'adaptation en extrémités de la chaîne de façon à dégrader le moins possible l'adaptation globale. L'ordre ainsi obtenu est donné en Figure 3.12. Comme pour la partie précédente, ce n'est pas l'unique agencement possible.

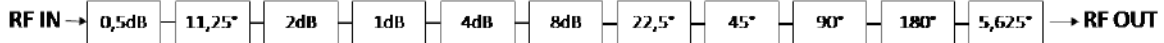


FIGURE 3.13 – Ordre de cellules placées par niveau d'adaptation

Les performances de cet agencement sont données en Figure 3.14. Nous obtenons une erreur RMS de phase de 3.6° et une erreur RMS d'amplitude de 1.1 dB. Concernant la constellation, nous obtenons des pertes d'insertion légèrement inférieures à -14dB. La zone de recouvrement est moindre comparée à celles des agencements initiaux précédents et il n'y a pas de zones non couvertes. C'est la constellation la plus dense parmi les 3 présentées, même si les erreurs RMS sont légèrement supérieures à celles de la configuration précédente (classement par pertes) qui procurait une couverture dégradée à 35 GHz.

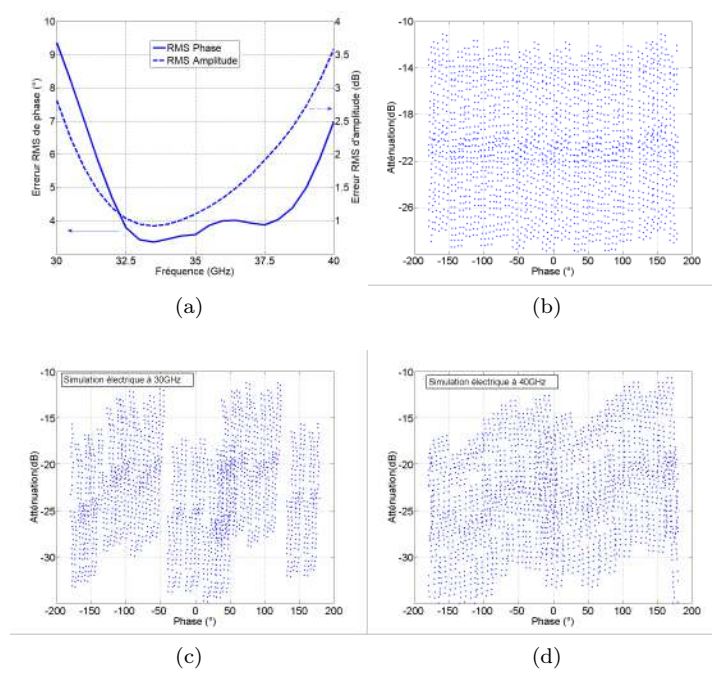


FIGURE 3.14 – Simulations électriques du core-chip dont les cellules sont classées par ordre de pertes d'insertion en termes de (a) erreurs fonctionnelles RMS sur la bande 30-40 GHz (b) constellation phase/amplitude à 35 GHz. Les cas (c) et (d) correspondent aux constellations en bords de bande, 30 et 40GHz.

Pour résumer, les deux approches par classement binaire et par niveau d'adaptation semblent plus efficaces pour obtenir de bonnes performances. En termes d'erreurs RMS fonctionnelles à 35 GHz nous pouvons dire que le classement par pertes est plus compétitif, mais cela ne se traduit pas par une meilleure constellation, au contraire.

Finalement, le critère d'adaptation semble avoir plus d'effets bénéfiques sur la constellation, puisque le tri selon l'adaptation propose la constellation avec la meilleure couverture. Il est néanmoins difficile de tirer des conclusions définitives concernant le rôle de l'adaptation sur la constellation. Nous allons maintenant présenter les performances obtenues avec l'aide de l'algorithme d'ordonnancement, et pouvoir comparer les différentes approches à privilégier.

3.2 En utilisant l'algorithme

Les cellules utilisées pour ces simulations assistées par l'algorithme sont rigoureusement les mêmes que celles utilisées lors de la partie précédente. En effet nous avons injecté les cellules dans l'algorithme, et celui-ci nous a retourné les différents agencements de cellules possibles avec leur EVM, leurs erreurs d'amplitude et de phase associées. Idéalement il aurait fallu injecter les 11 cellules dans l'algorithme pour trouver l'ordre optimal de cellules. Malheureusement, les temps de calcul impliqués pour dénombrer les 11! cas possibles (soit environ $40 \cdot 10^6$) sont beaucoup trop importants (estimés à plusieurs mois pour une simulation totale des 11 cellules) pour que nous puissions nous permettre ces simulations. Pour contourner cette difficulté, nous avons décidé d'appliquer l'algorithme sur les cellules de déphasage et d'atténuation séparément. Dans notre cas, cela représente 6 cellules de déphasage (720 cas possibles) et 5 d'atténuation (120 cas possibles) ce qui nous permet de réduire les temps de

simulation de l'algorithme à environ une à deux minutes. Une fois les résultats obtenus, nous avons classé les différents agencements selon 3 critères : l'EVM moyen, l'erreur de phase moyenne et l'erreur d'amplitude moyenne, par niveau de performance décroissante (du pire cas vers le meilleur cas). Puis nous avons commencé à idéaliser les paramètres S_{ii} et S_{ij} de chaque cellule, et ce les unes après les autres. En comparant les gains potentiels sur les 3 critères évalués, nous avons pu déterminer si une retouche de la cellule idéalisée était nécessaire dans le but de faire tendre les paramètres S_{ii} et S_{ij} vers l'idéalité après retouche du design de ladite cellule.

Nous allons donc présenter les opérations effectuées, que nous avons découpées en 3 parties correspondant respectivement à chaque critère.

Evaluation selon le critère EVM

Evaluation des déphaseurs

Au regard du critère de l'amplitude du vecteur erreur (EVM) moyen, l'ordre retenu pour les déphaseurs est donné en Figure 3.15. Plus généralement, les dix premiers agencements en termes d'EVM sont donnés dans le Tableau 3.2. Considérer un nombre d'agencements comme celui-ci nous permet de voir émerger des blocs de cellules responsables des bonnes performances du système. Ces différents blocs sont surlignés en vert dans le Tableau 3.2.

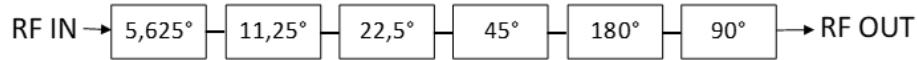


FIGURE 3.15 – Ordre de cellules placées par niveaux d'adaptation

Nous pouvons identifier deux blocs apparaissant sur tous les premiers agencements : l'association de la cellule 11,25° et 22,5° (quel que soit l'ordre) et l'association de la cellule 45° et 180° (quel que soit l'ordre) avec respectivement 9 occurrences et 8 occurrences sur les 10 meilleurs cas. Nous pouvons remarquer que l'association des cellules 5,625° et 90° apparaît régulièrement, cependant cette paire est aussi présente au niveau des pires agencements au regard de l'EVM moyen. Nous n'attribuons donc pas de bénéfices (ni de désavantages) à ce couple de cellules. En fait, sur 6 cellules, si deux couples se dégagent, les deux autres cellules seront régulièrement associées par voie de conséquence, avec 7 occurrences, mais le bémol quant à la présence de cet arrangement en situations de pire cas.

	Cellule 1	Cellule 2	Cellule 3	Cellule 4	Cellule 5	Cellule 6	Facteur de mérité EVM
Cas 1	5,625°	11,25°	22,5°	45°	180°	90°	0,02604
Cas 2	180°	5,625°	90°	45°	22,5°	11,25°	0,02618
Cas 3	90°	5,625°	11,25°	22,5°	180°	45°	0,0262
Cas 4	11,25°	22,5°	45°	180°	5,625°	90°	0,02624
Cas 5	90°	180°	45°	22,5°	11,25°	5,625°	0,02631
Cas 6	22,5°	45°	180°	11,25°	5,625°	90°	0,02635
Cas 7	180°	5,625°	90°	45°	11,25°	22,5°	0,02643
Cas 8	11,25°	22,5°	5,625°	90°	180°	45°	0,02654
Cas 9	45°	180°	22,5°	11,25°	5,625°	90°	0,02655
Cas 10	5,625°	90°	180°	45°	22,5°	11,25°	0,02656

TABLE 3.2 – 10 premiers meilleurs agencements de déphaseurs selon le critère EVM moyen

Pour vérifier que ces combinaisons sont bien responsables des performances améliorées, nous analysons les pires cas selon l'EVM moyen, donnés en Tableau 3.2. De cette façon nous pouvons déterminer si les bonnes performances sont dues à des combinaisons de cellules particulières ou plutôt à l'absence

de certaines. En première analyse, nous ne trouvons jamais d'association telle que décrite en vert sur le Tableau 3.3 dans les configurations pire cas.

	Cellule 1	Cellule 2	Cellule 3	Cellule 4	Cellule 5	Cellule 6	Facteur de mérite EVM
Cas 1	22,5°	45°	5,625°	90°	180°	11,25°	0,0385
Cas 2	22,5°	45°	5,625°	90°	11,25°	180°	0,038
Cas 3	5,625°	45°	22,5°	90°	180°	11,25°	0,0379
Cas 4	5,625°	45°	22,5°	90°	11,25°	180°	0,0376
Cas 5	45°	22,5°	5,625°	90°	180°	11,25°	0,0372
Cas 6	11,25°	180°	90°	5,625°	45°	22,5°	0,0369
Cas 7	5,625°	22,5°	45°	90°	11,25°	180°	0,0368
Cas 8	22,5°	5,625°	45°	90°	11,25°	180°	0,0368
Cas 9	11,25°	180°	90°	5,625°	22,5°	45°	0,0368
Cas 10	11,25°	180°	90°	22,5°	45°	5,625°	0,0366

TABLE 3.3 – 10 pires agencements de déphaseurs selon le critère EVM moyen

Ces différents agencements sont tous composés de deux blocs distincts respectivement formés des cellules 22,5°- 45° et 5,625° et des cellules 90°- 180° et 11,25° (quel que soit l'ordre dans les deux cas). Nous nous apercevons que ces blocs de 3 cellules ne sont jamais présents dans les agencements meilleurs cas et réciproquement les blocs à favoriser dans les meilleurs cas ne sont jamais présents dans les pires cas. La conclusion que nous tirons est qu'il faudra éviter ces blocs de 3 cellules, ou alors retravailler les cellules de façon à neutraliser cet effet inter-cellules négatif. Etant donné que nos observations sont vraies dans 100% des 20 cas observés nous pouvons supposer que cela constitue une bonne base de travail, et qu'il existe véritablement des regroupements de cellules à éviter ou à privilégier. Comme le montre la Figure 3.16, les 10 agencement choisis (meilleurs ou pires), se situent toujours à l'intérieurs des 50 meilleurs et pires cas sur les 720 tirages. Notre échantillon de 10 agencements est donc significatif et nous permet de tirer des conclusions statistiques.

Evolution des erreurs EVM pour 720 agencements de déphaseurs

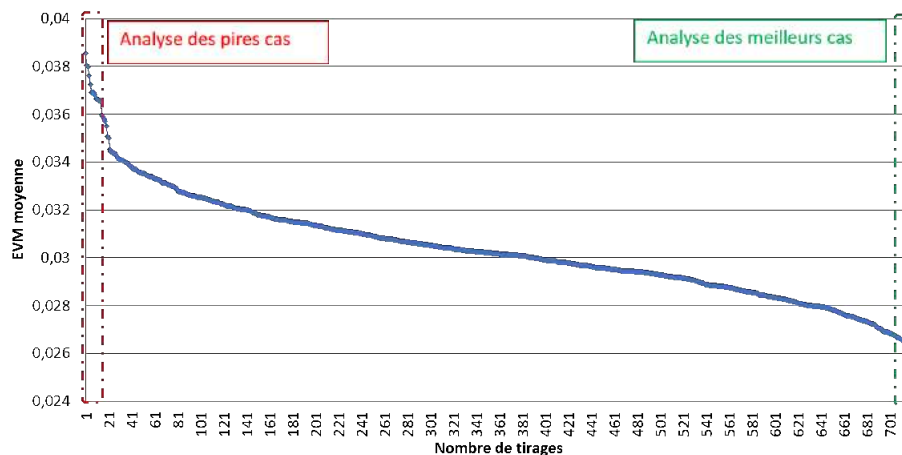


FIGURE 3.16 – Evolution de l'EVM moyen pour 720 agencements de 6 cellules de déphasage

Nous avons décidé de commencer notre processus d'optimisation à partir du meilleur ordre présenté

en Figure 3.15. Néanmoins, avant d'entamer celui-ci, nous allons réaliser le même tri d'organisation préalable, mais cette fois ci avec les cellules d'atténuateurs.

Evaluation des atténuateurs

De la même façon que pour les déphaseurs, nous sélectionnons le meilleur agencement de cellules selon le critère EVM moyen. Celui-ci est donné en Figure 3.17. Nous ne faisons pas figurer les dix premiers meilleurs/pires agencements puisque aucune tendance ou groupe de cellules particuliers n'a émergé de notre analyse. Ce n'est pas extrêmement surprenant puisque les cellules d'atténuation sont globalement mieux adaptées, ou alors présentent une atténuation permettant de baisser les niveaux de puissance incidente réfléchie.



FIGURE 3.17 – Agencement de cellules d'atténuation présentant l'EVM moyen la plus faible

Une fois l'ordre de cellule obtenu, nous pouvons les regrouper avec les cellules de déphasage.

Mise en commun en core-chip idéal

Il existe deux possibilités d'association des cellules, les cellules de déphasage avant ou après celles d'atténuation. Nous avons donc testé les deux possibilités au regard des constellations phase/atténuation et des erreurs RMS fonctionnelles. Les résultats sont présentés en Figure 3.17.

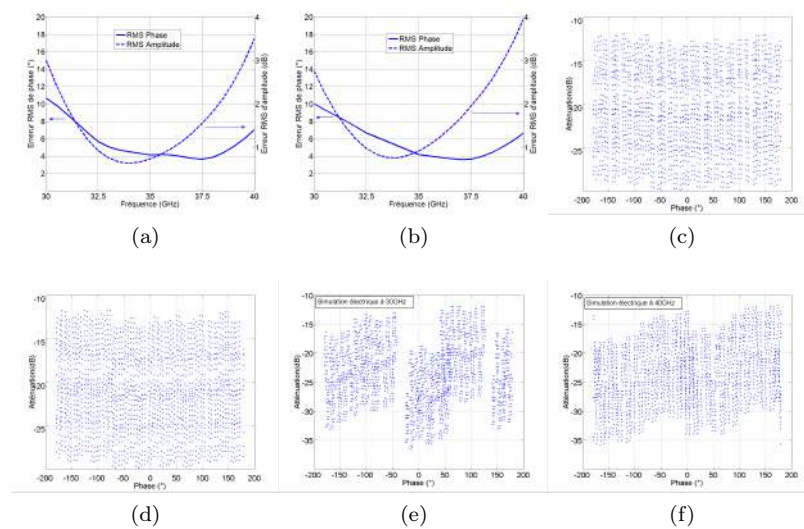


FIGURE 3.18 – simulations électriques présentant l'EVM moyenne la plus faible à 35 GHz en erreurs RMS fonctionnelles des agencements (a) atténuateurs puis déphaseurs (b) déphaseurs puis atténuateurs et constellations phase/amplitude des agencements (c) atténuateurs puis déphaseurs (d) déphaseurs puis atténuateurs. Les cas (e) et (f) correspondent aux constellations en bords de bandes, 30 et 40GHz pour l'ordre atténuateurs puis déphaseurs.

Même si nous avons fait figurer la bande de fréquence entière, nous rappelons que l'objectif de performances de ces cellules ciblait la fréquence centrale de 35 GHz. Nous pouvons voir que cet objectif est respecté puisque les minima des courbes sont quasiment toujours à 35 GHz. Les performances exactes

à 35 GHz sont des erreurs RMS de phase de 4.16° et 4.2° et d'amplitude de 0.74dB et 0.92dB respectivement pour les agencements atténuateurs puis déphaseurs et déphaseurs puis atténuateurs. Nous avons donc de légèrement meilleures performances pour l'agencement atténuateurs puis déphaseurs à 35 GHz, qui reste cependant peu sensible sur l'analyse des constellations. Cependant sur la totalité de la bande, les performances sont globalement similaires, avec quelques différences aux environs des fréquences limites. Malgré des performances assez similaires à celles des agencements précédents, nous pouvons voir que les constellations sont mieux couvertes et présentent peu de pertes par rapport aux méthodes d'organisation sans algorithme. De plus les différences de niveaux (ondulations d'amplitudes en fonction de la phase) en haut et bas de la constellation sont moins importants, cela donne une forme carrée au schéma. Nous pouvons donc pressentir que le critère de l'EVM moyen est pertinent pour traduire la qualité de couverture des constellations. Pour les constellations à 30 et 40 GHz, nous avons décidé de ne faire figurer que l'ordre atténuateurs puis déphaseurs puisque les erreurs RMS fonctionnelles à ces fréquences sont sensiblement égales. Nous allons maintenant réaliser la même étude mais en considérant l'erreur de phase comme critère discriminant.

Evaluation selon l'erreur de phase moyenne

Nous réalisons donc la même étude que précédemment. Nous ne ferons pas apparaître toutes les étapes présentées dans la partie ci-dessus. L'ordre présentant l'erreur de phase moyenne la plus basse est présenté en Figure 3.19a. Nous retrouvons le couple de cellule 22.5° - 11.25° présent dans les 10 meilleurs agencements selon l'EVM moyen. Ce couple de cellule peut donc présenter une organisation critique à privilégier. Pour les cellules d'atténuation, nous obtenons l'agencement illustré en Figure 3.19b.

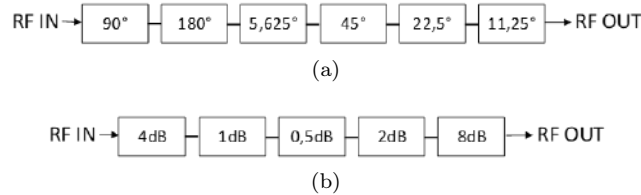


FIGURE 3.19 – Agencement de cellules (a) de déphasage (b) d'atténuation présentant l'erreur de phase moyenne la plus faible

Lors de la mise en commun de ces deux agencements nous obtenons les performances illustrées en Figure 3.20.

Nous pouvons voir dans les deux cas que l'erreur RMS de phase présente un optimum à 35 GHz. Ce minimum est égal à 1.9° et 2.1° respectivement pour les agencements atténuateurs puis déphaseurs et vice versa. En revanche pour l'erreur RMS d'atténuation le minimum est légèrement décalé vers la bande basse des fréquences, et se situe à un niveau d'atténuation de 1dB et 0.87dB respectivement aux deux ordres d'agencement. Bien que présentant des erreurs RMS satisfaisantes, les constellations présentent des couvertures à améliorer. Effectivement, même si celles-ci présentent une certaine densité, en comparaison avec la Figure 3.18, il existe des zones non couvertes. Il semblerait qu'il soit difficile d'optimiser les deux erreurs fonctionnelles simultanément (un écart de 0.3 dB sur l'erreur d'amplitude pouvant dégrader la constellation de façon non négligeable malgré une erreur de phase satisfaisante). Nous allons maintenant présenter le classement par erreurs d'amplitude moyenne puis conclure par une comparaison de ces trois critères pour choisir selon quel agencement le système final sera optimisé.

Evaluation selon l'erreur d'amplitude moyenne

Selon ce critère de tri des cellules, les meilleurs agencements choisis sont donnés en Figure 3.21.

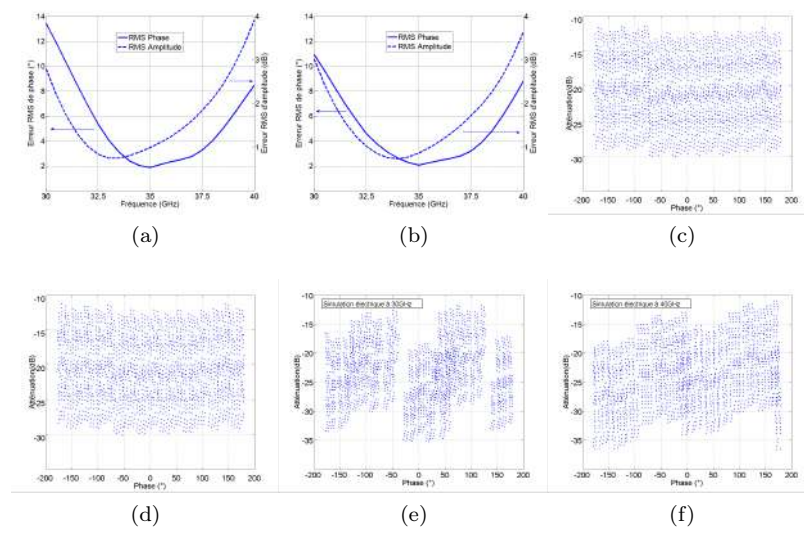


FIGURE 3.20 – simulations EM présentant l'erreur de phase moyenne la plus faible à 35 GHz en erreurs RMS fonctionnelles des agencements (a) atténuateurs puis déphaseurs (b) déphaseurs puis atténuateurs et constellations phase/amplitude des agencements (c) atténuateurs puis déphaseurs (d) déphaseurs puis atténuateurs. Les cas (e) et (f) correspondent aux constellations en bord de bandes, 30 et 40 GHz

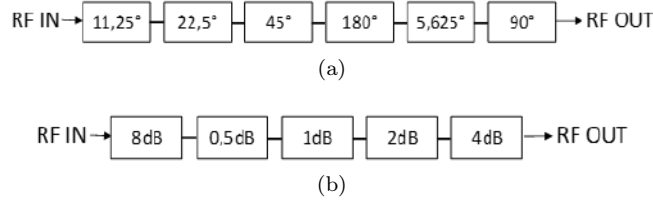


FIGURE 3.21 – Agencement de cellules (a) de déphasage (b) d'atténuation présentant l'erreur de phase moyenne la plus faible

Nous retrouvons les deux couples de cellules de déphasage identifiés en partie III.2.a (11.25° - 22.5° et 45° - 180°). Les performances obtenues lors de la mise en série des deux blocs sont données en Figure 3.22.

Nous obtenons des performances à 35 GHz égales à 3.6° et 3.2° pour l'erreur RMS de phase moyenne et de 0.73dB et 0.78dB respectivement aux ordres présentés en Figure 3.21. Les constellations présentent une meilleure couverture globale que celles des deux agencements précédents. Cependant les variations d'amplitudes aux bords des constellations sont plus importantes, ce qui peut réduire la zone exploitable de ces schémas. Nous pouvons remarquer que pour les classements suivants les 3 critères présentés, l'ordre atténuateurs puis déphaseurs présente des performances globalement meilleures que l'ordre inverse (dans des proportions peu significatives toutefois). Nous allons maintenant analyser plus en détails les performances de ces différents agencements.

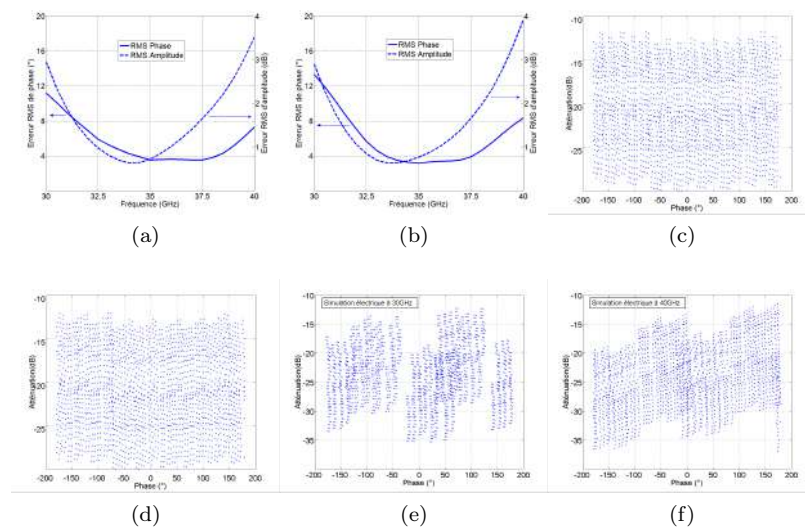


FIGURE 3.22 – Performances en erreurs RMS fonctionnelles des agencements (a) atténuateurs puis déphaseurs (b) déphaseurs puis atténuateurs et constellations phase/amplitude des agencements (c) atténuateurs puis déphaseurs (d) déphaseurs puis atténuateurs présentant l'erreur d'amplitude moyenne la plus faible à 35GHz. Les cas (e) et (f) correspondent aux constellations en bord de bandes, 30 et 40GHz.

Comparatif des performances

Avant de comparer les performances brutes il est important de rappeler que l'EVM moyen est fonction de l'erreur de phase et de l'erreur d'amplitude selon l'équation 3.1.

$$EVM = \sqrt{A_1^2 + A_2^2 - A_1 \cdot A_2 \cdot \cos(\phi)} \quad (3.1)$$

Avec A_1 et A_2 les amplitudes respectives du vecteur consigne et du vecteur réel, et ϕ l'erreur de phase. Nous pouvons en déduire que :

-l'erreur d'amplitude utilisée dans les parties précédentes correspond donc à la valeur RMS moyenne du ratio des deux amplitudes A_1 et A_2 , et ce exprimé en dB. D'une part ce ratio varie plus lentement que la phase, et en plus est traité après une opération logarithmique, ce qui veut dire qu'une erreur d'amplitude faible aura plus d'impact sur le calcul de l'EVM qu'une erreur de phase.

-En analysant la formule, nous pouvons aussi voir que l'erreur de phase varie selon une fonction cosinus mise à la racine carrée, alors que les amplitudes varient selon une fonction carrée. Cela confirme bien que l'erreur d'amplitude a un impact plus important sur le calcul de l'EVM.

Il faut aussi comprendre que les performances explicitées dans les parties précédentes sont données en valeurs RMS. Ainsi, l'information subit une opération de moyennage qui peut venir « lisser » les performances globales de telle sorte que l'on ne puisse plus voir les pires cas dont la valeur est assimilée au calcul de la moyenne. Il est donc important de ne pas se limiter à l'analyse des résultats des erreurs fonctionnelles RMS mais surtout de donner du crédit aux constellation phase/amplitude.

Sachant cela, nous pouvons analyser les performances des différents agencements explicités plus tôt. **Au regard des erreurs fonctionnelles RMS ce serait le classement selon l'erreur de phase qui donnerait les meilleures performances. Si nous analysons les constellations c'est le classement selon l'amplitude qui présente le meilleur résultat.** En effet, il présente la meilleure couverture globale phase amplitude, nous pouvons voir quelques zones de légers chevauchements mais

pas du tout de zones non couvertes comme pour les deux autres configurations. **En définitive, c'est l'ordre selon l'amplitude qui propose les meilleures performances selon les erreurs RMS et la constellation.** Une fois que le meilleur ordre a été identifié, il est possible d'étudier les gains de performances possibles grâce à des optimisations de cellules.

3.3 Optimisation des cellules

L'optimisation des cellules suit la même méthodologie, peu importe le mode de fonctionnement des cellules (*single-ended* ou différentielles). Nous allons présenter la façon générique selon laquelle nous avons optimisé les cellules. Nous commençons par idéaliser la performance que nous souhaitons optimiser ($S_{ii} = 0$ pour l'adaptation, $S_{ij} =$ « phase ou atténuation cible »). Puis nous injecterons la cellule avec le paramètre idéalisé dans l'algorithme. En analysant l'évolution des performances du système *core-chip* selon l'agencement idéalisé, nous pouvons juger si une amélioration peut être atteinte par un travail spécifique sur la cellule identifiée. Dans l'affirmative, une retouche de la cellule peut être envisagée (uniquement pour les cellules *single-ended* puisque l'algorithme fonctionne à une seule fréquence). Nous analysons ensuite les gains réels obtenus grâce à cette nouvelle conception de cellule. Si ces derniers sont conséquents, nous continuons ce processus pour les autres cellules.

Normalement, le cahier des charges initial stipulait des erreurs fonctionnelles nulles à 35 GHz. Cependant, les erreurs croisées des dispositifs (erreur de phase des atténuateurs et erreur d'atténuation des déphaseurs) ont pu mettre en défaut cette prérogative. En effet, même si c'est un peu moins vrai pour les atténuateurs où un effort a été fait pour limiter l'erreur de phase, l'erreur d'atténuation des déphaseurs n'a pas été sujette à des aménagements particuliers. C'est pour cela qu'une certaine liberté d'optimisation existe pour améliorer les performances globales du système. Pour réaliser l'optimisation, nous utilisons deux méthodes. La première consiste en l'utilisation d'isolateurs idéaux présentant des impédances de fermeture 50Ω dans le plan de charge que nous souhaitons idéaliser (approche par idéalisation de l'adaptation, également réalisable par action directe sur le fichier de paramètres S, paramètres $S_{ii} = 0$). La seconde permet d'idéaliser les paramètres de performances en transmission en modifiant directement les paramètres S en transmission et réflexion, sans toutefois modifier les niveaux d'adaptation réels (action uniquement sur les erreurs fonctionnelles). Nous avons appliqué ces améliorations pour les 3 meilleurs agencements trouvés dans la partie précédente.

Organisation des cellules de déphasage

Pour simplifier la mise en forme des résultats, nous avons choisi de faire figurer ces trois meilleurs agencements sous forme de tableau, et ce pour chacune des cellules optimisées. Chaque tableau est divisé en 3 parties en première colonne, une pour chaque meilleur agencement (*Best Case EVM*, *Best Case phase* et *Best Case amplitude*). Dans chacune de ces parties nous avons optimisé 3 paramètres ; l'adaptation d'entrée, l'adaptation de sortie et l'adaptation entrée-sortie. Les résultats sont visibles sur 3 colonnes selon chaque critère évalué (EVM moyenne, erreur de phase moyenne et erreur d'amplitude moyenne). Les chiffres en première colonne indiquent l'ordre des cellules étudié, les cellules étant numérotées dans l'ordre binaire, soit le numéro 1 pour la cellule 5.625° et le numéro 6 pour celle à 180° . Les résultats indiqués expriment la variation relative en pourcentage de la performance à la suite de l'optimisation. Le signe négatif indique une dégradation (colorée en rouge), les valeurs positives quant à elle indiquent une amélioration (colorée en vert). Nous allons donc présenter tous les résultats pour les différentes cellules de déphasage.

Nous commençons par la cellule 5.625° dans le Tableau 3.4.

Nous pouvons voir en Tableau 4 que les seules améliorations notables sont sur l'erreur de phase et ce sur l'ordre présentant la meilleure erreur d'amplitude moyenne (13 à 22%). Cette amélioration se

Best Case EVM 123465	EVM	Phase	Amplitude
Adaptation entrée (%)	0	0	0
Adaptation sortie (%)	-1,73	-0,35	-2,11
Adaptation E/S (%)	-1,73	-0,35	-2,11
Best Case Phase 561432	EVM	Phase	Amplitude
Adaptation entrée (%)	-3,71	-4,48	-2,62
Adaptation sortie (%)	-0,43	-8,13	0,43
Adaptation E/S (%)	-1,03	-5,14	0
Best Case Amplitude 234615	EVM	Phase	Amplitude
Adaptation entrée (%)	-1,61	13,21	-5,33
Adaptation sortie (%)	-1,38	17,96	-6,61
Adaptation E/S (%)	-1,70	22,11	-7,32

TABLE 3.4 – Evolution des performances de la cellule 5.625° lors de l'idéalisation de l'adaptation

fait au détriment d'une dégradation de l'erreur d'amplitude moyenne (5.3 à 7.3%). Cependant celle-ci a un impact plus important sur l'EVM moyenne puisqu'au final l'EVM moyenne associée à cette optimisation se dégrade malgré l'amélioration de l'erreur de phase. Un exemple de variation de l'erreur de phase suite à l'idéalisation de l'adaptation de sortie est donné en Figure 24. Les variations mises en jeu représentent respectivement des variations de 0%, 8% et 18% de l'erreur de phase par rapport à sa valeur initiale. Nous pouvons conclure que cette cellule ne se prête pas bien à l'amélioration de l'adaptation. Ce n'est pas très surprenant étant donné la topologie utilisée (ligne à retard), en effet celle-ci est centrée sur 35 GHz initialement, et ne laisse donc pas beaucoup de marge d'amélioration du fait d'une adaptation déjà correcte.

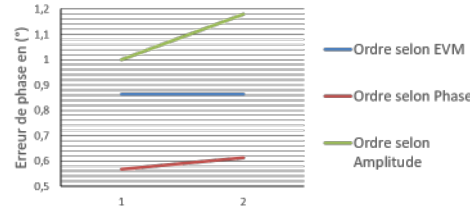


FIGURE 3.23 – Evolution des erreurs de phase à la suite de l'idéalisation de l'adaptation de sortie de la cellule 5,625°

Les résultats de la cellule 11.25° sont regroupés dans le Tableau 3.5. Nous pouvons voir qu'il n'y a pas d'amélioration notable possible des performances. Cette cellule illustre bien le fait qu'il existait des compensations d'erreurs grâce au niveau d'adaptation initial qui est médiocre. Nous pouvons le comprendre en observant la dégradation des performances avec l'idéalisation de l'adaptation. Intuitivement, ce changement devrait améliorer les performances cependant nous assistons à des dégradations allant jusqu'à 11%. Cela implique donc bien l'existence de compensations d'erreurs qui se voient perturber par cette nouvelle adaptation.

Les résultats pour la cellule 22.5° sont présentés en Tableau 3.6. Pour les meilleurs agencements selon l'EVM moyen et l'amplitude, nous pouvons voir qu'il existe une succession de cellules similaires dans les deux cas (11.25°-22.5°-45°-180°). Cette ressemblance est intéressante puisque nous pouvons voir les effets de cellules différentes sur les performances. En effet pour le meilleur ordre selon l'EVM nous obtenons une amélioration de l'erreur de phase contrairement à l'ordre selon l'amplitude. Nous

Best Case EVM 123465	EVM	Phase	Amplitude
Adaptation entrée (%)	-1,73	-0,35	-2,11
Adaptation sortie (%)	-4,05	4,17	-7,56
Adaptation E/S (%)	-7,24	0,61	-10,83
Best Case Phase 561432	EVM	Phase	Amplitude
Adaptation entrée (%)	-5,70	-10,85	-4,76
Adaptation sortie (%)	0	0	0
Adaptation E/S (%)	-5,70	-10,85	-4,76
Best Case Amplitude 234615	EVM	Phase	Amplitude
Adaptation entrée (%)	0	0	0
Adaptation sortie (%)	-8,14	-11,28	-10,96
Adaptation E/S (%)	-8,14	-11,28	-10,96

TABLE 3.5 – Evolution des performances de la cellule 11,25° lors de l'idéalisation de l'adaptation

pouvons nous questionner sur le rôle de la cellule 5.625° vis-à-vis des performances. Cette cellule, à moins qu'il existe des compensations internes, semble influencer les performances du système lors du changement de plan de charge autour de la cellule 22,5°. En effet cette dernière voit ses performances bouger uniquement quand la cellule qui lui est adjacente (11,25° dans les deux cas) est chargé en entrée sur la cellule 5.625°.

Best Case EVM 123465	EVM	Phase	Amplitude
Adaptation entrée (%)	-4,05	4,17	-7,56
Adaptation sortie (%)	-1,37	-0,41	-1,39
Adaptation E/S (%)	-5	9,48	-9,48
Best Case Phase 561432	EVM	Phase	Amplitude
Adaptation entrée (%)	1,67	-9,39	1,90
Adaptation sortie (%)	-5,70	-10,85	-4,76
Adaptation E/S (%)	-5,03	-10,09	-6,02
Best Case Amplitude 234615	EVM	Phase	Amplitude
Adaptation entrée (%)	-8,14	-11,28	-10,96
Adaptation sortie (%)	-2,42	-10,64	-1,94
Adaptation E/S (%)	-9,31	-17,63	-12,24

TABLE 3.6 – Evolution des performances de la cellule 22,5° lors de l'idéalisation de l'adaptation

Concernant la cellule 45°, nous avons regroupé les résultats dans le Tableau 3.7. Les améliorations sont plus nombreuses que pour les cellules précédentes. Une première explication peut venir du placement de la cellule dans les différents agencements. En effet, contrairement aux cellules précédentes, celle-ci est toujours placée en milieu de chaîne. De ce fait l'influence des plans de charge imposés par les autres cellules se fait plus sentir. Ainsi, en insérant un plan de charge idéal au milieu de la chaîne, les chances de replacer les autres cellules dans de meilleures conditions d'adaptation d'impédance. Nous pouvons aussi remarquer que malgré les améliorations notables de certaines erreurs de phase moyenne, l'erreur EVM moyenne n'est que légèrement modifiée. D'une part cela confirme le poids plus important de l'erreur d'amplitude dans le calcul de l'EVM, et d'autre part cela illustre le compromis récurrent

qui peut exister entre erreur de phase et erreur d'amplitude (il est difficile de réduire les deux simultanément). Il faut donc plutôt chercher à améliorer l'EVM plutôt qu'une des deux autres erreurs spécifiquement. Pour cette cellule par exemple, nous avons réussi à obtenir des améliorations importantes sur l'erreur de phase moyenne cependant. L'EVM n'a été que très légèrement impactée du fait de la dégradation de l'erreur d'amplitude.

Best Case EVM 123465	EVM	Phase	Amplitude
Adaptation entrée (%)	-1,37	-0,41	-1,39
Adaptation sortie (%)	-7,08	10,29	-9,68
Adaptation E/S (%)	-3,37	35,95	-7,37
Best Case Phase 561432	EVM	Phase	Amplitude
Adaptation entrée (%)	-0,43	-8,13	0,43
Adaptation sortie (%)	1,67	-9,39	1,90
Adaptation E/S (%)	2,70	9,65	2,33
Best Case Amplitude 234615	EVM	Phase	Amplitude
Adaptation entrée (%)	-2,42	-10,64	-1,94
Adaptation sortie (%)	0,17	20,32	-5,17
Adaptation E/S (%)	3,77	35,80	-1,84

TABLE 3.7 – Evolution des performances de la cellule 45° lors de l'idéalisation de l'adaptation

Nous avons rassemblé les résultats d'optimisation de la cellule 90° dans le Tableau 3.8. Il apparaît très rapidement que toutes les optimisations d'adaptation améliorent l'erreur de phase et ce quel que soit l'agencement optimisé. De plus dans le cas de l'agencement présentant les meilleures performances en erreur de phase moyenne, nous arrivons à obtenir une amélioration de l'EVM de 5%. Cela est très certainement dû à la présence de la cellule 180° après la cellule 90° puisque dans les autres configurations nous n'arrivons pas à reproduire ce gain de performances.

Best Case EVM 123465	EVM	Phase	Amplitude
Adaptation entrée (%)	-1,74	7,68	-4,32
Adaptation sortie (%)	0	0	0
Adaptation E/S (%)	-1,74	7,68	-4,32
Best Case Phase 561432	EVM	Phase	Amplitude
Adaptation entrée (%)	0	0	0
Adaptation sortie (%)	4,81	17,07	3,72
Adaptation E/S (%)	4,81	17,07	3,72
Best Case Amplitude 234615	EVM	Phase	Amplitude
Adaptation entrée (%)	-1,38	17,96	-6,61
Adaptation sortie (%)	0	0	0
Adaptation E/S (%)	-1,38	17,96	-6,61

TABLE 3.8 – Evolution des performances de la cellule 90° lors de l'idéalisation de l'adaptation

Enfin pour la dernière cellule de déphasage de 180° nous avons regroupé les performances dans le Tableau 3.9. Les résultats sont assez proches de ceux de la cellule précédente, c'est principalement dû au fait que les cellules sont adjacentes dans deux cas sur les trois. Nous obtenons donc la même

amélioration pour l'agencement présentant les meilleures performances en termes d'erreurs d'amplitude moyenne.

Best Case EVM 123465	EVM	Phase	Amplitude
Adaptation entrée (%)	-7,08	10,29	-9,68
Adaptation sortie (%)	-1,74	7,68	-4,32
Adaptation E/S (%)	-1	36,53	-6,84
Adaptation sortie (%)	0	0	0
Best Case Phase 561432	EVM	Phase	Amplitude
Adaptation entrée (%)	4,81	17,07	3,72
Adaptation sortie (%)	-3,71	-4,48	-2,62
Adaptation E/S (%)	3	12,83	2,40
Adaptation sortie (%)	0	0	0
Best Case Amplitude 234615	EVM	Phase	Amplitude
Adaptation entrée (%)	0,17	20,32	-5,17
Adaptation sortie (%)	-1,61	13,21	-5,33
Adaptation E/S (%)	-0,14	37,95	-8,20

TABLE 3.9 – Evolution des performances de la cellule 180° lors de l'idéalisation de l'adaptation

Les cellules 90° et 180° tirent donc le plus profit de l'idéalisation de l'adaptation, ce seront donc les cellules les plus susceptibles d'être redéfinies lors d'une refonte pour améliorer les niveaux d'adaptation. Cependant nous pouvons remarquer, pour les différentes cellules, qu'il est relativement compliqué d'améliorer l'EVM moyen et cela car les erreurs de phase et d'amplitude moyennes sont complexes à améliorer simultanément.

Nous ne ferons pas figurer cette idéalisation pour les atténuateurs, en effet pour les cellules *single-ended* l'apport de l'idéalisation pour les cellules d'atténuation que ce soit en adaptation ou en erreur fonctionnelle n'est pas suffisant et ne mérite donc pas d'être explicité. En effet l'erreur de phase des atténuateurs a été réduite à son maximum tout en gardant l'erreur d'amplitude minimale. L'existence de ce compromis rend l'idée d'idéalisation inutile. Nous allons maintenant réaliser le même cheminement mais en idéalisant les erreurs fonctionnelles, et plus spécifiquement l'erreur d'amplitude des déphaseurs puisque l'erreur de phase a été le point d'objectif de la conception, et est très proche de zéro à 35 GHz.

Nous commençons donc par la cellule 5.625°. Comme pour l'idéalisation de l'adaptation, nous ferons figurer les résultats sous forme de tableaux exprimant les variations en pourcentages des différents critères utilisés pour juger des performances. Cette fois-ci, les lignes représentent uniquement l'idéalisation de l'erreur d'amplitude. Pour la première cellule, le Tableau 3.10 présente les évolutions de performances. Nous n'obtenons pas d'amélioration très importante de l'EVM (3.3% maximum). Cette cellule présentait une erreur d'amplitude faible à 35 GHz (0.135dB), le gain potentiel de performances est donc réduit.

La cellule 11.25° quant à elle, présente une erreur d'amplitude de 0.5dB, le gain lors de l'idéalisation devrait donc être plus important. Le Tableau 3.11 explicite les performances obtenues. Contrairement à ce que nous pensions, l'idéalisation entraîne de fortes dégradations de l'erreur d'amplitude moyenne du système et donc de l'EVM moyenne (15%). Il y avait donc bien une compensation de l'erreur d'amplitude.

Nous pouvons faire le même constat sur la cellule 22.5° qui présente une erreur d'amplitude initiale de 0.4 dB. Ses performances sont répertoriées dans le Tableau 3.12. Nous pouvons questionner les

Best Case EVM 123465	EVM	Phase	Amplitude
Amplitude (%)	1,9	0	0,39
Best Case Phase 561432	EVM	Phase	Amplitude
Amplitude (%)	3,29	0,41	1,88
Best Case Amplitude 234615	EVM	Phase	Amplitude
Amplitude (%)	1,97	0,16	1,26

TABLE 3.10 – Evolution des performances de la cellule 5.625° lors de l'idéalisation de l'erreur d'amplitude à 35GHz

Best Case EVM 123465	EVM	Phase	Amplitude
Amplitude (%)	-15,5	-0,68	-14,04
Best Case Phase 561432	EVM	Phase	Amplitude
Amplitude (%)	-13,7	0	-10,26
Best Case Amplitude 234615	EVM	Phase	Amplitude
Amplitude (%)	-12,61	0	-13,02

TABLE 3.11 – Evolution des performances de la cellule 11.25° lors de l'idéalisation de l'erreur d'amplitude à 35GHz

interactions de cette cellule avec la 11.25° puisqu'elles sont toujours adjacentes dans les combinaisons. Il est donc possible qu'il existe une compensation d'erreur qui disparaît avec l'idéalisation de l'une ou l'autre des erreurs d'amplitude. Il faudrait peut-être corriger simultanément les erreurs d'amplitude des deux cellules pour confirmer cette théorie ; il faudrait pour cela faire un design commun des deux cellules. Cette approche (non mise en œuvre lors de ces travaux) est davantage détaillée à la fin du chapitre.

Best Case EVM 123465	EVM	Phase	Amplitude
Amplitude (%)	-11,9	-1,82	-12,44
Best Case Phase 561432	EVM	Phase	Amplitude
Amplitude (%)	-19,21	-1,59	-16,44
Best Case Amplitude 234615	EVM	Phase	Amplitude
Amplitude (%)	-12,2	-1,03	-12,99

TABLE 3.12 – Evolution des performances de la cellule 22.5° lors de l'idéalisation de l'erreur d'amplitude à 35GHz

La cellule 45° présentant une très faible erreur d'amplitude initiale (0.02dB), les améliorations présentées en Tableau 3.13 sont donc minimes.

Concernant la cellule 90°, étant donné que l'erreur d'amplitude était nulle, les performances restent les mêmes, et les résultats ne sont donc pas présentés.

Enfin pour la cellule 180°, nous obtenons des gains conséquents de performances (jusqu'à 19.5%). En effet avec une erreur d'amplitude initiale de 0.7dB nous optimisons grandement les performances globales du système. Celles-ci sont données en Tableau 3.14.

De cette analyse, nous pouvons observer que seule une retouche de la cellule 180° apporte un gain de performances notable en simulation. Nous avons donc décidé de tenter une refonte de la cellule pour obtenir des performances tendant vers celles obtenues après idéalisation.

Best Case EVM 123465	EVM	Phase	Amplitude
Amplitude (%)	0,40	0,14	-0,56
Best Case Phase 561432	EVM	Phase	Amplitude
Amplitude (%)	0,63	0,12	0,25
Best Case Amplitude 234615	EVM	Phase	Amplitude
Amplitude (%)	0,3	0,13	-0,14

TABLE 3.13 – Evolution des performances de la cellule 45° lors de l'idéalisation de l'erreur d'amplitude à 35GHz

Best Case EVM 123465	EVM	Phase	Amplitude
Amplitude (%)	10,52	0	5,64
Best Case Phase 561432	EVM	Phase	Amplitude
Amplitude (%)	22,96	0	17,97
Best Case Amplitude 234615	EVM	Phase	Amplitude
Amplitude (%)	7,16	4,89	1,87

TABLE 3.14 – Evolution des performances de la cellule 180° lors de l'idéalisation de l'erreur d'amplitude à 35GHz

La topologie de la cellule est une mise en parallèle d'un filtre passe bas avec un filtre passe haut. En diminuant la taille des transistors série du filtre passe haut de 5 μ m et en réduisant au minimum la longueur des lignes d'accès (comme représenté en Figure 3.24) nous arrivons à réduire l'erreur d'amplitude de 0.7 à 0.4dB.

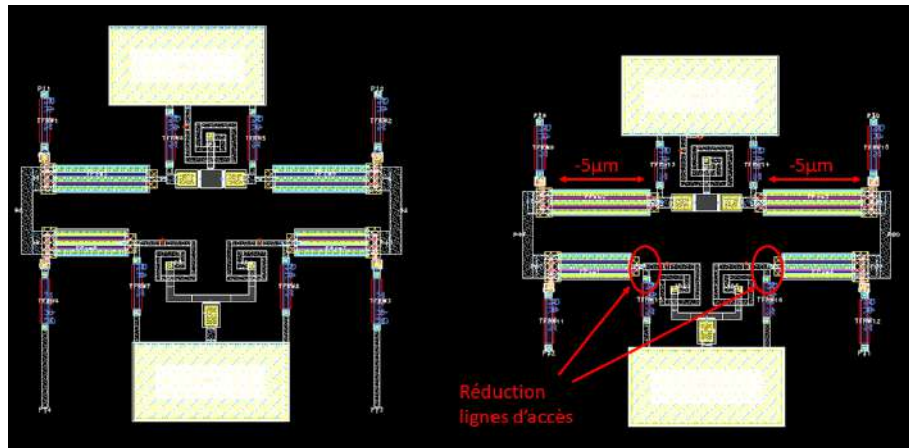


FIGURE 3.24 – Retouches opérées sur la cellule 180° single ended, à gauche la cellule initiale à droite la finale

Le Tableau 3.15 indique les valeurs d'erreurs de départ pour les meilleurs agencements.

Après la retouche de la cellule, nous obtenons les performances renseignées dans le Tableau 3.16. Contre toute attente, nous pouvons assister à une dégradation (très légère) des performances globales du système. Les seules erreurs ayant été améliorées sont l'erreur de phase pour l'ordre selon l'EVM, l'erreur d'amplitude et donc l'EVM pour l'ordre selon la phase. Ces améliorations sont sans conteste en dessous de nos prévisions, cependant en observant les nouveaux meilleurs agencements après ex-

Meilleurs cas	EVM	Erreur phase	Erreur amplitude
Ordre selon EVM : 123465	0,0260	0,862	0,050
Ordre selon phase : 561432	0,0279	0,567	0,058
Ordre selon amplitude : 234615	0,0262	1,00	0,049

TABLE 3.15 – Performances initiales des 3 différents agencements à 35GHz

exploitation de notre algorithme et après retouche de la cellule 180°, nous pouvons voir des améliorations non négligeables de performances puisque les meilleures erreurs passent à 0.0181, 0.625° et 0.0312dB respectivement pour l'EVM moyen, l'erreur de phase moyenne et l'erreur d'amplitude moyenne. Cependant il est difficile d'appliquer un tel raisonnement si nous ne suivons pas les évolutions suivant les agencements optimisés et que nous prenons uniquement les meilleurs nouveaux cas. En effet, les nouveaux meilleurs cas pourraient tirer profit de compensations que nous ignorons. Ainsi en analysant toujours le même ordre nous pouvons voir les gains en performances apportés par notre méthode. C'est pour cela que nous avons choisi de garder les performances initiales sans retouche de cellule. Nous les comparerons ensuite avec les performances obtenues sans l'aide de l'algorithme pour voir quel ordre sera choisi pour le système final.

Meilleurs cas	EVM	Erreur phase	Erreur amplitude
Ordre selon EVM : 123465	0,0266	0,829	0,0548
Ordre selon phase : 561432	0,026	0,866	0,0532
Ordre selon amplitude : 234615	0,0263	1,053	0,0515

TABLE 3.16 – Performances après retouche de la cellule 180° des 3 différents agencements à 35GHz

Pour comparer les performances nous avons choisi de suivre l'ordre selon l'EVM moyenne. En effet ce critère nous paraît le plus pertinent à comparer puisqu'il englobe les deux autres critères. Nous avons vu que dans certains cas, le classement selon l'erreur d'amplitude peut présenter de meilleures performances. Cependant, ces performances se font parfois au détriment de l'erreur de phase, qui même si elle n'a qu'une influence moindre sur l'EVM reste une caractéristique primordiale lors de la conception de *core-chips*. Nous avons donc trié les déphaseurs et atténuateurs selon l'EVM. Les deux différents ordres possibles sont donnés en Figure 3.25. Les performances en simulation électriques avaient déjà été présentées en Figure 3.17 et les simulations EM sont données en Figure 3.26.

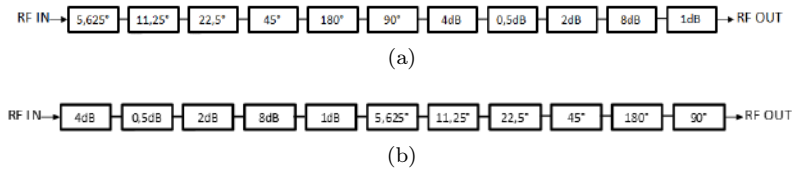


FIGURE 3.25 – Agencements des cellules présentant l'EVM moyenne la plus faible (a) déphaseurs puis atténuateurs (b) atténuateurs puis déphaseurs

A 35 GHz nous obtenons des erreurs RMS de phase de 4.05° et 3.34° et d'amplitude de 0.79dB et 0.75dB respectivement pour les ordres déphaseurs puis atténuateurs et vice versa. Concernant les constellations phase amplitude, la couverture semble correcte dans les deux cas. Nous pouvons distinguer de légers chevauchements au niveau d'atténuation -20dB, cependant ceux-ci ne sont pas préjudiciables aux performances puisqu'ils n'entraînent pas de zones non couvertes importantes. Pour les simulations EM dont les performances sont illustrées en Figure 3.26, nous assistons à une dégradation

des performances. A 35 GHz, nous obtenons respectivement des erreurs RMS de phase de 7.846° et 6.145° et d'amplitude de 1.662dB et 1.345dB pour les ordres déphaseurs-atténuateurs et atténuateurs-déphaseurs. C'est une augmentation de quasiment 4° pour l'erreur de phase RMS par rapport à la simulation électrique, ce qui une fois encore pose la question de la confiance accordée à la simulation électrique ou électromagnétique en première intention. Cela est bien visible sur les constellations qui présentent de nombreuses zones non couvertes et de recouvrement en simulation EM. Nous sélectionnons celui présentant les meilleures performances RMS des deux, malgré le fait qu'il ne présente pas une meilleure couverture au regard de la constellation correspondante.

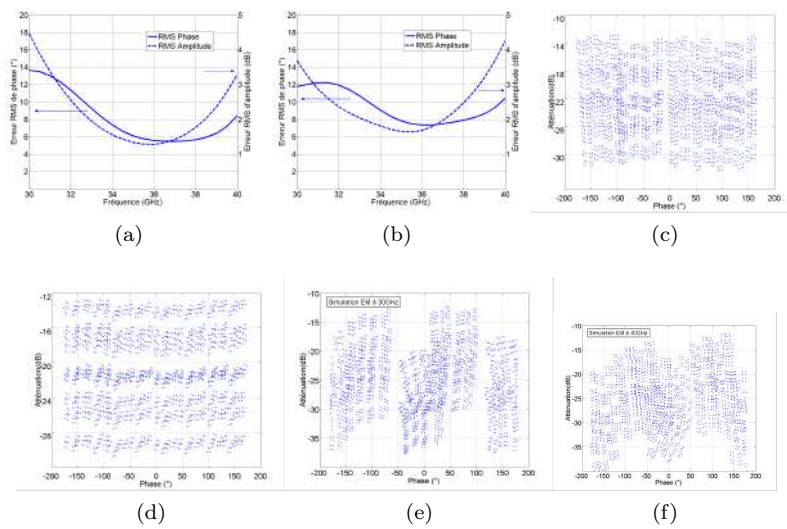


FIGURE 3.26 – Performances des simulations EM en erreurs RMS fonctionnelles des agencements (a) déphaseurs puis atténuateurs (b) atténuateurs puis déphaseurs et constellations phase/amplitude des agencements (c) déphaseurs puis atténuateurs et (d) atténuateurs puis déphaseurs présentant l'EVM moyenne la plus faible à 35GHz. Les cas (e) et (f) correspondent aux constellations en bord de bandes, 30 et 40GHz de l'agencement atténuateurs puis déphaseurs

Le Tableau 3.17 regroupe les performances des deux différents *core-chips* simulés jusqu'alors en EM. En comparant les performances nous voyons que l'utilisation de l'algorithme n'a pas apporté l'effet escompté sur les performances du système complet en version *single-ended*. En effet même si nous avons pu réduire grandement l'erreur d'amplitude, les performances appréciables sur la constellation démentent l'apport escompté par l'utilisation de l'algorithme étant donné les trop nombreuses zones non couvertes. Cependant ces propos restent mitigés puisque nous n'avons pas pu tirer réellement profit de l'algorithme au complet du fait que nous avons simulé les atténuateurs et déphaseurs séparément. En effet, il est vraisemblable que grâce à une optimisation manuelle appliquée à toutes les cellules, nous trouvons un ordre présentant des performances supérieures à celles obtenues dans notre cas de figure. Ainsi malgré le gain de temps obtenu grâce à l'algorithme, cette limitation due aux nombres de cellules utilisées apparaît. Loin de vouloir enterrer l'approche utilisant l'algorithme, il faudrait trouver une solution algorithmique permettant de manipuler des matrices aussi importantes (11!) tout en gardant des temps de calculs acceptables. Malheureusement par manque de temps nous n'avons pas pu explorer plus avant ce sujet, qui reste pertinent dans une optique future de pleine exploitation de l'algorithme d'ordonnancement.

Pour donner suite à ces conclusions, nous avons décidé de ne pas utiliser l'algorithme pour la seconde version du *core-chip single-ended*. Cependant, nous avons modifié les cellules à la suite d'un problème

Performances à 35GHz	Erreur RMS de phase (°)	Erreur RMS d'amplitude (dB)
Core-chip assemblé sans algorithme	4.973°	2.07
Core-chip assemblé avec algorithme	6.145	1.345

TABLE 3.17 – Comparatif des performances des deux différents core-chips simulés en EM

que nous avons remarqué lors des mesures de la première puce. Cet aspect sera plus largement détaillé dans le chapitre suivant consacré aux mesures. Néanmoins pour la compréhension du lecteur nous allons brièvement expliquer les motifs et les changements opérés sur les cellules *single-ended*.

Lors des deux premiers design (différentiel et *single-ended* non optimisé par algorithme) nous avons commis une erreur lors de la conception. En effet, les modèles de transistors en commutation présentaient une « auto polarisation », c'est-à-dire que les tensions statiques de polarisation des accès étaient déclarées comme paramètres intrinsèques du modèle. Ainsi, quelle que soit la tension réelle appliquée, les transistors voyaient toujours, en simulation, la tension que nous déclarions. En réalité, le circuit de polarisation influençait ces tensions et rien ne nous garantissait d'une part les tensions grille source de 0, -3 et -10V et d'autre part la tension drain source nulle qui garantissait un fonctionnement en commutation sans dissipation de courant. De plus le modèle électrique du transistor en commutation n'affichait pas les courants et tensions aux différents accès du transistors. Après rétro-simulation à la suite de l'apparition en mesure de cellules qui ne commutaient pas, nous avons utilisé des transistors amplificateurs (modèles large-signal analogique) avec les courants et tensions statiques pour vérifier la bonne polarisation. Nous avons trouvé que la masse statique nécessaire à l'établissement des tensions source-grille et drain-source ne se propageait pas à travers les cellules ; cela étant dû à l'ajout de capacité d'adaptation RF qui bloquaient le signal continu. La seconde version des circuits (*single-ended* et différentiel) intègre donc des cellules corrigées par l'ajout de *via-holes* assurant la présence de la masse statique pour toutes les cellules. Dans le cas du *core-chip single-ended* nous n'avons pas changé l'ordre des cellules, les performances après cette correction n'étant pas sensiblement différentes en simulation électrique. Celles-ci sont présentées en Figure 3.27. Contre toute attente, nous pouvons voir que les performances sont bien inférieures à l'itération précédente l'erreur RMS de phase étant de 8.25° et l'erreur RMS d'amplitude de 1.82dB à 35 GHz. Cependant par manque de temps relatifs aux délais courts nous imposant d'envoyer les circuits pour pouvoir obtenir une dernière conception (*single-ended* corrigé et différentiel corrigé), nous avons dû nous contenter de ces performances. Nous espérons cependant que cette seconde version de *core-chip single-ended* corrigera les défauts de polarisation inhérents à l'itération précédente.

Nous avons résumé les performances obtenues pour les deux *core-chips single-ended* centrés à 35 GHz envoyés en fabrication dans le Tableau 3.18. Les dessins des masques correspondants sont présentés en Figure 3.28.

Performances à 35 GHz	Erreur RMS de phase (°)	Erreur RMS d'amplitude (dB)
Core-chip assemblé sans algorithme	4.973°	2.07
Core-chip assemblé avec algorithme	8.249	1.823

TABLE 3.18 – Comparatif des performances simulés en EM des deux différents core-chips single-ended envoyés en fabrication

Les deux puces mesurent respectivement 3.13*1 mm² et 3.2*1.1 mm². Sur les deux puces, les plots de gauches et droites sont les entrées et sorties RF. Les plots alignés horizontalement sont les plots de polarisation statique de chaque cellule (commande des commutateurs de chaque état).

Nous allons maintenant nous intéresser à l'approche différentielle et aux simulations des deux *core-chips* large bande.

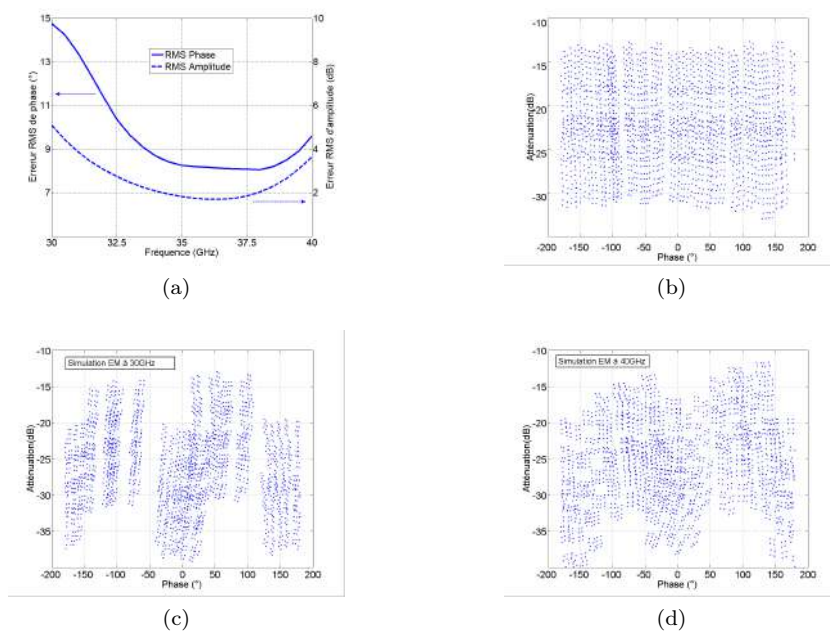
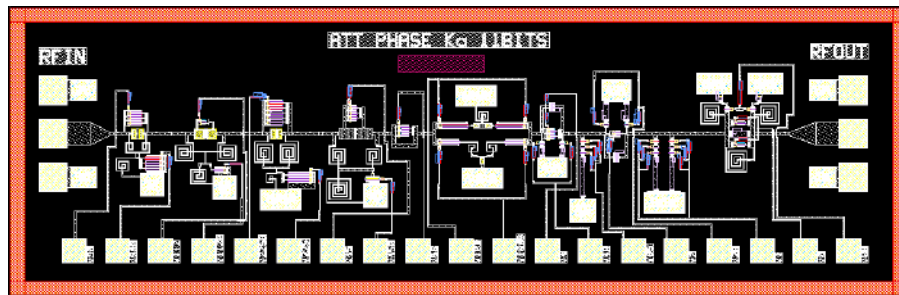
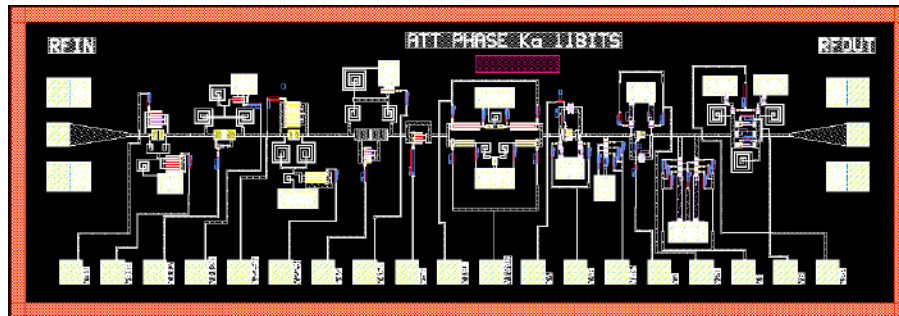


FIGURE 3.27 – Performances en simulation EM de la seconde itération du core-chip single ended centré à 35GHz envoyé en fabrication en (a) les erreurs fonctionnelles en (b) la constellation phase/amplitude à 35GHz et en (c) et (d) les constellations à 30 et 40 GHz



(a)



(b)

FIGURE 3.28 – Dessin des masques des deux core-chips single-ended envoyés en fabrication (a) première et (b) seconde version

4 Application pour les cellules différentielles

Les cellules différentielles ont été conçues selon un cahier des charges large bande. Ainsi, l'utilisation de l'algorithme est d'avantage limitée que pour les cellules *single-ended* centrées sur 35 GHz. En effet étant donné la largeur de bande mise en jeu, il est plus difficile d'utiliser l'algorithme pour cibler les paramètres à optimiser à une fréquence fixe. C'est pour cela que l'algorithme est plus utilisé pour trouver des bandes de fréquences réduites de fonctionnement où l'obtention d'agencements présentant de meilleures performances est envisageable. Nous allons donc présenter les différentes simulations avec ou sans algorithme.

4.1 Sans utilisation de l'algorithme

Nous avons commencé par tester quelques agencements, principalement inspirés par l'ordre binaire. En sélectionnant celui présentant les meilleures performances, nous avons ensuite cherché à améliorer les performances en permutant les cellules. Notre « optimisation manuelle » s'arrête quand nous n'arrivons plus à améliorer sensiblement les performances d'agencement des cellules. Pour obtenir un éventail important d'agencement, nous essayons de varier au maximum les différentes combinaisons de cellules. Une fois que certains emplacements de cellules sont identifiés comme critiques, des tendances vont se dégager et les permutations seront donc réduites. Pour cette raison, cette méthode d'optimisation a de grandes chances de tendre vers un optimum local. De plus, les cellules différentielles présentant deux lignes d'entrée et deux lignes de sortie, il est plus difficile de réaliser des accès standardisés pour connecter les cellules entre elles. En effet, les différences importantes de topologie rendent les distances entre les pistes de référence et déphasée extrêmement variable d'une cellule à l'autre. Pour les connecter il faut donc rajouter des longueurs de lignes plus ou moins importantes qui vont perturber l'adaptation et donc les performances des cellules mises en jeu. L'ordre des cellules choisi est celui donné en Figure 3.29.

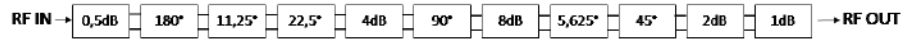


FIGURE 3.29 – Ordre des cellules différentielles final choisi

Cet agencement permet d'obtenir les performances données en Figure 3.30. Ces performances sont obtenues en simulation EM. Nous arrivons à conserver une erreur RMS de phase inférieure à 9.2° sur toute la bande et inférieure à 7° de 30 à 34,8 GHz. Pour l'erreur d'amplitude, elle est inférieure à 2.4 dB sur toute la bande, et inférieure à 2 dB de 30 à 37,5 GHz. La constellation présente de nombreuses zones non couvertes, et semble présenter des séquences de paquets de points comme si certains états de phase avaient été sautés. La qualité des constellations de bord de bande (30 et 40 GHz) illustrées en Figure 3.30c et en Figure 3.30d ne semblent pas suivre l'évolution des erreurs RMS fonctionnelles. En effet malgré une différence de 2° RMS et 0.8dB RMS entre les deux il ne semble pas qu'une des deux constellations soit sensiblement mieux couverte que l'autre.

C'est durant les mesures de ce circuit que nous nous sommes rendu compte du problème de polarisation évoqué dans la partie précédente sur le circuit *single-ended*. Pour pallier ce défaut, nous avons corrigé ces problèmes de polarisation avant d'appliquer l'algorithme d'ordonnancement. L'utilisation de ce dernier est détaillée dans la partie suivante.

4.2 En utilisant l'algorithme

Bien qu'initialement conçu pour optimiser l'ordre à une seule fréquence, nous avons choisi d'utiliser l'algorithme pour trouver l'ordre optimal des cellules différentielles à 35 GHz. Dans cette optique nous

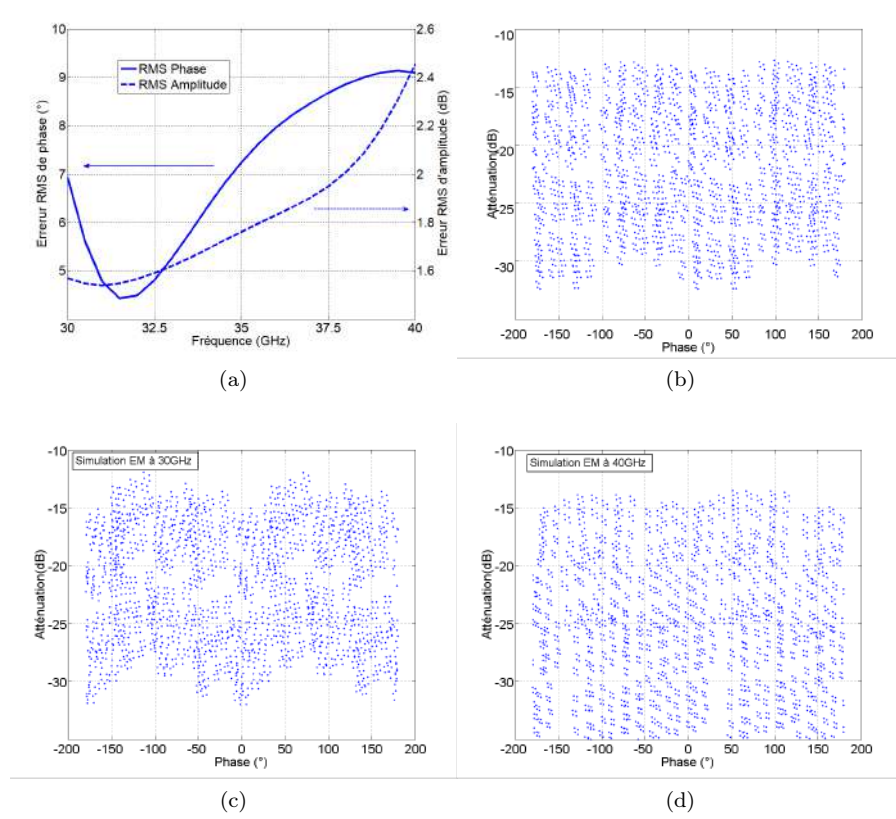


FIGURE 3.30 – Performances des core-chips différentiels à l’agencement de cellules optimisé manuellement (a) en RMS fonctionnelle et (b) la constellation phase/atténuation à 35 GHz et en (c) et (d) les constellations à 30 et 40 GHz

avons dans un premier temps trouvé les ordres présentant les meilleures performances à 35 GHz. Puis dans un second temps, nous avons observé l’évolution des performances aux fréquences basse (30 GHz) et haute (40 GHz). Comme nous l’avons notifié au chapitre 2, les variations d’erreur des cellules sont en général monotones autour d’une fréquence centrale. Quand cette fréquence est 35 GHz, nous obtenons des profils d’erreurs à peu près symétrique sur toute la bande ce qui nous permet de définir des sous bandes de fonctionnement où l’erreur est inférieure à un seuil fixé par nos soins. Nous avons utilisé quasiment les mêmes cellules que celles de la version non optimisée. La différence vient principalement d’un problème de polarisation statique que nous avons notifié lors des mesures du *core-chip* différentiel non optimisé par algorithme. Ce problème sera plus abondamment décrit dans le chapitre suivant consacré aux mesures. Les différences de dessin des masques sont donc majoritairement l’ajout de *via-holes* supplémentaires. Le problème de polarisation est d’autant plus critique étant donné que ce circuit, du fait de sa nature différentielle, n’est pas sensé avoir besoin de *via-holes* pour fonctionner.

Pour passer en revue les améliorations possibles grâce à l’algorithme nous avons regroupé les résultats sous forme de tableau figurant l’ensemble des optimisations par critères d’adaptation, d’erreur de phase et d’amplitude, et ce pour les trois meilleurs ordres au regard des critères étudiés (EVM moyenne et erreur de phase et d’atténuation moyennes). Ces trois meilleurs ordres sont présentés en Figure 3.31 pour les déphaseurs et en Figure 3.32 pour les atténuateurs.

Il n’y a pas vraiment de suite de cellules qui se répète parmi les déphaseurs ; l’enchaînement 45°-

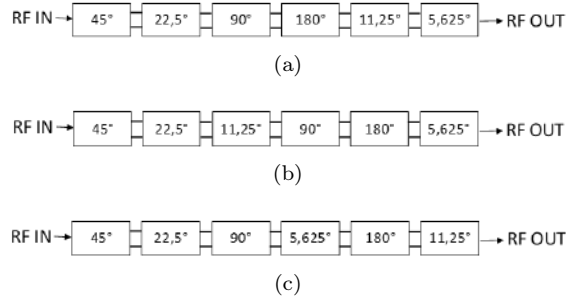


FIGURE 3.31 – Agencements de déphaseurs présentant les meilleures performances selon (a) l'EVM moyenne (b) l'erreur de phase moyenne et (c) l'erreur d'amplitude moyenne

22.5° est le seul à être présent dans les trois cas. Pour les atténuateurs cependant, il semble y avoir plus de réciprocité au sein des cellules puisque l'ordre selon l'EVM moyen correspond à l'ordre inverse de celui présentant l'erreur de phase la plus faible.

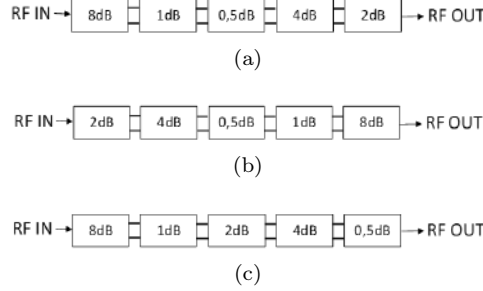


FIGURE 3.32 – Agencements de déphaseurs présentant les meilleures performances selon (a) l'EVM moyenne (b) l'erreur de phase moyenne et (c) l'erreur d'amplitude moyenne

Nous pouvons voir aussi que les atténuateurs ont été plus optimisés pour avoir des erreurs croisées moins, importantes puisqu'il n'y a pas le même compromis de conception que pour les déphaseurs où une erreur de phase faible s'obtient toujours au détriment d'une erreur d'amplitude forte.

Pour noter les améliorations, nous utiliserons des tableaux assez similaires à ceux-ci, excepté le fait que les colonnes correspondront à l'idéalisation effectuée et qu'il n'y aura qu'un seul tableau par ordre. Par exemple le Tableau 3.21 traite de l'idéalisation pour le meilleur ordre selon l'EVM moyenne des déphaseurs. Les numéros 1,2,...,6 correspondent aux cellules 5.625°,11.25°,..., 180°. Même chose pour les atténuateurs. Les valeurs négatives représentent une diminution de l'erreur en pourcentages et les valeur positives une augmentation, elles aussi en pourcentages par rapport aux erreur initiales présentées dans les Tableau 3.19 et Tableau 3.20.

Position	1	2	3	4	5	6	EVM	Phase	Amp
EVM	45°	22.5°	90°	180°	11.25°	5.625°	0,02339	2,8998	0,02894
Phase	45°	22.5°	11.25°	90°	180°	5.625°	0,02782	1,7746	0,06417
Amp	45°	22.5°	90°	5.625°	180°	11.25°	0,03684	5,2289	0,02400

TABLE 3.19 – Performances initiales des meilleurs agencements de déphaseurs

Position	1	2	3	4	5	EVM	Phase	Amp
EVM	8dB	1dB	0.5dB	4dB	2dB	0,02189	1,4666	0,2499
Phase	2dB	4dB	0.5dB	1dB	8dB	0,02198	1,4169	0,2513
Amp	8dB	1dB	2dB	4dB	0.5dB	0,02600	2,02689	0,2451

TABLE 3.20 – Performances initiales des meilleurs agencements de déphaseurs

Nous avons fait figurer les améliorations de performances en vert et les dégradations en rouge. Le Tableau 3.21 présente les résultats pour l'ordre des cellules de déphasage présentant la meilleure EVM moyenne. Nous pouvons voir qu'il n'y pas d'amélioration significative de l'EVM moyenne grâce à ces idéalizations.

Améliorations	Adapt IN (%)	Adapt OUT (%)	IN/OUT (%)	Amp (%)	Phase (%)
5,625°	-1,09	0	-1,11	52,30	0,48
11,25°	2,54	-1,10	-2,85	1,53	-4,13
22,5°	14,96	38,26	33,64	5,07	2,68
45°	0	14,92	14,87	2,98	-0,11
90°	37,83	9,21	50,16	69,12	2,83
180°	9,2	2,53	7,26	79,31	0

TABLE 3.21 – Evolutions des performances en termes d'EVM moyenne de l'ordre des cellules de déphasage présentant la meilleure EVM initiale

Toujours sur le même agencement de cellule de déphasage, nous avons effectué la même analyse, mais cette fois si en observant l'évolution de l'erreur de phase. Les résultats sont reportés dans le Tableau 3.22. Nous obtenons une amélioration de 7.2 % en adaptant idéalement la cellule 22.5° en entrée et sortie. Cependant il est visible dans le Tableau 3.21 que cela dégrade fortement l'EVM moyenne, certainement à cause de l'erreur d'amplitude moyenne qui doit augmenter.

Améliorations	Adapt IN (%)	Adapt OUT (%)	IN/OUT (%)	Amp (%)	Phase (%)
5,625°	3,95	0	3,95	0	1,1
11,25°	6,14	3,94	7,81	0,499	-3,4
22,5°	9,37	47,37	-7,19	0	3,49
45°	0	9,39	9,42	0	0,23
90°	2,19	5,44	16,41	63,17	-2,64
180°	5,41	6,14	6,61	0	

TABLE 3.22 – Evolutions des performances en termes d'erreur de phase moyenne de l'ordre des cellules de déphasage présentant la meilleure EVM initiale

Comme nous l'avions prévu l'adaptation entrée/sortie de la cellule dégrade l'erreur d'amplitude moyenne, le Tableau 3.23 illustre ces changements. Nous pouvons une nouvelle fois remarquer que les changements sur l'EVM moyenne semblent les plus pertinents à analyser puisqu'ils prennent en compte les deux autres critères d'évolution des performances.

Nous avons effectué les mêmes analyses pour l'ordre présentant l'erreur de phase moyenne la plus faible. Néanmoins, nous n'avons fait figurer que l'évolution de l'EVM puisque nous avons convenu que ce

Améliorations	Adapt IN (%)	Adapt OUT (%)	IN/OUT (%)	Amp (%)	Phase (%)
5,625°	-8,31	0	-8,38	-7,75	0
11,25°	-8	-8,31	-31,65	15,13	0,24
22,5°	22,1	97,97	120,84	5,6	0,11
45°	0	21,93	21,65	5,19	0
90°	121,46	18,24	143,65	44,89	17,84
180°	18,23	-8,01	6,87	23,06	

TABLE 3.23 – Evolutions des performances en termes d'erreur d'amplitude moyenne de l'ordre des cellules de déphasage présentant la meilleure EVM initiale

critère seul était suffisamment pertinent pour rendre compte d'une réelle amélioration de performances. Tableau 3.24 regroupe les différentes évolutions.

Améliorations	Adapt IN (%)	Adapt OUT (%)	IN/OUT (%)	Amp (%)	Phase (%)
5,625°	-2,27	0	-2,29	-5,58	0,81
11,25°	32,1	28,28	15,81	2,15	4,8
22,5°	13,51	55,38	20,1	-11,8	1,06
45°	0	13,44	13,36	-2,93	0,13
90°	27,97	14,36	45,31	25,48	5,46
180°	14,31	-2,27	11,64	68,4	0

TABLE 3.24 – Evolutions des performances en termes d'EVM moyenne de l'ordre des cellules de déphasage présentant la meilleure erreur de phase initiale

Améliorer l'erreur d'amplitude de la cellule 22.5° semble avoir un effet conséquent (12% d'amélioration) sur l'EVM moyenne. Cependant, comme nous pouvons nous en douter, cette amélioration est à mitiger puisque l'EVM moyenne initiale ne présentait pas une valeur très faible.

Enfin pour le meilleur ordre de cellules selon l'erreur d'amplitude, nous pouvons voir que nous avons de nombreuses améliorations, cependant l'EVM moyenne initiale était beaucoup trop importante pour se satisfaire de ces évolutions. Les résultats sont regroupés dans le Tableau 3.25.

Améliorations	Adapt IN (%)	Adapt OUT (%)	IN/OUT (%)	Amp (%)	Phase (%)
5,625°	-35,53	-15,33	-35,79	1,79	0,63
11,25°	-3,92	0	-3,92	-2,23	0,28
22,5°	-3,65	4,66	-6,04	6,81	-0,32
45°	0	-3,67	-3,69	3,36	-1,83
90°	-2,06	-35,5	-2,18	59,7	2,17
180°	-15,33	-3,92	-14,3	37,64	0

TABLE 3.25 – Evolutions des performances en termes d'EVM moyenne de l'ordre des cellules de déphasage présentant la meilleure erreur d'amplitude initiale

Nous avons donc décidé d'appliquer nos retouches de cellules sur l'agencement présentant le meilleur EVM moyenne initiale pour pouvoir atteindre les meilleures performances finales au regard de ce critère. Nous allons maintenant nous intéresser aux atténuateurs.

En se basant sur l'EVM, nous pouvons d'ores et déjà éliminer l'ordre selon l'amplitude qui présente une EVM moyenne initiale déjà trop élevée. Ensuite nous avons vu précédemment que les deux autres agencements de cellules n'était en réalité qu'un seul ordre qui avait été retourné selon une symétrie

verticale. Nous nous intéresserons donc qu'à l'ordre présentant le meilleure EVM moyenne initiale lors de notre étude des améliorations. Les résultats sont regroupés dans le Tableau 3.26.

Améliorations	Adapt IN (%)	Adapt OUT (%)	IN/OUT (%)	Amp (%)
0,5 dB	34,61	37,59	11,19	-0,25
1 dB	9,99	37,59	37,41	2,69
2 dB	28,88	0	28,88	0,25
4 dB	37,59	28,88	47,85	4,15
8 dB	0	9,56	9,54	-0,25-4,2

TABLE 3.26 – Evolutions des performances en termes d'EVM moyenne de l'ordre des cellules d'atténuation présentant la meilleure EVM initiale

Nous avons choisi de ne pas idéaliser l'erreur de phase puisque celle-ci a déjà été optimisée au sein des cellules individuelles. Les évolutions des performances ne renseignent aucune amélioration notable de l'EVM moyenne. Il n'apparaît donc pas pertinent d'effectuer une retouche des cellules. Nous choisirons l'agencement de cellules avec la meilleure EVM moyenne pour les intégrer dans le *core-chip* différentiel optimisé final.

Concernant les cellules de déphasage, les idéalizations ne permettent malheureusement pas d'améliorer l'EVM moyenne de façon significative. Après avoir tenté quelques retouches non fructueuses de cellules, sur les seuls paramètres permettant de gagner jusqu'à 3% sur l'EVM moyenne, nous avons choisi de laisser les cellules telles quelles et de les comparer avec les performances obtenues avec l'ordre utilisé dans la version sans algorithme.

La Figure 3.33 reporte les performances EM des meilleurs ordres d'atténuateurs et déphaseurs mis bout à bout. Nous pouvons d'ores et déjà éliminer l'agencement déphaseurs puis atténuateurs qui présente une erreur RMS de phase bien trop importante sur toute la bande ($>11^\circ$ RMS). En revanche le second agencement présente une bonne erreur de phase RMS ($>9^\circ$ de 30 à 37 GHz). Cependant l'erreur RMS d'amplitude est plus importante. En examinant les constellations, nous ne pouvons pas faire de lien direct entre les performances et les erreurs RMS fonctionnelles. En effet au regard des erreurs RMS fonctionnelles à 35 GHz la diminution de l'erreur de phase pour le second agencement n'apporte pas l'amélioration attendue sur la constellation du fait de l'erreur RMS d'amplitude qui a un poids plus important. Nous avons quand même décidé d'envoyer cette version de *core-chip* en fabrication car même si les erreurs RMS fonctionnelles sont importantes, les différentes constellations produites par cet agencement présentent des couvertures presque aussi bonnes que celles de la Figure 3.30. En plus de ce circuit nous avons ajouté une version corrigée du *core-chip* différentiel à l'organisation de cellule manuelle.

Nous avons donc choisi d'utiliser le même ordre que celui de la version sans algorithme mais cette fois ci avec le problème de polarisation corrigé. Les résultats de simulation EM obtenus sont reportés en Figure 3.34. Les performances sont moins bonnes que pour la première version (Figure 3.30) en termes d'erreurs fonctionnelles, avec une erreur de phase presque toujours supérieure à 7° RMS sur la bande. L'erreur d'amplitude est quant à elle légèrement supérieure à celle de la première version. Concernant la constellation, la différence de zones non couvertes et de recouvrement n'est pas flagrante cependant les niveaux d'amplitude les plus bas sont beaucoup moins bien alignés pour la dernière version. Cette différence doit donc logiquement se traduire par une ou plusieurs zones non couvertes. Pour les performances en bord de bande, encore une fois la couverture des constellations ne suit pas l'évolution des erreurs RMS fonctionnelles. La modification des cellules pour corriger les défauts de polarisation a donc impacté négativement les performances système de par une complexité de circuit

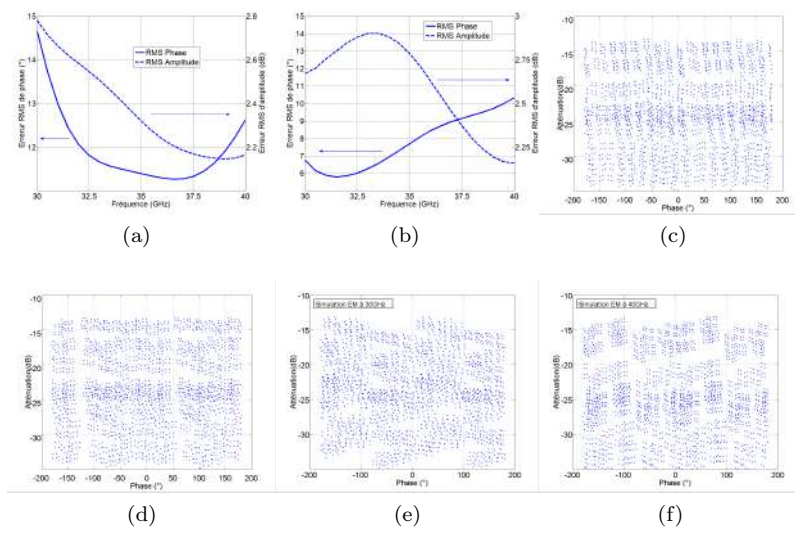


FIGURE 3.33 – Performances en simulations EM du core-chip lors de la mise en cascade des cellules en erreurs fonctionnelles RMS (a) déphaseurs puis atténuateurs (b) atténuateurs puis déphaseurs et en constellation phase amplitude à 35GHz des (c) déphaseurs puis atténuateurs et (d) atténuateurs puis déphaseurs. Les cas (e) et (f) correspondent aux constellations en bord de bandes, 30 et 40GHz de l'agencement atténuateurs puis déphaseurs

plus grande et une mise à défaut d'une symétrie de nos structures en version différentielle.

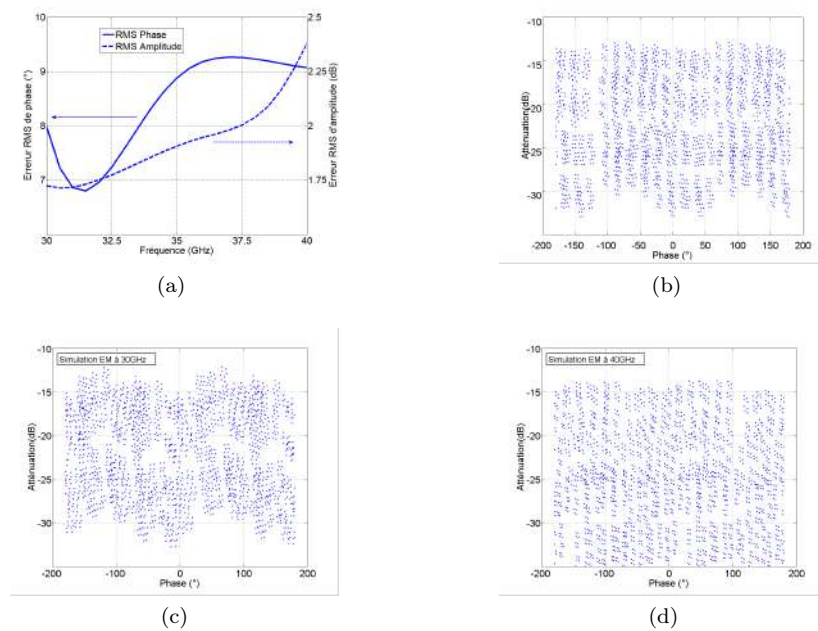


FIGURE 3.34 – Performances en simulations EM du core-chip corrigé envoyé en fabrication (a) en erreurs fonctionnelles RMS (b) en constellation phase amplitude à 35GHz et en (c) et (d) les constellations à 30 et 40GHz

Pour conclure ce chapitre, l'utilisation de l'algorithme en simulation ne nous a pas permis d'at-

teindre le niveau de performances obtenu grâce à l'agencement manuel des cellules que nous avons proposé. En effet les améliorations suggérées par l'algorithme n'étaient pas suffisantes ou alors impossibles à atteindre en retouchant les cellules. Cependant la plus grosse lacune de l'algorithme comme pour la partie *single-ended* vient de la limitation du nombre de cellules simulables du fait du temps de calcul impliqué. Cette limitation ne vient donc pas de l'algorithme en lui-même mais d'une réduction drastique du nombre d'agencements simulés ($11!$ à $5!$ et $6!$). Le dénombrement que nous faisons ne représente qu'une très petite partie de l'ensemble des solutions possibles. Les alternatives peuvent être par exemple d'appliquer l'algorithme aux cellules influençant le plus les performances systèmes et dans un second temps d'intégrer les cellules moins critiques comme l'explique la Figure 3.35. Malheureusement, nous n'avons pas eu le temps d'expérimenter cette possibilité.

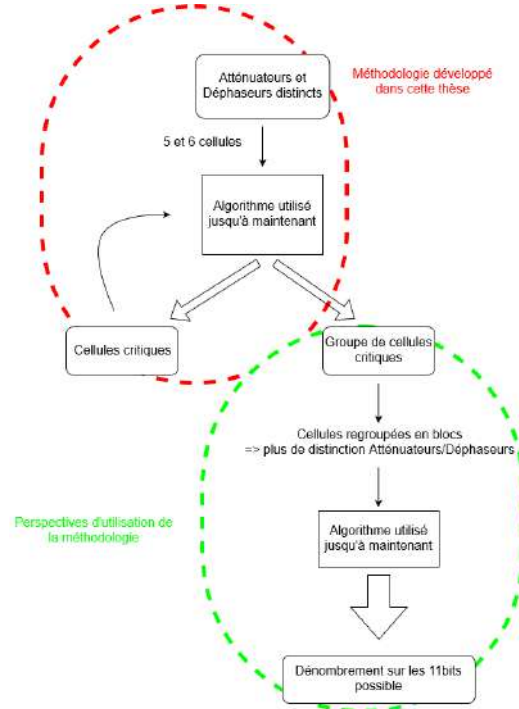
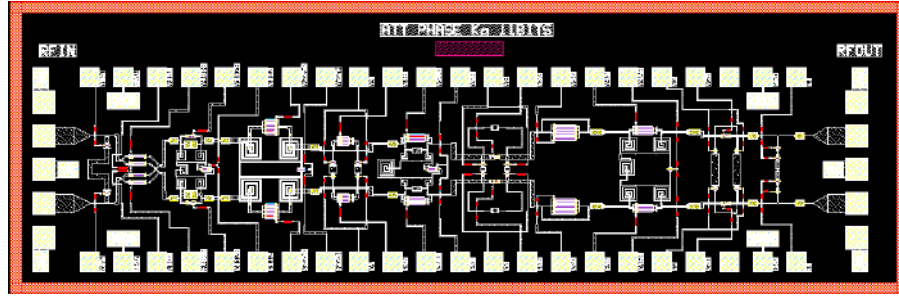


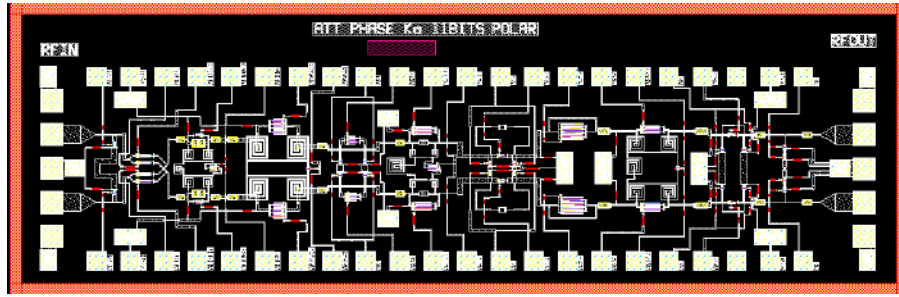
FIGURE 3.35 – Schématisation de la méthodologie utilisée et présentation des perspectives

La Figure 3.36 illustre les circuits différentiels envoyés en fabrication et comme nous pouvons le voir, ils sont assez similaires même si le second circuit est beaucoup plus dense du fait de l'ajout des nombreux *via-holes* assurant la correction de la polarisation des transistors. Les trois puces occupent toutes les deux une surface de $4 \times 1.3 \text{ mm}^2$.

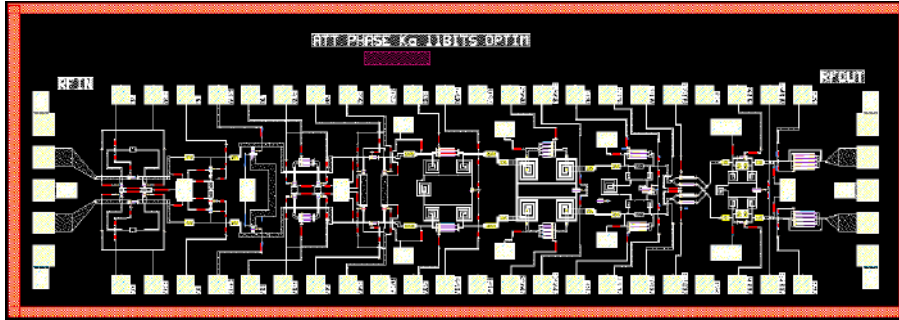
Pour conclure sur les performances des core-chips différentiels, nous pouvons dire que nous avons exploré toutes les possibilités de réalisation. En effet le premier core-chip réalisé a permis d'identifier les problèmes associés à la conception. Ce qui nous a permis d'une part de corriger les soucis de polarisation et d'autre part de comparer les performances atteignables avec ou sans l'utilisation de l'algorithme. Même si les résultats de simulation des deux derniers circuits ne sont pas à la hauteur de ceux du premier, les renseignements recueillis lors de ces différentes phases d'optimisation sont précieux. En effet nous avons pu mettre en évidence la pertinence du critère EVM moyenne choisi et sa bonne corrélation avec les performances en erreurs RMS traditionnellement utilisées pour caractériser les core-chips.



(a)



(b)



(c)

FIGURE 3.36 – Dessins des masques des trois core-chips différentiels envoyés en fabrication (a) première version (b) seconde version avec la polarisation corrigée et (c) troisième version avec l'ordre optimisé par algorithme

De plus face aux limitations de l'algorithme, nous avons pensé à une méthode d'intégration par blocs qui permettrait de simuler les 11 cellules simultanément. L'algorithme ayant montré ses potentialités, l'utiliser sans limitation garantirait sans conteste une hausse des performances.

Pour conclure de manière plus générale ce chapitre consacré à la mise en commun des cellules, plusieurs constats sont réalisés. D'une part, la cohérence entre les simulations électriques et EM nous a porté préjudice lors des étapes d'optimisation des cellules. En effet certaines performances obtenues en simulation électriques se révèlent inatteignables en EM. Il faudrait donc statuer sur la confiance à porter à chacune de ces simulations tout en proposant un processus itératif nous permettant de conserver des résultats de simulation plus constant tout au long des étapes de la méthodologie de conception suivies.

Concernant l'utilisation de l'algorithme, une approche plus astucieuse de regroupement des blocs

nous aurait très certainement permis d'obtenir des améliorations de performances notables. Ainsi nous ne pouvons pas juger de l'apport réel de ce dernier à la conception en regardant uniquement ce travail de thèse. Il faut plutôt le voir comme un outil d'aide à la conception permettant d'identifier rapidement les cellules ou groupes de cellules critiques et donc de cibler les aménagements topologiques nécessaires.

Concernant les deux approches utilisées (mono-fréquence et large bande) nous ne pouvons pas tirer de conclusions générales puisque comme nous avons pu le voir, la mise en commun des cellules à fortement influencée les performances que ce soit dans un mode ou l'autre. Peut être que plus de rigueur sur les erreurs croisées des cellules dès la première étape de conception aurait été bénéfique.

Après cette partie consacrée aux simulations, le chapitre suivant traitera des mesures des 5 différents circuits envoyés en fabrication. Il présentera donc l'environnement de mesures, les résultats obtenues et se terminera par une analyse globale des résultats et des différentes méthodes suivies durant ce travail de thèse et d'éventuelles perspectives pour compléter et parfaire ces travaux.

Bibliographie

- [1] Matthew A. Morton, Jonathan P. Comeau, John D. Cressler, Mark Mitchell, and John Papapolymmerou. Sources of Phase Error and Design Considerations for Silicon-Based Monolithic High-Pass/Low-Pass Microwave Phase Shifters. *IEEE Transactions on Microwave Theory and Techniques*, 54(12) :4032–4040, dec 2006.

Chapitre 4

Mesures des puces multifonctions

Dans ce chapitre nous allons présenter les résultats de mesures des 5 différents circuits fabriqués. Nous détaillerons tout d’abord l’environnement et les conditions de mesures. Puis nous séparerons les résultats en présentant les performances des deux *core-chips single-ended* en premier lieu pour finir sur les trois circuits différentiels. Enfin nous concluons par l’analyse de ces résultats et d’une brève comparaison entre *single-ended* et différentiel et entre circuit optimisé ou non.

1 Environnement de mesures

Classiquement la polarisation des *core-chips* est assurée par des module SIPO (Serial Input Parallel Output), ceux-ci reçoivent un mot binaire en entrée et répartissent des tensions parallèles pour polariser chaque cellule du circuit. Ce type de module sera utilisé dans les versions finales de *core-chips* cependant dans notre cas nous avons dû réaliser « manuellement » la polarisation pour pouvoir balayer tous les états possibles proposés par le circuit. Pour arriver à ce résultat, nous avons utilisé un Arduino Méga qui nous a permis d’effectuer la conversion série/parallèle. En effet un ordinateur écrit sur le port série de l’Arduino, ce dernier interprète la valeur de l’entrée et la traduit par des niveaux logiques sur ses sorties numériques. Les sorties de l’Arduino étant limitées à 0 et 5 V, l’ajout d’un circuit réalisant la transformation de ces tensions était nécessaire. En effet pour rappel, les tensions de polarisations nécessaires aux *core-chips* sont -10 V, -3V et 0V respectivement pour les circuits *single-ended* et différentiels. Pour réaliser la transformation de tension nous avons utilisé un circuit à base d’optocoupleurs, ceux-ci présentent en sortie la tension à laquelle ils sont polarisés en fonction de la valeur de la tension d’entrée. Ainsi si nous voulons transformer le couple de tensions [0,5] V en [0, -10] V, il suffit de polariser l’optocoupleur à -10V et selon que la tension d’entrée vaut 0 ou 5V il présentera en sortie respectivement 0 ou -10V. Cette solution est l’une des plus simple dans notre cas d’utilisation puisque nous n’allons pas faire de mesures rapides du fait de la vitesse d’acquisition de l’Analyseur de Réseaux Vectoriels (ARV). Le schéma de la Figure 4.1 explique le fonctionnement de ces deux éléments.

Le mot de N bits (N étant le nombre de cellule à piloter) est envoyé par un PC sur le port d’entrée série de l’Arduino. Comme nous avons pu voir sur les dessins des masques des *core-chips* précédemment présentés, les cellules ont généralement deux tensions de polarisation. Ces deux tensions sont complémentaires mais nécessitent chacune une sortie de l’Arduino. Ainsi, bien que le mot d’entrée soit composé de N bits, le nombre de sorties parallèles nécessaires pour la polarisation est supérieur à N et tend vers $N*2$, même si certaines cellules ne nécessitent qu’une seule tension (0.5dB, 1dB et

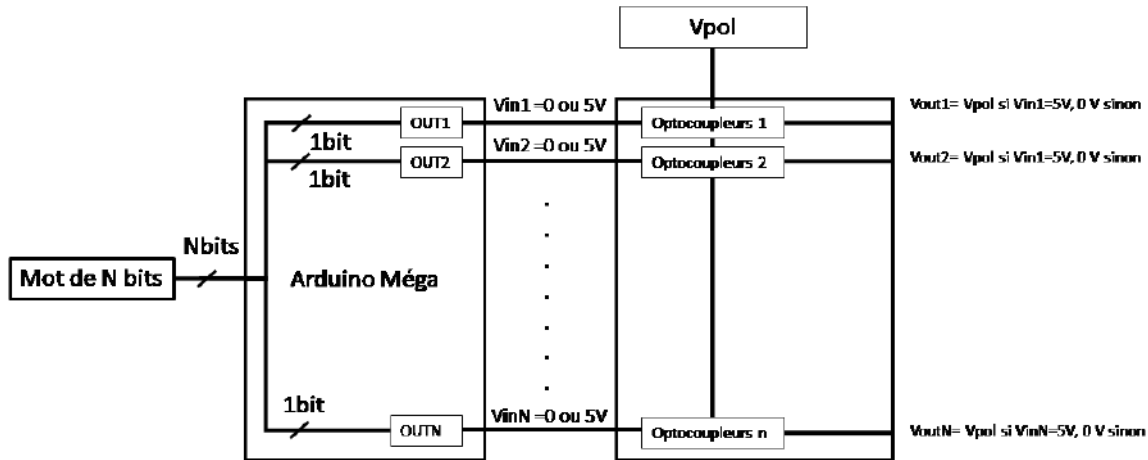


FIGURE 4.1 – Schéma de fonctionnement du système de polarisation

5.625°). A l'intérieur de l'Arduino, le code est simple et détaillé plus tard dans ce paragraphe. Le but est donc d'acquérir un nombre binaire entre 0 et $2N-1$ qui traduit l'état de chaque cellule du *core-chip*. Par suite de cette acquisition, le programme va venir scruter les états logiques bit par bit. Nous avons besoin de choisir une convention qui définit si l'état de référence d'une cellule (où elle ne modifie pas le signal) correspond à un bit haut ou bas. Une fois que cela est défini, nous pouvons écrire la condition d'activation des sorties numériques. Ainsi, il y a bien une entrée binaire associée à une ou deux sorties parallèles. Le code Arduino est réalisé pour chaque bit d'entrée comme dans l'exemple suivant :

```
if (bitRead(I,0) >0)
    {digitalWrite(43, HIGH); }
else
    {digitalWrite(43, LOW); }
```

Dans le code ci-dessus, nous regardons l'état du bit 0, s'il est supérieur à 0 (égal à 1), nous mettons la sortie 43 à l'état haut (5V) sinon nous la mettons à l'état bas (0V). Nous pouvons comprendre que ce bit correspond à une cellule n'ayant besoin que d'une seule tension pour fonctionner. Le code suivant explicite le cas d'une cellule nécessitant plusieurs tensions :

```
if (bitRead(I,1) >0)
    {digitalWrite(34, HIGH);
    digitalWrite(24, LOW);
    }
else
    {digitalWrite(34, LOW);
    digitalWrite(24, HIGH);
    }
```

Cette fois ci, nous regardons l'état du bit 1 et nous mettons les sorties 34 et 24 à l'état haut et bas en conséquence. Nous pouvons voir que ces sorties ne sont jamais simultanément dans le même état, elles sont donc bien complémentaires.

Nous effectuons ce procédé pour tous les bits du mot d'entrée. Il est très facile de vérifier le bon fonctionnement en utilisation le moniteur série disponible dans la bibliothèque Arduino. Avec ce dernier nous envoyons un mot sur l'entrée série de l'Arduino et pour vérifier le bon fonctionnement du code nous pouvons, soit afficher la valeur des sorties sur le moniteur, soit aller directement mesurer les niveaux de tensions de sortie avec un voltmètre.

Comme nous l'avons dit précédemment, les niveaux de sortie de l'Arduino sont de 0 et 5V. Il nous faut donc connecter l'Arduino à un PCB sur lequel sont répartis les optocoupleurs assurant la mise à niveau des tensions. Il y a autant d'optocoupleurs que de sortie actives de l'Arduino. Le PCB utilisé est illustré en Figure 4.2.

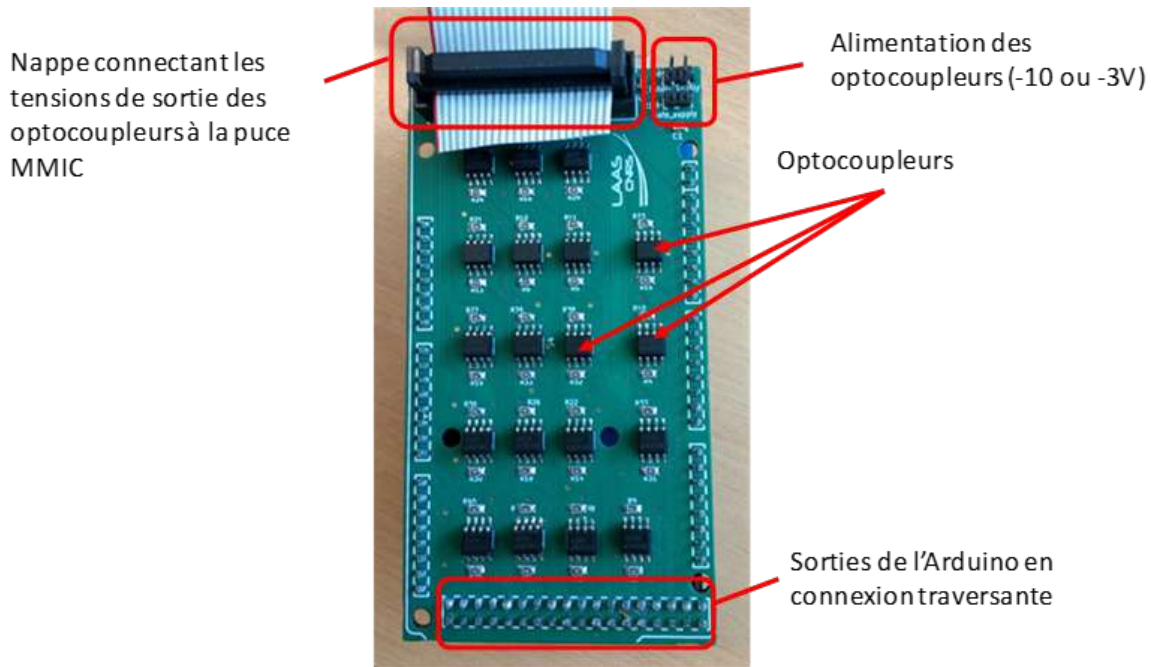


FIGURE 4.2 – Schéma du PCB assurant la mise à niveau des tensions de polarisation

Nous pouvons compter 19 optocoupleurs, nombre qui correspond au nombre maximal de tensions de polarisation de chaque core-chip. En effet, ce PCB doit être utilisable pour toutes les puces single-ended ou différentiel, première ou seconde version. La seule différence entre les deux modes de fonctionnement étant la tension de polarisation des optocoupleurs, réglable à l'aide d'une alimentation DC externe. La nappe de sortie du PCB est connectée à un second PCB sur lequel est collé la puce MMIC. Cette dernière repose sur un creuset de la résine de passivation de surface, l'adhérence est réalisée grâce à une colle conductrice, ainsi la face arrière du MMIC est connecté à la masse du PCB. Pour réaliser la connexion entre la nappe et les plots de polarisation du MMIC, nous avons tout d'abord approcher des pistes DC au plus près des accès du MMIC et la connexion électrique a été réalisée par fils de *bonding*.

Dans les conditions de mesures réelles, le mot binaire est envoyé par un PC. Celui-ci réalise l'interface logiciel entre l'Arduino et l'ARV. En effet grâce à un code CVI propre à l'ARV, nous avons réalisé un code permettant à la fois d'écrire sur le port d'entrée série de l'Arduino et de lancer les mesures en contrôlant l'ARV. Ainsi en réalisant une boucle au nombre d'itérations égal au nombre de mesures que nous voulons faire, nous pouvons mesurer tous les états qui nous intéressent. Nous ne présenterons pas

le code CVI de contrôle de l'ARV, d'une part car il a été réalisé par un ingénieur logiciel et d'autre part car ce n'est qu'un simple outil de mesure et nous ne considérons pas que sa présence soit pertinente dans ce manuscrit.

Le schéma récapitulatif de mesure est donné en Figure 4.3, nous pouvons bien y voir que le PC assure la synchronisation entre les mesures RF et la polarisation du MMIC via l'Arduino. Dans le cas de nos mesures, nous effectuons 2047 itérations correspondant chacune à un état du core-chip, le PC récupère ensuite les fichiers de paramètres S mesurés et les sauvegarde dans un espace dédié. Selon la temporisation entre itérations, la mesure de tous les états dure de 3 à 4h. L'ARV utilisé est le PNA-X 4 ports N5247A de Keysight. Concernant l'ARV nous avons utilisé une calibration SOLT (Short Open Load Thru). Nous avons obtenu une précision de phase de 0.5° et une précision d'atténuation de 0.1dB grâce à cette calibration. Comme nous l'avons dit précédemment c'est un ARV 4 ports, pour les mesures single-ended où nous n'utilisons que 2 ports. Cependant pour les mesures des circuits différentiels, nous n'avons pas réalisé de mesures purement différentielles. En effet nous avons réalisé des mesures single-ended dans un environnement différentiel et à partir des résultats nous avons calculé les paramètres S différentiels. Cette méthode d'extraction utilisant les paramètres S mixtes est détaillée dans [1].

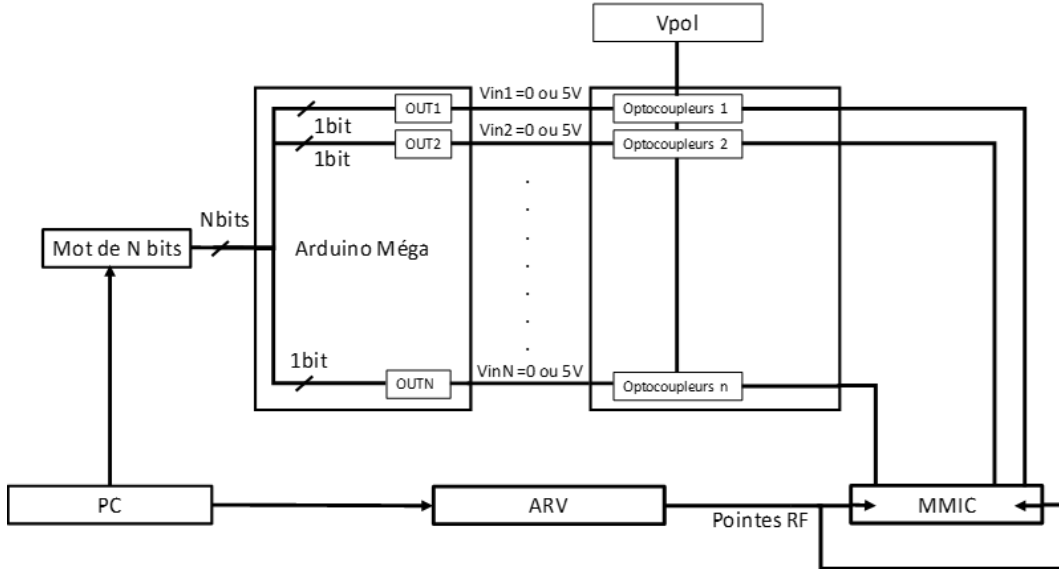


FIGURE 4.3 – Schéma de la procédure de mesures des core-chips

Les pointes utilisées sont respectivement des pointes GSG et GSGSG pour les mesures single-ended et différentielles respectivement. Ces pointes présentent toutes les deux un écartement de $150\mu\text{m}$.

Nous allons maintenant présenter les résultats de mesures des différents circuits réalisés.

2 Circuits single-ended

Durant ce travail de thèse nous avons réalisé deux circuits *single-ended*, une première version avec des performances centrées à 35GHz et une seconde version où les problèmes de polarisation du premier essai ont été corrigés. Les deux circuits fabriqués sont présentés en Figure 4.4. Les tensions de commutations statiques de cellules sont de 0V et -10V.

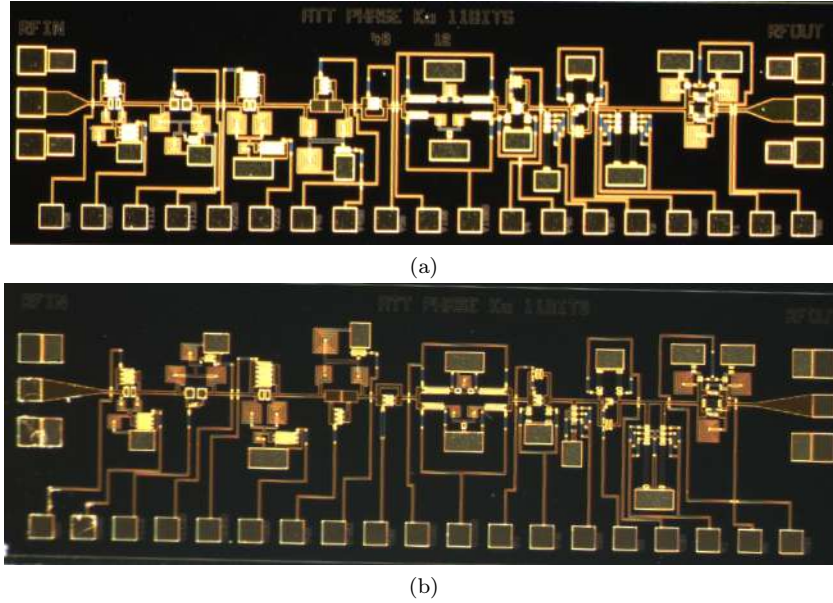


FIGURE 4.4 – Core-chips single-ended fabriqués (a) première version 3.13x1mm² (b) version avec la polarisation corrigée 3.2x1.1mm²

2.1 Première version

Les mesures sont présentées en Figure 4.5. Les résultats de simulations sont différents de ceux présentés dans le chapitre 3 car ils sont simulés avec tous les pads d'accès RF et DC.

Nous avons ensuite comparé les erreurs de mesures aux erreurs de simulation en Figure 4.6. Nous pouvons voir qu'à 35 GHz la zone non couverte double de la simulation (12.5° à 25°). Nous obtenons donc une dynamique d'atténuation de 20dB et un pas de phase de 25°.

Bien que les profils d'erreurs RMS et les constellations semblent indiquer un meilleur fonctionnement à 35GHz, les performances sont en dessous de nos attentes et complètement en dehors des exigences du cahier des charges. Les erreurs associées aux états principaux sont regroupées en Figure 4.7. Les erreurs sont identifiées par du vert pour celles qui sont satisfaisantes, en jaune celles qui sont acceptables et en rouge pour celles qui sont trop importantes par rapport à la résolution du système. Nous pouvons voir que la cellule 45° présente une erreur bien trop importante. Nous avons rétro-simulé cette cellule mais nous n'avons jusqu'alors pas trouvé l'origine de cette variation de phase. Cette version présente aussi des problèmes de polarisation statique même si toutes les cellules ont l'air de s'activer.

2.2 Deuxième version

Sur ce second circuit nous avons corrigé les problèmes de polarisation de la première version. Les mesures sont présentées en Figure 4.8. Nous les avons comparé sur la même figure aux performances de simulations. Malgré un meilleur fonctionnement à 35 GHz les différences entre simulations et mesures sont trop importantes pour tirer des conclusions sur les origines des différences. De plus cette seconde version ne semble pas corriger les erreurs du premier circuit puisque les performances sont au contraire moins bonnes.

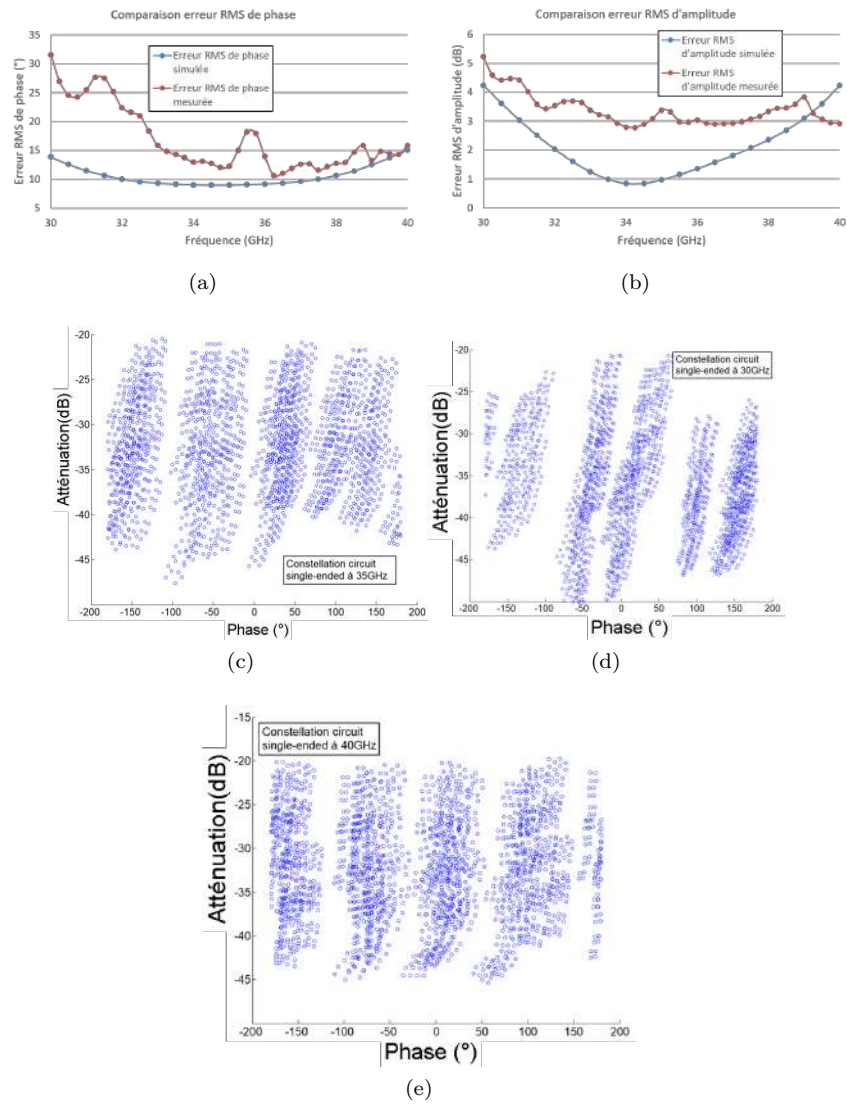


FIGURE 4.5 – Comparaisons erreurs RMS fonctionnelles (a) de phase (b) d'amplitude et constellation à (c) 35GHz (d) 30GHz et (e) 40GHz

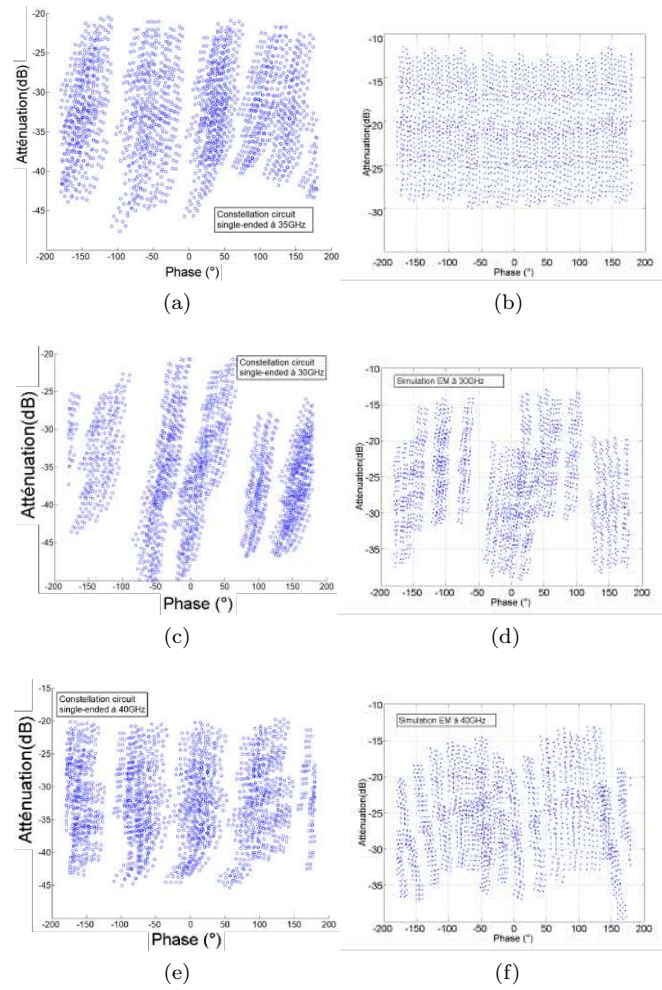


FIGURE 4.6 – Comparaisons des erreurs sur les constellations entre la simulation à droite et la mesure à gauche (a-b) 35 GHz (c-d) 30 GHz (e-f) 40GHz

Cellule	All OFF	5.625°	11.25°	22.5°	45°	90°	180°	0.5dB	1dB	2dB	4dB	8dB	All ON
Phase (°)	44	39.24	57	62.19	105.8	141.70	-142.2	44.64	40.26	37.48	39.98	39.37	22.5
Atténuation (dB)	-16.16	-15.70	-16.56	-15.54	-16.75	-17.61	-15.70	-17	-17.58	-19.19	-20.77	-24.18	-36.43
Erreur phase (°)	0	0.865	1.75	4.31	16.1	27	67	0.64	3.74	6.53	4.02	4.63	13.67
Erreur amp (dB)	0	0.46	0.5	0.62	0.64	1.45	0.46	0.34	0.42	1.03	0.62	0.02	2.77

FIGURE 4.7 – Tableau récapitulatif des erreurs associées aux états principaux de phase et d'atténuation

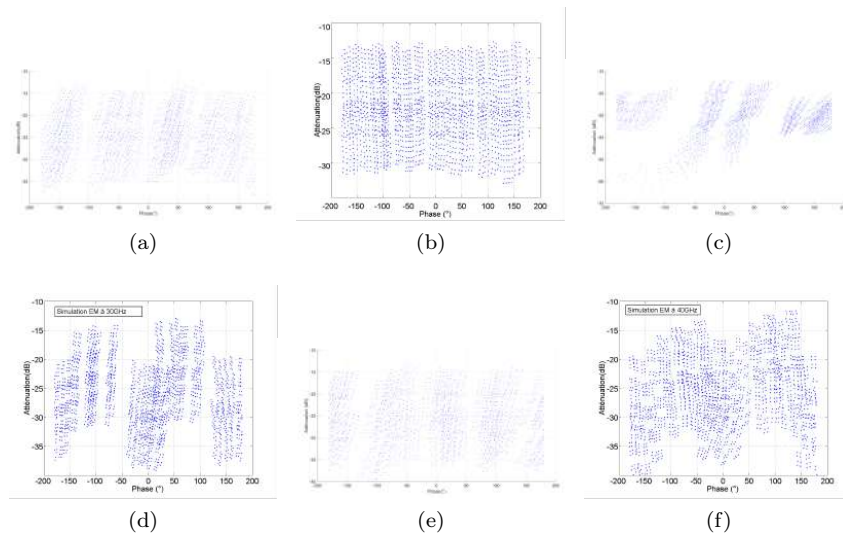
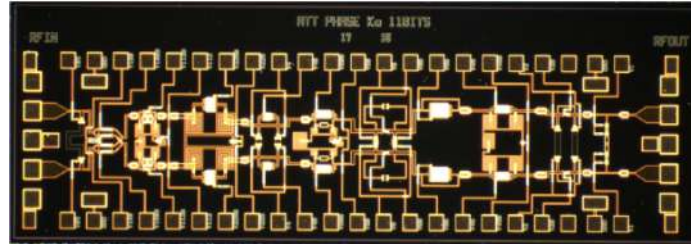


FIGURE 4.8 – Comparaisons des erreurs sur les constellations entre la simulation à droite et la mesure à gauche (a-b) 35 GHz (c-d) 30 GHz (e-f) 40GHz

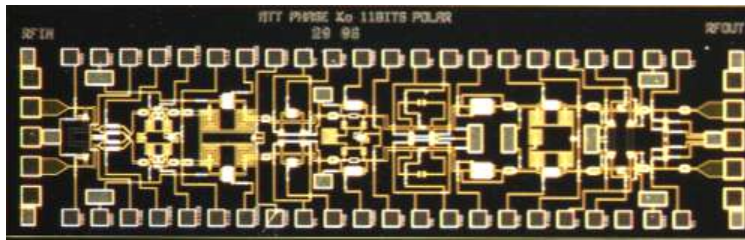
3 Circuits différentiels

Nous avons réalisé 3 circuits différentiels, tous les trois fonctionnant entre 30 et 40 GHz. Une première version où nous avons optimisé l'ordre des cellules manuellement, une seconde version du même circuit où les problèmes de polarisation ont été corrigés et enfin une dernière version à l'ordre des cellules optimisé par l'algorithme. Les conditions de mesures ne changent pas entre ces trois versions. Les 3 différents circuits sont présentés en Figure 4.9.

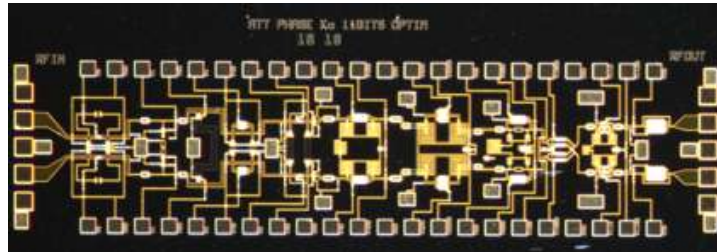
Comme expliqué précédemment nous n'avons pas effectué de vraies mesures en mode différentiel. Nous avons plutôt fait des mesures *single ended* 4 ports. Ce qui veut dire que seulement 2 ports de l'ARV étaient activés (envoi de puissance au port d'entrée/sortie du circuit) simultanément. Pendant ce temps d'activation, les deux autres ports présentaient des charges 50Ω au ports d'accès du circuits non activés. En effectuant ces mesures entre chaque port, nous obtenons une matrice 4×4 de paramètres S *single-ended*. A partir de là, comme l'explique Bockelman [1] nous pouvons reconstituer les paramètres S différentiels du circuit.



(a)



(b)



(c)

FIGURE 4.9 – Core-chips différentiels fabriqués (a) première version $4 \times 1.3 \text{ mm}^2$ (b) version avec la polarisation corrigée $4 \times 1.3 \text{ mm}^2$ (c) version optimisée par algorithme $4 \times 1.3 \text{ mm}^2$

3.1 Première version

Nous avons d'abord testé le fonctionnement du circuit en mesurant les états principaux de phase et d'atténuation. Les résultats sont donnés en Figure 4.10. Il semble y avoir des cellules qui ne s'activent pas (11.25° et 1dB). Cette hypothèse est confirmée sur les Figure 4.10c et Figure 4.10d. En rétro-simulant le circuit, nous nous sommes aperçu qu'il y avait un problème de polarisation statique. Pour confirmer ce dysfonctionnement, nous avons simulé le circuit en DC en remplaçant le modèle du transistor en commutation par le modèle en amplification. Ce dernier permet d'évaluer les potentiels statiques aux accès. Il s'est avéré que les capacités se situant sur la voie RF empêchaient la propagation de la masse statique entre les cellules. En effet les potentiels aux accès des transistors ne permettaient pas de commuter entre un V_{gs} de 0 et -3V. La Figure 4.11 montre les capacités de liaison en entrée et sortie des cellules 11.25° et 1dB.

Malgré la découverte de ce problème, nous ne pouvons pas pleinement expliquer pourquoi les cellules 11.25° et 1dB spécifiquement ne s'activent pas. En effet lors de la rétro-simulation ce n'étaient pas les seules cellules présentant des problèmes de polarisation. Il est quand même compliqué de croire que des compensations internes permettent d'atteindre les autres niveaux de phase et d'atténuation dans le cas où d'autres cellules ne commutent pas.

Nous avons ensuite fait la mesure sur les 2048 états de fonctionnement. Les résultats sont donnés en Figure 4.12.

Nous constatons de grosses différences entre les simulations et les mesures. Ces variations sont d'autant plus marquées pour l'erreur d'amplitude. Sachant que la cellule 1dB ne commute pas, nous pouvons considérer qu'elle influe sur l'augmentation de l'erreur d'amplitude. Nous pouvons faire le même raisonnement entre l'erreur de phase et la cellule 11.25°, même si la différence est moins marquée pour les fréquences les plus hautes de la bande. Nous n'avons pas réussi à retrouver ces niveaux d'erreurs RMS en rétro-simulation. Il nous est donc difficile de déterminer l'origine des erreurs de phase et d'atténuation. En effet, même si certaines cellules commutent, il est possible que les tensions de polarisation ne soient pas égales au -3V et 0V requis. Nous avons pu voir en simulation qu'une variation de tension pouvait entraîner des variations de déphasage ou d'atténuation. Certes ces variations à l'échelle d'une cellule ne sont pas de l'ordre des erreurs RMS trouvées mais la sommation aléatoire d'erreurs peut grandement influencer l'erreur finale du système.

Les secondes versions des circuits devraient nous renseigner sur ce point puisque les polarisations ont été corrigées.

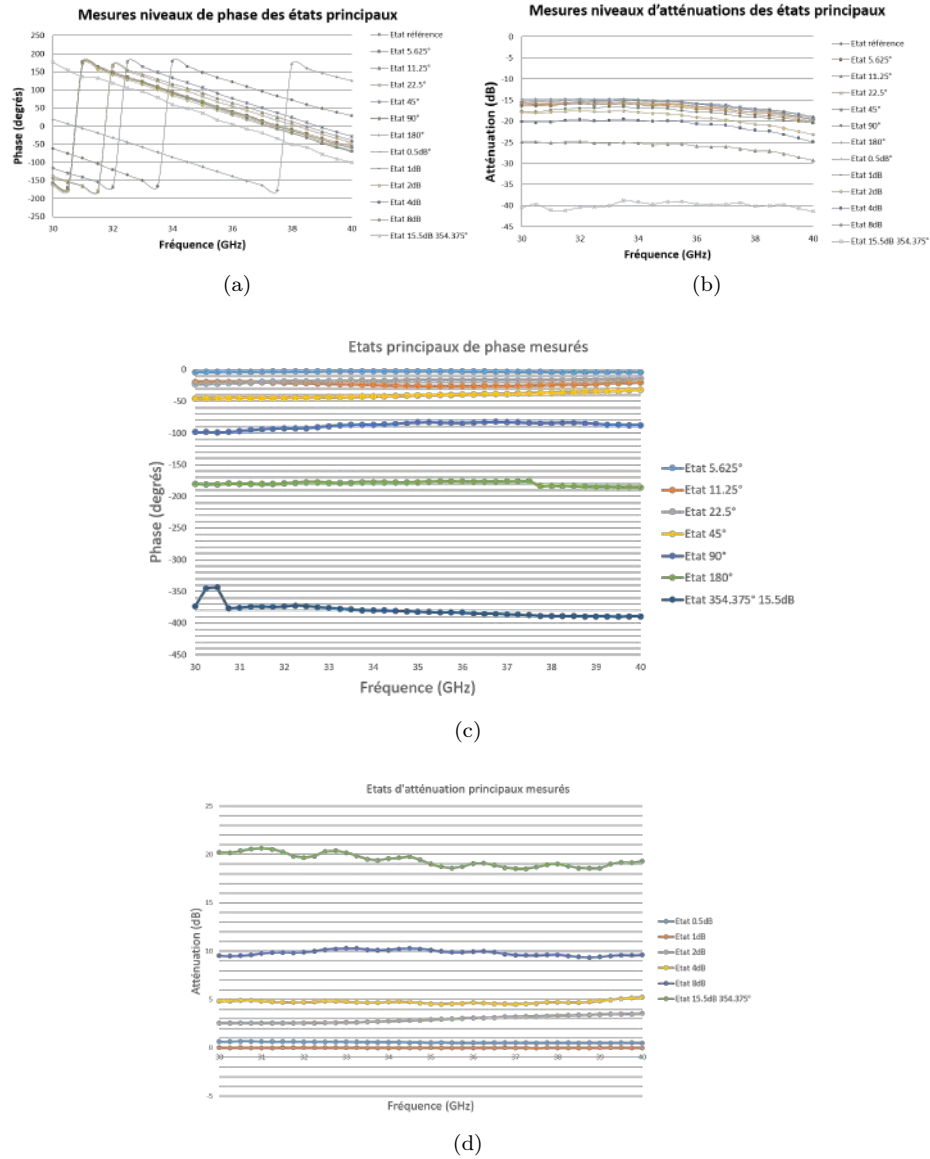


FIGURE 4.10 – Mesures des états principaux du core-chip (a) en phase (b) en amplitude et des états (c) de phase des déphaseurs et (d) d'amplitude des atténuateurs

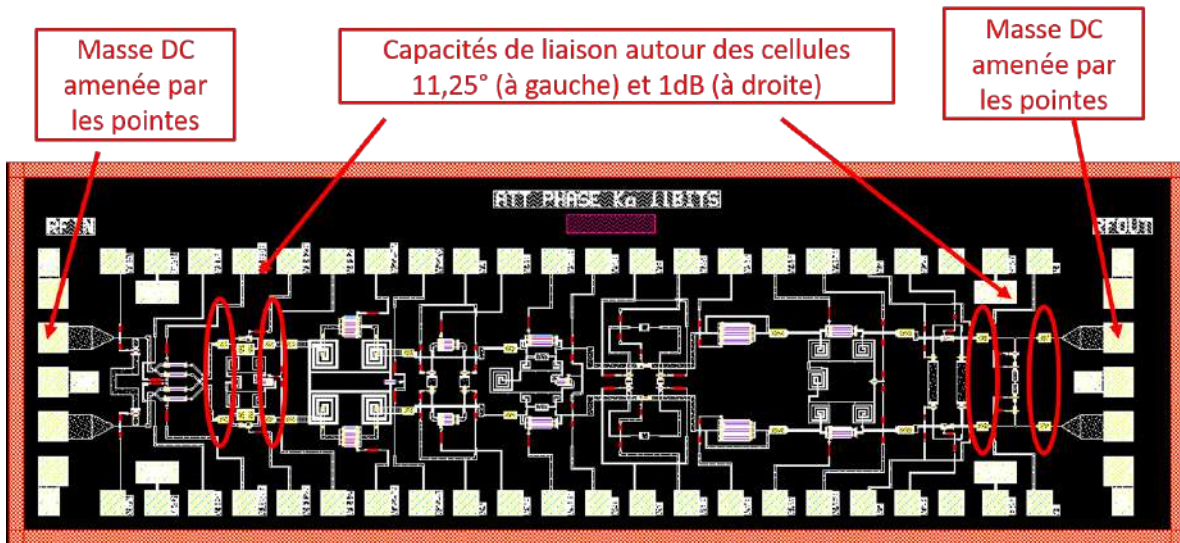


FIGURE 4.11 – Mise en évidence des défauts de polarisation

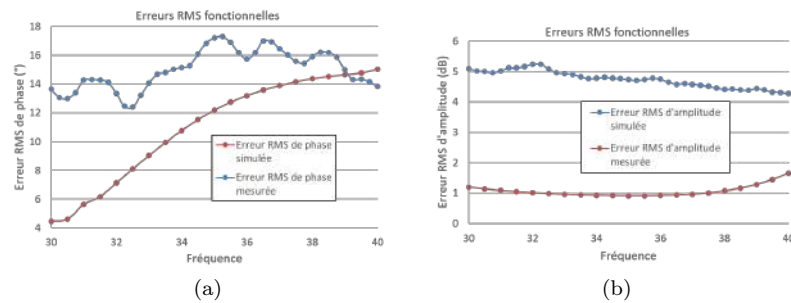


FIGURE 4.12 – Comparaisons erreurs RMS fonctionnelles (a) de phase (b) d'amplitude

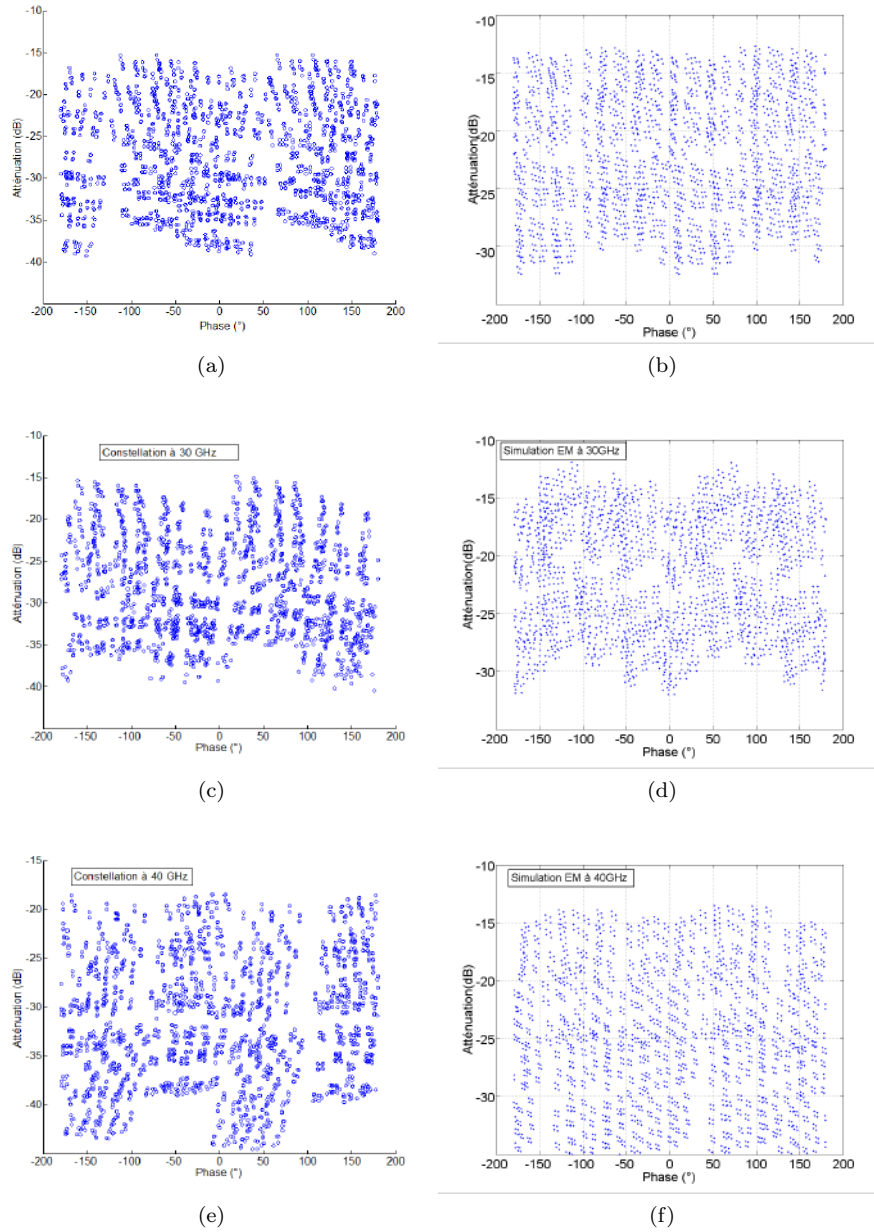


FIGURE 4.13 – Comparaisons des constellations en simulation à droite et mesure à gauche à (a-b) 35GHz (c-d) 30GHz et (e-f) 40GHz

Nous avons effectué des mesures en puissance sur ce circuit pour en déterminer la linéarité. Etant donné que nous mesurons les paramètres S petits signaux en utilisant les paramètres mixtes, nous avons du changer de méthode de mesure puisque ces derniers ne sont plus valables en grands signaux. Nous avons donc mesuré directement la linéarité sur une seule voie en terminant la seconde voie sur une charge 50Ω . Nous obtenons ainsi la une linéarité théorique que nous pouvons augmenter de 3dB puisque le signal sera divisé entre les deux voies en fonctionnement différentiel.

Nous avons ajouté un amplificateur sur la voie à mesurer. Celui ci présente les caractéristique présentées en Figure 4.14.

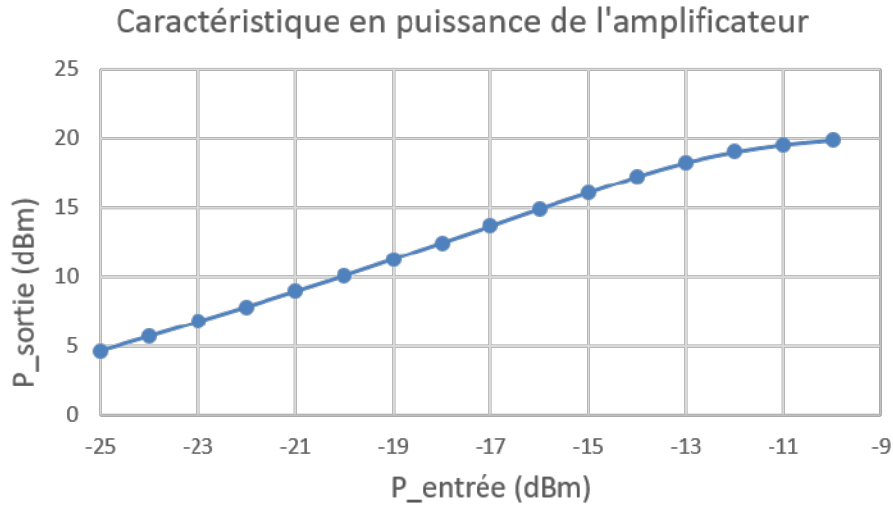


FIGURE 4.14 – Caractéristique en puissance de l'amplificateur utilisé

L'amplificateur présente un gain de 30dB jusqu'à une puissance de saturation de sortie mesurée de 20dBm. En retranchant les pertes des câbles et des pointes RF nous arrivons à une puissance injectable en entrée du circuit de 24dBm. Nous avons donc mesuré les pertes du circuit en fonction de la puissance d'entrée. Les résultats sont donnés en Figure 4.15. La bosse sur la courbe est due à une légère ondulation du gain de l'amplificateur. Nous pouvons voir que la linéarité est maintenue jusqu'à 24dBm, nous n'avons pas le matériel pour aller au delà en bande Ka. Nous pouvons donc supposer qu'en mode différentiel le circuit tolère une puissance d'entrée de 27dBm sans perte de linéarité. Ce seuil nous place au dessus des puissances de 20.5dBm tolérables en SiGe [2].

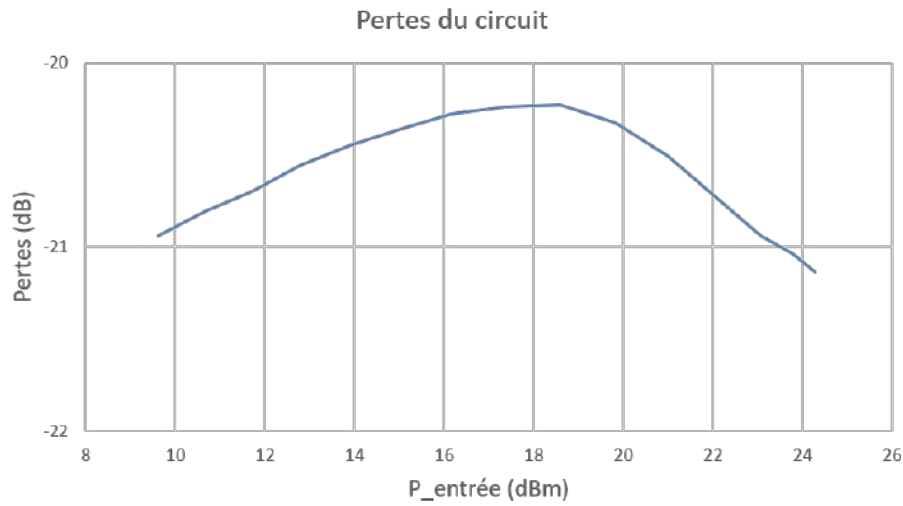


FIGURE 4.15 – Evolution des pertes du circuit selon la puissance d'entrée

3.2 Deuxième version et troisième version

La seconde version corrige les erreurs de polarisation du premier circuit différentiel. Les performances en mesure et comparées aux simulations sont présentées en Figure 4.17c. Cette correction n'apporte pas les améliorations attendues et présente même des performances en dessous de la première version. En effet les constellations sont complètement différentes des simulations puisque les points ont tendance à former des amas, tendance non visible en simulation. De plus le niveau de pertes supérieur à 25dB nous renseigne sur le non fonctionnement du système.

Pour la dernière version qui correspond à l'ordre optimisé par algorithme, les performances sont regroupées en Figure 4.17. Etant donné que la moitié des points sont dédoublés nous n'avons pas jugé pertinent de faire la comparaison aux simulations. La couverture reste correcte malgré une résolution en phase augmentée.

En Figure 4.18 est effectuée la comparaison entre les circuits *single-ended* et différentiel présentant les meilleures performances en mesures, ce sont pour chacun d'eux les premières versions.

Le circuit différentiel présente une meilleure couverture si l'on réduit la dynamique d'atténuation à 20dB et celle de phase à 22.5°. Pour le circuit *single-ended* cela tendrait plus vers des dynamiques de 20dB et 45°. Nous pouvons donc conclure au regard des performances obtenues sur les 5 différents circuits que le premier circuit différentiel offre les meilleures performances en termes de couverture.

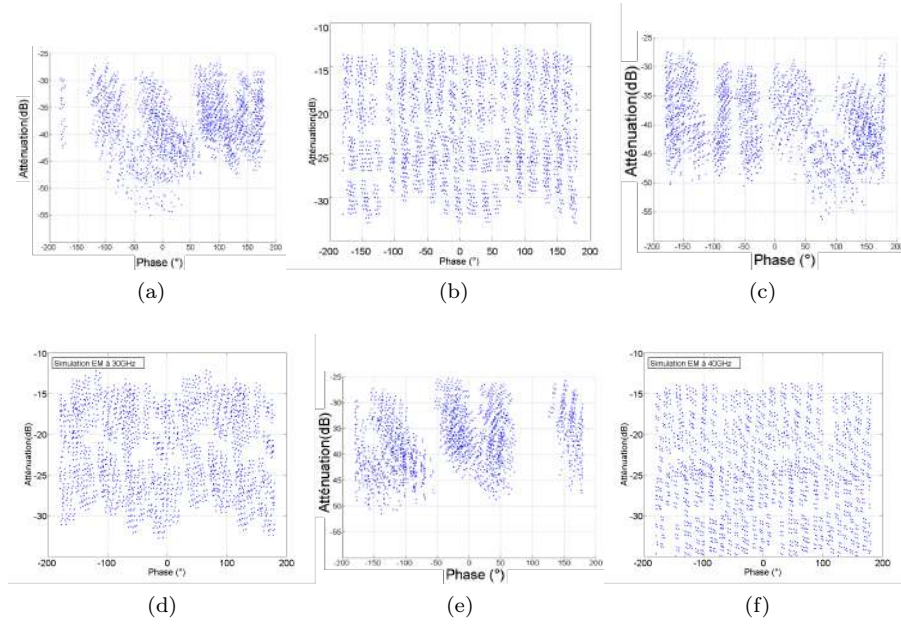


FIGURE 4.16 – Comparaisons des erreurs sur les constellations entre la simulation à droite et la mesure à gauche (a-b) 35 GHz (c-d) 30 GHz (e-f) 40GHz

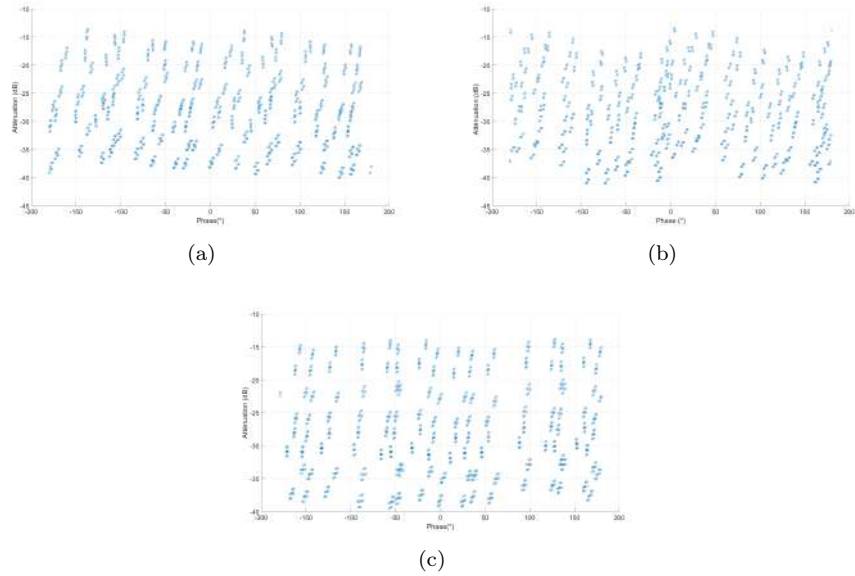


FIGURE 4.17 – Constellations du circuit différentiel optimisé par algorithme à (a) 35 GHz (b) 30 GHz (c) 40GHz

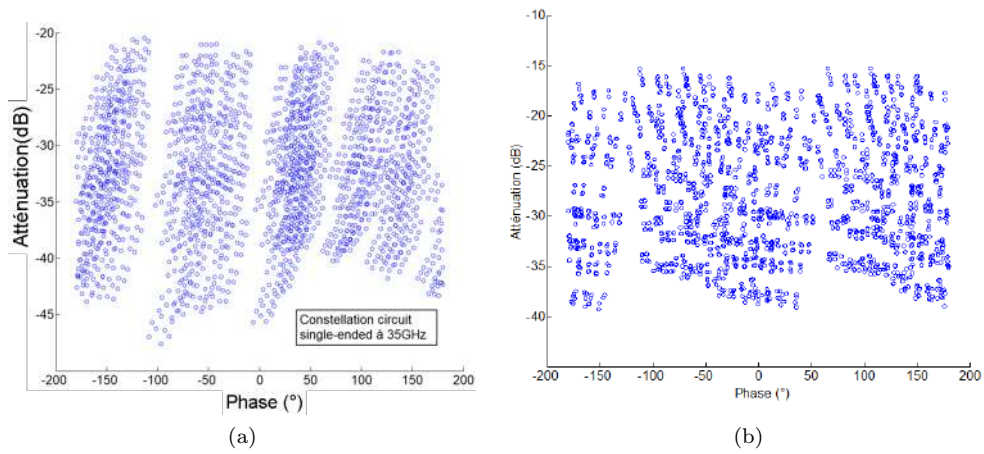


FIGURE 4.18 – Comparaisons des erreurs sur les constellations à 35 GHz entre les premières versions (a) single-ended et (b) différentielle

Pour conclure ce chapitre, nous pouvons dire que nos itérations successives n'ont pas porté leurs fruits. En effet les versions corrigées quelque soit le mode de fonctionnement présentent des performances inférieures à celles des versions initiales. Pour l'optimisation avec algorithme, nous supposons qu'il y a eu une erreur durant le processus puisque les simulations semblaient indiquer des performances bien meilleures. Les problèmes peuvent être attribuables à la maîtrise technologique du procédé de fabrication ou à un problème de modèle. Pour la première version nous avons inclus des masques de cellules élémentaires qui ont fonctionné en mesures, justifiant ainsi le questionnement sur la fiabilité des modèles et/ou de la maîtrise du procédé de fabrication. Malgré cela nous sommes optimistes pour la suite puisque nous avons identifié les verrous de conception et savons sur quels curseurs il faut jouer pour parvenir rapidement à un *core-chip* performant. Nous allons ensuite voir comment les circuits fabriqués peuvent être utilisés, la Figure 4.19 présente un exemple d'obtention d'une constellation 64 états à 35 GHz.

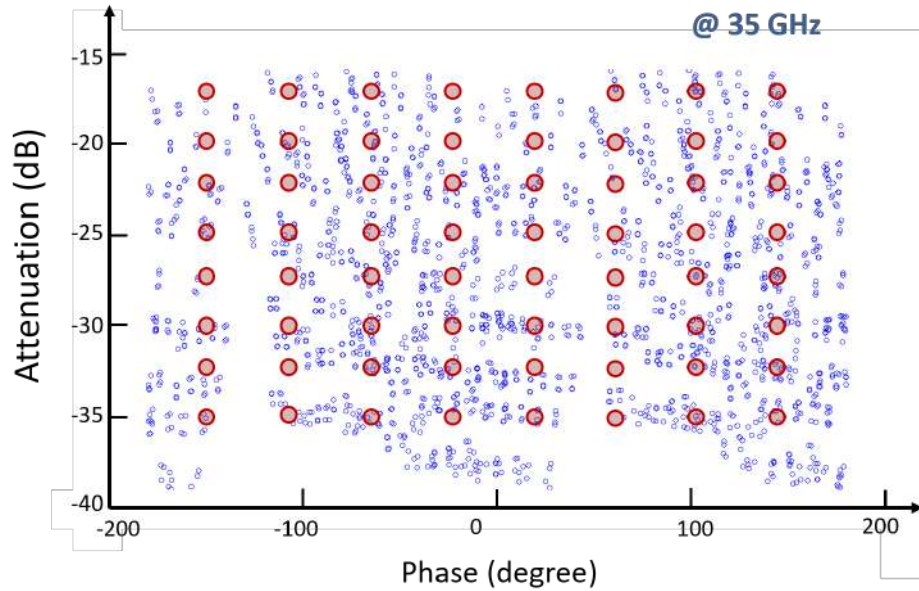


FIGURE 4.19 – Constellation 8x8 obtenue avec le circuit différentiel

Bibliographie

- [1] D.E. Bockelman and W.R. Eisenstadt. Combined differential and common-mode scattering parameters : theory and simulation. *IEEE Transactions on Microwave Theory and Techniques*, 43(7) :1530–1539, jul 1995.
- [2] Chao Liu, Qiang Li, Yihu Li, Xiao Dong Deng, Hailin Tang, Ruitao Wang, Haitao Liu, and Yong Zhong Xiong. A Ka-Band Single-Chip SiGe BiCMOS Phased-Array Transmit/Receive Front-End. *IEEE Transactions on Microwave Theory and Techniques*, 64(11) :3667–3677, nov 2016.

Conclusion générale

Cette thèse avait pour double objectif de démontrer la faisabilité de *core-chips* en GaN et d'identifier les topologies les plus adaptées à ce type de fonctionnalités. Les objectifs ont été atteints même si le travail de recherche s'est recentré vers l'élaboration d'une méthodologie de conception et d'intégration des cellules. Deux approches de conception sont proposées, une *single-ended* centrée à 35 GHz et l'autre différentielle sur la bande 30-40 GHz.

Dans le premier chapitre nous avons dressé un état de l'art des *core-chips* existant. L'absence du GaN dans cette revue de littérature confirme le côté novateur de ce projet. En effet l'utilisation de ce dernier pour réaliser des puces multi-fonctions MMIC permet de réaliser à plus longs termes des puces d'émissions réceptions tout en GaN, ce qui améliorerait grandement l'intégration. Nous explicitons ensuite notre méthodologie de conception dans sa globalité. Méthodologie nécessaire d'une part à cause des différences entre les éléments GaN et SiGe qui imposent des spécificités de conception et d'autre part pour exploiter pleinement les performances obtenues lors de la conception des cellules.

Le second chapitre présente les performances des cellules *single-ended* et différentielles obtenues. Nous pouvons déjà y voir plus clair concernant les différences de fonctionnement et de performances inhérentes à ces deux modes. Les performances de ces cellules sont comparées aux objectifs du cahier des charges puis utilisées comme briques de base dans l'utilisation de notre méthodologie d'intégration en vue de l'obtention des *core-chips* finaux.

Le troisième chapitre traite du plus gros du travail de recherche effectué durant cette thèse. Nous y présentons le processus d'optimisation de l'ordre des cellules pour les deux modes de fonctionnement. Ensuite nous évaluons les gains potentiels de performances suite à l'optimisation de certains paramètres S fonctionnels ou d'adaptation d'impédance. Dans l'affirmative nous tentons des retouches aux niveaux des cellules pour tendre vers ces performances idéalisées. Malgré le champ d'action réduit de notre algorithme (seulement 6 cellules simulables simultanément) nous avons tiré des conclusions pertinentes sur certaines combinaisons de cellules à favoriser ou à éviter. Ces hypothèses restent à confirmer lors des mesures des circuits fabriqués.

Le dernier chapitre présente les mesures des 5 circuits réalisés (2 *single-ended* et 3 différentiels). Malgré la fabrication de nombreux circuits (3 runs durant la thèse), nous pouvons justifier uniquement d'un circuit performant, la première version différentielle. En effet, les autres versions présentent des performances bien en dessous de nos attentes et ne nous permettent à l'heure actuelle pas de valider l'utilisation de notre algorithme (dans ces conditions). Cependant comme il l'est noté dans les perspectives nous avons bon espoir d'améliorer le champ d'action de l'algorithme et de pouvoir ainsi obtenir des circuits en GaN démontrant des performances à la hauteur de nos espérances. De cette façon il serait possible de réaliser la comparaison des deux modes de fonctionnement puisqu'ici, le différentiel se démarque non pas par ses performances mais plutôt par les mauvaises performances du *single-ended*. Un des axes d'études étant la comparaison des deux modes il serait intéressant de les comparer suivant un cahier des charges identiques (même largeur de bande et même marges d'erreur).

Les perspectives de ce travail sont nombreuses, que ce soit en termes de conception de cellules ou d'exploitation de l'algorithme. En premier lieu, la réalisation d'un *core-chip* avec des cellules d'amplification inter-étages assurant à la fois la compensation des pertes et l'adaptation d'impédance inter-cellules serait une bonne solution. En effet cela permettrait à la fois d'exploiter les *core-chips* GaN sur le terrain de la moyenne puissance mais aussi permettrait de proposer des performances à la hauteur de celles des puces SiGe.

Concernant l'algorithme, comme expliqué précédemment, nous n'avons fait qu'un usage restreint de celui ci. En effet, l'adapter non pas à une seule cellule mais lui permettre de traiter des groupes de cellules préalablement identifiés comme favorables ou critiques permettrait de repousser les performances du système en s'affranchissant des problèmes liés au temps de calcul. Comme nous avons pu le voir lors de l'optimisation, certains groupes de cellules ressortaient des tirages, que ce soit dans les meilleurs ou les pires cas. Dans ce travail, nous nous sommes arrêtés à cette conclusion sans en exploiter les implications, limitant ainsi l'utilisation de l'algorithme à un outil d'ordonnancement.

Résumé

La réduction de taille des technologies actives permet d'envisager des applications vers des fréquences toujours plus élevées. Cependant, l'exploitation de ces bandes de fréquences nécessite une redéfinition fondamentale des architectures en raison de pertes plus conséquentes qu'aux basses fréquences ; c'est de cette assimilation de nouvelles architectures à base de réseaux d'antennes directives programmables, conjointement à l'utilisation de technologies performantes que les fréquences supérieures à 30 GHz peuvent être pleinement exploitées pour des applications télécoms, de défense et commerciales. Dans les architectures notamment adoptées pour la génération de systèmes de communications 5G, les besoins concernent tout aussi bien une forte intégration hardware qu'un besoin en puissance élevé dans les bandes Ku, K et Ka. De plus, le gain en maturité technologique des filières à grande bande interdite GaN les rend éligibles pour prétendre à la conception des modules de puissance à ces fréquences. En effet, les circuits Tx-Rx réalisés traditionnellement en technologie SiGe ou GaAs se voient de plus en plus associés à (voire sont remplacées par) des éléments en technologie GaN. Les tendances d'utilisation de cette technologie ne la limitent plus au segment de la puissance pour l'amplification haute fréquence. Des récepteurs robustes ont prouvé l'intérêt de cette technologie en étage de réception (LNA), et d'autres travaux font état de performances avantageuses pour la synthèse d'oscillateurs stables hautes fréquences. C'est dans cette logique que le fondeur OMMIC a souhaité compléter la gamme des produits déjà disponibles, en motivant une étude sur la conception de puces multi-fonctions (atténuateurs-déphaseurs programmables) MMIC en technologie GaN. Ce travail a pour but de démontrer à terme les avantages que l'on peut tirer de modules totalement intégrés GaN d'une part, et de réaliser les premiers travaux de core-chip en bande Ka sur cette technologie GaN d'autre part. En effet, il n'existe pas à notre connaissance de circuit de contrôle du signal (core-chip) réalisé dans cette technologie.

Cette thèse a donc pour double objectif de démontrer la faisabilité de tels circuits, et de proposer une méthodologie de conception pour tendre vers les meilleures performances possibles. Pour atteindre le premier objectif, après une étude bibliographique approfondie, nous savons que la technologie à effet de champ GaN ne présente pas les propriétés intrinsèques les plus favorables pour réaliser de telles fonctions ; une analyse spécifique de chaque circuit est réalisée sur deux types de réalisations (cellules single ended et différentielles). Cela représente donc un challenge de répondre à ce premier objectif. C'est pour cela que, dans un deuxième temps, nous avons porté un effort spécifique sur la méthodologie de conception, qui se décline en deux parties : la conception des cellules individuelles du core-chip puis la mise en commun de celles-ci. Nous nous sommes rendu compte qu'une mauvaise maîtrise de cette étape d'assemblage des cellules pouvait aboutir à des performances dégradées. Une étude systématique et analytique du rôle des différents paramètres (adaptation en impédance et erreur fonctionnelle en amplitude et en phase) à l'échelle de chaque cellule de déphasage (6 cellules) ou d'atténuation (5 cellules) permet d'apprécier leur impact sur le comportement global du système. Par conséquent, nous avons développé un algorithme de dénombrement capable d'identifier les agencements de cellules

induisant les meilleurs niveaux de performances. Cette méthodologie a été appliquée à la conception de core-chips single-ended et différentiels centrés sur la bande 30 – 40 GHz. Une comparaison de ces deux approches est également menée.

Mots clés : Ka-band Corechip, GaN, MMIC, 5G, beamforming.

Liste des publications

Conférence internationale :

B. Berthelot, J. Tartarin, C. Viallon, R. Leblanc, H. Maher and F. Boone, "GaN MMIC Differential Multi-function Chip for Ka-Band Applications," 2019 IEEE MTT-S International Microwave Symposium (IMS), Boston, MA, USA, 2019, pp. 1399-1402.

Conférence nationale :

Puce multifonctions MMIC GaN pour applications en bande Ka, **B. Berthelot**, J. Tartarin, C. Viallon, R. Leblanc, H. Maher and F. Boone, Journées Nationales Micro-ondes 2019, Caen