



Contribution à la conception de convertisseurs de fréquence. Intégration en technologie arséniure de gallium et silicium germanium

David Dubuc

► To cite this version:

David Dubuc. Contribution à la conception de convertisseurs de fréquence. Intégration en technologie arséniure de gallium et silicium germanium. Micro et nanotechnologies/Microélectronique. Université Paul Sabatier - Toulouse III, 2001. Français. NNT: . tel-00132419

HAL Id: tel-00132419

<https://theses.hal.science/tel-00132419>

Submitted on 21 Feb 2007

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.

Année 2001

Thèse

Préparée au
Laboratoire d'Analyse et d'Architecture des Systèmes du CNRS

En vue de l'obtention du
Doctorat de l'Université Paul Sabatier de Toulouse

Spécialité : **Electronique**

Par
David DUBUC

CONTRIBUTION A LA CONCEPTION DE CONVERTISSEURS DE FREQUENCE. INTEGRATION EN TECHNOLOGIE ARSENIURE DE GALLIUM ET SILICIUM GERMANIUM

Soutenance le 19 décembre 2001 devant le jury :

Rapporteurs Jean Luc GAUTIER
 Robert-Alain PERICHON

Examineurs Jean François SAUTEREAU
 Jacques GRAFFEUIL
 Jean Louis CAZAUX
 Cyrille BOULANGER

Directeur Thierry PARRA

à Katia,

à mon père, à ma mère

à tous ceux que j'aime...

AVANT PROPOS

Le travail présenté dans ce mémoire a été effectué au sein du groupe Composants et Intégration des Systèmes Hyperfréquences pour Télécommunications (CISHT) du Laboratoire d'Analyse et d'Architecture des Systèmes (LAAS) du CNRS de Toulouse.

Je tiens tout d'abord à remercier Monsieur Jean-Claude LAPRIE, directeur du LAAS, pour la confiance qu'il m'a témoignée en m'accueillant dans le laboratoire.

Je remercie sincèrement Monsieur Jacques GRAFFEUIL, Professeur à l'Université Paul Sabatier de Toulouse III et responsable du groupe CISHT, pour m'avoir accueilli dans son groupe de recherche et soutenu durant ces années contre vents et marrées.

Je remercie vivement Monsieur Jean François SAUTEREAU, Professeur à l'Université Paul Sabatier de Toulouse III, pour m'avoir fait l'honneur de présider le jury de ma thèse.

J'adresse également mes remerciements à Messieurs Robert-Alain PERICHON, Professeur au LEST, et Jean Luc GAUTIER, Professeur à l'ENSEA, qui ont bien voulu me faire l'honneur de juger ce travail, en acceptant, malgré leurs nombreuses activités, d'être rapporteurs de ce mémoire.

De même, je tiens à remercier les membres du jury, Messieurs Jean Louis CAZAUX, ingénieur à ALCATEL SPACE à Toulouse et Cyrille BOULANGER, ingénieur au CNES à Toulouse, pour leur précieuse participation.

Enfin, je remercie Thierry PARRA, maître de conférences au LAAS et directeur de ma thèse pour toute l'aide qu'il a su m'apporter au cours de ces trois années.

Mes remerciements vont également à tous les membres permanents du groupe CISHT pour leur soutien et aide, à savoir Messieurs Robert PLANA, Eric TOURNIER, Jean-Guy TARTARIN, Laurent ESCOTTE, Olivier LLOPIS, ainsi que Jacques RAYSSAC.

J'adresse particulièrement mes remerciements les plus tendres à Katia. Tes encouragements et ton soutien infaillibles et permanents ont réellement nourri ces travaux, tout autant que le bonheur et le plaisir de t'avoir comme compagne de bureau ainsi que notre relation dans la vie. Je te dédie donc ce manuscrit avec toute mon affection pour tous les sourires radieux qui ont ensoleillé ces années de thèse.

Je tiens de plus à saluer l'amitié échangée avec certains et certaine : Jean-Brice (Bonjour chez vous) un compagnon de bureau irremplaçable *et d'une morale irréprochable*, Brigitte pour toutes nos discussions *purement philosophique* et Christian pour son humour, sa grande gentillesse et sa passion des grands vins.

Une pensée pour tous les doctorants et ex-doctorants du groupe : Myrienne (dont le rire inégalable retentit encore dans les couloirs), Anthony, Sabine, Jessica, Laurent, Christophe, Giana, Gilles, Mathilde, Jérôme, Wah, Abdel.

Enfin merci à toutes les personnes que j'ai pu oublier et qui ont participé de près ou de loin à l'aboutissement de ces travaux.

SOMMAIRE

Introduction générale	1
Chapitre 1 : Les mélangeurs micro-ondes	3
I. Définitions	4
I.1. La transposition de fréquence	4
I.1.1. Principe	4
I.1.2. Conventions.....	6
I.2. Les mélangeurs	7
I.2.1. Introduction	7
I.2.2. Principe des mélangeurs à non-linéarité	9
a) Définitions	9
b) Mélangeur à non-linéarité	11
c) Génération de fréquences	12
d) Fonction mélangeur.....	14
I.2.3. Mise en Œuvre de mélangeurs à non-linéarité	14
a) Mélangeur : système linéaire ou non-linéaire	15
b) Composants Non-linéaires	16
c) Etude des mélangeurs à non-linéarité.....	17
II. Caractéristiques des mélangeurs	19
II.1. Gain de conversion et puissance OL nécessaire.....	19
II.2. Facteur de bruit.....	20
II.2.1. Facteur de bruit en simple bande (SSB noise)	21
a) Définition de l'IEEE.....	21
b) Autre définition	22
II.2.2. Facteur de bruit en double bande (DSB noise)	22
II.2.3. Cascade de quadripôle : formule de Friis	22
II.3. Linéarité.....	23
II.3.1. Les méfaits du quatrième degré du développement de la non-linéarité ..	23
II.3.2. L'évaluation de la linéarité	24
a) L'indicateur de linéarité : C/I_3	25
b) L'indicateur de linéarité : Point d'interception d'ordre 3 ou PI3	25
c) L'indicateur de linéarité normalisé : Q_3	26
i - Association de quadripôles	26
ii - Définition de Q_3	27
d) Les autres indicateurs de linéarité	27
II.4. Isolation.....	27
II.5. Raies parasites	28
II.6. Adaptation des ports.....	29
II.7. Consommation.....	30

III. Classification des topologies.....	30
III.1. Les différents types de mélangeur.....	30
III.1.1. Les mélangeurs actifs	30
a) Injection de l'OL sur la grille.....	31
b) injection de l'OL sur le drain	32
c) Injection de l'OL sur la source	33
d) Mélangeur à 2 transistors : structure classique	34
III.1.2. Les mélangeurs passifs	35
a) Mélangeur à diode	35
b) Mélangeur à transistor série : interrupteur analogique.....	36
c) Mélangeur à transistor parallèle : mélangeur résistif	38
III.1.3. Les structures équilibrées	39
a) Principe, avantages et inconvénients.....	39
i - <i>Structure simple équilibrée</i>	40
ii - <i>Structure double équilibrée</i>	41
iii - <i>Avantages et inconvénients</i>	43
b) Mélangeur en anneau	43
c) Mélangeur à TEC froids	44
d) Mélangeur de Gilbert	44
III.2. Synthèse et état de l'art	45
IV. Conclusion.....	46
Références bibliographiques du chapitre 1	47

Chapitre 2 : Conception d'une cellule de mélange originale. Etude de la stabilité non-linéaire 49

I. Cellule Originale de mélange	50
I.1. Principe de la topologie.....	50
I.1.1. Topologie et principe de fonctionnement	51
I.1.2. Arrangement de la structure double équilibrée	53
I.2. Etude du mélangeur	54
I.2.1. Calcul du gain de conversion	54
I.2.2. Optimisation des dimensions des transistors et de leurs polarisations.....	59
I.3. Les impédances de fermetures	61
I.3.1. Contraintes à la fréquence f_{OL}	62
I.3.2. Contraintes aux fréquences f_{RF} et f_{FI}	62
a) Particularités du mélangeur sur source.....	62
b) contraintes de charges des mélangeurs à transconductance	62
<i>i - Influence de C_{gd} et solution</i>	<i>62</i>
<i>ii - Influence de G_{d1} et solution.....</i>	<i>63</i>
<i>iii - Influence de G_{m0} et solution.....</i>	<i>63</i>
<i>iv - Adaptation en puissance.....</i>	<i>64</i>
I.3.3. Synthèse et conclusion	64
II. Les matrices de conversion	67
II.1. Principe de la méthode	67
II.2. Représentation des signaux.....	70
II.3. Formation des matrices de conversion	72
II.4. Autre type d'éléments non-linéaire ou linéaire.....	74
II.4.1. Eléments linéaires	74
II.4.2. Capacité, trans-capacité	74
a) Capacité	74
b) transcapacité.....	75
II.4.3. Transistor à effet de champ.....	76
II.5. Mise en œuvre des matrices de conversion.....	77
II.5.1. Processus de calcul des matrices de conversion	77
II.5.2. Détermination des matrices de conversion	79
II.6. Application des matrices de conversion à l'optimisation de la cellule de mélange	81
II.6.1. Validation de la méthode	81
II.6.2. Optimisation	82

III. Etude de la stabilité non-linéaire.....	85
III.1. Définition de la stabilité.....	86
III.2. Stabilité linéaire	87
III.2.1. Stabilité dans le domaine spectral	87
a) Etude de la stabilité par une fonction de transfert équivalente	88
b) Etude de la stabilité par une fonction caractéristique.....	90
III.2.2. Application aux circuits micro-ondes.....	91
III.3. Stabilité non-linéaire : méthode mise en œuvre	93
III.3.1. Méthodes existantes.....	93
III.3.2. Méthodes mise en œuvre pour l'étude de la stabilité de mélangeurs. ...	94
a) Stabilité des circuits multi-dimensionnels.....	94
b) Recherche d'une fonction d'étude de la stabilité	95
c) Variables à observer	97
d) Conclusion	98
III.4. Application : stabilisation de la cellule de mélange	98
III.4.1. Etude de la stabilité non-linéaire de mélangeur	99
III.4.2. Optimisation de Lstab.....	100
IV. Conclusion.....	102
Références bibliographiques du chapitre 2	104

Chapitre 3 : Intégration monolithique de convertisseurs de fréquence.....107

Chapitre 3 Partie A : Mélangeur double équilibré sur GaAs en bande Ku109

I. Choix structurel.....	110
I.1. Cahier des charges.....	111
I.2. Etude système	111
II. Etude des coupleurs	113
II.1. Critères de sélection.....	115
II.2. Les coupleurs actifs à paire différentielle.....	116
II.2.1. Optimisation de Z_{mc}	120
II.2.2. Optimisation de Z_{grille}	122
II.2.3. Adaptation et stabilité	123
II.2.4. Conclusion	125
III. Intégration monolithique sur GaAs	125
III.1. La technologie PML/OMMIC	126
III.1.1. Les inductances intégrées.....	126
III.1.2. Eloignement minimal des composants	128
III.2. Intégration MMIC	129
III.3. Performances simulées	131
IV. Caractérisations des circuits MMIC.....	134
IV.1. Montage et caractérisation.....	134
IV.2. Cellules de base	135
IV.2.1. Mélangeur simple équilibré.....	135
IV.2.2. Diviseur de puissance.....	138
IV.3. Cellule de mélange complète.....	139
V. Conclusions	142
Références bibliographiques du chapitre 3 partie A.....	143

Chapitre 3 Partie B : Mélangeur simple sur BiCMOS-SiGe en bande Ka.....145

I. La technologie BiCMOS-SiGe de ST-Microelectronics146

I.1. Présentation de la technologie.....146

I.1.1. Les composants actifs 147

I.1.2. Les composants passifs 147

a) Les différents niveaux métalliques..... 148

b) Les capacités 148

c) Les inductances 149

I.2. Inductance dans la gamme millimétrique.....151

I.2.1. Le dimensionnement 151

I.2.2. Les différents types de pertes 153

I.2.3. Réalisation d'une inductance millimétrique..... 154

I.3. Conclusion158

II. Intégration de la cellule de mélange..... 159

II.1. Le schéma électrique.....159

II.2. Les performances mesurées161

II.2.1. Caractérisation du circuit 161

II.2.2. Discussion des caractérisations..... 164

III. Conclusion..... 165

Références bibliographiques du chapitre 3 partie B.....166

Conclusion générale169

Introduction générale

Ces dernières années ont vu le fort développement des systèmes de communication spatiaux et avec lui le changement des contraintes de conception de ces systèmes. En effet, la multiplicité des applications ainsi que leur ouverture au domaine grand public ont entraîné la nécessité de concevoir des systèmes performants à bas coût.

Ces nouvelles contraintes ont accéléré la migration des réalisations de la technologie hybride à la technologie monolithique qui permet l'intégration de circuits performants compatibles avec une production de masse. Réduisant simultanément coût, masse et volume, les technologies monolithiques ont fait l'objet de développements importants et de nouvelles technologies sont apparues telles que la technologie à transistors bipolaires à hétérojonction sur silicium. Les enjeux économiques et la compétition avec le segment terrestre ont de même conduits à une évolution de la conception des systèmes de communication vers l'assemblage de fonctions génériques afin de réduire le temps et le coût de la conception.

Au cœur de tous ces systèmes de télécommunication se trouve le mélangeur dont la fonction consiste à transposer la fréquence des signaux traités. C'est sur ce dernier composant que nous avons axé nos travaux de recherche et plus particulièrement sur sa réalisation en technologies monolithiques sur Arséniure de Gallium et sur Silicium-Germanium pour des applications de répéteurs spatiaux en bande Ku et Ka.

Nous avons de plus cherché à définir une méthodologie générique pour la conception monolithique des circuits mélangeurs à ces fréquences.

Le premier chapitre rappelle tout d'abord les définitions relatives aux mélangeurs ainsi que les principales caractéristiques nécessaires à leurs évaluations. Puis, nous présentons les différentes topologies de mélangeurs existantes en montrant, pour chacune, ses avantages et ses inconvénients. Ce chapitre conclue par la synthèse et l'état de l'art de ces différentes topologies de mélange permettant ainsi, suivant les contraintes imposées, d'en effectuer le choix judicieux.

S'appuyant sur cette classification, le second chapitre présente, tout d'abord, la définition d'une cellule originale de mélange. Puis, nous décrivons une méthode d'étude analytique des circuits mélangeurs qui nous a permis d'une part de définir des règles génériques de conception de mélangeur et d'autre part d'optimiser notre cellule de mélange suivant des critères fixés. Enfin, une partie importante de ce chapitre est dédiée à l'étude de la stabilité non-linéaire, point sensible de la topologie retenue.

Sur la base de ce travail théorique, le troisième chapitre est consacré à l'intégration monolithique de ce mélangeur performant pour deux technologies concurrentes : l'une basée sur des transistors à effet de champ à haute mobilité électronique sur Arséniure de Gallium et l'autre fondée sur des transistors bipolaires à hétérojonction Silicium/Silicium-Germanium. Nous exposerons les différents résultats de caractérisations de ces circuits qui ont prouvé qualitativement le bien fondé des études théoriques précédentes.

Enfin, nous concluons par la synthèse globale de nos travaux en insistant sur leurs aspects innovants et nous envisagerons les perspectives de développement de cette étude.

Chapitre 1 : Les mélangeurs micro-ondes

L'augmentation des performances et de l'intégration des systèmes de télécommunication alliée aux contraintes de fiabilité, de coût et de production de masse oriente les conceptions vers des circuits génériques capables d'être utilisés à la demande sans pour autant que le système ne perde en performances.

Malheureusement, contrairement aux amplificateurs, les circuits mélangeurs ne présentent pas de méthodologie de conception limpide permettant tout d'abord le choix d'une structure parmi l'ensemble des solutions existantes puis l'application de procédures de conception aboutissant à une conception optimale.

L'objectif premier de nos travaux de recherche fut donc la définition claire des différentes topologies de mélange, la compréhension de leur fonctionnement ainsi que l'interprétation de leurs performances, afin de permettre le choix d'une topologie suivant des contraintes fixées.

Nous présenterons tout d'abord la description du processus de transposition de fréquence ainsi que le mode opératoire des mélangeurs à non-linéarité. Puis, nous décrirons les caractéristiques principales nécessaires à l'évaluation des performances des différents mélangeurs existants dont nous présenterons les principes de fonctionnement au paragraphe III. Une synthèse des performances des différentes cellules de mélange conclura ce chapitre.

I. Définitions

Nous allons dans un premier temps présenter le principe de la transposition de fréquence, ainsi que les conventions utilisées. Puis nous décrirons plus en détail la fonction mélangeur, élément central du système de conversion de fréquence et objet de nos travaux de recherche. Nous verrons que sa réalisation dans le domaine micro-onde est basée sur la non-linéarité de composants semi-conducteur et nous décrirons un moyen simple d'appréhender qualitativement et quantitativement les phénomènes de conversion de fréquence.

I.1. La transposition de fréquence

Breveté en 1917 par Edwin Howard Armstrong, le récepteur superhétérodyne est un des premiers systèmes à mettre en œuvre le principe de transposition de fréquence. Cette invention est indiscutablement la plus importante dans l'histoire des télécommunications car son principe est toujours utilisé dans de nombreux systèmes quelques 80 ans après sa découverte.

I.1.1. Principe

Le changement de fréquence consiste à transposer le spectre d'un signal centré sur une fréquence initiale vers une autre bande de fréquence sans altération. D'un point de vue spectral l'opération de transposition de fréquence se résume donc à :

$$E(f) \rightarrow S(f)=E(f+f_0) \quad (\text{I-1})$$

où $E(f)$ et $S(f)$ sont respectivement les spectres du signal d'entrée et de sortie. L'équation (I-1) peut se mettre sous la forme suivante :

$$S(f) = E(f) * \delta(f+f_0) \quad (\text{I-2})$$

Où le signe $*$ correspond au produit de convolution et $\delta(f)$ désigne la fonction de Dirac. L'équation (I-2) se traduit, dans le domaine temporel par :

$$s(t) = e(t) \bullet e^{j2\pi f_0 t} \quad (\text{I-3})$$

L'équation (I-3) montre que l'opération est impossible car il faudrait multiplier au signal d'entrée un signal complexe donc irréalisable. L'opération de transposition de fréquence *stricto sensu* est donc impossible. Afin de réaliser cependant cette opération, la solution consiste à rajouter son conjugué au signal précédent. Le signal de sortie, décrit par l'équation (I-4), correspond alors au produit du signal d'entrée avec un signal sinusoïdal réel donc réalisable.

$$s(t) = e(t) \bullet \left(e^{j2\pi f_0 t} + e^{-j2\pi f_0 t} \right) = e(t) \bullet 2\cos(2\pi f_0 t) \quad (\text{I-4})$$

Ceci amène, dans le domaine spectral, à la transformation suivante :

$$E(f) \rightarrow S(f) = E(f+f_0) + E(f-f_0) \quad (\text{I-5})$$

Comme l'indique la Figure I-1, la transformation (I-5) réalise bien la transposition de fréquence puisque le spectre du signal d'entrée, centrée sur f_m , est translaté en sortie de la fréquence f_0 et est alors centré sur f_m+f_0 : $E(f+f_0)$. Par contre, nous avons de plus dupliqué ce spectre à la fréquence f_m-f_0 : $E(f-f_0)$, en raison de la multiplication du signal d'entrée par $e^{-j2\pi f_0 t}$.

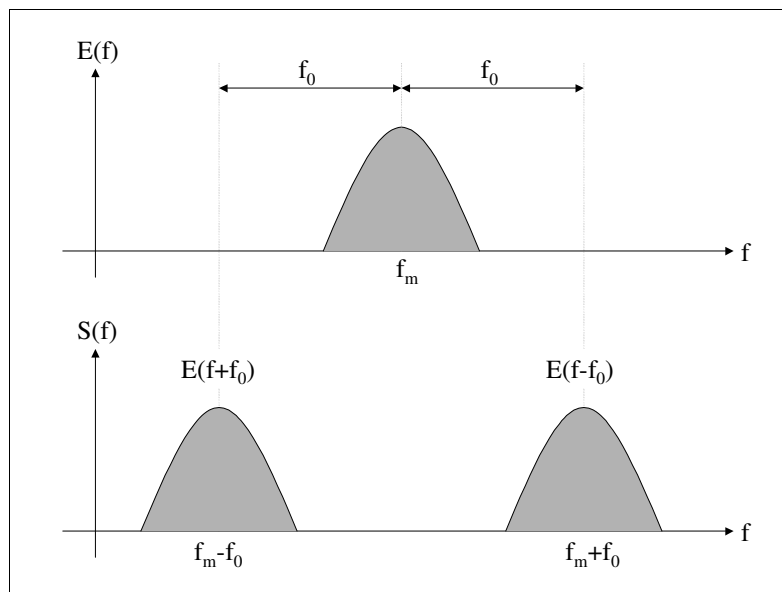


Figure I-1 : Transposition de fréquence

Si le spectre d'intérêt est centré sur $f_m + f_0$, la conversion obtenue est dite supradyne (up-conversion en anglo-saxon). Le spectre centré sur $f_m - f_0$ est alors sans intérêt et peut être éliminé par filtrage. Par contre, si le spectre d'intérêt est centré sur $f_m - f_0$, la conversion est alors infradyne (down-conversion en anglo-saxon), c'est alors le spectre centré sur $f_m + f_0$ qui est éliminé par filtrage. Enfin, dans le cas où $f_m = f_0$, la conversion est dite homodyne ou à « zéro IF » pour indiquer que la fréquence intermédiaire est alors nulle.

Dans la suite de ce manuscrit, nous ne considérerons que les conversions vers une fréquence inférieure car c'est sur ce type de transposition de fréquence que porte notre travail (voir le paragraphe II). Les principes sont rigoureusement identiques pour les conversions supradynes. Aussi le travail présenté dans ce manuscrit sera aisément transposable.

I.1.2. Conventions

La transposition de fréquence se résume donc à l'équation (I-4) dont la Figure I-2 donne la description fonctionnelle. Dans ce système, l'entrée est usuellement notée RF (pour signal Radio Fréquence) : il s'agit du signal dont on souhaite transposer de fréquence. Le signal de sortie se nomme quant à lui FI (pour signal à Fréquence Intermédiaire) : c'est le signal résultant de la transposition fréquentielle du signal RF.

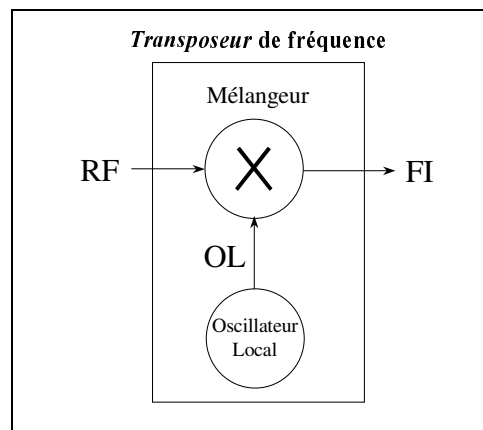


Figure I-2 : Description fonctionnelle d'un *transposeur* de fréquence

La Figure I-2 montre que le transposeur de fréquence est donc constitué de deux éléments :

1. Un oscillateur local qui est en charge de générer une onde sinusoïdale de fréquence correspondant à la fréquence de translation désirée : $f_{OL} = f_{FI} - f_{RF}$.
2. Un mélangeur dont la fonction est de réaliser la multiplication entre les signaux RF et OL.

Notre travail porte sur l'étude, la conception et la caractérisation d'une fonction mélangeur qui s'intègre dans un système de conversion de fréquence.

Bien évidemment, il existe également de nombreuses autres applications aux mélangeurs parmi lesquelles nous pouvons citer l'amplification ou atténuation variable, la modulation et la démodulation d'amplitude, la démodulation FM, la modulation et la

démodulation BPSK (Binary Phase Shift Keying), la détection de phase et le doublement de fréquence.

La suite de ce premier chapitre est donc dédiée à la description des principes et à l'énoncé des caractéristiques principales des mélangeurs. Nous aborderons ensuite les aspects plus pratiques de la mise en œuvre des mélangeurs dans le domaine micro-onde.

I.2. Les mélangeurs

Ce paragraphe traite des principes des mélangeurs à non-linéarité. Après une brève introduction, présentant deux mélangeurs classiques de l'électronique basse fréquence, nous exposons le principe des mélangeurs à non-linéarité, puis nous traiterons de leurs réalisations. Enfin, nous introduirons une étude simplifiée permettant la compréhension des mécanismes de conversion de fréquence.

I.2.1. Introduction

Comme nous l'avons vu dans le paragraphe précédent, un mélangeur idéal doit réaliser la multiplication des deux signaux d'entrée ; c'est donc naturellement que les **multiplieurs de Gilbert** réalisent la fonction mélangeur. Ces structures, du nom de leur inventeur Barrie Gilbert [I.1], sont constituées de trois paires différentielles comme le présente la Figure I-3.

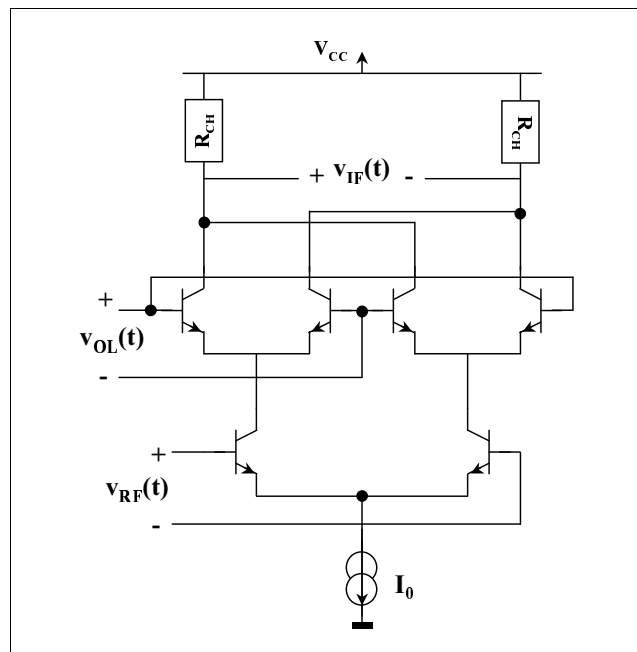


Figure I-3 : Multiplieur de Gilbert

La relation d'entrée-sortie est la suivante :

$$V_{FI} = R_{CH} I_0 \operatorname{th}\left(\frac{V_{RF}}{2U_T}\right) \operatorname{th}\left(\frac{V_{OL}}{2U_T}\right) \quad (I-6)$$

où $\tanh$ représente la fonction tangente hyperbolique et U_T la tension thermodynamique.

Dans le cas petits signaux, c'est à dire lorsque les expressions (I-7) sont vérifiées, la relation d'entrée-sortie se simplifie comme indiqué par (I-8) :

$$|V_{RF}| \ll 2U_T \text{ et } |V_{OL}| \ll 2U_T \quad (\text{I-7})$$

$$V_{FI} = \frac{R_{CH} I_0}{4(U_T)^2} V_{RF} V_{OL} \quad (\text{I-8})$$

Cette dernière relation démontre que la structure de Gilbert réalise une multiplication 4 quadrants et peut donc s'utiliser comme mélangeur. Nous verrons par la suite (III.1.3.d)) que lorsque les conditions de petits signaux (I-7) ne sont plus valables, la cellule de Gilbert ne réalise plus parfaitement une multiplication (voir la relation (I-6)) mais peut toujours faire office de mélangeur. La cellule résultante se nomme alors mélangeur de Gilbert et non plus multiplieur de Gilbert.

Néanmoins, la mise en œuvre d'une cellule de Gilbert peut parfois s'avérer lourde étant donnée qu'elle nécessite 6 transistors. Par simplicité il est possible de réaliser un **modulateur** qui consiste en un simple interrupteur, soit série, soit parallèle, commandé par le signal OL. Prenons le cas d'un interrupteur série décrit par la Figure I-4

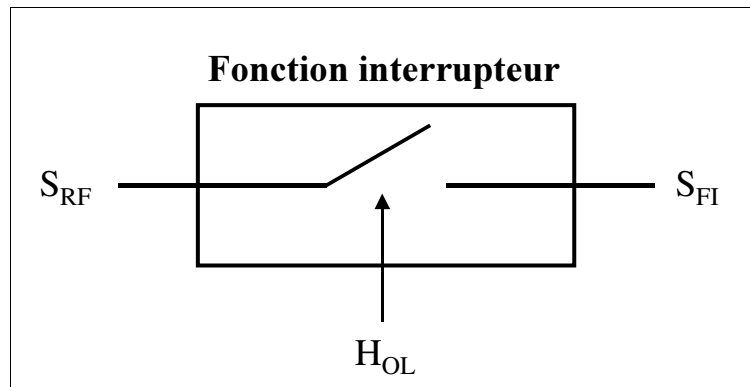


Figure I-4 : Modulateur : interrupteur série.

Le signal de sortie est donc l'image du signal d'entrée pendant la fraction de la période du signal OL pour laquelle l'interrupteur est fermé et vaut 0 le reste de cette même période. Le signal de sortie est donc le produit du signal d'entrée par un créneau de période OL dont les états sont 0 et 1.

Nous obtenons :

$$S_{FI} = S_{RF} \bullet H_{OL} \quad (\text{I-9})$$

avec H_{OL} représentant la fonction créneau à la fréquence f_{OL} . La multiplication entre deux signaux donc la fonction mélangeur est ainsi obtenue. Par contre, le signal H_{OL} n'est pas monochromatique à la fréquence f_{OL} et a le développement en série de Fourier suivant :

$$H_{OL}(t) = \frac{1}{2} + \frac{2}{\pi} \cos(\omega_{OL} t) + \sum_{k>1} a_{2k+1} \cos[(2k+1)\omega_{OL} t] \quad (\text{I-10})$$

En supposant le signal RF monochromatique, la sortie est de la forme :

$$S_{FI}(t) = \frac{1}{2} \cos(\omega_{RF} t) + \frac{1}{\pi} \cos(\omega_{FI} t) + \frac{1}{\pi} \cos[(\omega_{RF} + \omega_{OL}) t] \\ + \sum_{k>1} \frac{a_{2k+1}}{2} \left\{ \cos[\omega_{RF} + (2k+1)\omega_{OL} t] + \cos[\omega_{RF} - (2k+1)\omega_{OL} t] \right\} \quad (I-11)$$

Le second terme de l'équation (I-11) démontre que le modulateur réalise la transposition de fréquence tout en générant beaucoup d'autres raies parasites qu'il conviendra de filtrer.

Les deux types de mélangeur évoqués précédemment reposent sur des caractéristiques non-linéaires de composant : les mélangeurs de Gilbert sont basés sur la non-linéarité $I_c(V_{be})$ décrivant le courant de collecteur en fonction de la tension base émetteur, alors que les modulateurs utilisent des composants actifs (diode ou transistor) en régime de commutation.

D'autres types de mélangeurs existent et tous reposent sur des caractéristiques non-linéaires de composants. Nous allons donc détailler l'étude théorique des **mélangeurs à non-linéarité** dans le paragraphe suivant.

I.2.2. Principe des mélangeurs à non-linéarité

Tout circuit et, plus généralement tout système, est non linéaire. La linéarité des circuits n'est qu'une hypothèse nécessitant certaines conditions.

Deux types de circuits peuvent être différenciés [I.2].

1. Ceux qui sont faiblement non-linéaires (on parle alors de quasi-linéarité), par exemple les amplificateurs petit signal fonctionnant en classe A. Leur étude est réalisée en supposant un fonctionnement linéaire, les non-linéarités (faibles) étant dans ce cas considérées comme dégradant les performances.

2. Ceux dont le fonctionnement repose sur une ou plusieurs non-linéarité(s), comme les mélangeurs et les multiplieurs de fréquence. Dans ces cas là, les non-linéarités sont prises en compte et nécessaires pour réaliser la fonction désirée. Souvent il sera même nécessaire de minimiser les phénomènes linéaires de ces circuits. L'étude de ces circuits est considérablement plus complexe que celle des circuits où l'hypothèse de linéarité peut être faite.

Le fonctionnement des mélangeurs repose sur leur caractère non-linéaire. Nous verrons en effet, après quelques définitions relatives aux systèmes non-linéaires, comment ces derniers réalisent la conversion de fréquence.

a) Définitions

La définition du terme "**linéarité**" bien qu'intuitive mérite d'être signalée : un système est linéaire si et seulement si il satisfait à la condition suivante :

$$\left. \begin{array}{l} e_1(t) \rightarrow s_1(t) \\ e_2(t) \rightarrow s_2(t) \end{array} \right\} \Rightarrow a.e_1(t) + b.e_2(t) \rightarrow a.s_1(t) + b.s_2(t) \quad (\text{I-12})$$

a et b étant des constantes réelles. Cette définition est simplement l'expression du principe de superposition. La notion de linéarité fait référence à un système caractérisé par des équations différentielles linéaires.

Mais cette définition à elle seule ne suffit pas. A cela s'ajoutent les propriétés de variance ou d'invariance temporelle, de système dynamique ou non et de non-linéarité à mémoire ou non.

Ainsi, rappelons la notion d'**invariance temporelle**. Un système est invariant dans le temps si et seulement si :

$$e(t) \rightarrow s(t) \Rightarrow e(t-\tau) \rightarrow s(t-\tau) \quad (\text{I-13})$$

Un système linéaire et invariant dans le temps est régi par des équations différentielles linéaires à coefficients invariants dans le temps.

Remarque : d'un point de vue spectral, un système linéaire variant dans le temps peut générer des composantes spectrales de fréquence qui n'existent pas en entrée. Seuls les systèmes linéaires invariants dans le temps ne génèrent que les composantes spectrales qui sont présentes en entrée.

La notion de système **dynamique** est employée lorsque les variables de sorties dépendent de leurs propres passés **et** de celui des variables d'entrées.

Dans le cas des systèmes non-dynamiques, il est important de savoir si le **système** est **à mémoire**, c'est à dire si ses variables de sorties dépendent du passé de ses variables d'entrées.

Dans le cas général, la fonction de transfert d'un système est du type : $s(t)=F(e(\tau \leq t), t)$, ou F est appelée fonctionnelle du système.

Sans passer tous les cas en revue, nous proposons quelques exemples pratiques :

- ❖ Un système **linéaire sans mémoire** possède une relation d'entrée sortie du type : $s(t)=\alpha(t) \bullet e(t)$, si, de plus, il est **invariant dans le temps** alors $\alpha(t) = \text{constante}$.
- ❖ Un système **linéaire dynamique** est caractérisé par : $s(t)=h(t, \tau < t) * e(t)$, si de plus il est **invariant dans le temps**, alors $h(t, \tau)=h(t)$.
- ❖ Un système **non linéaire sans mémoire** est caractérisé par : $s(t)=F(e(t), t)$, si de plus il est **invariant dans le temps**, alors $s(t)=F(e(t))$.

Le Tableau I-1 donne la synthèse des fonctionnelles des systèmes selon les caractéristiques précédemment abordées.

Fonctionnelle Des systèmes	Système linéaire	Système non-linéaire
Dynamique Variant Dans le temps (=VT) Donc à mémoire	$s(t) = \int_{-\infty}^{+\infty} h(t, \theta) e(\theta) d\theta$	$s(t) = \int_{-\infty}^{+\infty} h_1(t, \theta) e(\theta) d\theta + \int_{-\infty}^{+\infty} h_2(t, \theta_1, \theta_2) e(\theta_1) e(\theta_2) d\theta + \dots$
Invariant Dans le temps (=IT)	$s(t) = (h * e)(t)$	$s(t) = \int_{-\infty}^{+\infty} h_1(t - \theta) e(\theta) d\theta + \int_{-\infty}^{+\infty} h_2(t - \theta_1, t - \theta_2) e(\theta_1) e(\theta_2) d\theta + \dots$
Non Dynamique À mémoire (Donc VT)	$s(t) = \int_{-\infty}^t h(t, \theta) e(\theta) d\theta$	$s(t) = \int_{-\infty}^t h_1(t, \theta) e(\theta) d\theta + \int_{-\infty}^t h_2(t, \theta_1, \theta_2) e(\theta_1) e(\theta_2) d\theta + \dots$
Vt Sans mémoire	$s(t) = \alpha(t) \bullet e(t)$	$s(t) = \sum_k \alpha_k(t) \bullet [e(t)]^k$
IT	$s(t) = \alpha \bullet e(t)$	$s(t) = \sum_k \alpha_k \bullet [e(t)]^k$

Tableau I-1 : Fonctionnelle des systèmes

Une dernière définition, un peu moins nette, concerne les systèmes faiblement ou fortement non linéaires : un système faiblement non linéaire peut être caractérisé par un développement limité en série de Taylor autour d'un point de repos.

Ces diverses définitions sont importantes car elles permettent la classification des différents types de non-linéarité suivant la complexité de leurs fonctionnelles.

Dans la pratique, pour des raisons de simplicité, nous nous limiterons à une description des caractéristiques non-linéaires des circuits électriques par des fonctionnelles non-dynamiques sans mémoire et invariantes dans le temps. Cette modélisation est en fait une hypothèse qu'il est nécessaire d'effectuer afin d'aboutir à des circuits exploitables pour les simulateurs électriques.

Nous allons voir dans les prochains paragraphes, comment une non-linéarité peut faire office de mélangeur. Nous découvrirons aussi, que ce type de mélangeur génère de nombreuses raies spectrales parasites.

b) Mélangeur à non-linéarité

L'explication du processus de conversion peut se faire sur l'exemple du mélangeur présenté en Figure I-5. Les deux signaux à mélanger sont additionnés, leur somme traverse un élément non-linéaire puis un filtre éliminant les fréquences indésirables.

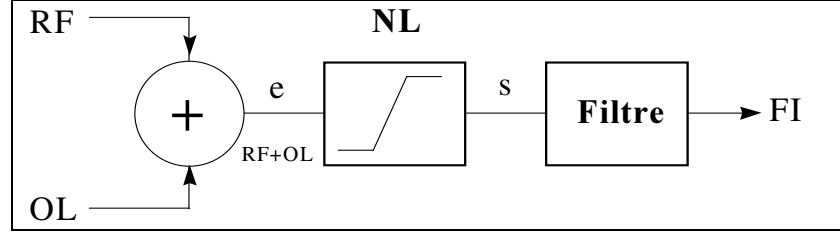


Figure I-5 : Synoptique de base d'un mélangeur

Nous considérons ici une non-linéarité non-dynamique sans mémoire et invariante dans le temps qui se caractérise donc par une équation d'état du style : $s=F(e)$. ou s et e sont des fonctions temporelles. Cette non-linéarité, supposée faible, peut être développée en série de puissance limitée, dans notre cas, à l'ordre 4 :

$$s = k_0 + k_1 e + k_2 e^2 + k_3 e^3 + k_4 e^4 \quad (\text{I-14})$$

Ce qui donne :

$$\begin{aligned} s &= k_0 + k_1 (OL + RF) + k_2 (OL + RF)^2 + \dots \\ &= k_0 + k_1 OL + k_2 OL^2 + \dots + k_1 RF + k_2 RF^2 + \dots + 2k_2 RF \bullet OL + \dots \end{aligned} \quad (\text{I-15})$$

L'équation (I-15) montre que la sortie se compose respectivement d'un terme continu, de puissances entières du signal OL, de puissances entières du signal RF, du terme qui nous intéresse, à savoir le produit du signal RF et du signal OL, et d'un certain nombre d'autres termes dits d'intermodulation.

Le terme d'intérêt est $2k_2 RF \bullet OL$: il démontre qu'une non-linéarité invariante dans le temps et sans mémoire réalise la multiplication des signaux additionnés et injectés en son entrée. Ceci prouve qu'une non-linéarité peut réaliser la fonction mélangeur. Par contre, comme dans le cas des modulateurs, la non-linéarité génère bon nombre d'autres signaux indésirables qu'il conviendra de filtrer.

c) Génération de fréquences.

Lorsque les accès RF et OL sont excités par des signaux monochromatiques, le signal d'entrée de la non-linéarité s'exprime par :

$$e = V_{RF} \cos(\omega_{RF} t + \varphi_{RF}) + V_{OL} \cos(\omega_{OL} t + \varphi_{OL}) \quad (\text{I-16})$$

La sortie est alors la somme des termes générés par chaque puissance de la non linéarité, mise sous la forme suivante : $s = s_0 + s_1 + s_2 + s_3 + s_4$. L'expression de chaque terme est décrite par les relations (I-17) à (I-21) qui suivent :

i) terme continu

$$s_0 = k_0 + \frac{k_2}{2} (V_{RF}^2 + V_{OL}^2) \quad (\text{I-17})$$

Ce terme correspond à la composante continue du signal de sortie, il est généré par les termes de degré pair de la non-linéarité. Outre le terme k_0 , qui correspond à la polarisation de la non-linéarité du mélangeur, les autres termes contribuent à son auto-polarisation.

Remarquons que le développement limité a été effectué autour d'un point de repos qui peut être modifié par l'auto-polarisation, entraînant la remise en cause de la validité du degré maximal de ce développement limité. Cependant, l'hypothèse d'une non-linéarité *faible* implique une faible influence de l'auto-polarisation sur la précision de cette modélisation.

ii) terme linéaire

$$s_1 = k_1 V_{RF} \cos(\omega_{RF} t + \varphi_{RF}) + k_1 V_{OL} \cos(\omega_{OL} t + \varphi_{OL}) \quad (\text{I-18})$$

Ces termes correspondent à la partie linéaire, ils sont générés par le terme de degré 1 de la non linéarité, et sont donc nommés termes d'ordre 1.

iii) produits d'intermodulation d'ordre 2

$$\begin{aligned} s_2 = & k_2 V_{RF} V_{OL} \cos((\omega_{RF} - \omega_{OL})t + \varphi_{RF} - \varphi_{OL}) \\ & + k_2 V_{RF} V_{OL} \cos((\omega_{RF} + \omega_{OL})t + \varphi_{RF} + \varphi_{OL}) \\ & + \frac{k}{2} V_{RF}^2 \cos(2\omega_{RF} t + 2\varphi_{RF}) \\ & + \frac{k}{2} V_{OL}^2 \cos(2\omega_{OL} t + \varphi_{OL}) \end{aligned} \quad (\text{I-19})$$

Ces termes sont l'expression du caractère non linéaire de la fonction de transfert. Ils sont générés par le degré 2 de la non-linéarité et sont donc nommés produits d'intermodulation d'ordre 2. Le premier terme est le terme utile dans le cas d'un abaissement de fréquence tandis que le second correspond à une élévation de fréquence. Les deux derniers correspondent aux harmoniques d'ordre deux des signaux d'entrée.

iv) produits d'intermodulation d'ordre 3

$$\begin{aligned} s_3 = & \frac{1}{4} k_3 V_{RF}^3 \cos(3\omega_{RF} t) + \frac{1}{4} k_3 V_{OL}^3 \cos(3\omega_{OL} t) \\ & + \frac{3}{4} k_3 V_{RF}^2 V_{OL} \cos(2\omega_{RF} + \omega_{OL}) t + \frac{3}{4} k_3 V_{RF}^2 V_{OL} \cos(2\omega_{RF} - \omega_{OL}) t \\ & + \frac{3}{4} k_3 V_{RF} V_{OL}^2 \cos(\omega_{RF} + 2\omega_{OL}) t + \frac{3}{4} k_3 V_{RF} V_{OL}^2 \cos(\omega_{RF} - 2\omega_{OL}) t \\ & + \frac{3}{2} k_3 (V_{RF}^3 + V_{RF} V_{OL}^2) \cos(\omega_{RF}) t + \frac{3}{2} k_3 (V_{OL}^3 + V_{RF}^2 V_{OL}) \cos(\omega_{OL}) t \end{aligned} \quad (\text{I-20})$$

Nous avons supposé les déphasages φ_{RF} et φ_{OL} nuls. Ces termes sont générés par le troisième degré de la non-linéarité et s'appellent donc : produits d'intermodulation d'ordre 3. Tous ces signaux ne sont pas utiles à la conversion de fréquence, il conviendra donc de les filtrer.

v) produits d'intermodulation d'ordre 4

Enfin, parmi les termes d'ordre 4, deux méritent une attention particulière :

$$s_4 = \frac{3}{4} k_4 (V_{OL} V_{RF}^3 + V_{RF} V_{OL}^3) \cos((\omega_{rf} - \omega_{ol})t) \quad (I-21)$$

Ces produits ont une fréquence d'intermodulation identique à celle du signal utile en sortie. Nous verrons par la suite l'importance du premier terme.

D'une manière générale, le k-ième degré d'une non-linéarité génère des produits d'intermodulation d'ordre k qui se caractérisent par des fréquences d'intermodulation telles que : $f = m f_{RF} \pm n f_{OL}$ avec $|m| + |n| = k$. De plus l'amplitude de ces termes est proportionnelle à : $V_{RF}^m V_{OL}^n$

Remarque : Considérons un système global qui regroupe un élément non linéaire et au moins un autre élément linéaire ou non. Si la non linéarité de l'élément non linéaire est de degré k, le système ne peut pas générer des produits d'intermodulation d'ordre supérieur à k.

d) Fonction mélangeur.

Parmi tous les termes générés par la non-linéarité décrite par l'équation (I-14), seul le premier terme de l'expression (I-19) nous intéresse. Ce terme démontre que l'on a bien réalisé la fonction mélangeur. L'inconvénient de ce type de mélangeur est la génération de signaux indésirables qui peuvent ne pas être gênants si ce n'est la dispersion de puissance qu'ils entraînent.

Enfin, remarquons que la fréquence d'intérêt $f_{FI} = f_{RF} - f_{OL}$ est générée par le terme d'ordre 2 de la non linéarité (en fait, par tous les termes d'ordre pair). Un mélangeur performant sera donc un système dans lequel le degré deux de sa non-linéarité sera maximisé.

Notons enfin qu'un système ne fonctionne en régime linéaire ou non-linéaire que suivant l'amplitude relative des signaux appliqués. Un mélangeur est donc un système pour lequel au moins un signal appliqué est d'amplitude suffisante pour qu'au minimum la non-linéarité de degré deux puisse être considérée. Le signal RF est généralement de faible amplitude car transportant l'information, c'est donc au signal OL d'être de fort niveau. Ce dernier est alors nommé signal de pompe car il est en charge de pomper la ou les non-linéarités du mélangeur.

I.2.3. Mise en Œuvre de mélangeurs à non-linéarité

Ce paragraphe décrit une approche simplifiée des mélangeurs à non-linéarité permettant la compréhension des phénomènes de conversion de fréquence.

Nous verrons tout d'abord que la transposition de fréquence n'est pas forcément réalisée à l'aide d'un système non-linéaire. En effet, un système linéaire mais variant dans le

temps vis à vis des signaux RF et FI permet la réalisation de cette fonction. Cette propriété est importante pour la compréhension de la méthode envisagée. Puis nous introduirons les composants non-linéaires constituant classiquement les mélangeurs, à savoir les diodes et les transistors. Nous classerons leurs non-linéarités et décrirons les hypothèses faites pour leur modélisation.

Enfin, nous exposerons une méthode simple pour l'évaluation des performances des mélangeurs qui nous a servi lors de la classification des divers types de mélangeur effectuée au paragraphe III.

a) Mélangeur : système linéaire ou non-linéaire

La notion de système est liée à la notion de frontière.

Nous modélisons un mélangeur selon la Figure I-6. Le mélange s'organise autour d'un multiplieur. Sur la voie OL le composant référencé par NL caractérise les non-linéarités introduites par le signal OL qui est de forte amplitude. Le signal RF est, quant à lui, considéré de faible amplitude et ne génèrera pas de non-linéarité (pratiquement, cela se traduit par $V_{RF}^2, V_{RF}^3, \dots = 0$ dans les expressions (I-17) à (I-21))

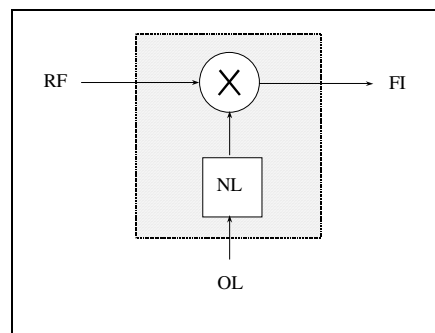


Figure I-6: un modèle de mélangeur.

La Figure I-7 diffère de la Figure I-6 par la localisation de la frontière du système. Dans cette configuration, seuls les accès RF et FI sont considérés, l'entrée OL est englobée dans le système : on retrouve le transposateur de fréquence. Dans ce cas, le système est linéaire mais variant dans le temps. Cette constatation est importante car elle démontre qu'un système de transposition de fréquence est linéaire lorsque l'on considère les accès RF et FI. C'est sur cette propriété qu'ont été inventés les mélangeurs résistifs qui ont la particularité d'être les mélangeurs les plus linéaires parmi toutes les topologies existantes (voir le paragraphe III.1.2.c)).

A titre de comparaison, considérons maintenant le système présenté en Figure I-8. Dans ce système, c'est l'entrée RF qui est englobée au système. Dans ce cas, le système est non linéaire à cause de la présence du bloc nommé NL.

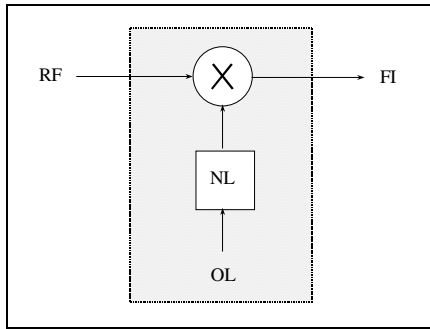


Figure I-7 : système n°1.

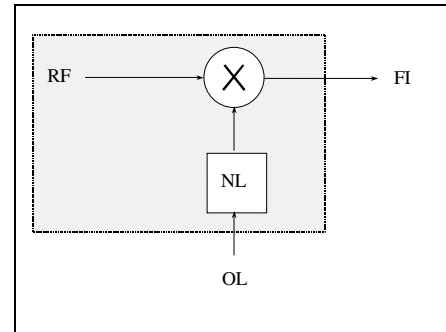


Figure I-8 : système n°2.

En conclusion, la notion de linéarité dépend de la frontière du système considéré ou bien, ce qui est équivalent, des variables d'entrée-sortie considérées. De plus, nous avons vu que la transposition de fréquence est une opération linéaire mais variante dans le temps. La linéarité de ce système est la propriété qui garantit la non dégradation de la bande passante, et c'est le caractère variant dans le temps qui permet la transposition de fréquence.

Nous présentons dans le prochain paragraphe les composants et leurs non-linéarités réalisant la fonction mélangeur.

b) Composants Non-linéaires

Tout système électronique se modélise par un schéma équivalent dans lequel apparaissent uniquement les composants de base suivants : résistance, inductance, capacité, sources de courant et de tension commandées. Par exemple, la Figure I-9 montre les schémas équivalents simplifiés d'une diode Schottky et d'un transistor à effet de champ (TEC), qui sont, avec le transistor bipolaire, les composants de base de tout mélangeur.

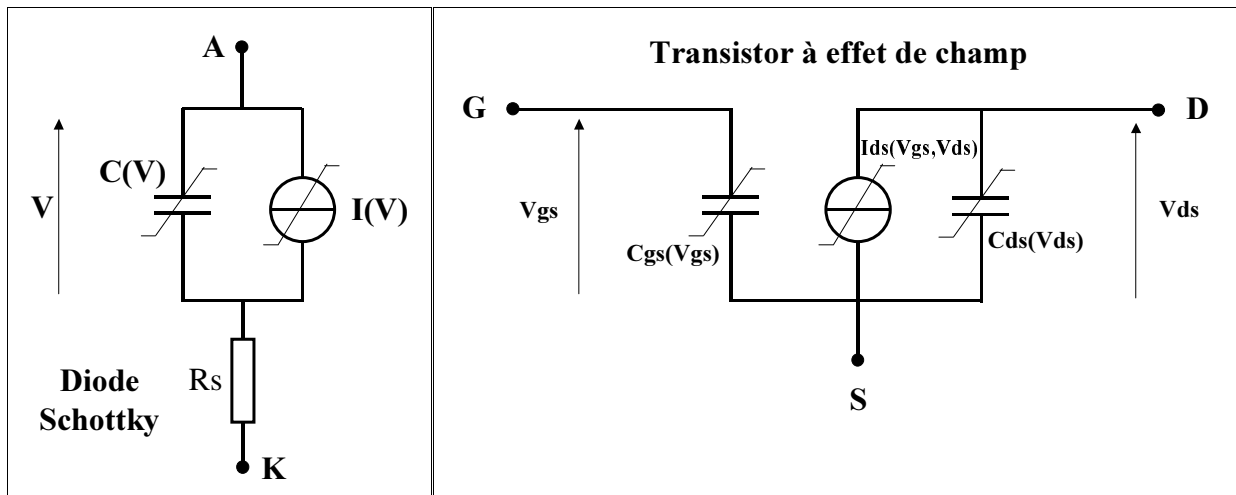


Figure I-9 : Modèle de la diode Schottky et d'un TEC

On supposera toujours les éléments non-linéaires **quasi-statiques**, c'est à dire ayant des caractéristiques (R pour une résistance, C ou Q pour une capacité, ...) ne dépendant que des variables d'états du système et non du temps.

Ces éléments non-linéaires peuvent se regrouper en deux types :

- 1) Les non-linéarités dipolaires regroupant tous les dipôles ne faisant intervenir que leurs propres variables d'état. Cette classe regroupe les résistances, inductances et capacités. Rentrent dans cette catégorie : $C(V)$, $I(V)$ pour la diode Schottky et $C_{gs}(V_{gs})$ dans le cas d'un TEC.
- 2) La seconde classe regroupe les dipôles dont l'équation d'état fait intervenir au moins une variable d'état non caractéristique du dipôle. Cette classe regroupe par exemple les transconductances des transistors ainsi que les capacités dont la charge est fonction de deux tensions. Les éléments $I_{ds}(V_{gs}, V_{ds})$ et $C(V_{gs}, V_{ds})$ du TEC rentrent dans cette classe.

Quel que soit le mélangeur envisagé, la conversion de fréquence résulte toujours de l'excitation d'une non-linéarité d'un composant actif. L'étude suivante permet d'appréhender les phénomènes de conversion de fréquence et de les quantifier approximativement.

c) Etude des mélangeurs à non-linéarité

Lors de la détermination précise des performances électriques d'un mélangeur, toutes les non-linéarités des composants le constituant doivent être considérées. Néanmoins, une et une seule non-linéarité est nécessaire à la conversion de fréquence, les autres ne jouant qu'un rôle secondaire, bénéfique ou néfaste.

Pour interpréter aisément les principes des mélangeurs, mais aussi pour évaluer leurs performances, seule la non-linéarité utile à la conversion de fréquence sera considérée.

Toutes les cellules de mélange actuelles sont basées sur la non-linéarité d'une source de courant (les mélangeurs paramétriques exploitant la non-linéarité d'une capacité sont en effet obsolètes). Dans le cas d'une diode, cette source de courant est $I(V)$ décrite en Figure I-9 ; pour un transistor il s'agit de $I_{ds}(V_{gs}, V_{ds})$.

D'une manière générale considérons une source de courant de la forme $I(V_1, V_2)$ pour laquelle V_1 correspond à la tension V_{OL} qui est de forte amplitude et V_2 correspond à V_{RF} qui est de faible amplitude. Dans ces conditions, on peut développer au premier ordre cette non-linéarité par rapport à la variable $V_2 = V_{RF}$.

$$I(t) = I(V_{OL}(t), V_{RF}(t)) = I(V_{OL}(t), 0) + \underbrace{\left. \frac{\partial I(V_1, V_2)}{\partial V_2} \right|_{\substack{V_1 = V_{OL}(t) \\ V_2 = 0}}}_{G(t)} \bullet V_{RF}(t) \quad (\text{I-22})$$

- ❖ Le premier terme est de même périodicité que le signal OL. Comme la source de courant est non-linéaire, ce signal n'est pas sinusoïdal et sa décomposition en série de Fourier exprime les amplitudes des harmoniques du signal OL en sortie.
- ❖ Le second terme est le terme utile car il est le produit du signal d'entrée RF par un coefficient dépendant du temps noté $G(t)$.

Ce coefficient est périodique, de même période que le signal OL, et non sinusoïdal. La décomposition en série de Fourier donne :

$$G(t) = \sum_k G_k \cos(k \omega_{OL} t + \varphi_k) \quad (I-23)$$

En considérant que le signal RF est un signal sinusoïdal de fréquence f_{RF} , le second terme de (I-22) s'écrit :

$$I(t) = \sum_k G_k V_{RF} \cos[(\omega_{RF} \pm k \omega_{OL})t + \varphi_k] \quad (I-24)$$

La raie spectrale d'intérêt du courant de sortie est de fréquence $f_{RF} - f_{OL}$. La présence de cette raie (premier terme de la relation (I-24)) démontre que l'on a effectivement transposition de fréquence. Il est important de noter que l'amplitude de cette composante spectrale est proportionnelle à G_1 .

Deux conclusions fondamentales peuvent être dégagées de cette étude :

- 1) Les mélangeurs reposent sur l'utilisation de la non-linéarité d'une conductance (ou d'une transconductance) issue d'une source de courant commandée. Cette conductance lie les variables d'entrée RF et de sortie FI par une relation linéaire mais variant dans le temps suivant le signal OL (relation (I-22)) cause de la transposition de fréquence.
- 2) L'amplitude de la raie d'intérêt en sortie est proportionnelle à l'amplitude du terme fondamental (à la fréquence f_{OL}) de la décomposition en série de Fourier de la conductance considérée : G_1 . Maximiser cette composante spectrale revient donc à maximiser la valeur de G_1 . Empiriquement, il faudra pour cela que l'excursion de $G(t)$ soit la plus importante possible.

La relation (I-22) suppose que le signal RF est de très faible amplitude. Lorsque cela n'est plus le cas, viennent se rajouter à $I(t)$ les termes suivants :

$$\left. \frac{\partial^2 I(V_1, V_2)}{(\partial V_2)^2} \right|_{\substack{V_1 = V_{OL}(t) \\ V_2 = 0}} \cdot \frac{(V_{RF}(t))^2}{2} + \left. \frac{\partial^3 I(V_1, V_2)}{(\partial V_2)^3} \right|_{\substack{V_1 = V_{OL}(t) \\ V_2 = 0}} \cdot \frac{(V_{RF}(t))^3}{6} + \dots \quad (I-25)$$

La présence de ces termes entraîne une relation non-linéaire entre les signaux RF et FI. Ceci provient de la non-linéarité de la conductance vis à vis du signal RF. Nous verrons au paragraphe II.3 l'impact de cette non-linéarité.

Les principes des mélangeurs micro-ondes étant exposés, nous allons présenter dans le prochain paragraphe leurs principales caractéristiques électriques.

II. Caractéristiques des mélangeurs

Afin de définir les caractéristiques des mélangeurs, il convient d'introduire des notations communes à toutes les définitions. Un mélangeur comporte trois accès : deux entrées notées ici RF et OL et une sortie notée FI. Comme le décrit la Figure II-1, les deux entrées sont attaquées par des signaux monochromatiques de fréquence f_{RF} et d'amplitude faible pour l'accès RF ; de fréquence f_{OL} et de forte amplitude pour l'accès OL. Le signal de sortie comporte, outre la raie d'intérêt à la fréquence $f_{FI}=f_{RF}-f_{OL}$, de nombreuses composantes spectrales parasites. Nous avons de plus fixé arbitrairement les fréquences en respectant : $f_{OL} \ll f_{FI} < f_{RF}$.

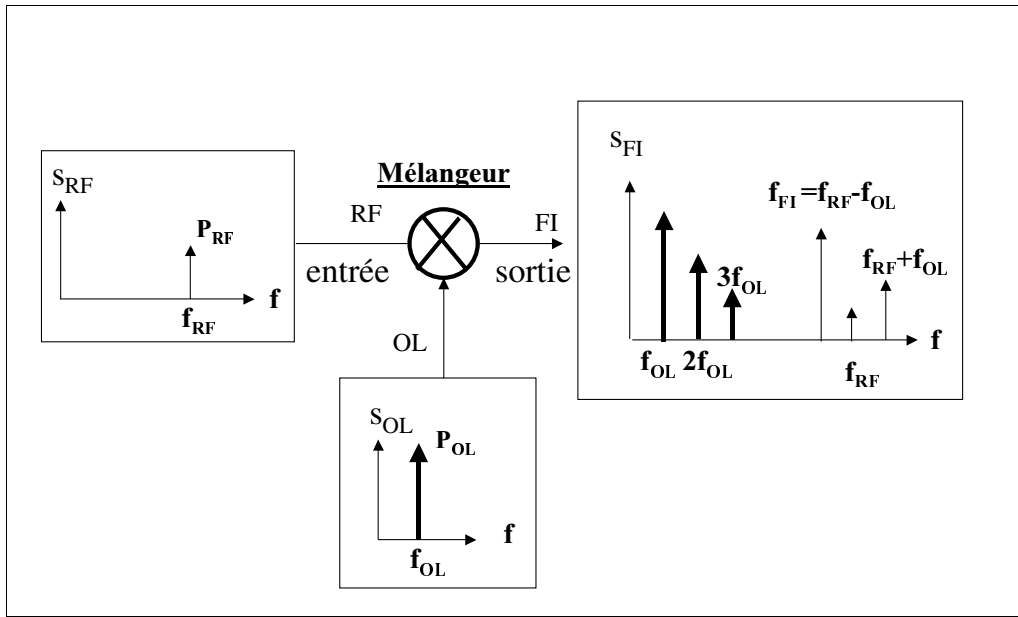


Figure II-1 : Notations

On notera la puissance à l'accès X à la fréquence f_0 : $P_X(f_0)$. Par souci de clarté on pose $P_{RF}(f_{RF})=P_{RF}$ et $P_{OL}(f_{OL})=P_{OL}$.

Nous exposons dans les paragraphes suivants les caractéristiques principales des mélangeurs permettant leurs évaluations.

II.1. Gain de conversion et puissance OL nécessaire

Une caractéristique importante de tout système est son gain. Pour les mélangeurs, on définit le gain de conversion **composite**, rapport de la puissance du signal de sortie à la fréquence d'intérêt (il s'agit de la puissance sur l'accès FI à la fréquence f_{FI} : noté $P_{FI}(f_{FI})$) sur la puissance disponible à l'entrée à la fréquence d'entrée (dans notre cas : accès RF pour la fréquence f_{RF} : $P_{RF}(f_{RF}) = P_{RF}$).

$$G_{c, composite} = \frac{P_{FI}(f_{FI})}{P_{RF, disponible}} \quad (II-1)$$

Dans la suite on omettra l'adjectif **composite**.

Le premier terme de l'équation (I-19) montre que l'amplitude du signal en sortie à la fréquence f_{FI} est : $k_2 V_{RF} V_{OL}$. Dans ces conditions, la puissance à la fréquence f_{FI} peut être évaluée par la relation :

$$P_{FI}(f_{FI}) \propto P_{RF} \bullet P_{OL} \quad (\text{II-2})$$

Dans cette relation le signe $\propto$ signifie « proportionnel à ». Si l'on reporte la relation (II-2) dans l'équation (II-1), on montre que :

$$G_c \propto P_{OL} \quad (\text{II-3})$$

Les deux dernières relations négligent les produits d'intermodulation d'ordre supérieur ou égal à 4. Cette hypothèse n'est donc valable que si les puissances P_{RF} et P_{OL} ne sont pas trop importantes. Dans le cas contraire, il est bien évident que $P_{FI}(f_{FI})$ n'augmente pas indéfiniment lorsque P_{OL} ou P_{RF} augmentent : il y a phénomène de saturation :

- ❖ Par rapport à P_{RF} . A faible puissance RF, le gain de conversion est constant. Au-delà d'un seuil de P_{RF} , la puissance $P_{FI}(f_{FI})$ sature et devient constante, le gain diminue alors lorsque la puissance RF augmente.
- ❖ Par rapport à P_{OL} . A faible puissance OL, le gain de conversion est proportionnel à la puissance OL (voir la relation (II-3)), au delà d'un certain seuil de puissance OL, le gain de conversion devient constant et peut même diminuer lorsque P_{OL} augmente. C'est ce seuil de puissance OL, qui correspond au maximum de gain de conversion, qui est nommé puissance OL nécessaire.

Le gain d'un système est une performance importante mais qui ne suffit pas à l'évaluation de la qualité de ce dernier. Il faut, en effet, de plus considérer les caractéristiques permettant de mesurer la dégradation de l'information qu'engendre ce système. Ces caractéristiques sont :

- Le facteur de bruit qui évalue le bruit propre qu'apporte le système au signal qui le traverse.
- La linéarité qui évalue le non-déformation du signal informatif.

II.2. Facteur de bruit

La définition du facteur de bruit des mélangeurs n'est pas aussi aisée que dans les systèmes linéaires du fait de la pluralité des fréquences mises en jeu.

Tout d'abord, il ne faut considérer que les accès RF et FI : le système étudié englobe alors l'oscillateur local dont on spécifiera les caractéristiques de bruit. Usuellement, seul le bruit thermique de ce dernier est considéré et non son bruit de phase ni d'amplitude.

De plus, pour la référence de bruit en entrée, ne sont considérées qu'une ou deux sources de bruit : une à la fréquence f_{RF} et l'autre à la fréquence image (voir paragraphe II.5). Les autres sources de bruit en entrée (bruit du canal d'entrée) aux fréquences différentes de f_{RF} et f_{image} , qui contribuent à l'augmentation du bruit à l'accès FI à la fréquence f_{FI} , sont intégrées au système étudié.

Afin de préciser les différentes définitions utilisées pour l'évaluation du facteur de bruit d'un mélangeur, nous allons considérer le mélangeur comme un système linéaire multipolaire (pour cette modélisation, se reporter au chapitre 2 sur les matrices de conversion) constitué d'une sortie représentant la puissance disponible sur l'accès FI à la fréquence f_{FI} et de deux entrées représentant les puissances disponibles sur l'accès RF aux fréquences f_{RF} et f_{image} . La Figure II-2 montre une telle représentation dans laquelle $G_{RF \rightarrow FI}$ et $G_{image \rightarrow FI}$ sont des gains en puissance correspondant aux deux mécanismes de conversion de fréquence considérés.

La pluralité des définitions provient de l'état du bruit en entrée à la fréquence image.

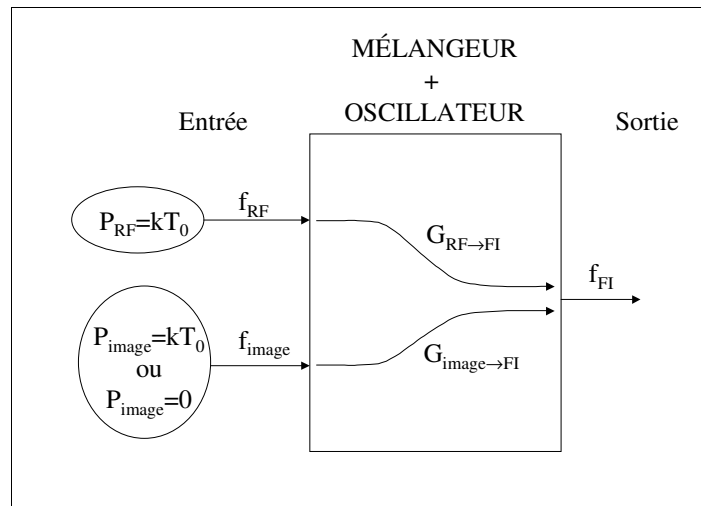


Figure II-2 : Représentation en bruit

II.2.1. Facteur de bruit en simple bande (SSB noise)

Dans cette définition, l'accès d'entrée à la fréquence f_{image} est englobé dans le système, seuls subsistent deux accès : l'entrée RF à la fréquence f_{RF} et la sortie FI à la fréquence f_{FI} .

a) Définition de l'IEEE

Pour cette définition, la source de bruit thermique du générateur sur l'accès d'entrée à la fréquence image est éliminée ($P_{image}=0$). Tout se passe donc comme s'il n'y avait aucun bruit incident sur l'entrée à la fréquence f_{image} . Ceci est loin d'être une simple hypothèse de travail car un filtre est souvent placé à l'entrée des systèmes de conversion de fréquence permettant la réjection de la fréquence image. Ces filtres présentent soit un court-circuit soit un circuit ouvert à la fréquence f_{image} , éliminant tout bruit du générateur à cette fréquence.

L'expression du facteur de bruit est basée sur la définition de North [I.3] :

$$F_{SSB}^{IEEE} = 1 + \frac{B_d}{G_{RF \rightarrow FI} * kT_0} \quad (\text{II-4})$$

Dans cette expression, B_d est la puissance disponible de bruit en sortie dans une bande de 1Hz ajouté par le système considéré (mélangeur+oscillateur).

b) Autre définition

Lorsque le filtrage de la fréquence image n'est pas possible, il faut tenir compte du bruit en entrée à cette fréquence. La définition suivante prend en compte ce bruit en supposant que les sources de bruit d'entrée ont la même température T_0 .

$$F_{SSB} = 1 + \frac{G_{RF \rightarrow image}}{G_{RF \rightarrow FI}} + \frac{B_d}{G_{RF \rightarrow FI} * kT_0} \quad (\text{II-5})$$

En comparant les équations (II-5) et (II-4), on s'aperçoit que la prise en compte du bruit à la fréquence image augmente le facteur bruit d'autant plus que ce bruit sera amplifié.

II.2.2. Facteur de bruit en double bande (DSB noise)

Dans ce cas, on considère les deux entrées excitées par des générateurs de bruits thermiques identiques et décorrélés. L'accès d'entrée à la fréquence f_{image} n'est plus intégré au système et la référence de puissance de bruit correspond alors à la somme des puissances des deux sources en entrée. L'expression du facteur de bruit est alors :

$$F_{DSB} = 1 + \frac{B_d}{(G_{RF \rightarrow image} + G_{RF \rightarrow FI}) * kT_0} \quad (\text{II-6})$$

II.2.3. Cascade de quadripôle : formule de Friis

Lorsque plusieurs quadripôles sont associés en série, il est possible, connaissant les gains disponibles et les facteurs de bruit de chaque quadripôle, de calculer le facteur de bruit du système global par la formule de Friis qui porte le nom de son inventeur [I.4].

Dans le cas de l'association série de trois quadripôles on a :

$$F_{total} = F_1 + \frac{F_2 - 1}{G_1} + \frac{F_3 - 1}{G_1 G_2} \quad (\text{II-7})$$

Dans cette expression, G_n et F_n sont le gain disponible et le facteur de bruit de l'étage n lorsque les étages sont numérotés de l'entrée vers la sortie.

Une interprétation importante de cette formule permet de dégager quelques règles de conception.

1. Le premier étage d'une chaîne de transmission de signal est en principe un étage préamplificateur à gain suffisant (G_1 grand) pour que le facteur de bruit total se réduise au facteur de bruit de ce premier étage, les second et troisième termes de (II-7) étant négligeables. Le système sera alors d'autant plus performant que le facteur de bruit de son premier étage sera faible. Il conviendra alors d'optimiser le premier étage d'une chaîne non seulement vis à vis de son facteur de bruit mais aussi de son gain.
2. Il faudra veiller à ce que le second étage ne présente pas des pertes comparables au gain du premier sous peine d'augmenter la contribution du troisième terme de (II-7) sur le facteur de bruit total de la chaîne.

II.3. Linéarité

Comme pour les amplificateurs, la linéarité est l'une des trois caractéristiques (avec le gain et le facteur de bruit) jugeant la qualité des mélangeurs.

Après la présentation des conséquences de l'intermodulation d'ordre 4, nous présenterons plusieurs indicateurs de linéarité.

II.3.1. Les méfaits du quatrième degré du développement de la non-linéarité

Rappelons les termes générés par le degré quatre du développement d'une non-linéarité qui ont pour fréquence f_{FI} :

$$\diamond k_2 V_{RF} V_{OL} \cos(\omega_{RF} - \omega_{OL})t \text{ introduit dans (I-19).}$$

$$\diamond \frac{3}{4} k_4 (V_{OL} V_{RF}^3 + V_{RF} V_{OL}^3) \cos((\omega_{rf} - \omega_{ol})t) \text{ introduit dans (I-21).}$$

Les termes dont l'amplitude est proportionnelle à V_{RF} contribuent à la translation de fréquence désirée : le spectre du signal d'entrée est translaté sans déformation de la fréquence f_{RF} vers la fréquence f_{FI} . Par contre, le terme en V_{RF}^3 contribue à la translation de fréquence tout en déformant l'amplitude du spectre du signal RF à cause de l'élévation à la puissance 3. Cette déformation va donc altérer le signal informatif véhiculé, voire détruire une grande partie de l'information qu'il contient.

La Figure II-3 est une illustration de la dégradation de l'information engendrée par le terme en V_{RF}^3 .

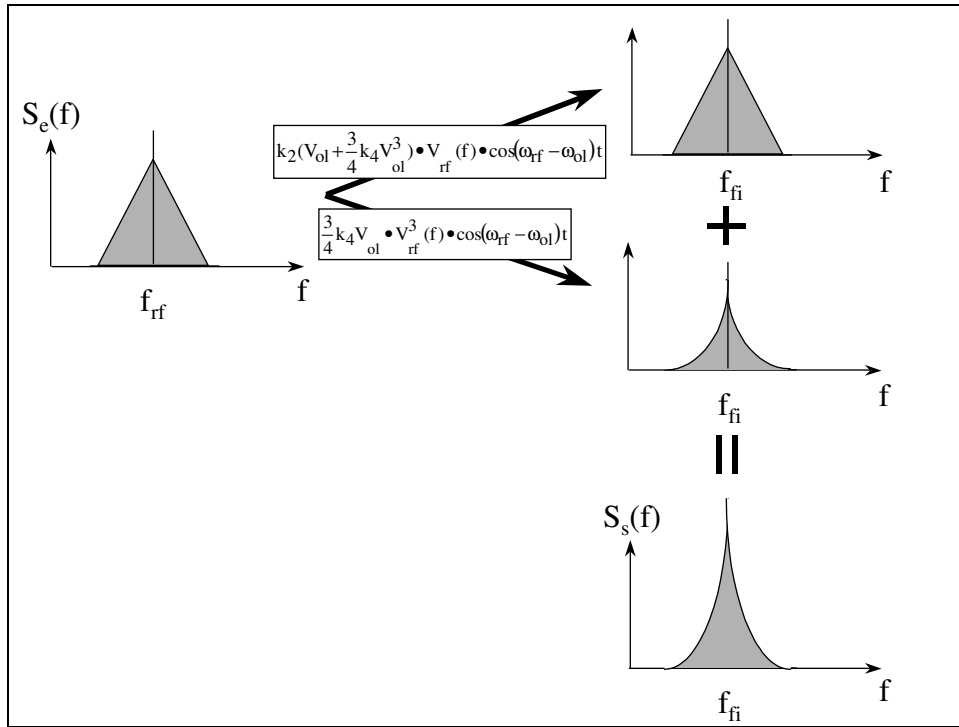


Figure II-3 : Effet du terme s4

Etant donnée que ce terme est proportionnel à V_{RF}^3 , il est nommé produit d'intermodulation d'ordre 3 (ordre 3 vis à vis du signal RF) bien qu'il soit la conséquence du degré quatre du développement de la non-linéarité.

Etant donné la dégradation qu'il engendre sur l'information, le terme en V_{RF}^3 devra être minimisé. Le paragraphe suivant décrit les indicateurs usuels d'évaluation de la linéarité des mélangeurs.

II.3.2. L'évaluation de la linéarité

L'évaluation de la linéarité est identique à celle d'un amplificateur : deux raies spectrales de fréquences voisines $f_{RF1}=f_0 + \delta$ et $f_{RF2}=f_0 - \delta$ et de mêmes puissances sont appliquées sur l'accès RF.

Note : δ doit être petit devant le rapport de f_0 par le plus grand des coefficients de qualité du circuit étudié, ceci afin de garantir que les deux raies subiront les mêmes effets de filtrage.

Aux alentours de la fréquence f_0 , la sortie est constituée des raies suivantes (se reporter Figure II-4):

- 1) Les deux raies d'intérêts aux fréquences $f_{FI1}=f_0-f_{OL}+ \delta$ et $f_{FI2}=f_0-f_{OL}-\delta$ qui résultent du premier terme de l'expression (I-19). Ces raies ont des puissances identiques notées P_{IM1} pour Puissance d'InterModulation d'ordre 1 car l'ordre

d'intermodulation sur le signal RF est 1. La puissance P_{IM1} est proportionnelle à P_{OL} et à P_{RF} .

- 2) Deux raies supplémentaires aux fréquences : $f_1=f_0 - f_{OL}+ 3\delta$ et $f_2=f_0-f_{OL}-3\delta$, prévues par le premier terme de l'expression (I-21). Ces raies ont des puissances identiques notées P_{IM3} pour Puissance d'InterModulation d'ordre 3 car l'ordre d'intermodulation sur le signal RF est 3 (voir note ci-avant). La puissance P_{IM3} est proportionnelle à P_{OL} et à $(P_{RF})^3$.

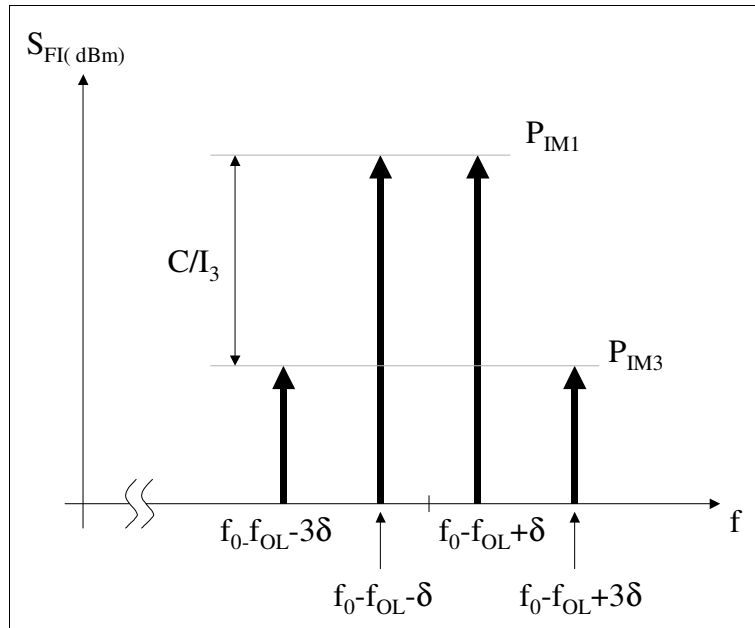


Figure II-4 : Spectre du signal FI

Plusieurs indicateurs de linéarité existent parmi lesquels :

a) L'indicateur de linéarité : C/I_3

Un moyen d'évaluer la linéarité des mélangeurs est de mesurer le rapport de puissance entre la puissance d'intermodulation d'ordre 1 et celle d'ordre 3, est alors introduit le rapport :

$$C/I_3 = \frac{P_{IM1}}{P_{IM3}} \quad (\text{II-8})$$

C/I_3 signifiant « Carrier / Intermodulation of 3rd order ». Notons que ce paramètre est indépendant de la puissance OL mais inversement proportionnel au carré de P_{RF} . Il faut donc veiller à spécifier le niveau de puissance RF lors de l'utilisation de C/I_3 .

b) L'indicateur de linéarité : Point d'interception d'ordre 3 ou $PI3$

Un autre indicateur de la linéarité d'un mélangeur est le point d'interception d'ordre 3 noté $PI3$ (IP3 en anglo-saxon pour Intercept Point of 3rd order). On introduit alors la puissance d'interception d'ordre 3 en entrée notée P_{I3e} (IIP3 en anglo-saxon) ou en sortie notée P_{I3s} (OIP3 en anglo-saxon).

La Figure II-5 illustre la détermination des paramètres précédemment introduits sur les caractéristiques évaluant les puissances du signal de sortie PIM1 et PIM3 en fonction de la puissance d'entrée P_{RF} .

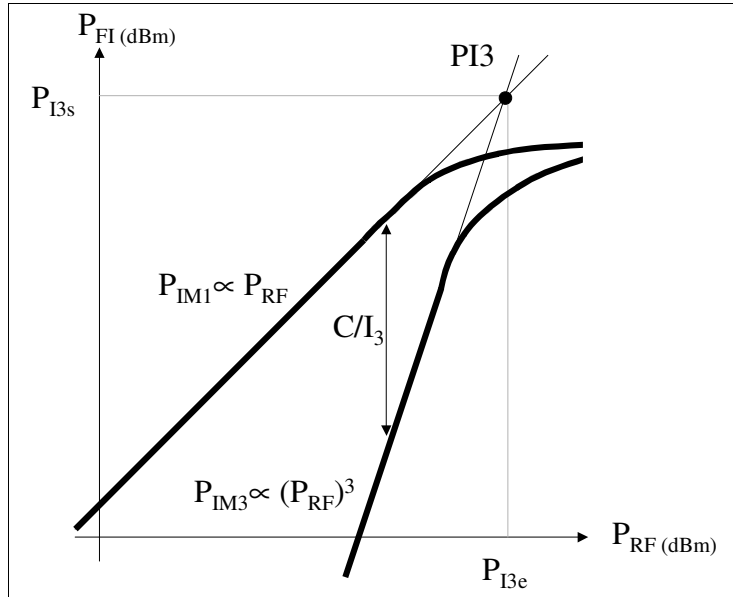


Figure II-5 : Illustration du point d'interception

P_{I3e} est nommée puissance d'interception d'ordre 3 en entrée et correspond à la puissance qu'il faudrait fournir au circuit pour que les puissances P_{IM1} et P_{IM3} soient égales. Pour cette définition, il ne faut pas considérer les saturations des puissances d'intermodulation d'ordre 1 et 3. En pratique, le point d'interception est déterminé par extrapolation des parties linéaires des courbes (faibles puissances P_{RF}).

L'avantage de ce paramètre est d'être une caractéristique absolue du mélangeur indépendante du niveau de puissance RF injectée.

On peut, connaissant le point d'interception, déduire la puissance d'intermodulation d'ordre 3 à partir de la puissance d'intermodulation d'ordre 1. L'expression (II-9) décrit cette relation pour des puissances exprimées en dBm.

$$P_{IM3s} = 3 P_{IM1s} - 2 P_{I3s} \quad (\text{II-9})$$

Remarquons de plus que la puissance P_{I3e} est indépendante de P_{OL} . Par contre, la puissance P_{I3s} qui résulte de la multiplication de P_{I3e} par le gain de conversion petit signal est proportionnel à la puissance P_{OL} injectée.

c) L'indicateur de linéarité normalisé : Q_3

i - Association de quadripôles

Comme dans le cas de la formule de Friss pour le facteur de bruit, il existe une expression liant la puissance d'interception d'ordre 3 en sortie d'un système en fonction des puissances d'interception d'ordre 3 en sortie des étages qui le composent et de leurs gains.

Par exemple, dans le cas de trois étages, numérotés de l'entrée vers la sortie, on a :

$$\frac{1}{P_{I3S}|_{total}} = \frac{1}{P_{I3S}|_3} + \frac{1}{G_3 \bullet P_{I3S}|_2} + \frac{1}{G_2 G_3 \bullet P_{I3S}|_1} \quad (II-10)$$

L'interprétation de cette formule est la suivante : lorsque le dernier étage d'une chaîne de transmission (ici le troisième) présente une valeur de gain suffisante (en principe c'est un étage de puissance) et que les étages en amont ne présentent pas de pertes trop importantes, la linéarité du système global se réduit à la linéarité de son dernier étage.

ii - Définition de Q3

Le paragraphe précédent montre qu'il est plus judicieux de raisonner sur le point d'interception en sortie. Cependant, cette donnée n'est pas absolue car elle est proportionnelle à la puissance P_{OL} . Un moyen efficace d'évaluer la linéarité, afin d'établir des comparaisons précises entre mélangeurs, est alors d'introduire le facteur de qualité Q_3 qui s'exprime par [I.5] :

$$Q_3 = \frac{P_{I3S}}{P_{OL}} \quad (II-11)$$

L'intérêt important de ce facteur est qu'il est indépendant à la fois de P_{RF} et de P_{OL} .

d) Les autres indicateurs de linéarité

D'autres indicateurs de linéarité existent parmi lesquels la puissance de compression à 1 dB qui correspond à la puissance d'entrée (ou de sortie) qui entraîne la diminution de 1 dB du gain de conversion par rapport à sa valeur en petit signal. La référence [I.6] démontre que la puissance de compression à 1 dB est environ 10 dB en dessous de la puissance d'interception d'ordre 3.

Il existe de plus, d'autres indicateurs de linéarité plus adaptés aux systèmes de communication sans fils terrestres tels que le NPR (pour Noise Power Ratio), l'ACPR (Adjacent Channel Power Ratio).

Parmi les autres caractéristiques qui peuvent, suivant le cahier des charges, être primordiales, nous allons décrire : l'isolation, la réjection des raies parasites, l'adaptation des différents ports et la consommation statique des mélangeurs.

II.4. Isolation

La sortie FI comporte de nombreuses raies spectrales et notamment les composantes spectrales aux fréquences d'entrée f_{RF} et f_{OL} . Afin d'évaluer l'importance de ces raies en sortie on exprime leurs puissances normalisées par rapport aux puissances de ces raies en entrée : on parle alors d'isolations $RF \rightarrow FI$ et $OL \rightarrow FI$. Les équations (II-12) et (II-13) définissent les expressions de ces isolations.

$$I_{RF \rightarrow FI} = \frac{P_{RF}(f_{RF})}{P_{FI}(f_{RF})} \quad (\text{II-12})$$

$$I_{OL \rightarrow FI} = \frac{P_{OL}(f_{OL})}{P_{FI}(f_{OL})} \quad (\text{II-13})$$

Le signal OL étant de forte amplitude, on retrouve à tous les accès une raie spectrale de puissance non négligeable à la fréquence f_{OL} . On définit donc de même la quantité de puissance à la fréquence OL que l'on retrouve sur l'entrée RF :

$$I_{OL \rightarrow RF} = \frac{P_{OL}(f_{OL})}{P_{RF}(f_{OL})} \quad (\text{II-14})$$

II.5. Raies parasites

Comme l'évoquent les expressions (I-18) à (I-21), sur l'accès FI, un grand nombre de signaux parasites sont présents. Ces produits de mélange parasites se classent en deux familles.

- Les signaux qui peuvent se retrouver dans la bande FI d'intérêt et qui ne peuvent donc pas être filtrés.
- Les signaux qui ont des fréquences qui ne coïncident pas avec la bande FI mais qui doivent néanmoins avoir de faibles amplitudes pour ne pas consommer de puissance ni perturber les éléments en aval du mélangeur.

Les raies prépondérantes sont bien évidemment les harmoniques de la fréquence f_{OL} , dont les amplitudes sont évaluées par rapport à la puissance P_{OL} injectée dans le mélangeur.

D'autres signaux gênants se situent aux fréquences f_{RF} , $f_{RF}-2f_{OL}$, ..., leur puissance est alors évaluée par rapport à celle injectée sur l'accès RF (P_{RF}), pour un niveau de puissance P_{OL} donné.

La Figure II-6 présente les produits d'intermodulation amenant des raies parasites dans le canal FI pour les spécifications de plan de fréquences suivantes :

- Frf : 12.75 - 14.80 GHz
- Ffi : 10.70 - 12.75 GHz
- Fol : 1.25 - 3.85 GHz

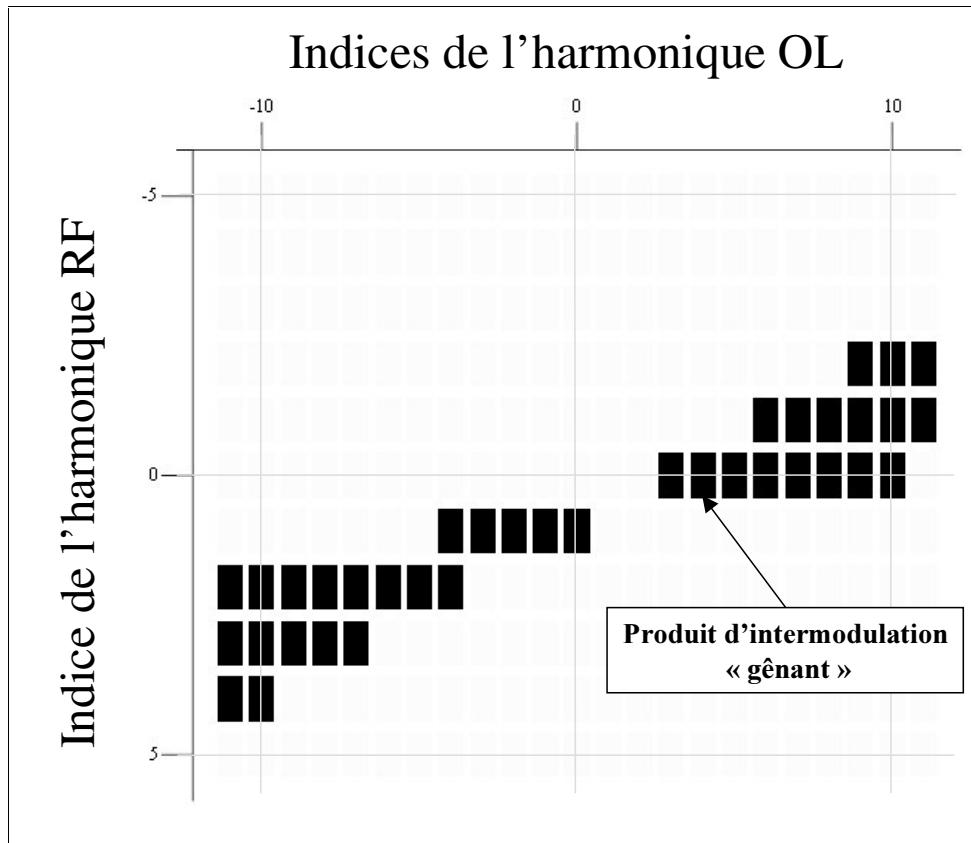


Figure II-6 : Indices d'intermodulation conduisant dans la bande FI

Outre les produits d'intermodulation des fréquences f_{RF} et f_{FI} qui se retrouvent en sortie du mélangeur et qui détériorent le signal informatif, nous allons présenter les méfaits de certains signaux d'entrée aux fréquences différentes de la fréquence f_{RF} .

La fréquence image est la fréquence d'un signal d'entrée, différente de f_{RF} , dont la puissance se retrouve translatée à la fréquence f_{FI} par les mécanismes de conversion décrit en (I-19). On a pour la fréquence f_{RF} la relation de conversion infradyne : $f_{FI} = f_{RF} - f_{OL}$, et pour la fréquence image la relation de conversion supradynne $f_{FI} = f_{image} + f_{OL}$. Un signal à la fréquence $f_{image} = f_{FI} - f_{OL}$ va donc se retrouver par conversion supradynne à la fréquence f_{FI} .

Ainsi, si un signal quelconque (le canal d'une application adjacente, du bruit, ...) se trouve à la fréquence f_{image} , il va directement se trouver ajouté au signal utile en sortie de mélangeur et parasitera l'information qu'il contient. Il convient donc, en entrée du mélangeur, de placer un filtre afin d'éviter la présence de tout signal à la fréquence image.

II.6. Adaptation des ports

Comme pour un amplificateur, on peut définir le coefficient de réflexion pour chaque port. Paramétré par la puissance P_{OL} , le coefficient de réflexion se définit toujours par le rapport de la puissance réfléchie à la fréquence f_{OL} (f_{RF} ou f_{FI}) sur la puissance incidente à la même fréquence.

Comme pour un système linéaire, une mauvaise adaptation entraînera la diminution du gain de conversion composite.

II.7. Consommation

Enfin, une donnée importante pour la classification et le choix d'un mélangeur est sa consommation statique. Cette caractéristique est primordiale pour les systèmes spatiaux embarqués et pour les systèmes portables, pour lesquels le dimensionnement des alimentations et la dissipation thermique représentent des contraintes importantes.

III. Classification des topologies.

Dans ce paragraphe, nous présenterons toutes les topologies de mélangeur existantes ainsi que leurs principe, avantages et inconvénients. Puis nous traitons le principe des structures équilibrées et présentons trois types classiques de mélangeurs doubles équilibrés. Enfin, nous proposons une synthèse de la classification des topologies élémentaires de mélangeur ainsi qu'un état de l'art des deux mélangeurs double équilibrés le plus souvent rencontrés dans la littérature.

III.1. Les différents types de mélangeur

On distingue deux grandes classes de mélangeurs :

- 1) Les mélangeurs actifs qui mettent en œuvre des composants actifs (transistors polarisés en amplificateur) dont on utilise les potentialités d'amplification. Ces mélangeurs peuvent avoir des gains de conversion supérieurs à 0dB.
- 2) Les mélangeurs passifs qui mettent en œuvre soit des composants passifs (diodes), soit des éléments actifs dont ne sont pas exploitées les capacités d'amplification (transistors non-polarisés sur le drain, par exemple). Ces mélangeurs ont des gains de conversion inférieurs à 0dB.

III.1.1. Les mélangeurs actifs

Les mélangeurs actifs mettent en œuvre des transistors pour lesquels on souhaite bénéficier des capacités d'amplification. Pour un transistor à effet de champ, le signal d'entrée RF doit donc, comme dans le cas d'un amplificateur, être appliqué à la grille du transistor tandis que le signal de sortie FI est recueilli sur le drain. La non-linéarité utilisée est la transconductance du transistor et les mélangeurs résultants se nomment « mélangeur à transconductance ». Le signal OL a pour fonction de moduler cette non-linéarité et peut, pour ce faire, être injecté soit sur la grille, soit sur le drain ou soit sur la source du transistor.

Les paragraphes suivants décrivent ces différents cas. Nous n'avons considéré que le cas des mélangeurs à transistor à effet de champ. Bien évidemment, ces structures peuvent s'implémenter de la même manière à l'aide de transistors bipolaires.

a) *Injection de l'OL sur la grille*

Cette topologie, décrite sur la Figure III-1(a), fut la première à mettre en œuvre un transistor à effet de champ pour réaliser la fonction de mélange. La configuration se nomme habituellement « mélangeur sur grille » pour rappeler que le signal OL est injecté sur la grille. Ce signal va moduler la transconductance via la tension V_{gs} , la tension V_{ds} étant maintenue constante : la non-linéarité exploitée est donc $G_m(V_{gs})$. La Figure III-1(b) décrit la variation de la transconductance en fonction du signal OL ou « cycle de pompage ». Le pompage optimal correspond au maximum d'excursion de G_m , le transistor doit donc fonctionner en régime saturé ($V_{ds} > 1V$) et la tension V_{gs} doit décrire au minimum la plage $[V_p ; 0V]$ où V_p est la tension de pincement [I.7et 8].

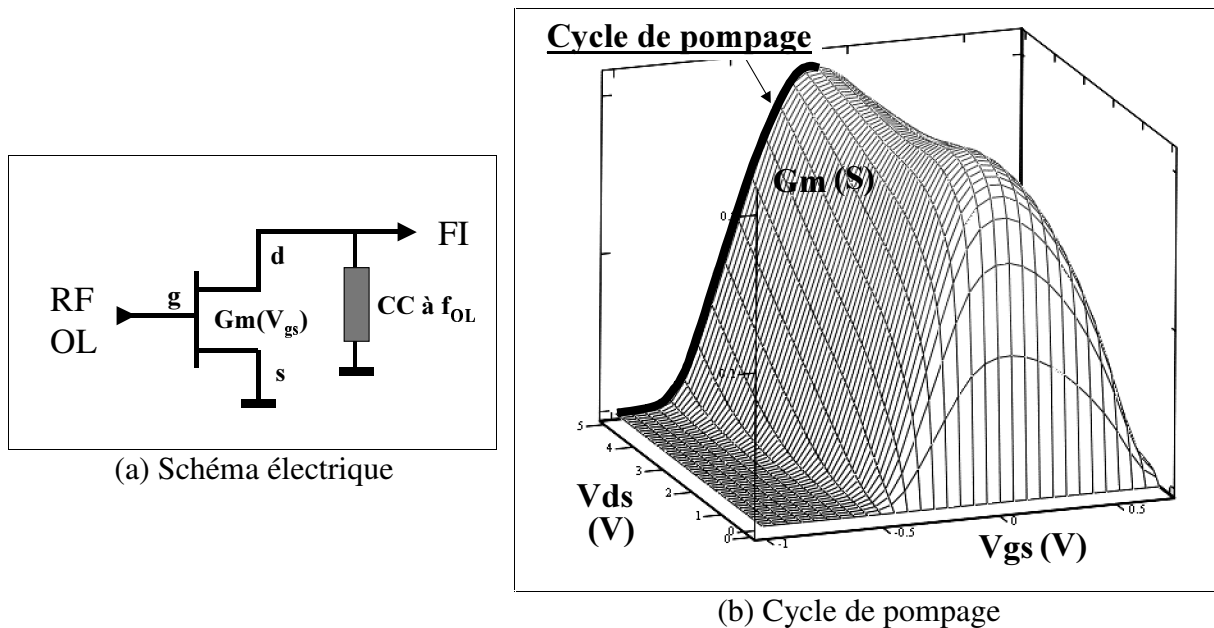


Figure III-1 : Mélangeur sur Grille.

Ce type de mélangeur présente de bonnes performances en gain et nécessite des puissances P_{OL} modérées par rapport aux autres configurations. Ses caractéristiques serviront donc de référence pour évaluer les autres types de mélangeur.

Une condition essentielle pour le bon fonctionnement de cette topologie est de garantir que la tension V_{ds} reste constante pendant le pompage par le signal OL. Ceci est assuré par la présence d'un filtre sur le drain du transistor qui présente un court-circuit à la masse pour la fréquence f_{OL} .

L'inconvénient majeur du mélangeur sur grille est qu'il nécessite l'utilisation d'un combineur de puissance pour l'application des signaux RF et OL sur la grille du transistor. Cet élément est soit réalisé par un coupleur hybride lorsque les fréquences f_{OL} et f_{RF} sont voisines soit par des filtres diplexeurs ([I.9]) dans le cas contraire. Son utilisation engendrera donc des pertes, du bruit et réduira la bande passante du mélangeur tout en augmentant la complexité et les dimensions du circuit résultant.

b) *injection de l'OL sur le drain*

Cette topologie, qui se nomme habituellement « mélangeur sur drain », est décrite sur la Figure III-2(a). Le signal OL va moduler la transconductance via la tension V_{ds} , la tension V_{gs} étant maintenue constante grâce à la présence d'un filtre sur la grille présentant un court-circuit à la masse pour la fréquence f_{OL} . La non-linéarité exploitée est donc cette fois $G_m(V_{ds})$.

Le cycle de pompage décrit sur la Figure III-2(b) montre que l'excursion optimale de G_m est obtenue pour V_{gs} voisin de 0V et pour des variations de V_{ds} comprises dans la zone de fonctionnement ohmique du transistor ($V_{ds} < 1V$). Dans ces conditions, l'excursion de G_m est identique à celle obtenue avec un mélangeur sur grille ce qui se traduit par des gains de conversion identiques pour les deux configurations ([I.10]).

Notons que la Figure III-2(b) ne représente pas la transconductance pour les valeurs négatives de V_{ds} . Pour une valeur nulle de V_{gs} , la transconductance est une fonction impaire de V_{ds} : un pompage symétrique par rapport à une tension de polarisation nulle, entraîne donc une excursions double de la transconductance de $-G_{m,max}$ à $+G_{m,max}$ et donc une augmentation de 6 dB du gain de conversion. C'est pour cette raison, que le gain des mélangeurs sur drain peut, sous très forte puissance OL, être supérieur à celui des mélangeurs sur grille.

D'une manière générale, le gain des mélangeurs sur drain est principalement limité par les impédances de fermeture aux différents accès et par les non-linéarités autres que la transconductance, en particulier, la non-linéarité de la conductance de sortie.

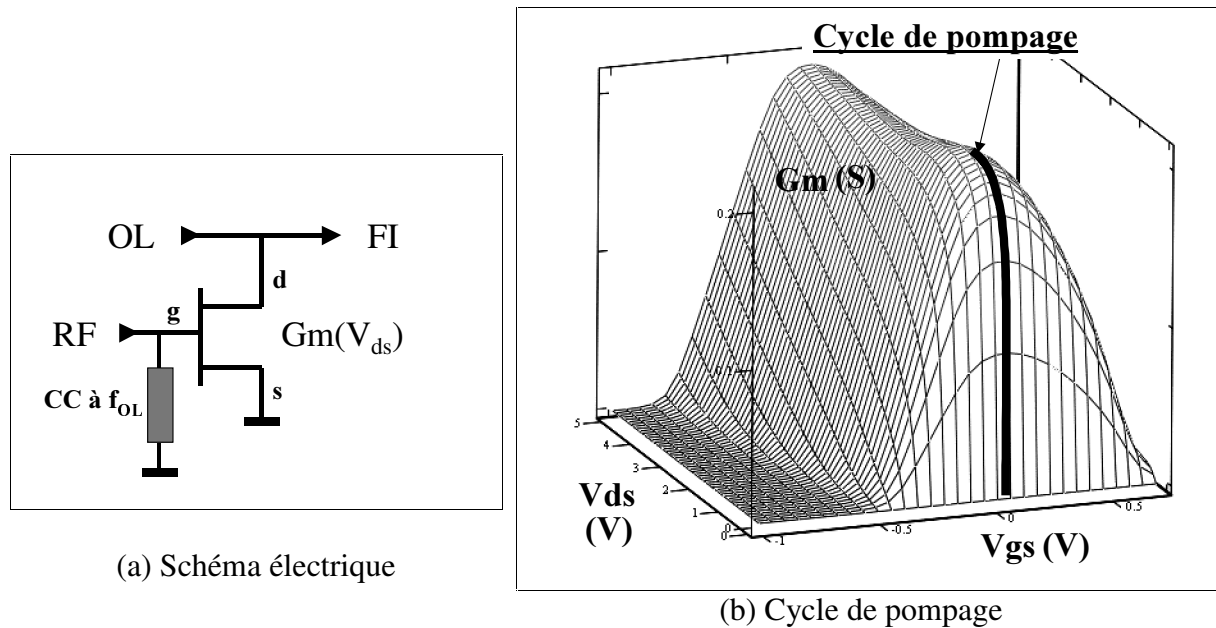


Figure III-2 : Mélangeur sur Drain.

Par contre, le signal OL ne bénéficie plus de l'amplification du transistor comme dans le cas du mélangeur sur grille. La puissance P_{OL} nécessaire est donc plus importante. La référence [I.11] montre un écart de 10 dB entre les deux configurations.

Enfin, le bruit des mélangeurs sur grille est légèrement plus fort que celui des mélangeurs sur drain [I.12]. Dans ce dernier cas, le fonctionnement du transistor se situe principalement en régime linéaire, ce qui lui confère un bruit propre plus faible que dans le cas d'un fonctionnement en régime saturé.

Comme dans le cas des mélangeurs sur grille, cette topologie nécessite toujours l'emploi de filtres diplexeurs pour la séparation des signaux OL et FI.

Cet inconvénient peut être cependant surmonté si le signal OL est injecté sur l'électrode inutilisée du transistor : la source ([I.13]), comme le décrit le prochain paragraphe.

c) Injection de l'OL sur la source

La topologie résultante est présentée sur la Figure III-1(a). Elle se nomme usuellement « mélangeur sur source ». Le signal OL module la transconductance conjointement par les tensions V_{gs} et V_{ds} . Néanmoins, comme le décrit la Figure III-1(b), le cycle de pompage de G_m est essentiellement la conséquence de la variation de V_{gs} , la non-linéarité exploitée est donc, comme pour le mélangeur sur grille, $G_m(V_{gs})$.

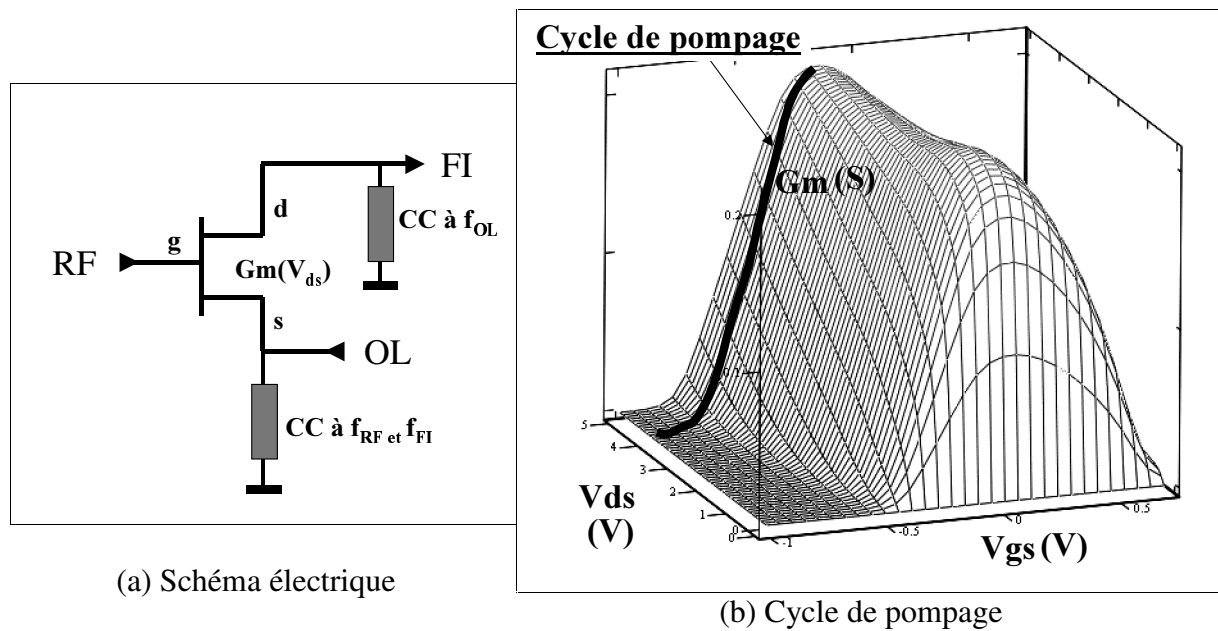


Figure III-3 : Mélangeur sur Source.

Cette topologie fonctionne de manière rigoureusement identique à celle du mélangeur sur grille et présente donc les mêmes caractéristiques ([I.14]).

Les précautions à prendre pour un fonctionnement optimal sont au nombre de deux :

1. Il faut présenter, comme dans le cas des mélangeurs sur grille, un court-circuit sur le drain du transistor de mélange à la fréquence f_{OL} .
2. Pour assurer un fort gain de conversion et une stabilité non-linéaire inconditionnelle, il faut éliminer toute contre-réaction potentielle sur la source du transistor de mélange. Ceci est

réalisé à l'aide d'un filtre placé sur cet accès et présentant un court-circuit à la masse pour les fréquences f_{RF} et f_{FI} . Ce point particulier fera l'objet d'une partie du second chapitre.

L'avantage de cette structure réside dans le fait qu'elle ne nécessite pas la mise en œuvre de coupleur de signaux (coupleur hybride ou filtre diplexeur). Elle nécessite cependant, pour assurer un fonctionnement correct, la mise en œuvre d'un filtre présentant un court-circuit entre la source du transistor et la masse aux fréquences f_{RF} et f_{FI} , ceci afin de garantir à la fois un fort gain de conversion mais aussi une stabilité inconditionnelle.

d) Mélangeur à 2 transistors : structure classique

Comme le mélangeur sur source, le mélangeur décrit dans ce paragraphe représente une amélioration du mélangeur sur drain à un transistor. Il met en œuvre deux transistors connectés comme indiqués sur la Figure III-4(a). Cette configuration est appelée transistor double grille car, historiquement, cette topologie ne nécessitait qu'un seul transistor intégrant une seconde grille [I.15].

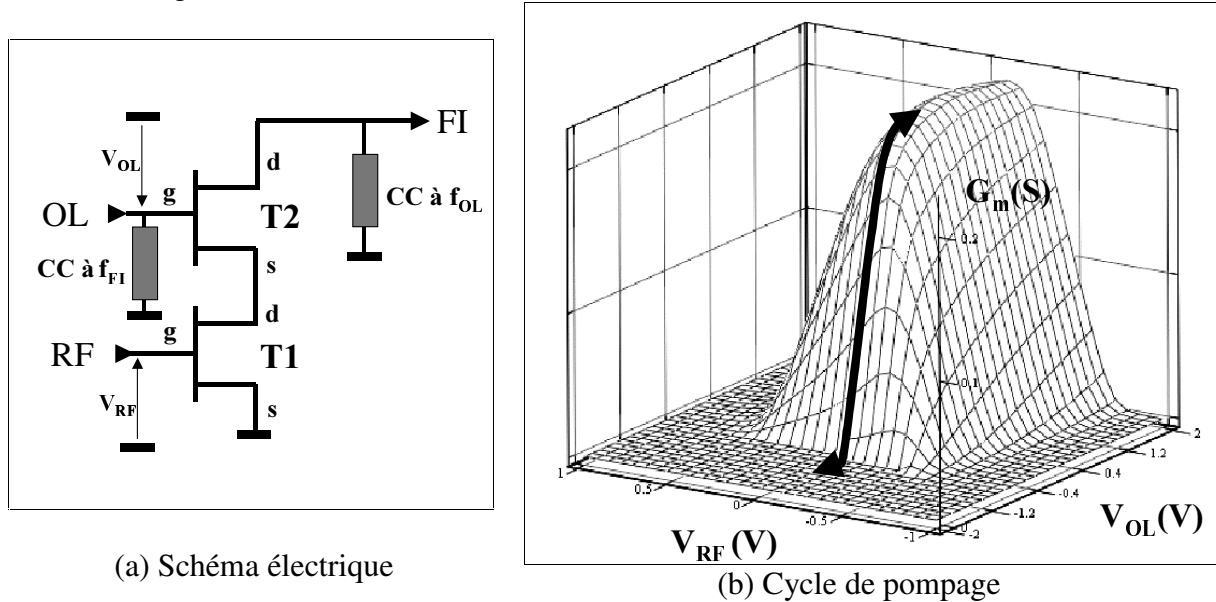


Figure III-4 : mélangeur à 2 transistors : structure classique

Le mélange est réalisé par la transconductance du transistor T1. Comme dans le cas des mélangeurs sur drain, cette transconductance est modulée par la tension drain-source du transistor T1 via le transistor T2 qui fonctionne en amplificateur drain-commun [I.16]. Le drain de ce dernier est en effet court-circuité à la masse à la fréquence f_{OL} par le filtre présenté sur la Figure III-4 (a).

La Figure III-4 (b) montre que la variation de la tension V_{OL} engendre la modulation souhaitée de la transconductance du transistor T1.

Enfin, le signal transposé à la fréquence f_{FI} et issu du mélange par le transistor T1 est amplifié en puissance par le transistor T2 qui est alors en configuration grille commune, sa grille étant connectée à la masse à la fréquence f_{FI} (cf. Figure III-4 (a)).

Les avantages de cette topologie à deux transistors sont tous basés sur l'ajout du transistor T2. Par rapport aux mélangeurs sur drain, cette topologie se caractérise par:

- ❖ Une nette diminution de la puissance OL nécessaire, grâce au transistor T2 réalisant l'amplification du signal OL.
- ❖ Une bonne isolation OL-RF, due à la faible valeur de la capacité C_{gd} de T1.
- ❖ Enfin, la présence de T2 dispense de l'utilisation de coupleur de signaux rendant cette topologie adaptée à l'intégration MMIC.

Par contre, le gain de conversion de cette structure est nettement inférieur à celui des structures à un transistor et peut même être comparable à celui des mélangeurs passifs [I.17]. Ce faible gain s'explique par la forte désadaptation entre le transistor T1 fonctionnant en mélangeur et le transistor T2 en configuration grille commune à la fréquence f_{FL} . Cette désadaptation entraîne de plus l'augmentation du facteur de bruit ainsi que la dégradation de la linéarité du mélangeur.

III.1.2. Les mélangeurs passifs

Les mélangeurs passifs mettent en œuvre des éléments non-linéaires fonctionnant en commutation au rythme du signal OL. Ce mode de fonctionnement ne permet pas d'obtenir un gain de conversion supérieur à 0dB, et engendre souvent des pertes de conversion comprises entre 6 et 10dB.

Les non-linéarités se modélisent idéalement par des interrupteurs faisant des mélangeurs passifs de simples modulateurs. Dans le domaine micro-ondes, cette représentation devient fortement imprécise et il est préférable soit de modéliser les éléments non-linéaires par des conductances variables suivant le signal OL, soit de modéliser les diodes et transistors par leurs schémas équivalents complets présentés Figure I-9.

Nous allons décrire les trois types de cellules de mélange passives couramment rencontrées. Il est à noter que ces topologies ne sont pas utilisées telles quelles dans la pratique mais que leur mise en œuvre est usuellement faite au sein de structures équilibrées (voir paragraphe III.1.3) pour lesquelles la simplicité des cellules de mélange est un élément important.

a) Mélangeur à diode

Historiquement ce fut les premiers mélangeurs micro-ondes, du fait des bonnes performances fréquentielles des diodes dans ce domaine de fréquences.

La topologie de base est décrite sur la Figure III-5. Comme dans toutes les cellules jusqu'ici exposées, un pompage optimal de l'élément non-linéaire requiert la présence d'un filtre. Dans ce type de mélangeur ce filtre présentera un court-circuit entre la cathode de la diode et la masse pour la fréquence f_{OL} , permettant à la tension OL en entrée de se retrouver totalement appliquée au borne de la diode.

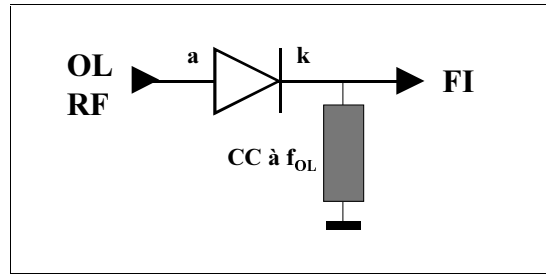


Figure III-5 : Mélangeur à diode

L'étude de cette cellule peut se faire simplement en considérant la diode passante pendant les alternances positives du signal OL et bloquée le reste de la période. Le schéma équivalent est donc celui d'un modulateur décrit en Figure I-4 dont la valeur du gain de conversion se déduit de l'expression (I-11) :

$$G_c = \left(\frac{1}{\pi} \right)^2 = -10 \text{ dB} \quad (\text{III-1})$$

Cette étude simplifiée, permettant la compréhension du phénomène physique de la conversion de fréquence dans ce type de mélangeur, ne permet qu'une estimation au premier ordre du gain de conversion.

Pour plus de précision dans cette évaluation, il faut considérer la non-linéarité de la conductance de la diode. Cette conductance, issue de la source de courant non-linéaire du schéma équivalent de la Figure I-9, varie en fonction du signal OL : elle est infinie lorsque la tension aux bornes de la diode est négative et très faible pour des valeurs de tension supérieures à la tension de seuil. Remarquons que la fonction interrupteur est un passage à la limite du cas précédent en considérant que la conductance varie entre une valeur nulle et une valeur infinie.

Parmi les inconvénients de cette topologie, le plus marqué est sa très mauvaise linéarité due au caractère fortement non-linéaire de la diode.

Enfin, malgré de fortes pertes de conversion ainsi qu'une forte puissance OL nécessaire, ce type de mélangeur est employé lorsqu'une large bande passante est désirée et/ou lorsque les fréquences de fonctionnement élevées ne permettent la mise en œuvre que de diodes [I.18], [I.19].

b) Mélangeur à transistor série : interrupteur analogique

Le principe de fonctionnement de ce type de mélangeur est identique à celui des mélangeurs à diode. Le transistor fonctionne idéalement en commutation, le signal OL étant appliqué sur la grille de ce transistor comme le présente la Figure III-6 [I.9].

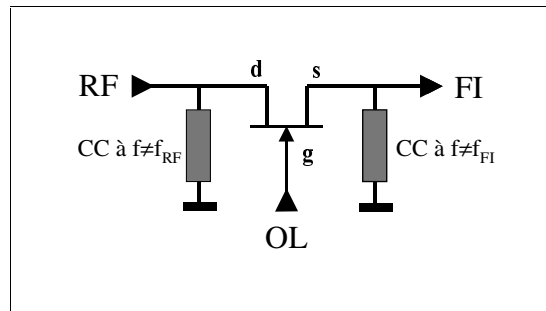


Figure III-6 : Mélangeur à transistor série

On qualifie souvent le transistor utilisé au sein de cette topologie de transistor froid, car aucune polarisation continue n'est appliquée entre le drain et la source de celui-ci, ce qui rend la consommation de cette cellule idéalement nulle.

Par contre, la valeur du gain de conversion est le plus souvent de l'ordre de -10dB et la puissance OL nécessaire reste élevée car elle doit être suffisante pour assurer le fonctionnement en commutation de la structure.

Le facteur de bruit est supérieur à celui des cellules actives à cause des fortes pertes engendrées. Par contre, l'absence de bruit de grenaille et la seule présence du bruit thermique rend le facteur de bruit du mélangeur à transistor série un peu plus faible que son homologue à diode.

Pour une évaluation plus fine des caractéristiques du mélangeur, il faut considérer la non-linéarité de la conductance de sortie du transistor en fonction de la tension V_{gs} , la tension V_{ds} étant maintenue constante.

Un fonctionnement optimal de cette cellule nécessite ainsi les conditions de charge suivantes :

- ❖ Un court-circuit à la masse pour la fréquence f_{OL} sur le drain et sur la source, afin de garantir le maintien de V_{ds} à 0V pour un pompage optimal.
- ❖ Un court-circuit à la masse sur la source à la fréquence f_{RF} , pour éliminer tout signal RF que l'on pourrait retrouver sur l'accès FI en raison de la présence de la capacité C_{ds} dans le schéma équivalent du TEC. Ce couplage éventuel conduirait à la fois à un fonctionnement potentiellement instable et à une augmentation de la distorsion d'intermodulation.
- ❖ Un court-circuit à la masse pour la fréquence f_{FI} sur le drain afin d'empêcher toute contre-réaction néfaste pour la stabilité et le gain du mélangeur.

Une caractéristique remarquable de ce mélangeur par rapport aux autres types présentés est sa très bonne linéarité vis à vis des signaux RF et FI. Cette particularité provient de la grande linéarité de la conductance du canal drain-source sous polarisation nulle [I.9].

Enfin, l'application des signaux se faisant sur trois accès différents, les inconvénients de l'emploi de coupleurs de signaux sont évités.

c) Mélangeur à transistor parallèle : mélangeur résistif

Cette dernière topologie est présentée en Figure III-7(a).

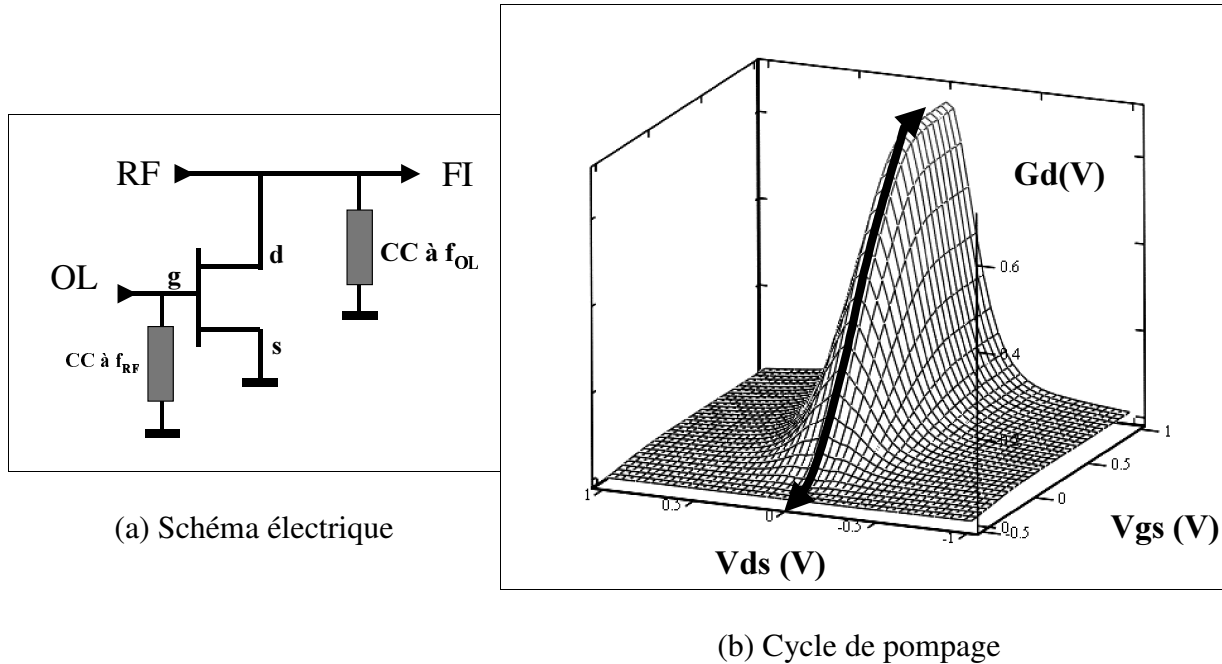


Figure III-7 : Mélangeur résistif

Cette cellule peut être étudiée comme précédemment en considérant le transistor en commutation. Néanmoins, pour ce mélangeur, nous allons directement considérer que la non-linéarité exploitée est la conductance de canal drain-source, en fonction de la tension grille-source c'est à dire du signal OL. La Figure III-7(b) décrit le cycle de pompage optimal obtenu pour $V_{ds} = 0V$. Le maintien de cette tension durant le pompage est assuré par un filtre présentant un court-circuit entre le drain et la masse pour la fréquence f_{OL} .

Les avantages de choisir une polarisation nulle pour V_{ds} ont déjà été cités :

- 1) Une consommation nulle.
- 2) Une linéarité accrue car la résistance du canal drain-source possède un troisième terme du développement en série de Taylor est relativement faible (c'est ce terme qui génère les produits d'intermodulation d'ordre trois).

Comme le mélangeur précédent, cette topologie se distingue donc par sa très grande linéarité [I.20].

Comme dans le cas des mélangeurs à transconductance, le gain de conversion est d'autant plus important que l'excursion de la conductance de canal est élevée. La tension OL sur la grille doit donc balayer la plage $[-V_t, V_{gs,max}]$. Bien évidemment, aucune amplification n'est générée et ce type de mélangeur présente des pertes voisines de 6 dB. La puissance OL nécessaire est inférieure à celle requise pour les mélangeurs à diode mais supérieure à celle des mélangeurs sur grille.

Une attention particulière est à porter aux effets néfastes de la forte valeur de la capacité grille-drain engendrée par une valeur nulle de V_{ds} . Des filtres présentant des courts-circuits à la masse doivent donc être placés [I.21] :

- ❖ Sur le drain, afin de maintenir V_{ds} à 0V durant le pompage.
- ❖ Sur la grille, pour éviter que le signal RF ne module la tension V_{gs} et n'engendre pas de la sorte des produits d'intermodulation.

Le bruit n'est d'origine que thermique, on observe donc des facteurs de bruit inférieurs à ceux obtenus avec les topologies mettant en œuvre des diodes à jonction et des transistors bipolaires dans lesquelles le bruit de grenaille est à prendre en considération.

III.1.3. Les structures équilibrées

La mise en œuvre d'une seule non-linéarité pour réaliser le mélange, comme dans le cas de toutes les structures que nous venons de décrire, présente l'inconvénient majeur d'engendrer un nombre important de signaux parasites qui vont affecter le signal de sortie.

La solution consiste alors à utiliser plusieurs non-linéarité au sein de structures équilibrées qui vont, de façon interférométrique, éliminer un grand nombre de signaux indésirables.

Nous allons ici décrire le principe des structures équilibrées, puis nous présenterons les trois topologies doubles équilibrées les plus largement rencontrées dans les systèmes de conversion de fréquence actuels.

a) Principe, avantages et inconvénients

Nous avons vu au paragraphe II.5 que de nombreuses raies parasites apparaissent en sortie des mélangeurs. Elles engendrent de nombreux inconvénients parmi lesquels :

- ❖ Il est impossible de filtrer les raies parasites apparaissant dans le canal FI.
- ❖ Les filtres atténuant les raies indésirables en dehors du canal FI présentent des pertes et occupent beaucoup de surface.

La solution est de mettre en œuvre des structures équilibrées qui éliminent naturellement ces raies parasites sans nécessité de filtrage.

Considérons le mélangeur à non-linéarité décrit par la Figure III-8 et dont la relation entre les signaux est décrite par un développement limité à l'ordre 3 (III-2).

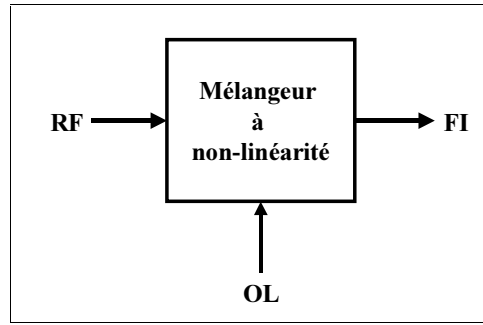


Figure III-8 : Mélangeur simple

$$FI = k_1(OL+RF) + k_2(OL+RF)^2 + k_3(OL+RF)^3 \quad (III-2)$$

$$\Downarrow$$

$$FI = k_1(OL+RF) + k_2(OL^2 + RF^2 + 2OL \bullet RF) + k_3(OL^3 + 3OL^2 \bullet RF + 3OL \bullet RF^2 + RF^3)$$

Seul le terme $2k_2OL \bullet RF$ est utile, les autres termes génèrent des raies spectrales indésirables qu'il est nécessaire d'éliminer.

Se basant sur des propriétés de symétrie / anti-symétrie électrique, l'équilibrage de structure va de façon naturelle, à l'image des interférences destructives, rejeter certaines raies spectrales parasites.

Nous allons tout d'abord décrire la structure simple équilibrée qui met en œuvre deux non-linéarité et rejète de la sorte certaines raies parasites. Puis nous présenterons la structure double équilibrée, de loin la plus usitée, qui rejète un grand nombre de raies indésirables au prix d'un accroissement de la complexité lié à la mise en œuvre de quatre non-linéarité.

i - Structure simple équilibrée

Considérons la structure simple équilibrée constituée de deux mélangeurs simples présentés précédemment et arrangés comme le présente la Figure III-9.

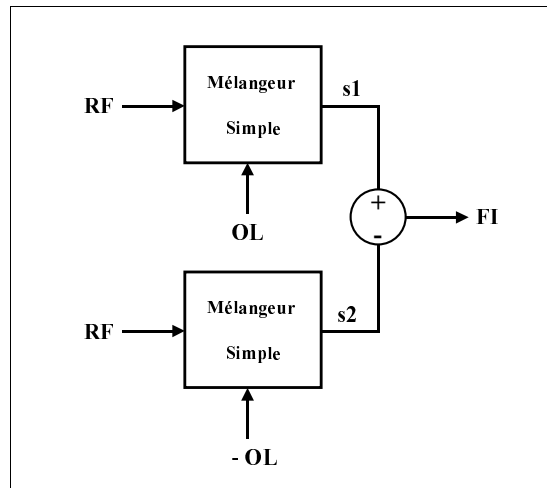


Figure III-9 : Structure simple équilibrée

L'expression de la sortie FI peut alors être calculée aisément :

$$s1 = k_1(OL + RF) + k_2(OL^2 + RF^2 + 2OL \bullet RF) + k_3(OL^3 + 3OL^2 \bullet RF + 3OL \bullet RF^2 + RF^3) \quad (III-3)$$

$$s2 = k_1(-OL + RF) + k_2(OL^2 + RF^2 - 2OL \bullet RF) + k_3(-OL^3 + 3OL^2 \bullet RF - 3OL \bullet RF^2 + RF^3) \quad (III-4)$$

$$FI = 2k_1(OL) + 2k_2(2OL \bullet RF) + 2k_3(OL^3 - 3OL \bullet RF^2) \quad (III-5)$$

Si l'on compare les sorties du mélangeur simple (relation (III-3)) et du mélangeur simple équilibré (relation (III-5)), on s'aperçoit que la structure équilibrée rejette de façon naturelle :

- ❖ toutes les fréquences harmoniques du signal RF,
- ❖ les fréquences harmoniques paires du signal OL,
- ❖ les produits d'intermodulation suivant : $nf_{RF} \pm mf_{OL}$ avec m pair.

Remarque : Nous avons présenté une structure constituée de deux mélangeurs excités par les signaux (RF, OL) et (RF, -OL) et rejetant toutes les harmoniques RF et les harmoniques paire OL. Nous pouvons tout aussi bien envisager une structure similaire rejetant toutes les harmoniques OL et les harmonique paires RF pour laquelle les excitations seraient (RF, OL) et (-RF, OL). En pratique, le choix du type d'équilibrage est déterminé en fonction de la nature des raies parasites fortement indésirables dont on souhaite la réjection par la structure.

ii - Structure double équilibrée

Considérons maintenant la structure double équilibrée constituée de quatre mélangeurs simples et arrangés comme le présente la Figure III-10.

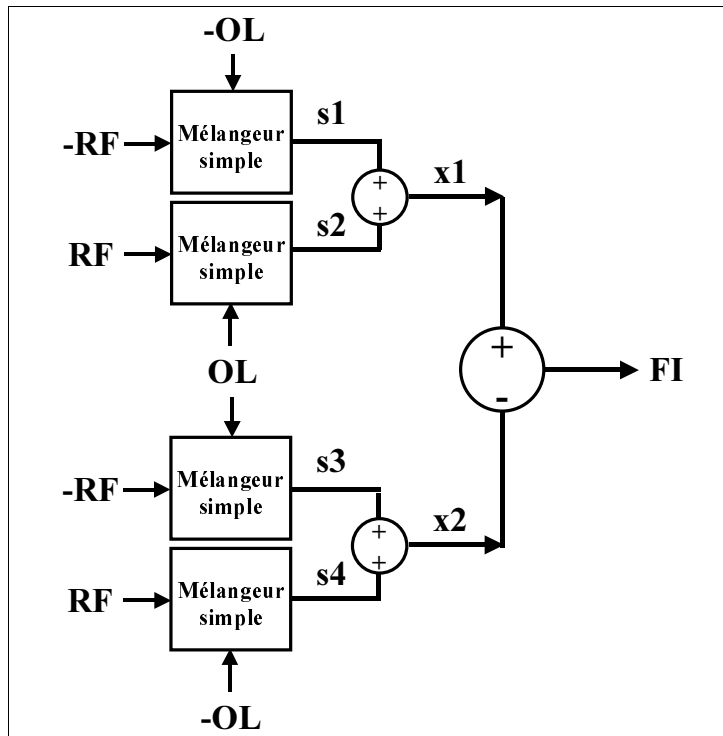


Figure III-10 : Structure double équilibrée

Le calcul suivant démontre que la sortie ne comporte plus que le terme utile.

$$s_1 = k_1(-OL-RF) + k_2(OL^2 + RF^2 + 2OL \bullet RF) + k_3(-OL^3 - 3OL^2 \bullet RF - 3OL \bullet RF^2 - RF^3) \quad (\text{III-6})$$

$$+ s_2 = k_1(OL+RF) + k_2(OL^2 + RF^2 + 2OL \bullet RF) + k_3(OL^3 + 3OL^2 \bullet RF + 3OL \bullet RF^2 + RF^3) \quad (\text{III-7})$$

$$x_1 = \frac{+2 k_2(OL^2 + RF^2 + 2OL \bullet RF)}{\quad} \quad (\text{III-8})$$

et

$$s_3 = k_1(OL-RF) + k_2(OL^2 + RF^2 - 2OL \bullet RF) + k_3(OL^3 - 3OL^2 \bullet RF + 3OL \bullet RF^2 - RF^3) \quad (\text{III-9})$$

$$+ s_4 = k_1(-OL+RF) + k_2(OL^2 + RF^2 - 2OL \bullet RF) + k_3(-OL^3 + 3OL^2 \bullet RF - 3OL \bullet RF^2 + RF^3) \quad (\text{III-10})$$

$$x_2 = \frac{+2 k_2(OL^2 + RF^2 - 2OL \bullet RF)}{\quad} \quad (\text{III-11})$$

d'où

$$\boxed{FI = x_1 - x_2 = 4k_2(2OL \bullet RF)} \quad (\text{III-12})$$

La structure double équilibrée rejette donc de façon naturelle :

- ❖ Toutes les harmoniques du signal RF.
- ❖ Toutes les harmoniques du signal OL.
- ❖ Les produits d'intermodulation suivant : $nf_{RF} \pm mf_{OL}$ avec m pair ou n pair.

Ainsi, un grand nombre de raies parasites sont rejetées de façon naturelle, sans filtrage par les structures double équilibrées. Parmi les composantes spectrales restantes, il en est deux importantes :

- 1) La raie à la fréquence $f_{RF} + f_{OL}$ dont la présence est justifiée au paragraphe I.1.1. Cette raie n'est pas forcément gênante en elle même, mais sa présence démontre celle de la conversion de fréquence supradynne et donc la nécessité de filtrer la fréquence image en entrée. Notons qu'il existe une structure dite « à réjection de fréquence image » [I.22] qui permet l'élimination de la raie à la fréquence image au prix d'un accroissement de complexité.
- 2) La raie à la fréquence $3f_{RF} - f_{OL}$, issue de l'ordre 4 du développement limité de la non-linéarité et conséquence de l'intermodulation d'ordre 3. La présence de cette raie montre que l'équilibrage de structures n'améliore pas leur linéarité.

Remarque : Il est important de noter que les sommateurs éliminent les harmoniques RF et OL impaires ainsi que les produits d'intermodulation tels que m+n soit impair tandis que le soustracteur élimine les harmoniques RF et OL paires ainsi que les produits d'intermodulation tels que l'indice sur l'OL soit pair. Cette constatation permet la localisation des lieux d'élimination des raies spectrales, et donc la détermination des éléments à mettre en cause lors d'une mauvaise réjection de certaines raies en fonction de leur fréquence.

iii - Avantages et inconvénients

L'inconvénient majeur des structures doubles équilibrées est leur difficulté de mise en œuvre car elle nécessite trois coupleurs de signaux (2 diviseurs 180° pour les voies RF et OL et un combineur 180° pour la sortie FI) et quatre cellules simples de mélange.

En considérant des coupleurs passifs idéaux, si l'on compare la structure double équilibrée au mélangeur simple qui la constitue, pour un gain et un facteur de bruit identiques, elle présente :

Pour les inconvénients :

- ❖ Une consommation multipliée par quatre.
- ❖ Une complexité de mise en œuvre accrue.
- ❖ Une surface de réalisation au moins multipliée par 5 ou 6.

Pour les avantages :

- ❖ Une réjection naturelle de nombreuses raies parasites.
- ❖ Une augmentation du point d'interception de 6 dB.

Quelle que soit l'importance des inconvénients que l'on peut trouver aux structures doubles équilibrées, elles seront systématiquement utilisées dans les systèmes imposants de fortes contraintes de réjection des raies parasites.

Nous allons décrire maintenant les structures doubles équilibrées classiques utilisant trois types de cellules de base présentées précédemment.

b) Mélangeur en anneau

Cette structure met en œuvre quatre mélangeurs à diode et trois transformateurs à point milieu pour les coupleurs 180° comme le présente la Figure III-11.

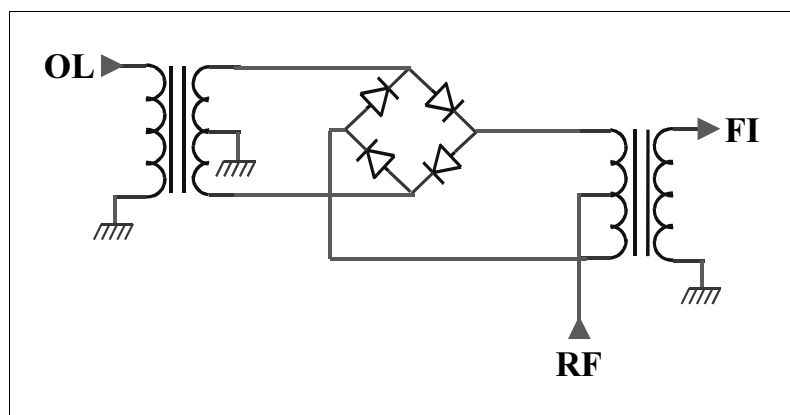


Figure III-11 : Mélangeur en anneau

Ce mélangeur présente beaucoup de pertes de conversion mais est de réalisation très facile. Sa gamme de fréquence est évidemment limitée par les transformateurs à quelques GHz. A plus hautes fréquences, on peut trouver d'autres types de coupleurs, mais les mélangeurs à diodes souffriront toujours de leur faible gain de conversion, de leur forte puissance OL nécessaire et surtout de leur très mauvaise linéarité.

c) Mélangeur à TEC froids

Cette structure repose sur l'emploi de topologies de mélange à transistor série utilisé en interrupteur analogique profitant ainsi de sa grande linéarité.

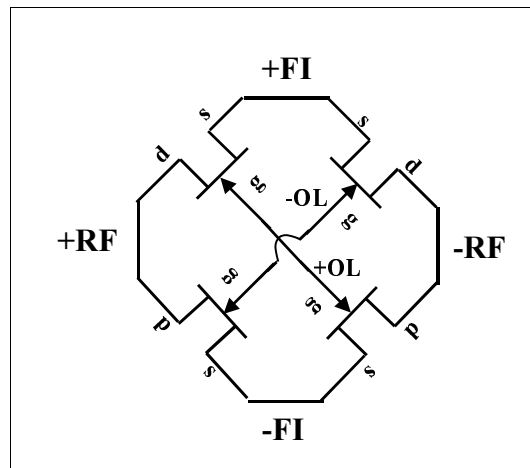


Figure III-12 : Mélangeur double équilibré à TEC froids

Comme pour la cellule de base, le gain de conversion est de l'ordre de -10dB , pour une puissance OL de l'ordre de 0 à 10dBm . Ce type de mélangeur est utilisé lorsque de fortes contraintes de linéarité sont requises [I.6]

d) Mélangeur de Gilbert

La Figure I-3 présente la cellule de Gilbert. Son fonctionnement diffère de celui du multiplieur de Gilbert car le signal OL est de forte amplitude. Les deux paires différentielles supérieures fonctionnent ainsi en commutation, les cellules de mélange de base sont alors des transistors séries utilisés en interrupteurs analogiques. La paire différentielle inférieure fonctionne en amplificateur du signal RF pouvant conférer à cette cellule un gain de conversion supérieur à 0dB .

Longtemps confinés aux applications radio fréquence, les mélangeurs de Gilbert deviennent de plus en plus utilisés dans le domaine micro-ondes grâce au développement des technologies des transistors bipolaires (TBH SiGe) [I.23]. La mise en œuvre de ces topologies peut aussi s'effectuer à l'aide de transistors à effet de champ et même conjointement avec les deux types de transistors dans une technologie BiCMOS [I.24].

Cependant, les mélangeurs de Gilbert possèdent intrinsèquement les inconvénients des mélangeurs passifs qui présentent des performances en gain et en bruit médiocres par rapport à celles des mélangeurs actifs.

III.2. Synthèse et état de l'art

La comparaison qualitative de toutes les topologies de mélange n'est pas aisée, voire impossible, car les bandes de fréquences et les technologies diffèrent d'une référence à l'autre.

Cette étude est d'autant plus difficile que les topologies de base ne sont plus utilisées telles quelles mais toujours mises en œuvre au sein de structures double équilibrées. Il est alors difficile d'obtenir les performances des cellules de base avec les technologies actuelles.

Le Tableau III-1 tente cependant une classification qualitative des différentes topologies de mélange vues précédemment. Ce tableau s'appuie d'une part sur quelques références relativement anciennes [1.7, 10, 14, 20], et d'autre part sur des simulations que nous avons effectuées. Ces simulations ont été effectuées à l'aide d'une même technologie à base de transistors PHEMT, d'un même cahier des charges, et des conditions de fermeture idéales adaptées à chaque structure.

Topologies de base	Topologies	Gc	Pol	OPI3	NF	Pcons	Mise en oeuvre
	Mélangeur sur grille	😊	😊	😐	😊	😞	Nécessite 1 coupleur
	Mélangeur sur drain	😊	😞	😐	😊	😊	Nécessite 1 coupleur
	Mélangeur sur source	😊	😐	😐	😊	😐	Pas de coupleur
	Mélangeur à 2 transistors	😐	😊	😐	😐	😐	Pas de coupleur
	Mélangeur à diode	😞	😐	😞	😞	😊	Nécessite 1 coupleur
	Mélangeur à transistor série	😞	😊	😐	😐	😊	Pas de coupleur
	Mélangeur à transistor parallèle	😞	😊	😊	😐	😊	Nécessite 1 coupleur
	Mélangeur de Gilbert	😊	😊	😐	😐	😐	Structures double équilibrées
	Mélangeur double équilibré à FET froids	😞	😐	😊	😐	😊	

Tableau III-1 : Classification des cellules de mélange

Cette étude met en avant les qualités et les défauts de chaque topologie, notons que, parmi les topologies de base, seuls le mélangeur sur source ne présente pas d'inconvénient majeur. Parmi les topologies double équilibrées, présentées aux deux dernières lignes du Tableau III-2, le mélangeur de Gilbert sera utilisé lorsque du gain de conversion est nécessaire, tandis que le mélangeur à transistors froids sera mis en œuvre lorsqu'une faible distorsion d'intermodulation est requise.

Comme nous l'avons déjà signalé, les topologies de base ne sont plus utilisées seules mais au sein de structures doubles équilibrées, les publications récentes concernent principalement deux types de mélangeurs double équilibrés :

- ❖ Les cellules de Gilbert, principalement mise en œuvre à l'aide de Transistor Bipolaire à Hétérojonction (TBH) en Silicium-Germanium (SiGe).
- ❖ Les topologies double équilibrées à transistors froids à base de transistor à effet de champ à haute mobilité électronique, par exemple.

Même le classement de ces deux types de mélangeurs est difficile, car les performances publiées sont souvent celles de systèmes dans lesquels le mélangeur est encadré d'amplificateurs. Le Tableau III-2 présente deux types de mélangeurs, reportés par [I.25] et [I.26], les plus représentatifs possible de l'ensemble des réalisations opérants aux alentours de 10 GHz.

Topologies	Gc(dB)	Pol(dBm)	NF(dB)	OPI3(dBm)	Q3(dB)
Mélangeur de Gilbert	16	-2	9.4	2	4
Mélangeur double équilibré à FET froids	-9	7	10	11	5

Tableau III-2 : Etats de l'Art des mélangeurs micro-ondes

Les mélangeurs de Gilbert ont la particularité de présenter un gain de conversion supérieur à 0dB pour une puissance OL nécessaire relativement faible. Les deux types de mélangeur sont comparables au niveau bruit et le mélangeur à FET froid présente un facteur de qualité de linéarité (défini au paragraphe II.3.2.c) ii -) Q3 légèrement supérieur.

IV. Conclusion

Nous avons, dans ce chapitre présenté le principe de fonctionnement des mélangeurs dans le domaine micro-ondes puis défini les caractéristiques permettant d'évaluer leurs performances.

Nous avons de plus recensé toutes les cellules de mélange de bases existantes, expliqué leur principe de fonctionnement et ainsi justifié leurs performances. Sur cette base, nous avons établi une classification de ces diverses topologies de laquelle est ressorti le mélangeur sur source. Cette topologie présente en effet de bonnes performances en gain de conversion et de par la séparation physique des signaux RF, OL et FI est facilement intégrable.

Nous avons ainsi répondu au premier besoin de toute conception qui débute par le choix de la structure envisagée conformément à des contraintes fixées et nous allons aborder, dans le chapitre suivant, la définition d'une méthodologie de conception de mélangeur qui débutera par le choix justifié de sa topologie.

REFERENCES BIBLIOGRAPHIQUES DU CHAPITRE 1

- [I.1] : B. Gilbert, « A precise four quadrant multiplier with subnanosecond response », IEEE – Journal of Solid-State Circuits, vol.3, december 1968, pp 365-373.
- [I.2] : B. Razavi, « RF Microelectronics », PRENTICE HALL, 1998.
- [I.3] : G. Vasilescu, « Bruits et signaux parasites », DUNOD, 1999.
- [I.4] : H.T. Friis, « Noise figure of radio receivers », Proceedings of IRE, vol. 32, n°7, 1944, pp. 419-422.
- [I.5] : S. Weiner, D. Neuf, S. Spohrer, « 2 to 8 GHz double balanced MESFET mixer with +30 dBm input 3rd intercept point », IEE- MTT Digest Symposium, 1988, pp. 1097-1100.
- [I.6] : D. Prieto, « Conception et caractérisation de circuits intégrés micro-ondes monolithiques (MMICs) en technologie d'interconnexions uniplanaires. Application à la conception d'un convertisseur de fréquence en bande Ku », Thèse de doctorat de l'Université Paul Sabatier de Toulouse, Janvier 1999.
- [I.7] : R. Pucel, D. Massé, R. Bera, « Performance of GaAs MESFET mixers at X band », IEEE MTT, vol. 24, n° 6, June 1976, pp. 351-360.
- [I.8] : S.A. Maas, « Design and performances of a 45-GHz HEMT mixer », IEEE MTT, vol. 34, n° 7, July 1986, pp. 799-803.
- [I.9] : S.A. Maas, « The RF and microwave circuit design cookbook », Artech House, 1998.
- [I.10] : I. Angelov, H. Zirath, « On the performance of different types of MESFET-mixers », Microwave and optical technology letters, vol. 4, n° 12, november 1991.
- [I.11] : P. Bura, R. Dikshit, « FET mixers for communication satellite transponders », IEEE MTT-Symposium, 1976, pp.90-92.
- [I.12] : P. Bura, R. Dikshit, « FET mixer with the drain L.O. injection », Electronics Letters, Vol. 12, n° 20, September 1976, pp.536-537.
- [I.13] : B. Lloriu, J.C. Leost, « Design and performance of low noise C- and X-band GaAs FET mixers », 7th European Microwave Conference, 1977, pp. 95-100.
- [I.14] : P. Harrop, R. Dessert, P. Baudet : « Performance of some GaAs MESFET mixers », 6th European Microwave Conference, 1976, pp. 8-13.

- [I.15] : S.A. Maas, « Microwave mixers », Artech House, 2nd edition, 1993.
- [I.16] : C. Tsironis, R. Meierer, R. Stahlmann, « Dual-gate MESFET mixers », IEEE MTT, vol. 32, n° 3, March 1984, pp. 248-255.
- [I.17] : C.-H. Lee, S. Han, J. Laskar, « GaAs MESFET dual-gate mixer with active filter design for Ku-band applications », IEEE- MTT Symposium Digest, 1999, pp. 841-844.
- [I.18] : H. Gu, K. Wu, « A novel uniplanar balanced subharmonically pumped mixer for low-cost broadband millimeter-wave transceiver design », IEEE-MTT Symposium Digest, 2000, pp. 635-638.
- [I.19] : C-Y Chang, C-C Yang, D-C Niu, « A multioctave bandwidth rat-race singly balanced mixer », IEEE Microwave and guided wave letters, vol. 9, n° 1, January 1999, pp. 37-39.
- [I.20] : S.A.Maas, « A GaAs MESFET mixer with very low intermodulation », IEEE-MTT, vol. 35, n° 4, April 1987, pp. 425-429.
- [I.21] : J.C. Cayrou, « Modélisation, conception et caractérisation de convertisseurs de fréquence micro-ondes à transistor à effet de champ à grille schottky sur aséniure de Gallium », Thèse de doctorat de l'Université Paul Sabatier de Toulouse, Juillet 1993.
- [I.22] : P.F. Combes, J. Graffeuil, J.F. Sautereau, « Composants, dispositifs et circuits actifs en micro-ondes », Dunod 1985.
- [I.23] : J. Glenn, M. Case, D. Harane, B. Meyerson, R. Poisson, « 12-GHz Gilbert mixers using a manufacturable Si/SiGe epitaxial-base bipolar technology », BCTM, 1995, pp. 186-189.
- [I.24] : S. Colomines, « Conception et caractérisation de mélangeurs radiofréquences en technologie BiCMOS pour applications de téléphonie cellulaire », Thèse de doctorat de l'Université Paul Sabatier de Toulouse, Mars 1999.
- [I.25] : M.C. Tsai, M.J. Schindler, W. Strubler, M. Ventresca, R. Binder, R. Waterman, D. Danzilio, « A compact Wideband Balanced Mixer », IEEE Microwave and Millimeter-Wave Monolithic Circuits Symposium, 1994, pp. 135-138.
- [I.26] : W. Dürr, U. Erben, A. Schüppen, H. Dietrich, H. Schumacher, « Low-power low-noise active mixers for 5.7 and 11.2GHz using commercially available SiGe HBT MMIC technology », Electronics Letters, vol. 34, n° 21, 15th October 1998, pp. 1994-1996.

Chapitre 2 : Conception d'une cellule de mélange originale.

Etude de la stabilité non-linéaire.

Introduction

Nous avons présenté au premier chapitre toutes les cellules de mélange existantes parmi lesquelles nous devons faire un choix. Nous allons baser notre choix sur une étude système qui va lier les performances globales d'un système de conversion de fréquence à celles de chaque sous système, dégagant ainsi les caractéristiques importantes que doit présenter le mélangeur.

Tout en prenant en compte l'aspect intégration monolithique future, nous allons introduire une cellule originale de mélange répondant aux caractéristiques désirées. Nous présentons ensuite le fonctionnement de cette cellule de mélange ainsi que le choix des transistors et de leur polarisation. Puis nous présentons les diverses précautions à prendre pour garantir un fonctionnement correct et optimal. Cette étude démontrera l'importance des impédances de fermeture et justifiera la mise en œuvre du formalisme des matrices de conversion permettant leurs optimisations. Cette méthode fera l'objet du paragraphe II qui présentera le principe du calcul et son application à l'optimisation de la cellule étudiée. Enfin, nous mettrons à profit les matrices de conversion pour l'étude de la stabilité non-linéaire, point délicat de la structure conçue. L'application du critère de Nyquist au déterminant de la matrice admittance de conversion permettra d'optimiser le circuit afin de garantir sa stabilité non-linéaire inconditionnelle.

I. Cellule Originale de mélange

Nous avons vu au chapitre précédent les avantages apportés par l'ajout d'un second transistor pour l'application du signal de pompe OL dans le circuit : abaissement de la puissance OL nécessaire et économie d'un coupleur de signaux. Ce principe fut mis en œuvre pour les mélangeurs sur drain conduisant à la structure présentée au paragraphe III.1.1 du premier chapitre et résolvant le problème de la forte puissance OL que nécessite le mélangeur sur drain (cf. tableau III-1 du premier chapitre). La structure résultante présente néanmoins un faible gain de conversion, un fort facteur de bruit et une mauvaise linéarité en raison de la mauvaise adaptation entre les deux transistors.

Sur la base de ce principe, nous avons étudié les apports d'un second transistor associé à une cellule de mélange sur source dont les performances sont plus attrayantes.

I.1. Principe de la topologie

Les mélangeurs sur source qui semblent réaliser le meilleur compromis de performances (cf. tableau III-1 du premier chapitre) présentent pour inconvénient de nécessiter une puissance importante du signal OL. Aussi, en vue d'aboutir à une cellule de mélange optimale, nous avons étudié l'influence de l'insertion d'un second transistor pour appliquer le signal de pompe.

I.1.1. Topologie et principe de fonctionnement

La structure initiale du mélangeur piloté par la source, présentée sur la Figure I-1, à laquelle est rajouté un deuxième transistor pour l'application du signal de pompe OL permet alors d'aboutir à la cellule originale dont la topologie est donnée sur la Figure I-2.

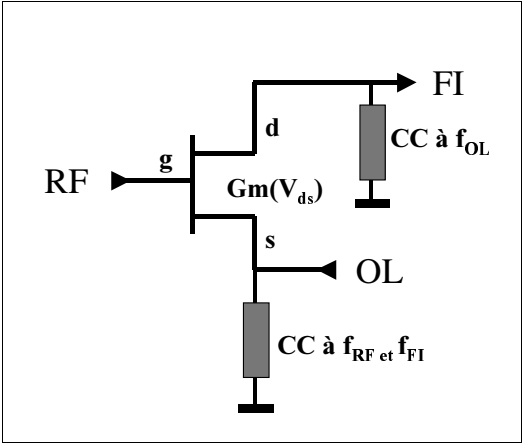


Figure I-1 : Mélangeur sur source

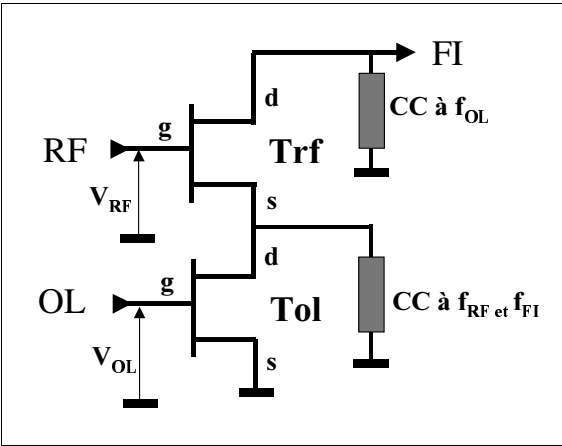


Figure I-2 : Mélangeur sur source modifié

L'insertion du transistor Tol ne dégrade pas les performances, comme cela est observé pour les cellules à deux transistors classiques. Les performances résultantes correspondent toujours à celles des mélangeurs sur source accompagnées d'une forte diminution de la puissance P_{OL} nécessaire. Afin de comparer cette cellule aux autres, nous présentons ci-dessous une ligne à rajouter à la synthèse présentée au tableau III-1 du premier chapitre 1 faisant de cette structure la plus performante parmi celles présentées.

Topologies	Gc	Pol	OPI3	NF	Pcons	Mise en oeuvre
Mélangeur original	☺	☺	☹	☺	☹	Pas de coupleur

Tableau I-1 : Caractéristiques de la nouvelle cellule de mélange

Précisons que le transistor T_{RF} fonctionne en régime saturé pour assurer un fort gain de conversion, tandis que nous avons choisi de placer la transistor T_{OL} dans sa zone linéaire afin de limiter la puissance consommée globale.

Le signal OL, appliqué sur la grille de T_{OL} , module la tension drain-source du transistor T_{RF} , rapprochant le fonctionnement de cette cellule à celui de mélangeur sur source présenté au paragraphe III.1.1.c) du premier chapitre. La Figure I-3 montre la dépendance de la transconductance en fonction des deux tensions appliquées sur les deux grilles des transistors. Le cycle de pompage optimal est indiqué en gras et correspond à une tension de polarisation V_{RF} de 0V et une excursion de la tension V_{OL} dans la plage $[V_p, 0V]$.

Comme pour les autres cellules, afin de garantir un pompage optimal de la transconductance, il est nécessaire de placer un filtre présentant un court-circuit à la masse pour la fréquence f_{OL} sur le drain du transistor T_{RF} .

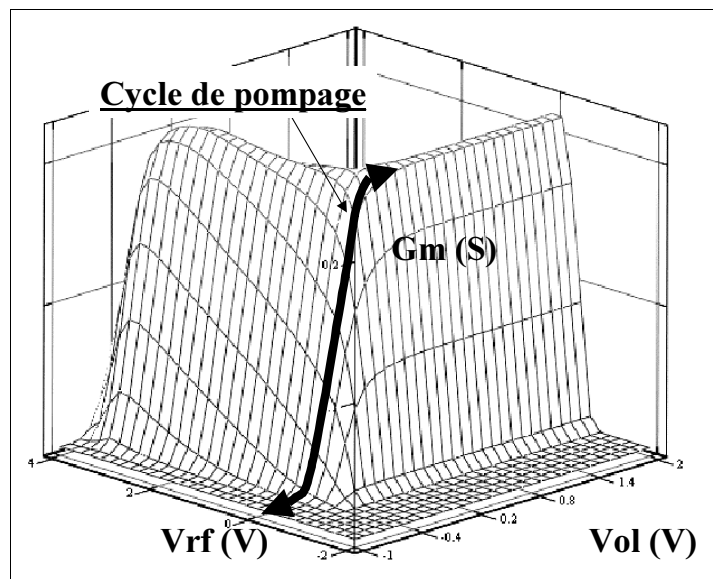


Figure I-3 : Cycle de pompage du mélangeur original

Malgré un principe simple et de bonnes performances, cette cellule reste marginale ou mal exploitée dans la littérature [II.1]. En effet, les fortes potentialités d'instabilité, étant donné que la source du transistor de mélange n'est pas directement connectée à la masse, ont le plus souvent limité l'étude de cette cellule.

Toute la stabilité repose sur la conception d'un filtre présentant idéalement un circuit ouvert à la fréquence f_{OL} et un court-circuit à toutes les autres fréquences placé au niveau de la connexion des deux transistors. Pratiquement, l'ensemble transistor T_{OL} et filtre présente une impédance non-nulle aux fréquences f_{FI} et f_{RF} qu'il est nécessaire d'optimiser pour garantir un fonctionnement stable. Malheureusement, il n'existe pas d'outils dans les simulateurs commerciaux permettant l'étude complète de la stabilité en régime non-linéaire.

L'originalité de notre travail repose donc à la fois sur la topologie du mélangeur, sa polarisation, ses impédances de fermeture mais aussi sur l'optimisation de l'impédance garantissant un fonctionnement stable.

I.1.2. Arrangement de la structure double équilibrée

Un autre intérêt important de cette nouvelle topologie réside dans son aptitude pour l'intégration monolithique au sein d'une structure double équilibrée.

La Figure I-4 montre en effet que la constitution de la topologie simple équilibrée se simplifie par la mise en commun du transistor T_{OL} . Cette simplification, qui facilite considérablement la conception MMIC tout en améliorant l'équilibrage de la cellule [II.2], rend possible d'envisager la conception MMIC d'une structure double équilibrée.

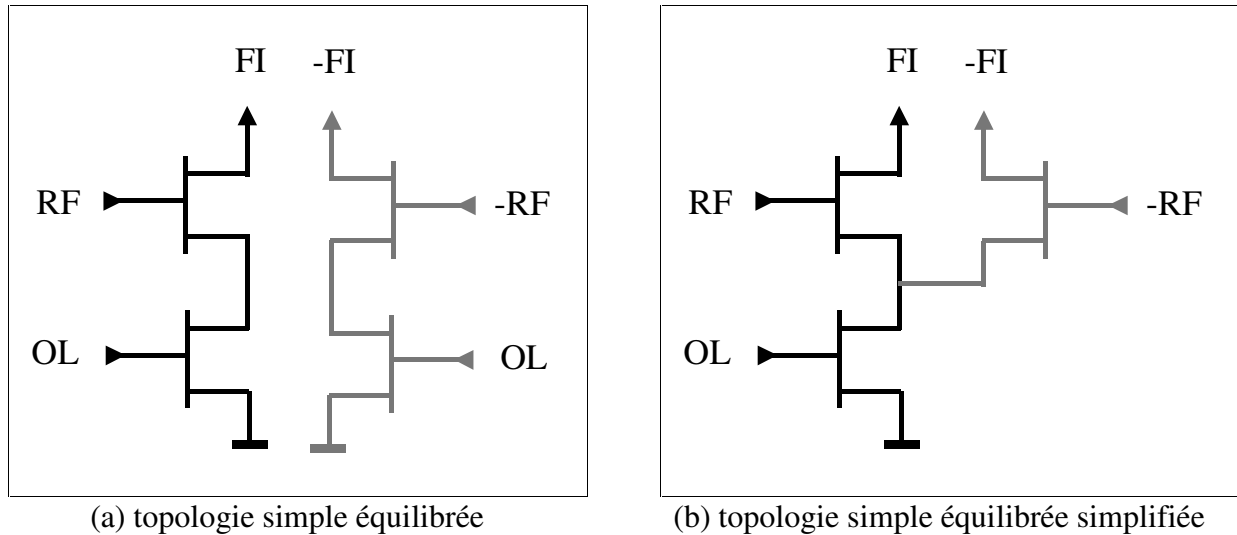


Figure I-4 : Arrangement en structure simple équilibrée

La cellule double équilibrée résultante, présentée sur la Figure I-5, possède une structure semblable à la cellule de Gilbert [II.3] mais présente un fonctionnement fondamentalement différent. En effet, dans notre cas, la paire différentielle supérieure fonctionne en régime saturé et non en commutation tandis que la paire différentielle inférieure fonctionne en régime linéaire.

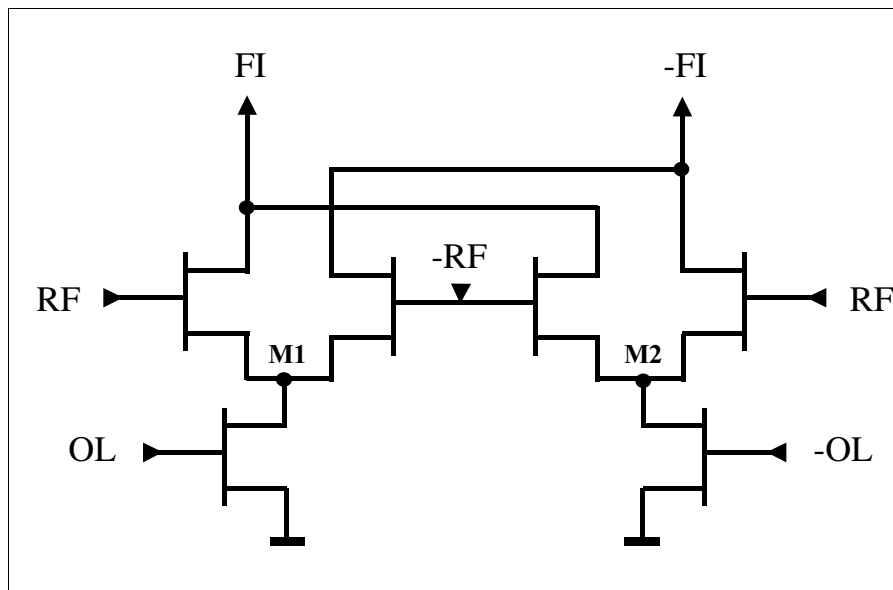


Figure I-5 : Structure double équilibrée originale

Par rapport à la cellule simple, les symétries et anti-symétrie électriques de la topologie double équilibrée apportent les avantages suivant [II.2] :

- Les accès +RF et -RF sont virtuellement court-circuités à la masse pour la fréquence f_{OL} , augmentant l'isolation OL vers RF.
- Les accès +FI et -FI sont court-circuités à la masse pour la fréquence f_{OL} , entraînant l'élimination des filtres correspondant.
- Les équipotentiels M1 et M2 sont court-circuités à la masse pour les fréquences f_{RF} et f_{FI} , garantissant un fort gain de conversion car éliminant les contre réactions.

Concernant la stabilité, le problème persiste toujours avec la structure double équilibrée, comme nous le verrons au paragraphe III de ce chapitre. Nous montrerons alors la nécessité de filtres entre M1 et la masse et entre M2 et la masse.

Nous allons maintenant détailler le fonctionnement de la cellule envisagée.

I.2. Etude du mélangeur

Nous présentons dans un premier temps le calcul analytique du gain de conversion de la cellule de mélange proposée. Outre l'obtention de l'expression au premier ordre du gain de conversion, ces développements présentent aussi l'intérêt de révéler le détail du fonctionnement de la cellule. Les informations qui ressortent de cette étude analytique pourront être judicieusement exploitées au paragraphe I.3 pour la détermination des impédances de fermetures optimales. L'expression du gain de conversion nous permettra, quant à elle, d'identifier les éléments les plus influents sur la valeur du gain et notamment de conduire une optimisation rapide des polarisations et du dimensionnement des transistors, première étape de conception de notre cellule originale.

I.2.1. Calcul du gain de conversion

Ce calcul se décompose en deux parties [II.4] :

- Dans un premier temps, le circuit est étudié uniquement sous l'excitation du signal de pompe OL. Ceci permet la détermination du fondamental à la fréquence f_{OL} de la transconductance responsable de la conversion de fréquence.
- Lorsque l'état de la non-linéarité fixée par le signal OL est connu, il est alors possible, dans un second temps, d'étudier la réponse du circuit uniquement aux fréquences f_{RF} et f_{FI} en ne considérant que l'application du signal RF.

Pompage de la transconductance par le signal OL :

La seule non-linéarité considérée est la transconductance du transistor T_{RF} en fonction de la tension grille-source. Cette non-linéarité possède la forme simplifiée reportée sur la Figure I-7 : elle vaut 0 lorsque V_{gs} est inférieur à la tension de pincement V_p et croît linéairement jusqu'à $g_{m,max}$ pour $V_{gs}=0V$.

Remarque : la plupart des autres éléments du schéma équivalent sont aussi non-linéaires et donc des fonctions du temps au rythme du signal OL. Ces variations ne sont cependant pas capitales pour la conversion de fréquence et ne seront donc pas prises en compte dans le calcul. Dans un souci de simplification des développements analytiques, chaque élément (à l'exception de la transconductance) sera donc maintenu constant et caractérisé par sa valeur moyenne sur une période OL (R sera noté $\bar{R}$).

Si l'on considère uniquement le signal à la fréquence f_{OL} , le circuit de la Figure I-2 se simplifie comme présenté en Figure I-6. Le transistor T_{OL} est considéré comme un amplificateur de la tension V_{OL} , et nous supposons donc : $V_{gs_{OL}}=G \times V_{OL}$. Cette hypothèse est d'autant plus vraie que le signal OL est de faible fréquence. Enfin, la tension de polarisation grille-source du transistor T_{RF} est égale à sa tension de pincement V_p .

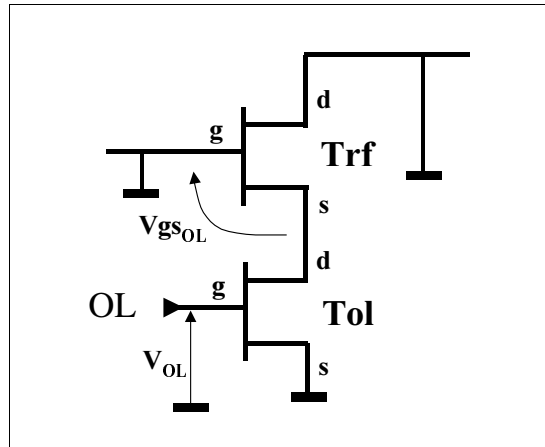


Figure I-6 : Mélangeur sous l'excitation OL

Ces hypothèses font que la variation de la transconductance est équivalente à une arche de sinusoïde d'amplitude $g_{m,max}$ de 0 à $T_{OL}/2$ et de valeur 0 de $T_{OL}/2$ à T_{OL} . La Figure I-7 montre l'allure de g_m en fonction du temps.

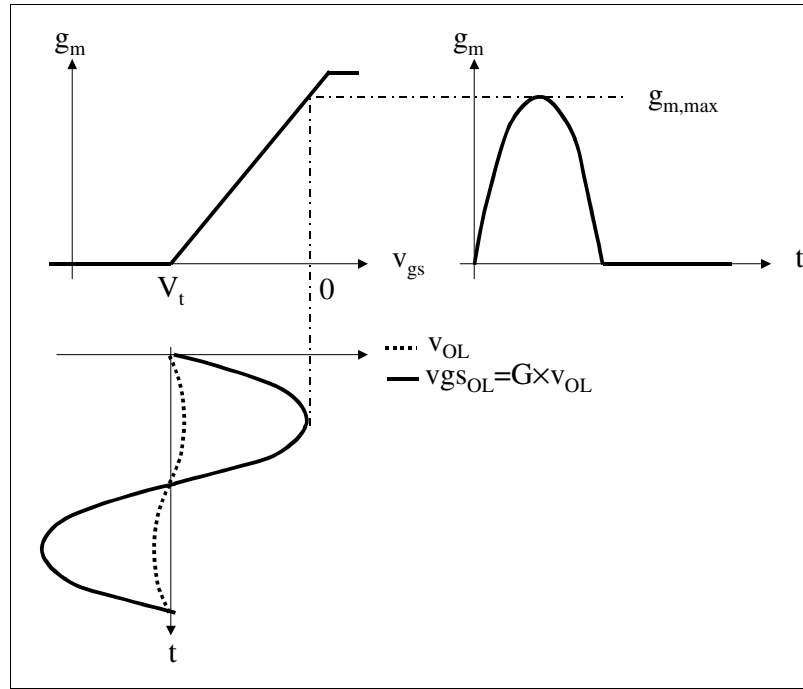


Figure I-7 : Pompage de la transconductance [II.5]

Seul nous intéresse le terme à la fréquence fondamentale f_{OL} de la décomposition en série de Fourier de la transconductance $g_m(t)$. Ce terme est d'amplitude :

$$g_{m,1} = \frac{2}{2\pi} \int_0^{2\pi} g_m(t) \sin(\omega_{OL} t) d(\omega t) = \frac{2}{2\pi} \int_0^{\pi} g_{m,max} \sin^2(\omega_{OL} t) d(\omega t) = \frac{g_{m,max}}{2} \quad (\text{I-1})$$

Maximiser le gain de conversion, donc $g_{m,1}$, revient à maximiser $g_{m,max}$. En pratique, le maximum de transconductance est obtenu, pour des PHEMTs, pour une tension grille-source $V_{gs,max} = 0V$. La tension de pompe V_{gsOL} , polarisée à V_p , doit ainsi atteindre $0V$, et est donc d'amplitude égale à V_p . Dans ces conditions, la puissance OL nécessaire à l'entrée du transistor T_{OL} vaut donc :

$$P_{OL} = \frac{1}{2} \left(\frac{V_t}{G} \right)^2 R_{gs}^{T_{OL}} \left(C_{gs}^{T_{OL}} \right)^2 \omega_{OL}^2 \quad (\text{I-2})$$

Où $R_{gs}^{T_{OL}}$ et $C_{gs}^{T_{OL}}$ sont la résistance d'entrée et la capacité grille-source du transistor

T_{OL} . G le gain en tension du transistor T_{OL} à la fréquence f_{OL} , et $\frac{|V_p|}{G}$ l'amplitude que doit avoir le signal V_{OL} pour que le signal V_{gsOL} atteigne la valeur $0V$.

Application du signal RF:

Comme nous l'avons développé au premier chapitre, le signal RF est considéré d'amplitude suffisamment faible pour que le régime de fonctionnement du circuit par rapport à ce signal soit linéaire.

On suppose que ne sont alors présentes en entrée que la composante spectrale à la fréquence f_{RF} et en sortie que la composante spectrale à la fréquence f_{FI} . Nous verrons plus loin que ceci n'est pas qu'une simple hypothèse de calcul mais doit être vérifiée en pratique. On placera pour ce faire des circuits résonnants aux accès RF et FI.

Dans ces conditions, le circuit équivalent pour les fréquences f_{RF} et f_{FI} est présenté sur la Figure I-8. Le transistor T_{RF} est en source commune tandis que le transistor T_{OL} devient non-opérationnel et ne sera plus considéré.

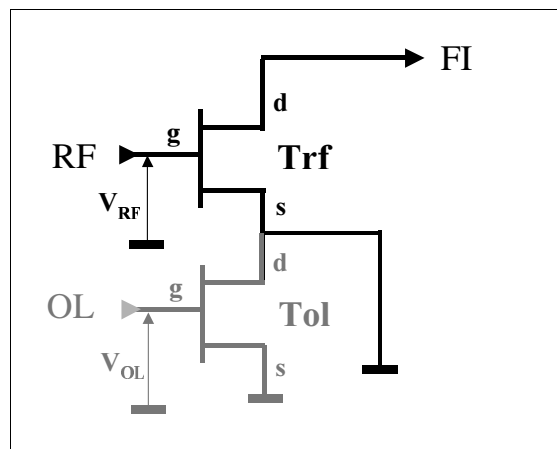


Figure I-8 : Mélangeur aux fréquences f_{RF} et f_{FI}

Le générateur et la charge sont adaptés en puissance au transistor respectivement aux fréquences f_{RF} et f_{FI} .

Les éléments du schéma équivalent du transistor T_{RF} sont, en entrée, une résistance ($\bar{R}_{gs}$) en série avec la capacité de canal ($\bar{C}_{gs}$), et en sortie, une source de courant commandé par la tension aux bornes de C_{gs} et de relation : $i_d(f_{FI}) = g_{m,1} \bullet V_{gs}(f_{RF})$ avec en parallèle, une capacité $\bar{C}_{ds}$ et une conductance $\bar{G}_{ds}$ constituant l'impédance de sortie notée Z_s .

La Figure I-9 synthétise les hypothèses faites en exposant le schéma équivalent du mélangeur.

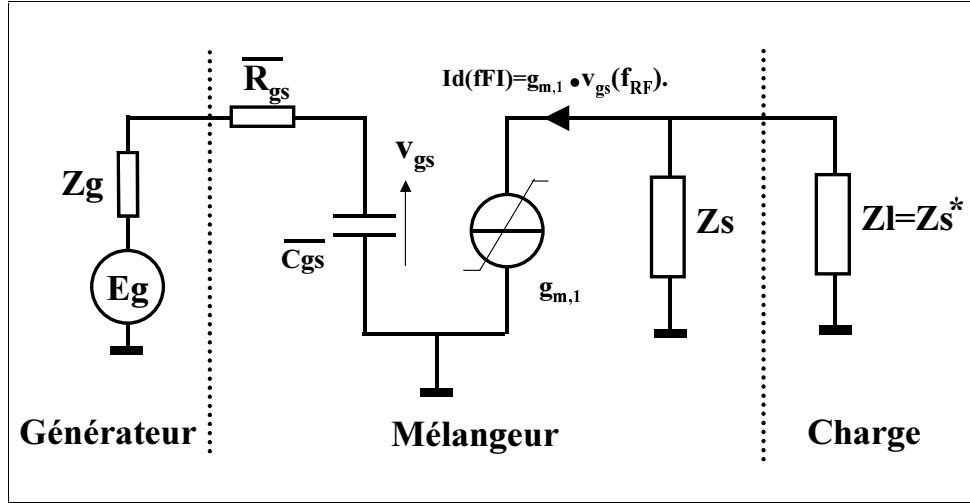


Figure I-9 : Schéma équivalent du mélangeur sur grille.

Dans ces conditions, la valeur du gain de conversion (disponible) est [II.6] :

$$G_c = \frac{g_{m,1}^2 R_l}{16 \omega_{RF}^2 C_{gs}^2 R_e} = \frac{g_{m,max}^2}{64 \omega_{RF}^2 C_{gs}^2 R_e} \times \frac{1}{g_{ds}} \quad (I-3)$$

Il est intéressant de comparer cette valeur à celle obtenue pour un amplificateur source commune, polarisé à $V_{gs} = V_p/2$ afin de réaliser un bon compromis gain-bruit-linéarité, dont la valeur est :

$$G_{ampli} = \frac{g_{m,max}^2}{16 \omega_{RF}^2 C_{gs}^2 R_e} \times \frac{1}{g_{ds}} \quad (I-4)$$

Cela conduit à l'expression du rapport des deux gains en puissance :

$$\eta = \frac{G_c}{G_{ampli}} = \frac{1}{4} \left(\frac{C_{gs}}{C_{gs}} \right)^2 \frac{R_e}{R_e} \frac{g_{ds}}{g_{ds}} \quad (I-5)$$

Une application numérique de l'expression (I-5) a été menée en se basant sur les hypothèses précédentes et sur le modèle d'un transistor PHEMT à 6 doigts de grille de largeur 50 μm fournit par PML. La résistance R_e , faiblement non-linéaire et peu pompée par le signal OL, a été supposé constante. L'évaluation de η en décibel donne, en respectant l'ordre des facteurs :

$$\eta = -6 + 5,4 + 0 + 1,9 = +1.3 \text{ dB} \quad (I-6)$$

Cette valeur démontre que le gain d'un mélangeur n'est pas fatalement inférieur à celui d'un amplificateur de même technologie, et peut même être supérieur selon les cas. En réalité, le gain de conversion est toujours inférieur à celui de l'amplificateur correspondant car les conditions de fermetures du mélangeur sont très délicates à réaliser et très sensibles sur le gain de conversion. Nous reviendrons sur cette remarque au paragraphe I.3 lorsque nous étudierons l'influence des impédances de fermeture.

Remarque : Plusieurs auteurs négligent la partie réelle de la conductance de sortie du transistor T_{RF} ce qui conduit à l'apparition d'un coefficient multiplicateur de valeur 4 dans l'expression du gain de conversion donné par la formule (I-3), ainsi qu'au remplacement de $\frac{1}{g_{ds}}$ par R_l , partie réelle de l'impédance de charge.

I.2.2. Optimisation des dimensions des transistors et de leurs polarisations

Nous avons tout d'abord choisis des transistors T_{OL} et T_{RF} de mêmes dimensions car étant parcourus par le même courant de drain. En effet, à taille de transistor T_{RF} constante, le pompage optimal est assuré lorsque le courant I_{ds} varie de 0 à I_{dss} (courant I_{ds} pour $V_{gs}=0V$). Le sous-dimensionnement de T_{OL} entraînerait le risque de ne pouvoir atteindre la valeur I_{dss} et résulterait donc un mauvais pompage du transistor de mélange. Tandis que le sur-dimensionnement de T_{OL} serait inutile.

Ceci étant posé, nous allons maintenant voir le dimensionnement de ces transistors.

Les relations (I-2) et (I-3) permettent un certain nombre d'optimisations portant d'une part sur la polarisation du transistor et sur son pilotage par le signal OL et d'autre part sur le dimensionnement des transistors réalisant le mélange. Dans les expressions citées précédemment, il n'est plus nécessaire de différencier les deux transistors qui sont supposés identiques.

Si l'on tient compte de la dépendance des éléments du schéma équivalent du transistor : R_e , C_{gs} et g_m vis à vis du nombre de doigts de grille (N_{bd}) et de la largeur de chaque doigt (W_u), on s'aperçoit que le gain de conversion et la puissance OL correspondante augmentent lorsque le nombre de doigts augmente. On constate de plus que la puissance OL nécessaire augmente lorsque W_u augmente, par contre le gain de conversion est maximal pour une valeur fixée de largeur de grille.

La Figure I-10 montre les variations relatives en dB du gain de conversion et de la puissance OL nécessaire en fonction des dimensions d'un transistor à effet de champ de type PHEMT issu de la technologie de PML.

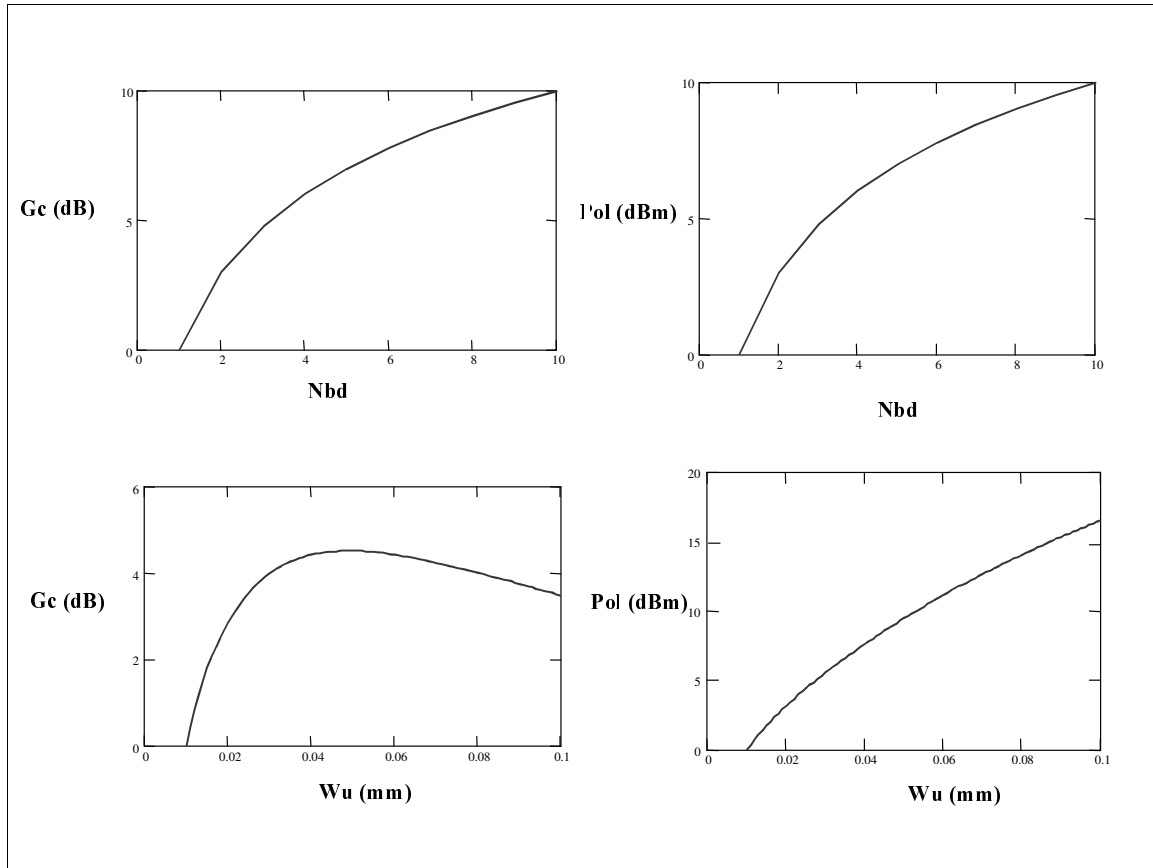


Figure I-10 : Dépendance de G_c et P_{ol} vis à vis de N_{bd} et W_u

Notre choix s'est porté tout d'abord sur une largeur de grille de $50\ \mu\text{m}$, car cette valeur maximise le gain de conversion. Quant au nombre de doigts, nous avons choisi une valeur de 6 réalisant le compromis entre un fort gain de conversion et une faible puissance OL correspondante. Il est à noter que ces valeurs correspondent de plus aux transistors de taille usuelle de cette technologie pour lesquels les modèles sont fiables.

La dernière étape consiste à choisir la polarisation du transistor ainsi que l'amplitude du signal OL optimale afin de maximiser le gain de conversion du mélangeur. En première approximation, seul le terme $g_{m,1}$ est sensible à ces paramètres. Maximiser le gain de conversion reviendra donc à maximiser la dynamique de $g_m(t)$ lorsque le transistor est pompé par le signal OL. Il faut donc que g_m atteigne ses valeurs extrêmes 0 et $g_{m,max}$, c'est à dire que V_{gs} doit atteindre les valeurs V_p et $0v$. Si la valeur de V_{gs} dépasse $0v$, la valeur de g_m va décroître, ce qui aura tendance à diminuer $g_{m,1}$; par contre, rien n'interdit à V_{gs} d'être inférieur à V_p . En conclusion, le point de polarisation peut être proche de V_p , ce point de polarisation est idéal car il minimise la puissance consommée du mélangeur ; et la dynamique du signal OL est telle que le signal V_{gs} ne dépasse pas $0v$, c'est à dire que l'amplitude du signal OL doit être proche de la valeur absolue de V_p .

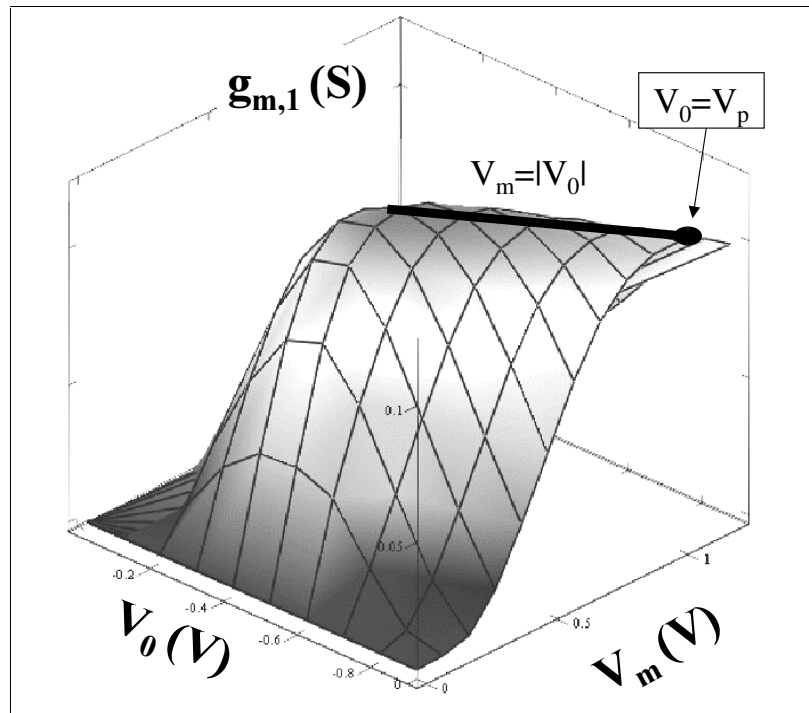


Figure I-11 : variation de $g_{m,1}$ en fonction de V_p et V_m

La Figure I-11 montre les variations de $g_{m,1}$ d'un transistor à effet de champ ($N=6$ et $W_u=50 \mu m$) en fonction des caractéristiques du signal de pompe : V_0 : tension de polarisation et V_m : amplitude du signal alternatif. Cette figure confirme l'analyse qualitative précédente et montre que l'amplitude optimale du signal alternatif est égale, en valeur absolue, à la tension de polarisation et que l'on obtient une valeur maximale de $g_{m,1}$ pour une polarisation égale à V_p .

I.3. Les impédances de fermetures

Comme nous l'avons déjà souligné, l'optimisation des réseaux présents sur les différents accès d'un mélangeur est bien plus délicate que celle menée dans le cas d'un amplificateur.

Cette difficulté rend une optimisation logicielle globale quasiment impossible. En effet, comme nous allons le présenter dans les paragraphes suivants, outre l'adaptation en puissance des différents accès (somme toute déjà délicate à optimiser compte tenu des trois fréquences à considérer dans le cas de mélangeurs) il est nécessaire que les impédances placées sur les accès satisfassent un certain nombre de règles afin d'assurer un fonctionnement optimum au mélangeur.

Pour ces raisons, il est préférable de mener cette optimisation manuellement afin de garantir à chaque étape le respect de ces règles.

Nous allons donc décrire les précautions à prendre lors de l'optimisation des impédances de fermetures de mélangeurs. Nous allons scinder cette présentation en deux parties, à l'image du calcul du gain de conversion :

- Nous présenterons tout d'abord les contraintes pour assurer un pompage optimal par le signal OL.
- Puis nous présenterons pour les fréquences f_{FI} et f_{RF} les contraintes assurant les meilleures caractéristiques possibles à la cellule de mélange.

I.3.1. Contraintes à la fréquence f_{OL}

Le pompage optimal du transistor, c'est à dire le maximum de variation de la transconductance, est assuré lorsque des court-circuits à la masse pour la fréquence f_{OL} et ses harmoniques sont placés sur le drain et la grille du transistor de mélange : T_{RF} [II.7]. Nous avons déjà présenté, à maintes reprises, la nécessité de ces filtres quelle que soit la topologie de mélange considérée.

I.3.2. Contraintes aux fréquences f_{RF} et f_{FI}

a) Particularités du mélangeur sur source

Nous avons déjà évoqué les fortes instabilités potentielles de la cellule de mélange choisie liées au fait que la source du transistor de mélange n'est pas directement connectée à la masse. La contre-réaction engendrée dégrade toutes les caractéristiques mais surtout peut se traduire par des oscillations parasites de la cellule.

Afin d'éviter ces disfonctionnements du transistor de mélange T_{RF} , il est nécessaire de placer une impédance entre la source et la masse présentant :

- un court-circuit idéalement à toutes les fréquences (à l'exception de la fréquence f_{OL}) et pratiquement au moins aux fréquences f_{RF} et f_{FI} .
- Un circuit ouvert à la fréquence f_{OL} et éventuellement à ses harmoniques.

Ces conditions sont particulières à la cellule de mélange choisie mais nous allons maintenant présenter d'autres conditions de charges qui s'appliquent à tous les mélangeurs à transconductance.

b) contraintes de charges des mélangeurs à transconductance

Les autres règles à observer sont toutes liées à la présence d'éléments supplémentaires dans le schéma équivalent de la Figure I-9. En effet, la prise en compte de la capacité grille-drain (C_{gd}) du transistor de mélange et du coefficient du premier harmonique à la fréquence f_{OL} de la conductance de sortie du transistor (G_{d1}) ainsi que de la valeur moyenne de la transconductance du transistor de mélange (G_{m0}), démontre la dégradation des performances de la cellule de mélange.

i - Influence de C_{gd} et solution

La capacité C_{gd} induit, quel que soit le circuit et la configuration, beaucoup de désagréments. Dans un amplificateur son principal impact est l'effet Miller qui réduit sa bande de fonctionnement. Dans un mélangeur, l'effet Miller n'intervient plus car les

fréquences d'entrée et de sortie sont différentes. Par contre la capacité C_{gd} entraîne une contre réaction parallèle qui se traduit par une chute significative du gain de conversion [II.8] mais aussi par une dégradation des isolations, de la linéarité [II.9] et surtout de la stabilité. Des applications numériques montrent, en effet, une diminution pouvant atteindre 8 dB du gain de conversion par rapport au gain obtenu pour les conditions de fonctionnement indiquées aux paragraphes I.2).

Pour les amplificateurs, le remède est le montage cascode ou le *neutrodynage*. Pour les mélangeurs, la solution est plus simple, il suffit de placer un court-circuit à la masse sur le drain du transistor de mélange à la fréquence f_{RF} , éliminant ainsi toute contre réaction de la sortie vers l'entrée à cette fréquence.

ii - Influence de G_{d1} et solution

Si l'on tient compte du coefficient à la fréquence fondamentale f_{OL} de la conductance de sortie du transistor de mélange (G_{d1}), le gain de conversion est multiplié par le coefficient suivant :

$$\left| 1 - \frac{1}{2} \frac{G_{d1}}{G_{d0}} \frac{G_{m0}}{G_{m1}} \right|^2 < 1 \quad (I-7)$$

Ce coefficient traduit la présence de deux modes de conversion de fréquence de la fréquence f_{RF} en entrée vers la fréquence f_{FI} en sortie :

- Le premier mode est celui que nous avons considéré jusqu'à présent, la conversion de fréquence se fait au travers de G_{m1} , le courant de sortie à la fréquence f_{FI} se partage entre G_{d0} et la résistance de charge.
- Le second mode apparaît lorsque l'on tient compte de G_{m0} et G_{d1} . Le courant de drain possède, via G_{m0} , une composante spectrale à la fréquence f_{RF} et c'est G_{d1} qui réalise la conversion de fréquence générant une composante spectrale de la tension à ses bornes à la fréquence f_{FI} .

Les deux modes de conversion se combinent, théoriquement, soit en phase soit en opposition de phase, mais pratiquement on observe toujours une chute du gain de conversion.

Nous avons estimé que ce coefficient induisait une diminution de 5 dB sur la valeur du gain de conversion. La solution est identique à celle reportée au paragraphe précédent. Un court-circuit à la masse pour la fréquence f_{RF} sur le drain du transistor de mélange empêche, en effet, la circulation d'un courant à cette fréquence dans la conductance de sortie du transistor. On évite de la sorte le mode de conversion de fréquence induit par G_{d1} .

iii - Influence de G_{m0} et solution

L'influence de G_{m0} porte essentiellement sur le facteur de bruit du mélangeur [II.10]. En effet, tout bruit présent en entrée à la fréquence f_{FI} se retrouve dans la bande FI d'intérêt en sortie par amplification par G_{m0} .

La solution consiste à éliminer tout bruit en entrée à la fréquence f_{FI} en plaçant un court-circuit à la masse à cette fréquence sur la grille du transistor de mélange. Le facteur de bruit correspond alors à la définition donnée par l'IEEE (chapitre 1, relation (II.4)) :

$$F_{SSB}^{IEEE} = 1 + \frac{B_d}{G_{RF \rightarrow FI} * kT_0} \quad (I-8)$$

Dans le cas où aucune précaution n'est prise, le bruit propre du mélangeur est additionné de la quantité suivante :

$$G_{FI \rightarrow FI} * kT_0 \quad (I-9)$$

Où $G_{FI \rightarrow FI}$ est le gain en puissance définit par :

$$G_{FI \rightarrow FI} = \frac{P_{FI}(f_{FI})}{P_{RF}(f_{FI})} \quad (I-10)$$

et kT_0 la puissance de bruit en entrée à la fréquence f_{FI} . Le facteur de bruit correspond alors l'expression suivante :

$$F_{SSB} = 1 + \frac{G_{FI \rightarrow FI}}{G_{RF \rightarrow FI}} + \frac{B_d}{G_{RF \rightarrow FI} * kT_0} \quad (I-11)$$

Ce qui représente une majoration de 1 à 2 dB par rapport à l'expression (I-8) suivant les valeurs relatives des différents gains en puissance intervenant dans les formules. Cet écart est important et justifie amplement la nécessité du court-circuit à l'entrée RF pour la fréquence f_{FI} .

iv - Adaptation en puissance

Comme pour la conception d'amplificateurs, il est nécessaire d'adapter en puissance chaque accès du mélangeur. En effet, le gain de conversion est d'autant plus important que le générateur est adapté au mélangeur pour la fréquence d'entrée f_{RF} et que la charge est adaptée à la sortie du mélangeur pour la fréquence f_{FI} . Comme pour les amplificateurs, les impédances d'entrée et de sortie peuvent être déterminées afin de privilégier un faible facteur de bruit au détriment du gain. Le cahier des charges imposera l'une ou l'autre des contraintes, et quoi qu'il en soit, le concepteur devra optimiser les impédances d'entrée et de sortie respectivement aux fréquences f_{RF} et f_{FI} .

I.3.3. Synthèse et conclusion

Le Tableau I-2 résume les règles à observer pour la conception d'un mélangeur de performances optimales suivant la nature des causes qui les affectent.

Conséquences Causes / Eléments	↘ Gain	↗ Instabilité	↗ Bruit	↘ Isolation	Règles de conception
Contre réaction sur source	↘	↘	↘	↘	CC à f_{RF} et f_{FI} sur la source T_{RF}
Mauvais Pompage	↘				CC à f_{OL} sur grille et drain de T_{RF}
Influence de C_{gd}	↘	↘		↘	CC à f_{RF} sur le drain T_{RF}
Influence de G_{dl}	↘				CC à f_{RF} sur le drain T_{RF}
Influence de G_{m0}			↘		CC à f_{FI} sur la grille T_{RF}
Adaptation	↘	↘			$\Gamma_{géné} = \Gamma_{in}^*$ $\Gamma_{charge} = \Gamma_{out}^*$

↘ : effet sensible de la cause sur une caractéristique du mélangeur

Tableau I-2 : Synthèse des règles de conception

Ainsi, pour réduire le facteur de bruit, outre l'élimination des contre-réactions sur source réalisée en court-circuitant à la masse la source du transistor de mélange aux fréquences d'intérêt (f_{RF} et f_{FI}), il est nécessaire d'annihiler l'effet de G_{m0} en court-circuitant à la masse la grille du transistor de mélange pour la fréquence f_{FI} .

Afin de valider ces conclusions, nous avons mené des simulations électriques sur la cellule envisagée. Le Tableau I-3 synthétise les résultats de simulations en quantifiant le gain de conversion et le facteur de bruit suivant différentes configurations des impédances de fermetures.

La première ligne de ce tableau correspond au cas où toutes les précautions seraient prises, c'est à dire :

- La grille du transistor T_{RF} est en court-circuit à la masse pour toutes les fréquences à l'exception de la fréquence f_{RF} .
- Le drain du transistor T_{RF} est en court-circuit à la masse pour toutes les fréquences à l'exception de la fréquence f_{FI} .
- La source du transistor T_{RF} est en court-circuit à la masse pour toutes les fréquences à l'exception de la fréquence f_{OL} .

Les lignes suivantes rendent compte des performances en gain de conversion et en facteur de bruit lorsqu'une de ces conditions n'est pas satisfaite.

Les filtres étant considérés sans pertes, les résultats de simulation sur le facteur de bruit sont donc sous-estimés. C'est la raison pour laquelle nous n'avons considéré que ses variations relatives par rapport au cas où toutes les précautions sont prises.

Nous vérifions ainsi que seule l'impédance présentant un court-circuit à la masse sur la grille pour la fréquence f_{FI} affecte le facteur de bruit en l'augmentant de 2dB. De plus, nous

constatons que le gain de conversion diminue fortement si l'impédance connectée au drain et présentant un court-circuit à la masse aux fréquences f_{OL} et f_{RF} est omise. Enfin, nous avons pu observer des problèmes de convergence de nos simulations lorsque l'impédance connectée à la source du transistor de mélange T_{RF} ne satisfaisait pas les conditions requises, laissant présager la présence d'instabilité.

Conditions de fermeture	Gain de conversion	Dégradation du facteur de bruit
Toutes précautions prises	15 dB	Référence
Toutes sauf CC à f_{RF} et f_{FI} sur la source	Instabilité	
Toutes sauf CC à f_{OL} sur la grille	15 dB	≈ 0 dB
Toutes sauf CC à f_{OL} sur le drain	8 dB	≈ 0 dB
Toutes sauf CC à f_{RF} sur le drain	4 dB	≈ 0 dB
Toutes sauf CC à f_{FI} sur la grille	10 dB	+2 dB

Tableau I-3 : Simulations électriques : évaluation des conditions de charges du transistor T_{RF}

Nous venons de poser les règles à observer lors de la conception de la cellule de mélange afin d'aboutir aux meilleures caractéristiques possibles. Nous avons considéré jusqu'à présent des filtres idéaux (impédances soit nulles soit infinies) qui donnaient, bien évidemment, les performances idéales.

Lors de la réalisation pratique de ces impédances passives, de nombreux éléments, fixant la sélectivité des diverses impédances par exemple, resteront à dimensionner et interviendront grandement sur les performances du mélangeur. Il sera donc nécessaire de mener leur optimisation.

De plus, la réalisation pratique du filtre placé sur la source du transistor T_{RF} et assurant la stabilité en éliminant les contre-réactions éventuelles ne permet pas de garantir la stabilité de la cellule à toutes les fréquences. Il est donc nécessaire d'étudier plus précisément la stabilité du mélangeur afin d'optimiser cette impédance.

C'est dans ce double objectif que nous avons mis en œuvre le formalisme des matrices de conversion. Cette méthode va nous permettre en effet d'optimiser, pas à pas, les différents éléments du circuit, avec l'avantage, par rapport à des simulations électriques, de permettre un contrôle de chaque étape d'optimisation. Enfin, les matrices de conversion vont nous permettre d'étudier la stabilité non-linéaire du mélangeur et d'en optimiser l'impédance sensible.

II. Les matrices de conversion

L'étude analytique simplifiée précédente nous a permis d'énoncer les contraintes théoriques nécessaires à un fonctionnement optimal du mélangeur. Cependant, en pratique, le strict respect de ces conditions sera très difficile à satisfaire, voire impossible, et il est nécessaire d'effectuer une dernière étape d'optimisation des diverses impédances du circuit. Devant la complexité de cette optimisation, les simulations électriques longues et fastidieuses peuvent s'avérer peu adaptées pour la convergence vers une solution optimale masquant de plus compréhension du fonctionnement du circuit..

Ainsi, la mise en œuvre d'une méthode analytique fine du type des matrices de conversion peut permettre de pallier aux inconvénients des outils classiques, notamment des simulateurs électriques [II.11]. En effet :

- Les matrices de conversion permettent la compréhension des phénomènes de conversion. Cette opportunité permet la localisation des éléments responsables des différentes conversions de fréquence sur lesquels le concepteur peut interagir.
- Concernant les adaptations, la méthode matricielle considérée permet l'application des formules d'hyperfréquence usuelles pour la synthèse des réseaux d'adaptation en puissance à deux ou trois éléments. Il est alors possible de s'affranchir des longues optimisations, souvent infructueuses, conduites par les simulateurs électriques.
- Enfin, l'implémentation de la méthode des matrices de conversion sur un logiciel de traitement numérique permet une optimisation pas à pas des différents filtres évoqués au paragraphe I.3.

Nous allons donc décrire, dans ce qui suit, la méthode des matrices de conversion et son application à l'optimisation de la cellule de mélange choisie.

II.1. Principe de la méthode

La difficulté de l'analyse des circuits non-linéaires du type mélangeur repose essentiellement sur deux points :

1. Le grand nombre de variables d'état. En effet, le nombre important de composantes spectrales à considérer multiplie d'autant plus le nombre de variables d'état du système.
2. L'impossibilité d'appliquer le principe de superposition. La décomposition du circuit aux différentes fréquences est alors impossible et l'on doit considérer simultanément tous les signaux.

Pour les simulateurs commerciaux utilisant la méthode de l'équilibrage harmonique, ce nombre élevé de variables d'état peut rapidement conduire à des temps de simulation prohibitifs et interdire la mise en œuvre d'optimisations. Ceci est d'autant plus rédhibitoire lors des études analytiques qui nécessitent alors la gestion d'énormes matrices perdant ainsi leur avantage qui réside dans la possibilité d'interpréter les résultats obtenus.

Afin de permettre ces études analytiques et la possibilité d'en dégager la compréhension du fonctionnement des circuits étudiés, nous avons introduit une hypothèse simplificatrice qui consiste à considérer le signal RF de faible amplitude [II.12] et qui correspond au fonctionnement conventionnel d'un mélangeur.

Afin d'explicitier le principe de l'analyse reposant sur cette hypothèse, nous allons l'illustrer en considérant une conductance non-linéaire (non-dynamique, sans mémoire et invariante dans le temps) décrite par la relation $i=f(v)$ et soumise à une différence de potentiel égale à $V_{OL} + \tilde{v}_{RF}$. Toute la méthode repose sur l'hypothèse petit signal de la tension $\tilde{v}_{RF}$ et sur le développement au premier ordre de la relation non-linéaire qui en découle :

$$i = f(V_{OL} + \tilde{v}_{RF}) = f(V_{OL}) + \left. \frac{\partial f}{\partial v} \right|_{V_{OL}} \times \tilde{v}_{RF} \quad (\text{II-1})$$

Le courant de sortie est donc composé de trois quantités :

- La quantité $f(V_{OL})$ dont la détermination nécessite une résolution de type « Equilibrage Harmonique » pour laquelle seul le signal OL est considéré. Cette étude est considérablement plus aisée que l'analyse globale qui considère les deux signaux OL et RF.
- La quantité $g(V_{OL}) = \left. \frac{\partial f}{\partial v} \right|_{V_{OL}}$ dont l'évaluation peut se faire de deux manières distinctes :
 - soit directement par l'étude décrite au point précédent qui résout les circuits en calculant le gradient des fonctions non-linéaires, donc cette quantité, pour l'application de la méthode de convergence de Newton-Raphson [II.13].
 - soit indirectement mais tout aussi simplement par un logiciel de traitement numérique en connaissant la fonction $f(v)$, issue de la modélisation du fondeur, et la fonction $V_{OL}(t)$, issue de l'« Equilibrage Harmonique » précédent.

Pour notre part, n'ayant pas directement accès au gradient des fonctions non-linéaires estimé par le simulateur électrique employé (HP-MDS), nous avons appliqué cette dernière méthode.

- La quantité $\tilde{v}_{RF}$ qui correspond alors à notre variable.
En considérant uniquement la réponse aux petits signaux RF, l'expression (II-1) devient :

$$\tilde{i}_{RF} = g(V_{OL}) \times \tilde{v}_{RF} \quad (\text{II-2})$$

Remarquons que le signal V_{OL} n'est plus une variable d'état mais est considérée comme un paramètre d'où le nom d'étude paramétrique.

La relation (II-2) a la propriété remarquable d'être linéaire et variante dans le temps. Cette propriété engendre les avantages suivants :

1. On peut associer à cette relation un schéma équivalent petit signal dont les éléments sont linéaires mais variants dans le temps.
2. On peut appliquer (vis à vis des petits signaux RF) le théorème de superposition.

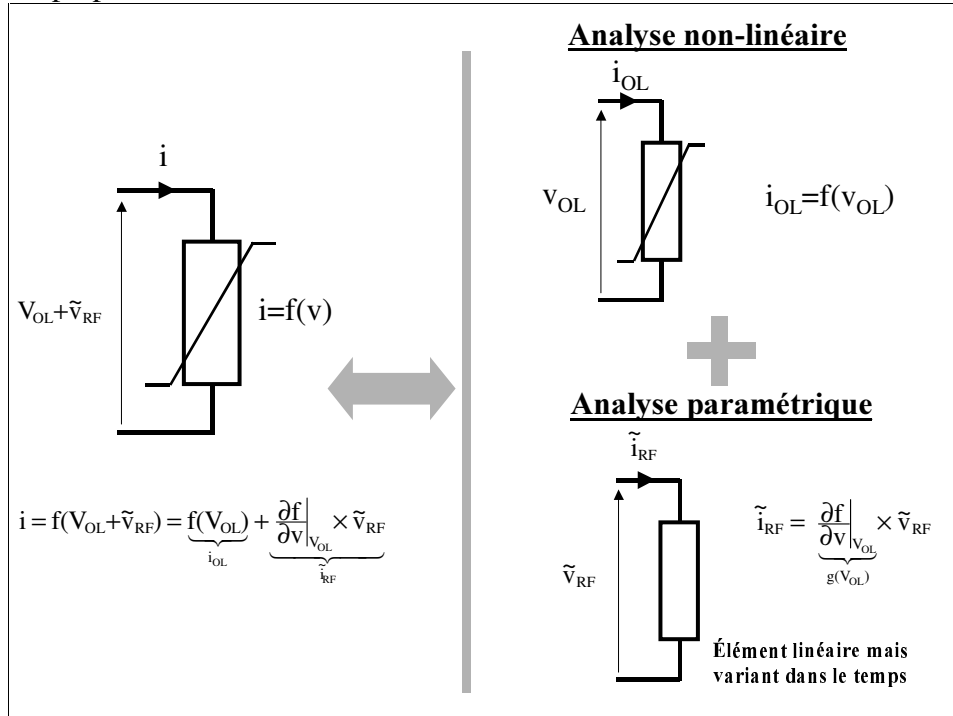


Figure II-1 : Principe de l'analyse paramétrique

La Figure II-1 résume le principe de la méthode exposée qui se scinde en deux étapes [II.14] :

- Une étude non-linéaire de type « Equilibrage Harmonique » qui ne considère que le signal OL et dont sont déduits i_{OL} et V_{OL} (ou $\left. \frac{\partial f}{\partial v} \right|_{V_{OL}}$ selon la méthode choisie).
- Une étude linéaire dont les éléments sont variants dans le temps dit étude paramétrique dont les principes de résolution sont simples et bien connus.

La dernière étape de l'analyse paramétrique consiste à décomposer en série de Fourier la quantité $g(t) = g(V_{OL}(t)) = \left. \frac{\partial f}{\partial v} \right|_{V_{OL}(t)}$. En effet, le signal V_{OL} étant périodique, il en est de même de $g(t)$ dont la décomposition en série de Fourier prend la forme suivante :

$$g(t) = \sum_k G_k \cos(k \times \omega_{OL} t + \varphi_k) \quad (\text{II-3})$$

La relation (II-2) s'écrit alors (en considérant $\tilde{v}_{RF} = V_{RF} \cos(\omega_{RF} t)$) :

$$\tilde{i}_{RF}(t) = G_0 V_{RF} \cos(\omega_{RF} t) + \sum_{k \geq 1} \frac{G_k}{2} V_{RF} \cos [\omega_{RF} \pm k \omega_{OL} t + \varphi_k] \quad (\text{II-4})$$

Ceci Conduit à la représentation spectrale de $\tilde{i}_{RF}(t)$ suivante :

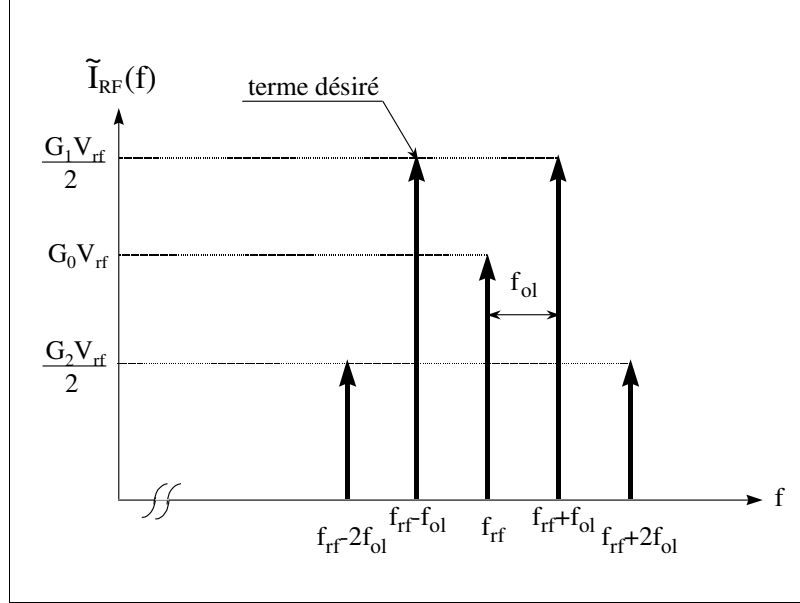


Figure II-2: Spectre monolatéral du signal en sortie du mélangeur.

La relation (II-4) et son interprétation spectrale présentée sur la Figure II-2 démontrent la parfaite adéquation de ce formalisme à l'étude de circuit mélangeur. En effet, la méthodologie exposée exprime directement les diverses conversions de fréquence qui sont engendrées par un élément non-linéaire.

Ce formalisme peut se généraliser quelle que soit la non-linéarité d'expression générale : $s = f(\xi = V_{OL}, \psi = \tilde{v}_{RF})$, où f est une fonction à deux variables. Cette fonction se décompose en effet en série de Taylor suivant la variable $\psi = \tilde{v}_{RF}$ qui est de faible amplitude :

$$s(t) = f(V_{OL}(t), 0) + \underbrace{\left(\frac{\partial f(\xi, \psi)}{\partial \psi} \right)_{\substack{\xi = V_{OL}(t) \\ \psi = 0}}}_{g(V_{OL}) = \sum_k G_k \cos(\omega_{OL} t + \varphi_k)} \times \tilde{v}_{rf} \quad (\text{II-5})$$

L'expression (II-5) est ainsi la généralisation de la relation (II-1).

II.2. Représentation des signaux

La dernière étape pour la formation des matrices de conversion consiste à définir la représentation des petits signaux.

A l'image de l'étude des circuits linéaires qui se réalise aisément à l'aide des phaseurs, nous allons introduire une manière de représenter les petits signaux aux fréquences $f_{RF} \pm kf_{OL}$ qui est particulièrement adaptée au système de conversion de fréquence.

Quelle que soit la variable d'état considérée, elle contient de nombreuses composantes spectrales aux fréquences $f_{RF} \pm kf_{OL}$ résultantes des phénomènes de conversion de fréquence induits par le signal de pompe OL. Il est particulièrement adapté de représenter de tels signaux de la manière suivante [II.14] :

$$x(t) = \underbrace{\sum_{k=-\infty}^{\infty} X_k e^{i(\omega_{RF} + k\omega_{OL})t}}_{x_1(t)} + \underbrace{\sum_{k=-\infty}^{\infty} X_k^* e^{i(-\omega_{RF} + k\omega_{OL})t}}_{x_2(t)} \quad (\text{II-6})$$

Les composantes spectrales aux fréquences $f_{RF} \pm kf_{OL}$ sont ainsi séparées de celles aux fréquences $-f_{RF} \pm kf_{OL}$. Comme les signaux traités sont réels, la connaissance d'un des deux groupes de composantes spectrales suffit à définir parfaitement le signal. Nous allons donc montrer qu'il est possible de se restreindre au premier groupe de composantes spectrales ($x_1(t)$) pour étudier les circuits mélangeurs.

Considérons une non-linéarité définie par une relation paramétrique de même type que celle décrite par l'expression (II-2) : $y(t)=g(t) \bullet x(t)$ et pour laquelle $g(t)$ se décompose en série de Fourier (exponentielle) de la façon suivante :

$$g(t) = \sum_{n=-\infty}^{\infty} G_n e^{in\omega_{OL}t} \quad \text{avec } G_n = G_{-n}^* \text{ car } g(t) \text{ est réel} \quad (\text{II-7})$$

La sortie se met alors sous la forme :

$$y(t) = \underbrace{\sum_n \sum_k X_k G_n e^{i[\omega_{RF} + (k+n)\omega_{OL}]t}}_{y_1(t)} + \underbrace{\sum_n \sum_k X_k^* G_n e^{i[-\omega_{RF} + (k+n)\omega_{OL}]t}}_{y_2(t)} \quad (\text{II-8})$$

Le premier membre se simplifie par le changement de variable : $n \rightarrow \mu=k+n$:

$$y_1(t) = \sum_{\mu} \underbrace{\left\{ \sum_k G_{\mu-k} X_k \right\}}_{Y_{\mu}} e^{i[\omega_{RF} + \mu\omega_{OL}]t} \quad (\text{II-9})$$

$y_1(t)$ prend alors la même forme que $x_1(t)$, et l'on déduit les coefficients Y_{μ} par la formule suivante :

$$Y_{\mu} = \sum_k G_{\mu-k} X_k \quad (\text{II-10})$$

Quant au second membre deux changements de variable sont nécessaires : $k \rightarrow m=-k$ et $m \rightarrow \mu=k+m$:

$$y_2(t) = \sum_{\mu} \underbrace{\left\{ \sum_m G_{\mu+m} X_m^* \right\}}_{\left(\sum_m G_{-\mu-m} X_m \right)^* = Y_{-\mu}^*} e^{i[\omega_{RF} + \mu\omega_{OL}]t} \quad (\text{II-11})$$

Il est remarquable que $y_2(t)$ prenne la même forme que $x_2(t)$ et cela prouve que représenter un signal par uniquement ses composantes spectrales aux fréquences $f_{RF} \pm kf_{OL}$ pour l'étude de circuit paramétrique est possible.

Nous allons donc utiliser cette représentation pour l'étude des circuits convertisseurs de fréquence dont la Figure II-3 résume le principe.

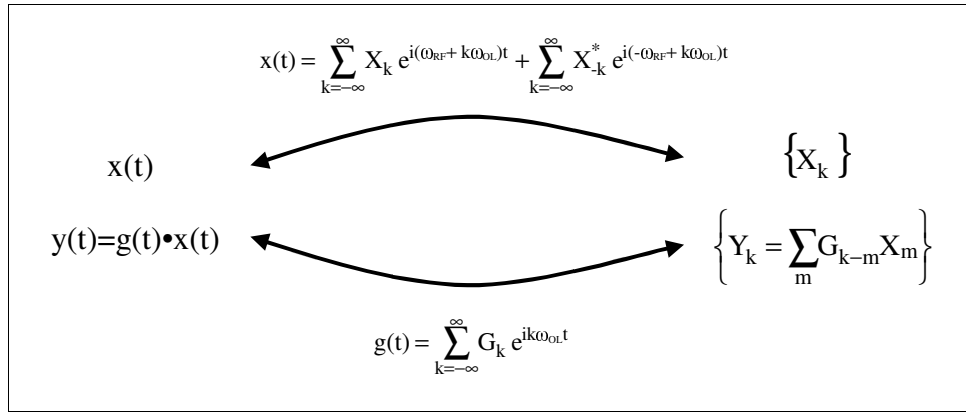


Figure II-3 : Représentation des signaux

Tout signal va donc se caractériser par un ensemble de coefficient complexe X_k . En théorie k varie de $-\infty$ à $+\infty$, mais en pratique, il est indispensable de limiter le nombre de variables à considérer. Cette limitation va dépendre du degré de développement en série de Fourier de $g(t)$ assurant une précision suffisante. Si le degré convenable du développement de $g(t)$ est N alors le nombre de coefficients X_k à considérer est alors : $2N+1$ (k variant de $-N$ à N).

II.3. Formation des matrices de conversion

L'intérêt du formalisme précédemment introduit réside dans la formule (II-10) qui en est déduite et qui devient la nouvelle relation d'état (ou ensemble de relations d'état) de tout élément linéaire ou non-linéaire.

Remarquons tout d'abord que, comme introduit au paragraphe II.1, la relation d'état non-linéaire mono-dimensionnelle initiale ($y=f(x)$) s'est transformée en $2N+1$ relations d'état linéaires ($2N+1$ étant le nombre de composantes spectrales considérées).

Il est évident que travailler sur un ensemble de relations, certes linéaires, n'est pas aisé. C'est pour cette raison qu'est introduit le formalisme matriciel qui permet de résumer en une relation un ensemble de relations linéaires. Ce formalisme est d'autant plus adéquate que la relation (II-10) se met sous la forme matricielle suivante :

$$\vec{Y}_n = [G]_{n \times n} \bullet \vec{X}_n \quad (\text{II-12})$$

Dans laquelle :

$$\vec{Y}_n = \begin{bmatrix} Y_{-N} \\ \vdots \\ Y_0 \\ \vdots \\ Y_N \end{bmatrix}, \quad \vec{X}_n = \begin{bmatrix} X_{-N} \\ \vdots \\ X_0 \\ \vdots \\ X_N \end{bmatrix} \quad \text{et} \quad [G]_{n \times n} = \begin{bmatrix} G_0 & G_1^* & \dots & G_{2N}^* \\ G_1 & G_0 & \ddots & \vdots \\ \vdots & \ddots & \ddots & G_1^* \\ G_{2N} & \dots & G_1 & G_0 \end{bmatrix} \quad (\text{II-13})$$

$[G]$ est la matrice de conversion admittance liant les deux variables d'état vectorielles formées par les $2N+1$ composantes spectrales complexes des signaux $x(t)$ et $y(t)$ aux fréquences $f_{RF} + kf_{OL}$ (k variant de $-N$ à N).

La Figure II-4 conjuguée à la Figure II-3 résume le formalisme des matrices de conversion. Ce formalisme donne une représentation linéaire multi-dimensionnelle d'un phénomène mono-dimensionnel non-linéaire. L'équivalence entre les deux représentations n'est pas mathématiquement rigoureuse à cause de la troncature des développements en série de Fourier.

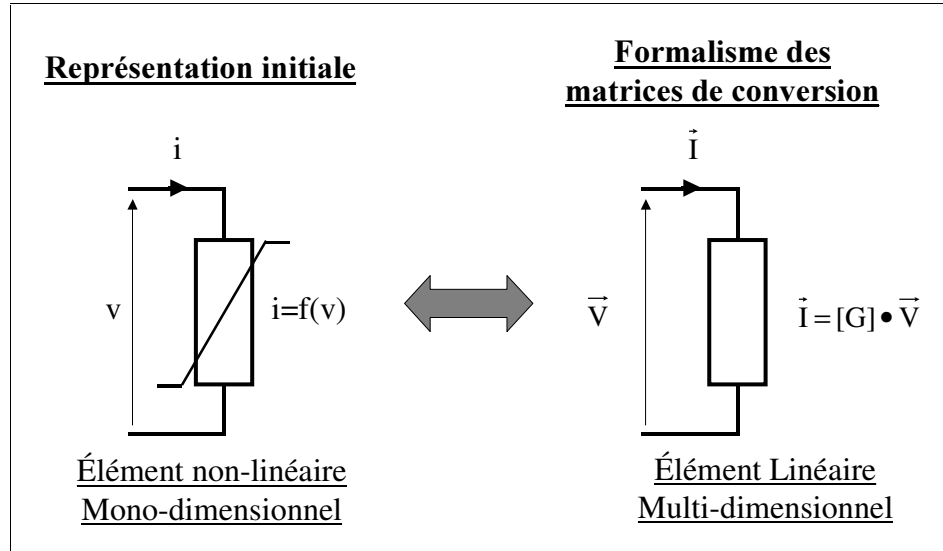


Figure II-4 : Formalisme des matrices de conversion

De plus, rappelons que ce formalisme ne rend pas compte des composantes spectrales aux fréquences kf_{OL} des signaux traités. Le concepteur désireux de connaître leur puissance se reportera à l'analyse non-linéaire préalable à la constitution des matrices de conversion.

Pour conclure, insistons sur le fait que ce formalisme aboutit à une représentation linéaire à laquelle les méthodes analytiques usuelles des circuits linéaires peuvent être appliquées. L'inconvénient majeur réside dans la multiplication de la dimension des variables d'état par le nombre de fréquences considérées.

II.4. Autre type d'éléments non-linéaire ou linéaire

Nous avons décrit la mise en œuvre des matrices de conversion en considérant l'exemple d'une conductance non-linéaire (non dynamique, sans mémoire et invariante dans le temps). Bien évidemment ce formalisme est adapté à tout élément non-linéaire quasi-statique (c'est à dire de caractéristiques R, G, C, L non-dynamiques) sans mémoire et invariant dans le temps. Nous allons traiter deux familles d'éléments particuliers : les éléments linéaires et les capacités, puis le cas d'un transistor à effet de champ qui se décomposera en éléments de base linéaires ou non-linéaires.

II.4.1. Éléments linéaires

Un élément linéaire et invariant dans le temps ne génère que les composantes spectrales présentes en son entrée. Cette propriété démontre que les matrices de conversion qui caractérisent de tels éléments sont diagonales interdisant le couplage entre deux variables d'état de fréquences différentes. La matrice de conversion impédance d'une inductance vaut donc :

$$\begin{bmatrix} V_{-N} \\ \vdots \\ V_0 \\ \vdots \\ V_N \end{bmatrix} = \begin{bmatrix} V(\omega_{RF}-N\omega_{OL}) \\ \vdots \\ V(\omega_{RF}) \\ \vdots \\ V(\omega_{RF}+N\omega_{OL}) \end{bmatrix} = i L \begin{bmatrix} \omega_{RF}-N\omega_{OL} & 0 & \cdots & \cdots & 0 \\ 0 & \ddots & 0 & & \vdots \\ \vdots & 0 & \omega_{RF} & & \vdots \\ \vdots & & & \ddots & 0 \\ 0 & \cdots & \cdots & 0 & \omega_{RF}+N\omega_{OL} \end{bmatrix} \bullet \begin{bmatrix} I_{-N} \\ \vdots \\ I_0 \\ \vdots \\ I_N \end{bmatrix} \quad (\text{II-14})$$

II.4.2. Capacité, trans-capacité

a) Capacité

Une capacité dérive toujours de l'interaction de deux conducteurs métalliques caractérisés par des charges opposées et fonctions de la différence de potentiel qu'il existe entre ces deux corps : Q(V).

Cette interaction est alors caractérisée par une capacité différentielle (à l'image de la résistance dynamique pour les résistances) définie par :

$$c(V_0) = \left. \frac{\partial Q(v)}{\partial v} \right|_{V_0} \quad (\text{II-15})$$

Dans le cas linéaire, on retrouve évidemment la formule bien connue : $Q=C \bullet V$, dans laquelle C est indépendant de V.

Dans le cas des composants semi-conducteurs, les variations des zones de charge d'espace en fonction des polarisations peuvent rendre les capacités non-linéaires. Reprenons alors le développement effectué au paragraphe II.1 :

$$q(t) = q(V_{OL}(t) + \tilde{v}_{RF}(t)) = \underbrace{q(V_{OL}(t))}_{q_{OL}(t)} + \underbrace{\frac{\partial q(v)}{\partial v} \bigg|_{V_{OL}(t)}}_{c(V_{OL}(t))} \times \tilde{v}_{RF}(t) \quad (\text{II-16})$$

En se restreignant à l'étude des petits signaux RF, on aboutit à :

$$\tilde{q}(t) = c(V_{OL}(t)) \times \tilde{v}_{RF}(t) \quad (\text{II-17})$$

De la même manière que précédemment, $c(t)$ est développé en série de Fourier et $\tilde{v}_{RF}(t)$ est mis sous la même forme que $x_1(t)$ dans l'expression (II-6). L'expression (II-17) s'écrit alors :

$$\tilde{q}(t) = \sum_n \sum_k V_n C_k e^{i(\omega_{RF} + (k+n)\omega_{OL})t} \stackrel{\mu=k+n}{=} \sum_n \sum_\mu V_n C_{\mu-n} e^{i(\omega_{RF} + \mu\omega_{OL})t} \quad (\text{II-18})$$

L'expression du courant circulant dans la capacité est alors déduite par dérivation de la charge par rapport au temps :

$$\tilde{i}(t) = \sum_n \underbrace{\left(\sum_\mu i(\omega_{RF} + \mu\omega_{OL}) V_n C_{\mu-n} \right)}_{I_n} e^{i(\omega_{RF} + \mu\omega_{OL})t} \quad (\text{II-19})$$

La forme matricielle de l'expression (II-19) est la suivante [II.15] :

$$\vec{I} = i \underbrace{\begin{bmatrix} \omega_{RF} - N\omega_{OL} & 0 & \dots & \dots & 0 \\ 0 & \ddots & & & \vdots \\ \vdots & & \omega_{RF} & & \vdots \\ 0 & \dots & \dots & 0 & \omega_{RF} + N\omega_{OL} \end{bmatrix}}_{[\Omega]} \cdot \underbrace{\begin{bmatrix} C_0 & \overline{C_1} & \dots & \dots & \overline{C_N} \\ C_1 & \ddots & \ddots & & \vdots \\ \vdots & \ddots & C_0 & \ddots & \vdots \\ \vdots & & \ddots & \ddots & C_1 \\ C_N & \dots & \dots & C_1 & C_0 \end{bmatrix}}_{[C]} \cdot \vec{V} = i \cdot [\Omega] \cdot [C] \cdot \vec{V} \quad (\text{II-20})$$

On prendra garde de ne pas commuter le produit entre les matrices $[\Omega]$ et $[C]$.

b) transcapacité

L'étude précédente considère que la charge q des deux conducteurs considérés n'est fonction que d'une variable d'état qui correspond à leur différence de potentiel. Dans le cas de transistors, bipolaires et unipolaires, les zones de charge d'espace sont fonctions de deux tensions ξ et ψ ($\xi = V_{gs}$ ou V_{be} et $\psi = V_{gd}$ ou V_{bc}), les charges sont alors aussi fonctions de ces deux variables ξ et ψ .

Considérons la charge de deux conducteurs métalliques dont ξ est leur différence de potentiel et ψ une tension secondaire qui influence cette charge. Le développement en série de Taylor de la charge par rapport aux deux variables s'écrit alors :

$$\begin{aligned}
 q(t) &= q(\xi = \xi_{OL} + \tilde{\xi}_{RF}, \psi = \psi_{OL} + \tilde{\psi}_{RF}) \\
 &= \underbrace{q(\xi_{OL}, \psi_{OL})}_{q_{OL}(t)} + \underbrace{\frac{\partial q(\xi, \psi)}{\partial \xi} \bigg|_{\xi = \xi_{OL}(t), \psi = \psi_{OL}(t)}}_{c(t)} \times \tilde{\xi}_{RF}(t) + \underbrace{\frac{\partial q(\xi, \psi)}{\partial \psi} \bigg|_{\xi = \xi_{OL}(t), \psi = \psi_{OL}(t)}}_{c_m(t)} \times \tilde{\psi}_{RF}(t)
 \end{aligned}
 \tag{II-21}$$

Finalement, la charge petit signal s'écrit :

$$\tilde{q}(t) = c(t) \times \tilde{\xi}_{RF}(t) + c_m(t) \times \tilde{\psi}_{RF}(t)
 \tag{II-22}$$

Où $c(t)$ est la capacité (variant dans le temps) différentielle liant la charge entre deux conducteurs avec leur différence de potentiel et $c_m(t)$ est une trans-capacité liant la charge entre ces deux mêmes conducteurs et une tension quelconque du circuit. Le nom de trans-capacité fait référence au terme trans-conductance qui lie un courant à une différence de potentielle géographiquement délocalisée.

Dans le cas des transistors à effet de champ, l'écriture matricielle du courant grille source est alors :

$$\vec{I}_{gs} = i[\Omega] \cdot [C_{gs}] \cdot \vec{V}_{gs} + i[\Omega] \cdot [C_{gs}^m] \cdot \vec{V}_{gd}
 \tag{II-23}$$

II.4.3. Transistor à effet de champ

Dans tous nos circuits, les non-linéarités seront générées exclusivement par les transistors. Afin de mener l'étude paramétrique, les transistors seront associés à un schéma équivalent constitué d'éléments de base linéaires ou non-linéaires tels que des résistances, capacités, inductances et sources commandées. La Figure II-5 montre le schéma équivalent d'un transistor à effet de champ.

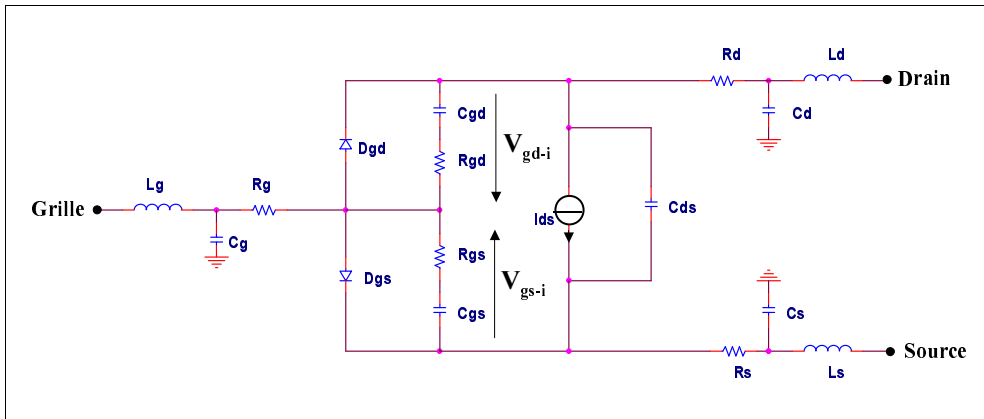


Figure II-5 : Schéma équivalent d'un transistor à effet de champ [II.16]

Il est important de remarquer que toutes les non-linéarités du schéma équivalent du transistor ont pour uniques variables d'état les tensions V_{gd-i} et V_{gs-i} .

Cette dernière remarque est primordiale pour la mise en œuvre des matrices de conversion que nous présentons au prochain paragraphe. En effet, quel que soit le circuit envisagé, l'ensemble des variables d'état de toutes les non-linéarités se limitera aux seules tensions de commande V_{gd-i} et V_{gs-i} de tous les transistors.

Ainsi, lors du processus de détermination des matrices de conversion pour l'étude d'un circuit mettant en œuvre M transistors, l'évaluation des quantités $\frac{\partial f}{\partial v} \Big|_{V_{gs-i}^1(t) V_{gd-i}^1(t) \dots V_{gd-i}^M(t)}$ ne nécessitera la connaissance que des variables $V_{gs-i}^1(t) V_{gd-i}^1(t) \dots V_{gd-i}^M(t)$ issues de l'étude non-linéaire du circuit soumis uniquement au fort signal de pompe OL.

II.5. Mise en œuvre des matrices de conversion

L'utilité du formalisme que nous venons introduire repose sur l'application rendue possible des méthodes de l'électrocinétique : lois de Kirchhoff et les formules d'associations série parallèle qui en sont déduites.

De plus, le formalisme des matrices de conversion donnant une représentation linéaire du circuit (non-linéaire) étudié, il est possible d'appliquer le théorème de superposition ainsi que les nombreuses formules analytiques des circuits linéaires : adaptation conjuguée, formules de gain, ...

L'unique précaution à considérer est liée à l'algèbre linéaire imposée par les matrices de conversion qui se différencie de l'algèbre « mono-dimensionnelle » par la non-commutativité du produit.

II.5.1. Processus de calcul des matrices de conversion

La Figure II-6 décrit le processus d'établissement des matrices de conversion que nous avons mis en place.

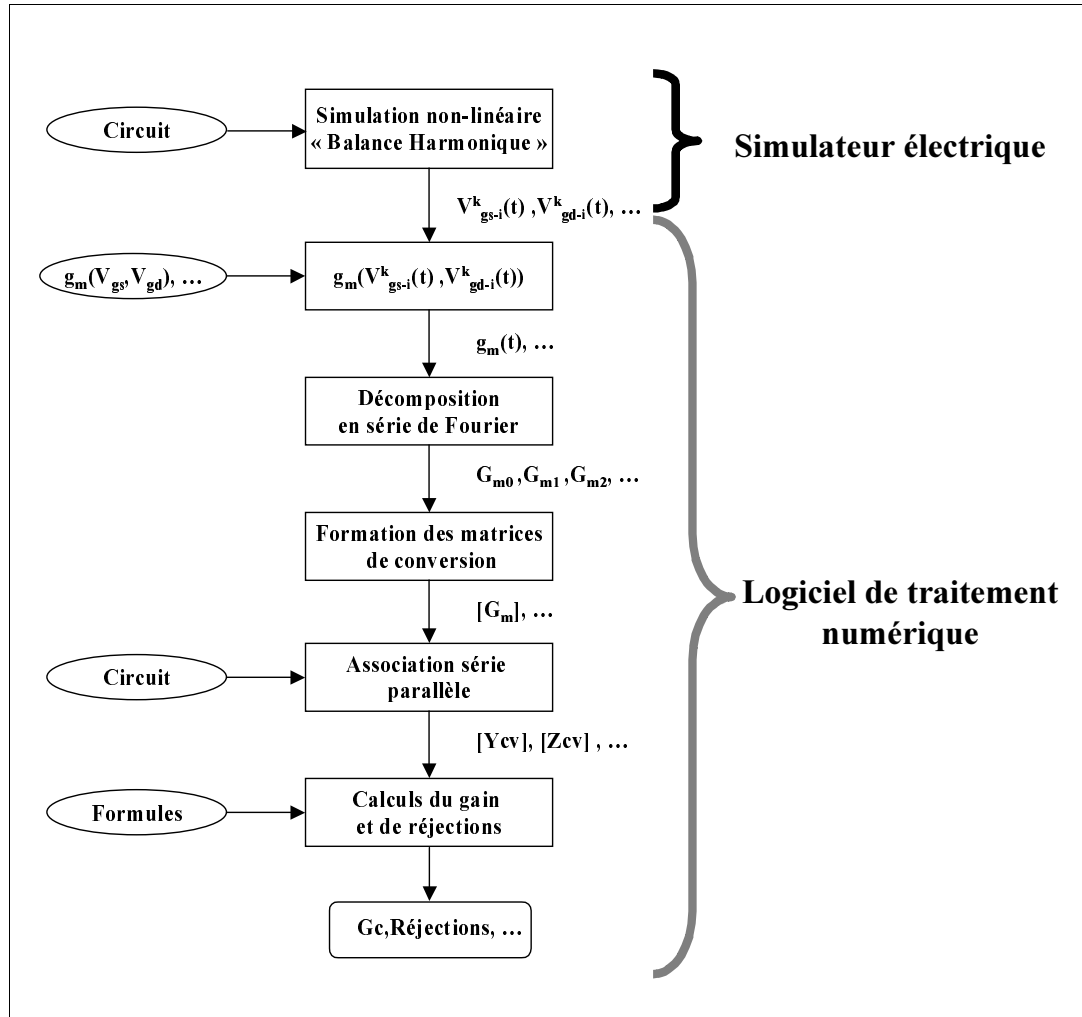


Figure II-6 : Processus de résolution du circuit

Tout d'abord, nous avons mené une étude du circuit uniquement soumis au fort signal de pompe OL à l'aide d'un simulateur électrique de type « Equilibrage Harmonique » (HP-MDS). Nous avons donc pu extraire les variables d'état $V_{gs-i}^k(t), V_{gd-i}^k(t), \dots$ intervenant en tant que commandes des non-linéarités. Ces variables, fonctions du temps, ont été calculées pour différentes puissances du signal OL : tous les résultats seront donc paramétrés par la puissance P_{OL} .

Ces variables d'état, ainsi que les expressions $(g_m(V_{gs}, V_{gd}), \dots)$ des différents éléments non-linéaires constituant le circuit ont été intégrées dans un logiciel de traitement numérique afin d'extraire leurs variations temporelles $(g_m(t), c(t), \dots)$. Nous avons ensuite procédé à la décomposition en série de Fourier de ces éléments $(G_{m0}, G_{m1}, \dots)$ afin de former les matrices de conversion qui leurs sont associées $([G_m], \dots)$ puis les matrices de conversion globales $([Y_{cv}], ([Z_{cv}], \dots))$.

Enfin, l'application de formules classiques de l'électrocinétique nous a permis de déterminer le gain de conversion et les réjections $(G_c, \text{Réjections}, \dots)$ du mélangeur étudié.

II.5.2. Détermination des matrices de conversion

La Figure II-7 présente le schéma simplifié de la cellule de mélange objet de notre étude. La finalité de la méthode mise en œuvre est la détermination des deux cellules d'adaptation et de filtrage aux accès RF et FI ainsi que de l'impédance de stabilisation Z_{stab} (la cellule de l'accès OL est déterminée par les simulations électriques non-linéaires préalables au formalisme des matrices de conversion).

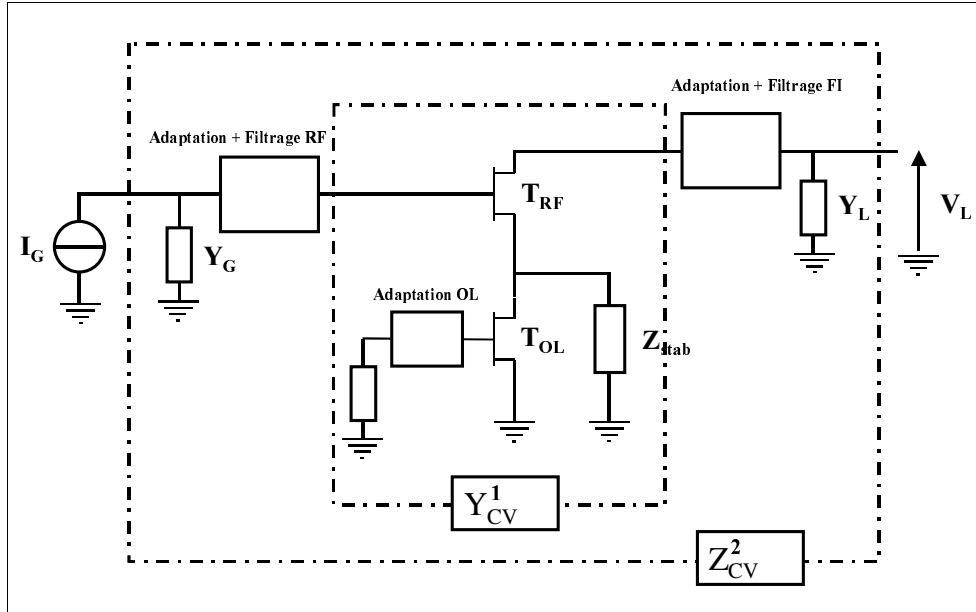


Figure II-7 : Matrices de conversion utiles

Le schéma présenté correspond à celui de l'étude paramétrique dans lequel les éléments sont caractérisés par des matrices. Le générateur OL est quant à lui éteint car seuls les petits signaux RF sont considérés.

A ce stade de l'étude, deux étapes restent encore à conduire afin d'aboutir à une cellule optimale de mélange :

1. **L'adaptation** de l'accès RF à la fréquence f_{RF} et de l'accès FI à la fréquence f_{FI} . Les cellules envisagées sont des circuits à trois éléments réactifs, en T ou en Π , permettant d'ajuster la fréquence et la bande passante de l'adaptation. L'évaluation de la matrice de conversion admittance Y_{CV}^1 (se reporter à la Figure II-7 pour la définition de cette matrice) ainsi que l'application des formules analytiques classiques de l'adaptation en puissance permettent la détermination rapide et optimale des réseaux d'adaptation. Ce travail est aussi possible par des simulations électriques mais les optimisations de type gradient et/ou aléatoire ne convergent pas toujours et souvent ne donnent pas les résultats optimaux.

Le calcul de l'adaptation est mené comme suit. Partant de la définition de Y_{CV}^1 :

$$I = \begin{bmatrix} I_e = \begin{bmatrix} \vdots \\ I_e(f_{RF}) \\ \vdots \end{bmatrix} \\ I_s = \begin{bmatrix} \vdots \\ I_s(f_{FI}) \\ \vdots \end{bmatrix} \end{bmatrix} = [Y_{CV}^1] \bullet \begin{bmatrix} V_e = \begin{bmatrix} \vdots \\ V_e(f_{RF}) \\ \vdots \end{bmatrix} \\ V_s = \begin{bmatrix} \vdots \\ V_s(f_{FI}) \\ \vdots \end{bmatrix} \end{bmatrix} \quad (\text{II-24})$$

On supposera les filtres d'entrée (accès RF) et de sortie (accès FI) idéaux, c'est à dire ne laissant passer que les fréquences d'intérêt et court-circuitant à la masse les autres : $V_e(f \neq f_{RF})=0$ et $V_s(f \neq f_{FI})=0$. En ne considérant que les courants aux fréquences d'intérêt, la relation (II-24) se réduit donc à :

$$\begin{bmatrix} I_e(f_{RF}) \\ I_s(f_{FI}) \end{bmatrix} = [Y_{CV}^1]_{2 \times 2} \bullet \begin{bmatrix} V_e(f_{RF}) \\ V_s(f_{FI}) \end{bmatrix} \quad (\text{II-25})$$

La nouvelle matrice est de dimension 2×2 , ses coefficients sont extraits de la matrice présente dans l'expression (II-24). Il est alors possible d'appliquer les calculs analytiques de l'adaptation conjuguée d'un quadripôle [II.17].

2. **Le filtrage** à l'entrée RF et la sortie FI. La nécessité de ces filtres a été longuement évoquée et justifiée au paragraphe I.3. Cette étude ultime est menée grâce à l'évaluation de la matrice Z_{CV}^2 (se reporter à la Figure II-7 pour la définition de cette matrice). Cette matrice englobe la charge et l'impédance du générateur. L'accès FI de sortie est alors en circuit ouvert : $I_s(\forall f)=0$ et l'accès RF d'entrée n'est excité que par le courant I_g à la fréquence f_{RF} : $I_e(f_{RF})=I_g$ et $I_e(f \neq f_{RF})=0$.

On a alors :

$$V = \begin{bmatrix} V_e = \begin{bmatrix} \vdots \\ V_e(f_{RF}) \\ \vdots \end{bmatrix} \\ V_s = \begin{bmatrix} \vdots \\ V_s(f_{FI}) \\ \vdots \end{bmatrix} \end{bmatrix} = [Z_{CV}^2] \bullet \begin{bmatrix} I_e = \begin{bmatrix} \vdots \\ I_e(f_{RF})=I_g \\ \vdots \end{bmatrix} \leftarrow k^{\text{ième}} \text{ ligne} \\ I_s = \begin{bmatrix} 0 \\ \vdots \\ 0 \end{bmatrix} \end{bmatrix} \quad (\text{II-26})$$

L'évaluation de tous les coefficients de Z_{CV}^2 est alors superflue car seule la $k^{\text{ième}}$ colonne (correspondant à $I_e(f_{RF})$) intervient pour l'évaluation du vecteur V . On a donc :

$$\vec{V} = \begin{bmatrix} \vec{V}_e = \begin{bmatrix} \vdots \\ V_e(f_{RF}) \\ \vdots \end{bmatrix} \\ \vec{V}_s = \begin{bmatrix} \vdots \\ V_s(f_{FI}) \\ \vdots \end{bmatrix} \leftarrow n^{\text{ième}} \text{ ligne} \end{bmatrix} = \begin{bmatrix} Z_{CV}^2|_{1,k} \\ \vdots \\ Z_{CV}^2|_{2N+1,k} \end{bmatrix} \bullet I_g \quad (\text{II-27})$$

Le gain de conversion (transducique) ainsi que les différentes rejections se calculent alors simplement à partir de l'expression (II-27).

Pour le gain de conversion :

$$G_c = 4 \frac{\Re(Y_L)}{\Re(Y_G) |Y_G|^2} \left| \frac{V_s(f)}{I_g} \right|^2 = 4 \frac{\Re(Y_L)}{\Re(Y_G) |Y_G|^2} \left| Z_{CV}^2 \right|_{n,k}^2 \quad (\text{II-28})$$

Pour les réjections de l'entrée RF à la fréquence f_{RF} vers l'accès X à la fréquence f (qui correspond à la $m^{i\text{ème}}$ ligne des vecteurs d'états), notées $R_{RF(f_{RF}) \rightarrow X(f)}$:

$$R_{RF(f_{RF}) \rightarrow X(f)} = 4 \frac{\Re(Y_X)}{\Re(Y_G) |Y_G|^2} \left| \frac{V_X(f)}{I_g} \right|^2 = 4 \frac{\Re(Y_L)}{\Re(Y_G) |Y_G|^2} \left| Z_{CV}^2 \right|_{m,k}^2 \quad (\text{II-29})$$

Un processus d'optimisation (manuel et/ou automatisé) peut être alors conduit afin d'évaluer la nécessité de certains filtre, d'optimiser certains éléments, et au final de maximiser le gain de conversion et les différentes réjections.

II.6. Application des matrices de conversion à l'optimisation de la cellule de mélange

II.6.1. Validation de la méthode

Nous avons dans un premier temps cherché à valider la méthode analytique des matrices de conversion mise en œuvre. Nous l'avons, pour cela, confrontée à des simulations électriques effectuées sur la cellule de mélange munie de filtres idéaux et sans aucune adaptation.

La Figure II-8 montre le gain de conversion en fonction de la puissance OL appliquée obtenu par les deux méthodes.

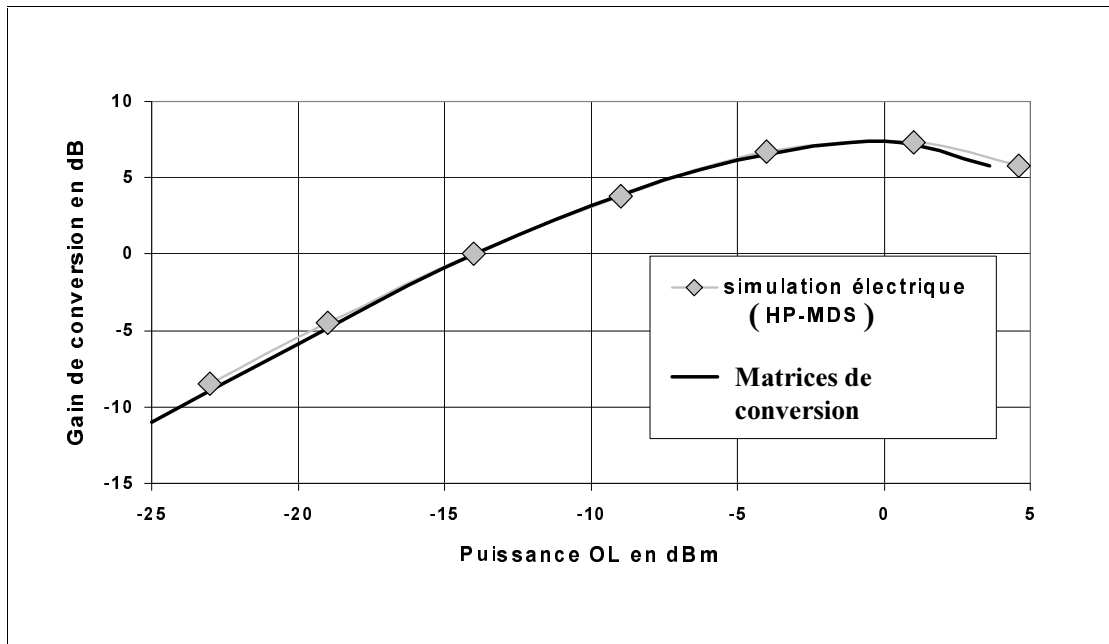


Figure II-8 : Gain de conversion en fonction de la puissance OL, validation de la méthode analytique.

Le faible écart observé (inférieur à 0.3dB) prouve que le formalisme des matrices de conversion permet d'aboutir à des résultats suffisamment fiables et valide le processus mis en œuvre.

Il est donc permis d'envisager l'utilisation des matrices de conversion pour l'optimisation des éléments passifs constituant filtres et cellules d'adaptation du mélangeur.

A ce stade, une dernière précaution est à observer. Il faudra veiller à ce que l'ajout et/ou l'ajustement de cellules passives ne modifie pas fortement le comportement du circuit à la fréquence OL et à ses harmoniques les plus significatives. Dans le cas contraire, le pompage serait alors modifié et il faudrait tenir compte de ces différences dans les matrices de conversion, obligeant à reprendre au début le processus décrit Figure II-6.

II.6.2. Optimisation [II.18]

Le circuit considéré est présenté Figure II-9. Compte tenu des règles de conception énoncées au Tableau I-2, il est constitué des éléments suivants :

- Un filtre de stabilisation constitué d'un circuit résonnant série : L_{stab} , C_{stab} .
- Un filtre de sortie (accès FI) constitué des composants : L_{s0} , C_s et L_{s1} dont le rôle est de présenter une faible impédance à f_{RF} tout en adaptant la partie réactive en sortie à la fréquence f_{FI} .
- Un circuit de stabilisation constitué de L_0 , C_0 et R_0 et dont la nécessité sera démontrée au chapitre 3. Ce circuit fera de plus office de filtre d'entrée en charge de limiter la puissance des composantes spectrales sur la grille de T_{RF} autres que celles à la fréquence f_{RF} .

- Un circuit d'adaptation en entrée (accès RF) constitué des éléments L_{e0} , C_{e0} et L_{e1} formant une cellule en T capable de réaliser une adaptation plus large bande que les cellules à deux éléments.

Remarque : Seule la partie réactive en sortie est adaptée car il s'est avéré qu'une adaptation conjuguée augmentait la complexité du circuit sans pour autant améliorer de façon significative la valeur du gain de conversion.

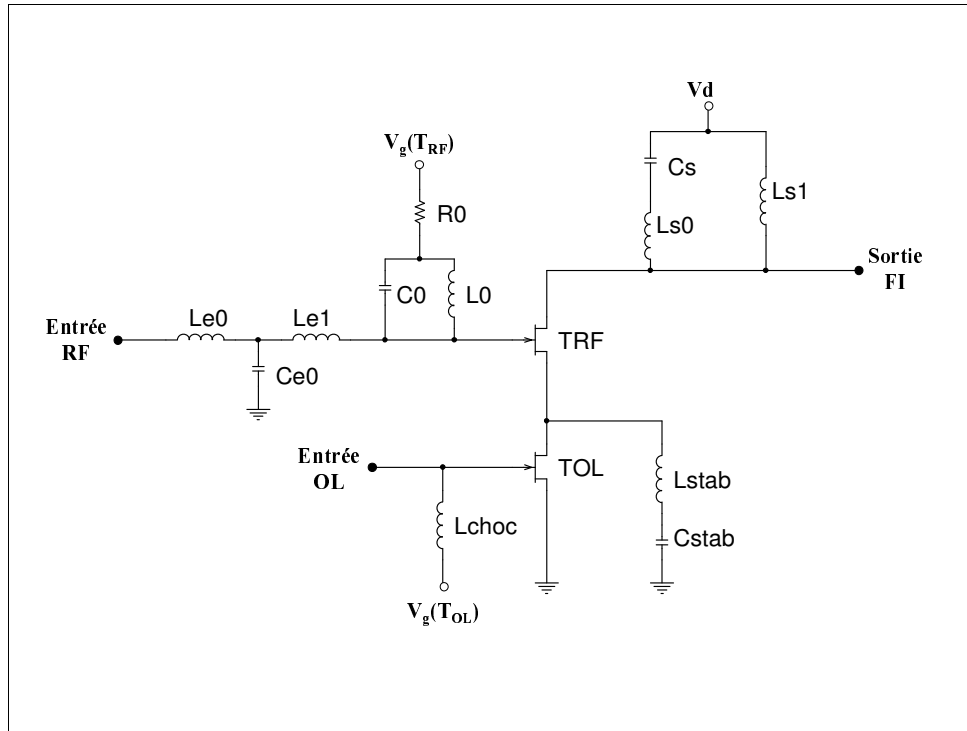


Figure II-9 : Schéma électrique de notre cellule originale de mélange

Tous les éléments du circuit ne sont pas à optimiser car certaines relations liant leur valeur sont imposées pour un fonctionnement correct de la cellule (se reporter au paragraphe I.3) :

- Le filtre de stabilisation (L_{stab} , C_{stab}) doit court-circuiter les fréquences f_{RF} et f_{FI} , sa résonance a donc été fixée au milieu de ces deux fréquences.
- Le filtre de sortie doit présenter une faible impédance à la fréquence f_{RF} fixant une de ses deux résonances.
- Le filtre de stabilisation en entrée doit résonner à la fréquence f_{RF} .

Il ne reste, tout de même, pas moins de cinq éléments à déterminer :

- L_{e0} , L_{e1} et C_{e0}
- L_0
- L_{s1}

Remarque : L_{stab} sera optimiser au paragraphe III afin de garantir la stabilité du circuit.

Le réseau d'adaptation en entrée (L_{e0} , L_{e1} et C_{e0}) a été déterminé par voie analytique tandis que les deux éléments restants (L_0 , L_{s1}) ont fait l'objet d'une optimisation, comme il est décrit au paragraphe II.5.2. Ces études ont été conduites afin de garantir un fonctionnement à fort gain tout en assurant de fortes réjections des signaux indésirables. Nous avons aussi veillé à obtenir une bande passante de valeur conforme à un cahier des charges fixé.

Enfin, cette méthode a permis l'établissement et l'interprétation de nombreuses règles de conception que nous avons exposées au paragraphe I de ce chapitre.

La cellule de mélange conçue présente des performances en gain optimales puisqu'une valeur de 11dB de gain de conversion est atteinte au centre de la bande passante, comme le présente la Figure II-10 qui décrit les variations du gain de conversion en fonction de la fréquence f_{RF} .

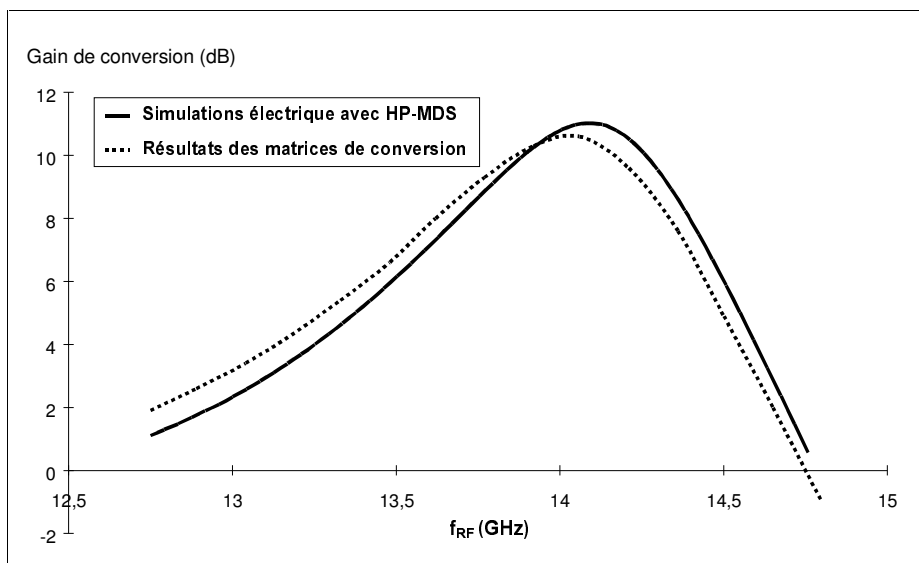


Figure II-10 : Performances en gain de conversion du mélangeur optimisé

En outre, nous avons mené des simulations électriques sur la cellule optimisée dont les résultats apparaissent aussi sur la Figure II-10. Un écart inférieur à 0.5 dB entre les résultats donnés par les deux méthodes valide, ici encore, l'approche des matrices de conversion.

La dernière étape avant l'intégration monolithique de cette cellule de mélange est l'étude de sa stabilité et l'optimisation de l'impédance constituée par les éléments L_{stab} et C_{stab} placée sur la source du transistor de mélange T_{RF} .

Cette étude fait l'objet du paragraphe III que nous allons maintenant aborder.

III. Etude de la stabilité non-linéaire

Les bases de la stabilité ont été établies dès la fin du 19^{ième} siècle successivement par trois mathématiciens [II.19].

En 1854, Charles Hermite écrivit une lettre dans laquelle il décrit l'intérêt de l'étude des pôles du système pour l'investigation de sa stabilité. La méthode qu'il mit en œuvre sur un exemple simple, en ne considérant que les pôles à partie imaginaire positive, fut généralisée aisément aboutissant au critère de stabilité des pôles à parties réelles négatives.

Suivant ce critère, Edward John Routh développa en 1877 un algorithme testant la stabilité des systèmes dont la simplicité et l'efficacité en ont fait le critère algébrique de stabilité utilisé par les automaticiens quelques 120 ans après son établissement.

En marge de ces découvertes, Aleksander Michailovich Lyapunov, élève de Chebyshev et intéressé par la stabilité du système solaire, fut le premier à généraliser les concepts de stabilité à tout système linéaire ou non-linéaire. Il définit en 1892, lors de sa thèse, la stabilité et introduisit des fonctions, qui portent son nom et qui s'apparentent à des énergies, dont l'étude permet celle de la stabilité. Comme pour le critère de Routh, ces concepts ont traversé plus d'un siècle et sont toujours usités.

Plus tard, ces concepts furent appliqués aux systèmes électroniques dans lesquels les concepts de contre réaction mis en œuvre nécessitèrent l'étude systématique de la stabilité. C'est dans l'objectif d'aboutir à une méthode efficace d'évaluation de la stabilité que le mathématicien Harry Nyquist des laboratoires Bell énonça en 1932 [II.20] le fameux critère de stabilité géométrique qui porte son nom. Ce critère facilita grandement l'étude de la stabilité des circuits électriques en évaluant celle-ci visuellement de façon graphique.

Enfin, mathématicien des mêmes laboratoires américains que H. Nyquist, Hendrik W. Bode, qui donna son nom aux diagrammes bien connus des électroniciens, fut le pionnier pour l'étude des circuits électroniques à contre réaction. Il publia en 1945 un ouvrage [II.21] dont les principes et théorèmes sont les références exhaustives de la stabilité des circuits linéaires.

Bien connue de nos jours, la stabilité des systèmes linéaires mono-variables et multi-variables ne pose aucune difficulté d'étude aux concepteurs de circuits électroniques hyperfréquences qui appliquent efficacement le critère géométrique de Nyquist à des fonctions spécifiques correctement choisies. Nous ferons le bilan de ces différentes méthodes au paragraphe III.2.

Cependant, les outils pour l'étude de la stabilité ne sont plus aussi développés lorsque les circuits sont non-linéaires, comme pour le cas des mélangeurs. Seuls les travaux de Lyapunov, et ceux qui en sont déduits, sont applicables et reposent sur une étude analytique des caractéristiques de transfert des systèmes. Dans le cas des circuits électroniques hyperfréquences actuels, les études analytiques sont impossibles et seules les méthodes numériques et géométriques sont applicables. Nous ferons le catalogue critique des différentes méthodes existantes puis nous présenterons la méthode que nous avons choisie et nous traiterons de son application au cas de notre mélange.

Avant la présentation des différentes méthodes d'analyse de la stabilité linéaire et non-linéaire, nous allons définir le concept de système stable et instable.

III.1. Définition de la stabilité

Cette définition fut posée en 1892 par A. Lyapunov. Elle constitue la définition la plus générale de la stabilité car elle concerne tout système qu'il soit linéaire ou non-linéaire.

Dans l'espace des variables d'état : $\vec{X}$, supposons que l'origine soit un point d'équilibre : $\dot{\vec{X}}(0)=0$ et que le système soit initialement dans cet état d'équilibre.

Supposons que le système soit soumis à une perturbation $\delta\vec{X}$ de faible amplitude ($\exists A>0$ tel que $|\delta\vec{X}| < A$) de cet état d'équilibre, trois conséquences sont alors possibles, comme illustré Figure III-1 :

1. Le système retourne dans la position d'équilibre initiale ($\vec{X}(\infty)=0$) sans diverger ($\exists R>0$ tel que $\forall t >0, |\vec{X}(t)| < R$). Le système est alors asymptotiquement stable.
2. Le système ne retourne pas dans sa position d'équilibre initiale ($\vec{X}(\infty) \neq 0$) mais ne diverge pas ($\exists R>0$ tel que $\forall t >0, |\vec{X}| < R$). Le système est stable.
3. Le système diverge, le système est alors instable.

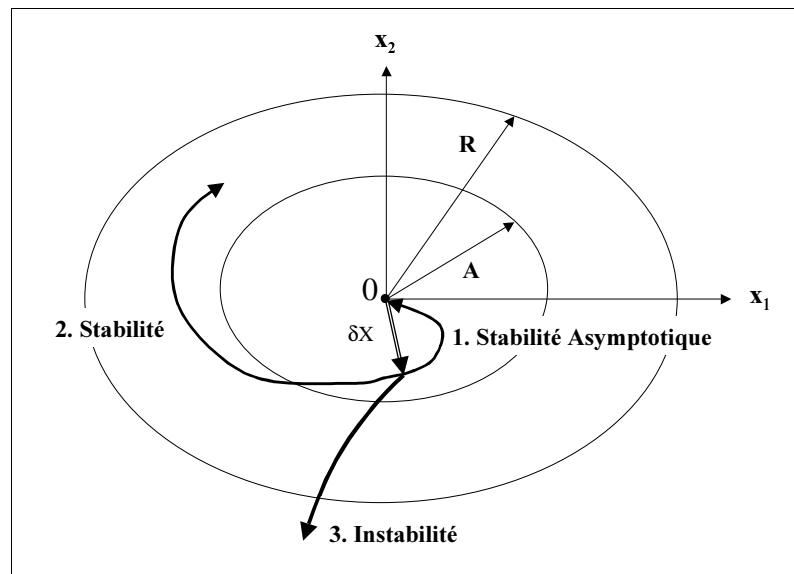


Figure III-1 : Stabilité de Lyapunov

L'ensemble des schémas de la Figure III-2 donne une illustration des différentes configurations de la stabilité au travers de l'étude d'une bille soumise à la gravitation. Ce diagramme permet de plus d'introduire la notion de bifurcation que nous ne traiterons pas dans la suite de ce manuscrit car nous nous limiterons à l'étude de la stabilité asymptotique des systèmes.

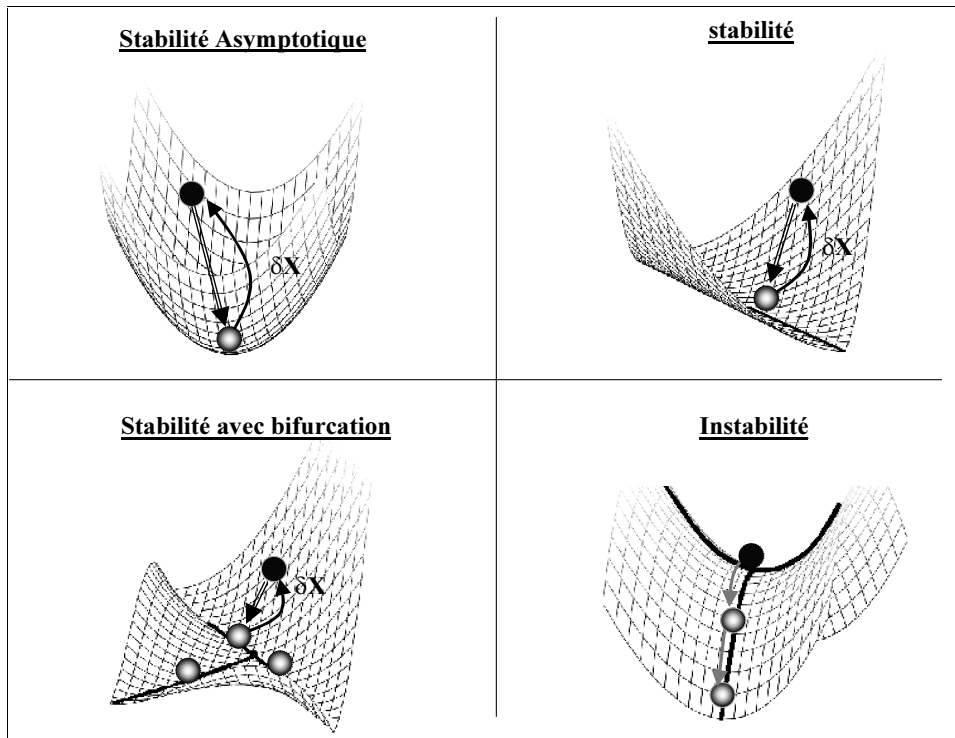


Figure III-2 : Illustration de la stabilité

Nous allons, dans les paragraphes suivants, appliquer ces concepts de stabilité aux systèmes linéaires puis non-linéaires afin d'en dégager des critères exploitables par les concepteurs de circuits électroniques.

III.2. Stabilité linéaire

La stabilité linéaire est la stabilité des circuits linéaires ou linéarisés autour d'un point de fonctionnement n'évoluant pas au cours du temps.

III.2.1. Stabilité dans le domaine spectral

Dans le cas des systèmes linéaires, l'étude de la stabilité se réalise plus aisément dans le domaine spectral. La traduction des trois configurations décrites par A. Lyapunov sont :

1. Le système est asymptotiquement stable si et seulement si les pôles de sa fonction de transfert sont à partie réelle **strictement** négative.
2. Le système est stable si et seulement si les pôles de sa fonction de transfert sont à partie réelle négative **et** ceux dont la partie réelle est nulle ont une multiplicité égale à un.
3. Dans le cas contraire au 2. le système est instable.

L'étude de la stabilité se mène donc par l'analyse des pôles de la fonction de transfert, dont la détermination n'est pas indispensable puisque seule la connaissance de l'existence de pôles à partie réelle positive est requise et non celle de leur valeur.

Algébriquement menée via le critère de Routh, l'étude de la présence de pôles à partie réelle positive peut être extrêmement simplifiée par l'application du critère géométrique de Nyquist.

En effet, le tracé du lieu de Nyquist d'une fonction dans le plan complexe, la fréquence variant de $-\infty$ à $+\infty$, et l'évaluation graphique du nombre d'encerclement de l'origine, compté positivement dans le sens trigonométrique ($=N_0^{\text{trigo}}$), permet de connaître la différence entre le nombre de pôles à partie réelle positive ($=P_{\text{Re} \geq 0}$) et le nombre de zéros à partie réelle positive ($=Z_{\text{Re} \geq 0}$) de cette fonction :

$$N_0^{\text{trigo}} = P_{\text{Re} \geq 0} - Z_{\text{Re} \geq 0} \quad (\text{III-1})$$

Appliquée à la fonction de transfert du système, le critère de stabilité de Nyquist est alors le suivant :

$$\text{Stabilité} \Leftrightarrow P_{\text{Re} \geq 0} = 0 \Leftrightarrow N_0^{\text{trigo}} = -Z_{\text{Re} \geq 0} \quad (\text{III-2})$$

Ainsi, le critère de Nyquist permet de conclure de la stabilité d'un système par l'évaluation graphique du nombre d'encerclement de l'origine du tracé de Nyquist de la fonction de transfert de ce système et de la connaissance du nombre de zéros à partie réelle positive de cette fonction.

Dans la suite de ce manuscrit, l'évaluation du nombre de tours d'un lieu sera toujours effectuée autour de l'origine et compté positivement dans le sens trigonométrique. De plus, P et Z désigneront toujours les nombres de pôles et de zéros à partie réelle positive d'une fonction.

Nous présentons, dans ce qui suit, l'étude de la stabilité par le critère de Nyquist d'une part sur une fonction de transfert équivalente (F.T.E.) que nous définirons, et d'autre part sur une fonction caractéristique (F.C.) image du dénominateur de la fonction de transfert. Nous nous attacherons à dégager les différences de ces deux cas d'étude.

a) Etude de la stabilité par une fonction de transfert équivalente

Une fonction de transfert équivalente (F.T.E.) est une fraction rationnelle (rapport de deux polynômes de la variable p) proportionnelle à la fonction de transfert (F.T.) du système considéré. Notons que la fonction de transfert équivalente peut se déduire de la fonction de transfert par multiplication au numérateur et au dénominateur par un même polynôme M(p) quelconque.

Appliqué à la fonction de transfert équivalente (indice F.T.E) précédemment définie, la relation (III-1) devient :

$$N_{\text{F.T.E}} = P_{\text{F.T.E}} - Z_{\text{F.T.E}} = (P_{\text{F.T.}} + Z_M) - (Z_{\text{F.T.}} + Z_M) = P_{\text{F.T.}} - Z_{\text{F.T.}} \quad (\text{III-3})$$

La stabilité asymptotique étant assurée si et seulement si $P_{F.T.}=0$, le critère de Nyquist devient donc :

$$\text{Stabilité} \Leftrightarrow P_{F.T.}=0 \Leftrightarrow N_{F.T.E} = -Z_{F.T.} \quad (\text{III-4})$$

La stabilité est ainsi testée par la détermination graphique du nombre d'encerclement de l'origine par le tracé de Nyquist de la fonction de transfert équivalente, mais requiert de plus la connaissance du nombre de zéros à partie réelle positive de la fonction de transfert. L'évaluation de ce dernier se fait à partir de la détermination du nombre de zéros à partie réelle positive de la fonction de transfert en boucle ouverte associée que nous allons définir.

La boucle ouverte associée à une fonction de transfert correspond au système dans lequel les contre-réactions des variables de sortie vers les variables d'entrée ont été éliminées, comme présenté sur la Figure III-3.

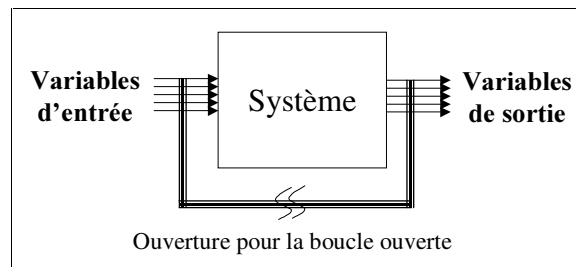


Figure III-3 : Obtention de la boucle ouverte

On démontre aisément que la fonction de transfert résultante, notée F.T.B.O pour fonction de transfert en boucle ouverte, possède les mêmes zéros que la fonction de transfert du système initial (F.T.) :

$$Z_{F.T.} = Z_{F.T.B.O.} \quad (\text{III-5})$$

Enfin, remarquons que la boucle ouverte associée à un système est fonction des variables d'entrée et de sortie considérées et donc de la fonction de transfert du système. Le terme associé fait donc référence à la fonction de transfert du système ou, de manière équivalente, aux variables d'entrée-sortie considérées. Ainsi à un système donné, peut correspondre plusieurs boucles ouvertes associées à des variables d'entrée-sortie différentes.

Ainsi, d'après (III-5), la condition de stabilité (III-4) devient :

$$\text{Stabilité} \Leftrightarrow N_{F.T.E} = -Z_{F.T.B.O.} \quad (\text{III-6})$$

L'avantage de cette nouvelle définition repose sur un théorème démontré par H. Bode [II.21] qui stipule qu'un circuit passif ne comporte aucun zéro à partie réelle positive ($Z_{F.T.B.O.}=0$).

Ainsi, dans le cas où la boucle ouverte associée à un système correspond à un circuit passif, dans lequel aucun sous-circuit ne peut être instable, la stabilité de ce système est facilement vérifiable car la condition de stabilité devient alors :

$$\text{Stabilité} \stackrel{\text{B.O. passive}}{\Leftrightarrow} N_{F.T.E} = 0 \quad (\text{III-7})$$

Dans le cas où la boucle ouverte associée n'est pas passive, la fonction de transfert en boucle ouverte peut comporter des zéros à partie réelle positive. Il est alors impossible de statuer sur la stabilité du système uniquement à partir du nombre d'encerclement de l'origine par le tracé de Nyquist de la fonction de transfert équivalente.

Il est donc important de considérer les variables d'entrée-sortie « sensibles » vis à vis de la stabilité afin que la boucle ouverte associée soit purement passive et qu'il soit ainsi possible de statuer graphiquement sur la stabilité. Nous reviendrons sur ce point au paragraphe III.3.2.c).

Pour conclure, la fonction de transfert équivalente ne semble pas optimale pour l'évaluation de la stabilité puisqu'une condition d'application délicate est à observer (boucle ouverte passive) afin de statuer sur la stabilité du système considéré.

Nous présentons dans le prochain paragraphe une fonction d'étude permettant de rendre compte de la stabilité d'un système tout en diminuant les contraintes.

b) Etude de la stabilité par une fonction caractéristique

Une fonction caractéristique (F.C.) est une fraction rationnelle dont les zéros sont égaux aux pôles de la fonction de transfert du système considéré.

Appliquée à cette fonction caractéristique, la relation (III-1) devient :

$$N_{F.C.} = P_{F.C.} - Z_{F.C.} = P_{F.C.} - P_{F.T.} \quad (\text{III-8})$$

Comme précédemment, le nombre de pôles instables du système ne se déduit pas directement du nombre d'encerclement de l'origine par le lieu de Nyquist. Il est en effet nécessaire, pour ce faire, de connaître le nombre de pôles à partie réelle positive de la fonction caractéristique.

Le concept de la boucle ouverte associée est ici aussi utilisé afin d'outrepasser cette difficulté. En effet, les pôles de la fonction de transfert en boucle ouverte sont égaux aux pôles de la fonction caractéristique. La condition de stabilité devient alors :

$$\text{Stabilité} \Leftrightarrow P_{F.T.} = 0 \Leftrightarrow N_{F.C.} = P_{F.T.B.O.} \quad (\text{III-9})$$

Lorsque la connaissance du nombre de pôles instables de la fonction de transfert en boucle ouverte n'est pas possible, la condition nécessaire et suffisante (III-9) peut être restreinte à une condition suffisante de stabilité. En effet, si la boucle ouverte associée est stable, son nombre de pôle à partie réelle positive est nul ($P_{F.T.B.O.} = 0$), la relation (III-9) devient alors :

$$\text{Stabilité} \stackrel{\text{B.O. stable}}{\Leftrightarrow} N_{F.C.}^{\text{trigo}} = 0 \quad (\text{III-10})$$

Cette condition est moins restrictive que celle de l'expression (III-7), car la boucle ouverte associée n'est pas ici forcément passive.

Historiquement, la première étude géométrique de stabilité a été effectuée par H. Bode à l'aide d'une fonction caractéristique habilement introduite qu'il nomma « Return Difference ».

L'application du critère de stabilité (III-10) a permis aux concepteurs de circuits électroniques basses-fréquences, radio-fréquences et hyperfréquences de vérifier la stabilité des systèmes conçus. Nous traitons dans le prochain paragraphe le cas des circuits hyperfréquences.

III.2.2. Application aux circuits micro-ondes

Toutes les méthodes appliquées aux circuits micro-ondes sont basées sur l'étude de la fonction caractéristique suivante : $1 - \Gamma_1(\omega)\Gamma_2(\omega)$ de point critique l'origine ou, plus souvent, sur la fonction $\Gamma_1(\omega)\Gamma_2(\omega)$ dont le point critique est alors 1. La Figure III-4 présente la définition de $\Gamma_1(\omega)$ et $\Gamma_2(\omega)$ qui sont les coefficients de réflexion des deux sous-circuits déduits du circuit à évaluer par coupure suivant un plan d'étude arbitraire ; ainsi que le schéma bloc modélisant le comportement de ces deux sous-circuits.

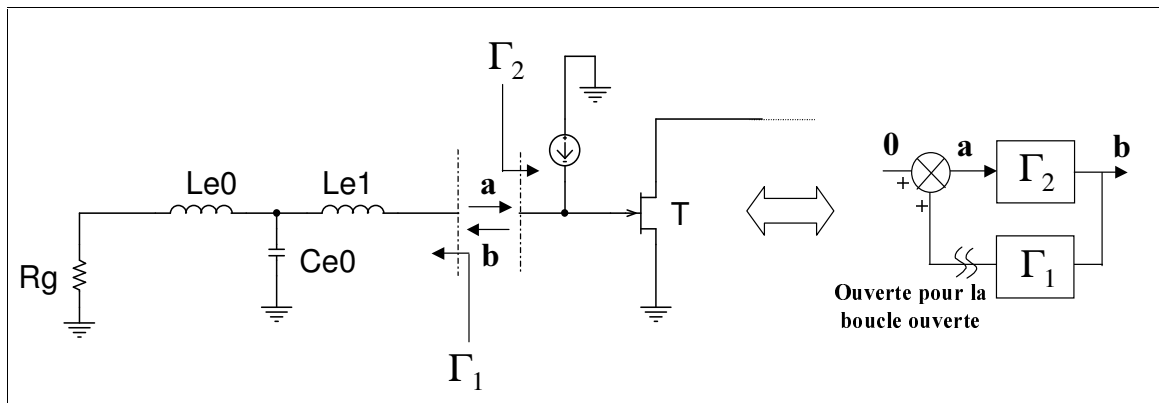


Figure III-4 : Etude de la stabilité de circuit micro-ondes

Reprenant la condition (III-10), la stabilité est donc assurée si (mais non seulement si car la condition est suffisante mais non nécessaire) :

- 1) Le circuit est stable en boucle ouverte : c'est à dire lorsque les deux sous-circuits sont connectés à des charges de coefficient de réflexion nul, éliminant les contre-réactions par réflexion.
- 2) Le lieu de Nyquist de la fonction caractéristique $\Gamma_1(\omega)\Gamma_2(\omega)$ n'entoure pas le point critique 1, assurant $N_{F.C.}=0$.

De façon à simplifier l'étude en évitant le tracé du lieu de Nyquist, la condition 2) peut être restreinte à la condition suivante :

$$2') \quad \forall \omega, |\Gamma_1(\omega)\Gamma_2(\omega)| < 1.$$

L'application de cette condition suffisante de stabilité aux circuits à un seul transistor engendra de nombreuses méthodes devenues classiques telles que celle de Rollett [II.22], de Linvill ou la méthode des cercles de stabilité.

Pour les circuits à un seul transistor, la condition 1) est toujours vérifiée, puisque la boucle ouverte associée s'affranchit de toute contre-réaction sur l'élément actif. Ce n'est plus le cas des circuits à plusieurs transistors dans lesquels la boucle ouverte peut ne pas être stable, empêchant l'application du critère de stabilité évoqué par la relation (III-10).

Le circuit de la Figure III-5 montre un exemple classique [II.23] infirmant les méthodes précédemment évoquées.

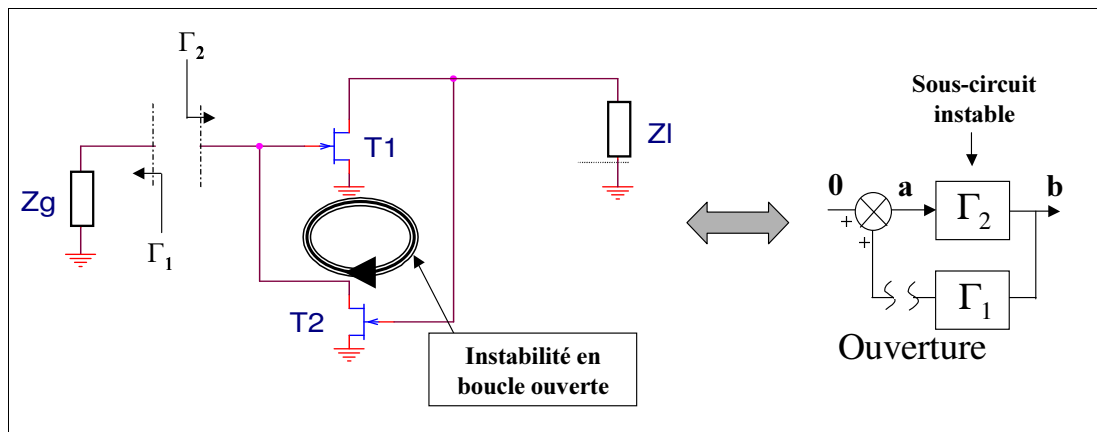


Figure III-5 : Circuit à deux transistors instable en boucle ouverte

La solution pour ce cas de figure fut prévue par H. Bode [II.21] dès 1945, puis formalisée par A. Platzker en 1993 [II.23]. Ce dernier introduisit une fonction caractéristique nommée « Normalized Determinant Function (NDF) » ne possédant aucun pôle à partie réelle positive. Le lieu de Nyquist de ce NDF permet la détermination sans condition de la stabilité du circuit puisque la condition nécessaire et suffisante de stabilité est alors :

$$\text{Stabilité} \Leftrightarrow N_{\text{NDF}} = 0 \quad (\text{III-11})$$

Cette méthode permet donc l'évaluation de la stabilité de circuits à plusieurs transistors mais nécessite, en contre partie, autant de simulations que de transistors dans le circuit. Nous verrons au paragraphe III.3.2.a) une autre méthode d'analyse de circuit multi-transistors.

Les méthodes évoquées précédemment sont des analyses de la stabilité des circuits linéaires. Lorsque les circuits sont non-linéaires, comme c'est le cas pour des mélangeurs, ces études ne peuvent plus être utilisées. Dans le prochain paragraphe, nous allons développer les moyens mis en œuvre pour effectuer l'étude de la stabilité des circuits non-linéaires.

III.3. Stabilité non-linéaire : méthode mise en œuvre

La stabilité non-linéaire est la stabilité de circuits non-linéaires en condition de fonctionnement nominal, c'est à dire soumis à des signaux de forte amplitude.

Les seules méthodes applicables pour l'étude de la stabilité des systèmes non-linéaires sont celles qui se déduisent de la méthode de Lyapunov. D'un énoncé simple, cette méthode est pourtant extrêmement difficile à mettre en œuvre car elle requiert au préalable la définition d'une fonction d'étude bien particulière.

Les méthodes actuelles tentent donc de contourner ces difficultés rédhibitoires.

Nous ferons, tout d'abord, le catalogue critique des méthodes existantes pour l'étude de la stabilité non-linéaire. Puis nous présenterons la méthode que nous avons mise en œuvre et qui nous a permis l'étude de la stabilité de la cellule de mélange.

III.3.1. Méthodes existantes

A notre connaissance, trois méthodes existent pour l'étude de la stabilité de circuits non-linéaires :

- 1) La méthode de la sonde qui consiste à insérer une source autonome dans le circuit, par exemple de tension, dont l'amplitude et la fréquence sont ajustées afin de rechercher l'annulation du courant débité par cette source [II.24]. Si une telle situation se produit, le système est instable et la fréquence d'oscillation correspondra à la fréquence de la source. Malheureusement, il n'existe pas de règle générale pour le positionnement de cette source dans le circuit. De plus, cette méthode pose des problèmes conséquents de convergence car deux variables sont à optimiser.
- 2) La méthode déduite de l'équilibrage harmonique. La fonction caractéristique pour l'application du critère de Nyquist est : $F(\omega) = \det(1 - A(\omega)U)$, où $A(\omega)$ est la matrice représentation du sous-circuit linéaire et U la matrice Jacobienne du sous-circuit non-linéaire. De plus, toutes les variables doivent être observées de façon à ce que le système en boucle ouverte associée soit stable. La condition de stabilité est alors celle de l'expression (III-10). Cette méthode est lourde car elle met en œuvre les matrices imposantes de l'équilibrage harmonique, conduisant à des temps de simulation importants. Enfin, les diagrammes de Nyquist sont souvent inexploitable [II.25].
- 3) Une méthode récente, publié par S. Mons de l'université de Limoges [II.26], qui applique le concept de fonction normalisée du déterminant (NDF) aux matrices de conversion. Cette méthode nécessite, comme celle de Platzker, autant de simulations que de transistors présents dans le circuit et requiert le traitement numérique des matrices de conversions issues de ces simulations sur un logiciel de traitement disjoint du logiciel de simulation électrique.

Quelle que soit la méthode envisagée, elle se résume à la superposition de perturbations de faibles amplitudes au régime de fonctionnement permanent fort signal du circuit et à l'étude de l'impact de ces perturbations sur l'état d'équilibre. La stabilité ainsi

déduite est donc nommée *stabilité locale* car son étude est conduite localement autour d'un point de fonctionnement dynamique. Seule l'étude du diagramme de bifurcation, qui dépasse le cadre de ce manuscrit, donnerait des indications sur la stabilité globale du circuit.

Le prochain paragraphe présente la méthode que nous avons choisie de mettre en œuvre et qui s'appuie, comme celle présenté par S. Mons, sur le formalisme des matrices de conversion.

III.3.2. Méthodes mise en œuvre pour l'étude de la stabilité de mélanges.

L'étude de la stabilité linéaire d'un mélangeur, ne donne aucune information sur la stabilité du circuit considéré. En effet, cette étude suppose que le circuit n'est soumis qu'à des perturbations de faibles amplitudes, et impose donc l'extinction du fort signal de pompe OL. Le circuit est alors dans un état d'équilibre statique différent de l'état d'équilibre dynamique imposé par le signal OL, et la stabilité déduite de l'état statique n'impose aucunement celle de l'état dynamique et vice-versa.

Les circuits mélangeurs nécessitent donc l'étude de la stabilité de l'état d'équilibre dynamique imposé par le signal OL dite stabilité non-linéaire.

Les matrices de conversion sont toutes indiquées pour mener cette étude puisqu'elles modélisent un circuit, soumis à un fort signal de pompe OL, vis à vis de signaux de faibles amplitudes. L'énorme avantage de ce formalisme est de donner une représentation linéaire matricielle du circuit (exprimée au paragraphe II.5.2) offrant la possibilité d'utiliser les méthodes linéaires d'étude de la stabilité.

Les perturbations du système seront donc les variables des matrices de conversion et la stabilité sera étudiée sur une fonction de transfert ou une fonction caractéristique non plus mono-dimensionnelle mais matricielle.

Nous présentons dans un premier temps l'étude de la stabilité de système multi-dimensionnel, puis nous définirons une fonction (F.T.E. ou F.C.) d'étude pour l'application du critère de Nyquist. Un paragraphe traitera d'un point crucial pour l'étude de la stabilité de système portant sur l'observation des variables « sensibles » du système. Enfin, nous conclurons ces paragraphes en résumant la méthode choisie pour l'étude de la stabilité non-linéaire du mélangeur que nous proposons.

a) Stabilité des circuits multi-dimensionnels

Soit $[Q(p)]$ une matrice représentative du fonctionnement du système étudié (ce peut être une matrice de conversion) et $\vec{E}(p)$, $\vec{S}(p)$ les vecteurs représentant les variables d'entrée et de sortie de ce système :

$$\vec{S}(p) = [Q(p)] \cdot \vec{E}(p) \quad (\text{III-12})$$

Soient $\lambda_i(p)$ les valeurs propres de la matrice $Q(p)$ et $\bar{u}_i(p)$ les vecteurs propres qui leur sont associés. Toute perturbation bornée ($|\vec{E}(t)| < A$) peut alors se décomposer comme suit :

$$\vec{E}(p) = \sum \alpha_i \bar{u}_i(p) \quad (\text{III-13})$$

Où α_i sont des constantes (bornées). Comme $\vec{E}(t)$ est bornée, pour tout i les vecteurs $\bar{u}_i(t)$ sont bornés : $|\bar{u}_i(t)| < U$. La sortie se met alors sous la forme :

$$\vec{S}(p) = \sum \alpha_i [Q(p)] \bar{u}_i(p) = \sum \alpha_i \bullet \lambda_i(p) \bullet \bar{u}_i(p) \quad (\text{III-14})$$

Afin d'appliquer la théorie de Lyapunov, revenons dans le domaine temporel :

$$\vec{S}(t) = \sum \alpha_i \bullet \lambda_i(t) * \bar{u}_i(t) \quad (* : \text{produit de convolution}) \quad (\text{III-15})$$

On en déduit que $\vec{S}(t)$ est borné si et seulement si toutes les valeurs propres sont bornées (car α_i et $\bar{u}_i(t)$ sont bornés), amenant aux critères de stabilité suivant :

$$\begin{array}{ccc} \text{Stabilité} & \begin{array}{c} \text{domaine temporel} \\ \Leftrightarrow \end{array} & \forall i : \left(\int_{-\infty}^{+\infty} |\lambda_i(t)| dt \right) < \infty \\ & \begin{array}{c} \text{domaine spectral} \\ \Leftrightarrow \end{array} & \forall i : \text{les pôles de } \lambda_i(p) \text{ sont à} \\ & & \text{partie réelle strictement négative} \end{array} \quad (\text{III-16})$$

Finalement, l'étude de la stabilité d'un système multi-dimensionnel revient à considérer la nature des pôles des valeurs propres $\lambda_i(p)$ de la matrice représentative du système, ces derniers devant être considérées comme des fonctions de transfert équivalentes (F.T.E.).

b) Recherche d'une fonction d'étude de la stabilité

La stabilité peut donc être menée en appliquant le critère généralisé de Nyquist [II.25] aux valeurs propres de la matrice représentant le système. Cette étude est néanmoins impossible car elle nécessite la détermination et l'étude de toutes les valeurs propres (fonctions de la variable complexe p), ce qui en pratique n'est pas réalisable.

Nous avons donc cherché à définir une fonction **mono-dimensionnelle** d'étude de la stabilité du système qui soit facilement accessible par le calcul.

Dans une publication de 1977, A. MacFarlane [II.27] démontra, au prix de nombreuses pages de calculs que :

- Le nombre de pôles instables de toutes les valeurs propres $\lambda_i(p)$ (noté P_{λ_i}) est égal à la somme du nombre de pôles à partie réelle positive du déterminant de $Q(p)$ (noté $P_{\det(Q)}$) et du nombre de zéros à partie réelle positive d'un polynôme $X(p)$ inconnu (noté Z_X) : $P_{\lambda_i} = P_{\det(Q)} + Z_X$.
- Le nombre de zéros à partie réelle positive des valeurs propres $\lambda_i(p)$ (noté Z_{λ_i}) est égal à la somme du nombre de zéros à partie réelle positive du déterminant de $Q(p)$ (noté $Z_{\det(Q)}$) et du nombre de zéros à partie réelle positive du polynôme $X(p)$ précédemment introduit (noté Z_X) : $Z_{\lambda_i} = Z_{\det(Q)} + Z_X$.

Cette démonstration fait du déterminant de la matrice $Q(p)$ une fonction de transfert équivalente pour laquelle l'application du critère de Nyquist entraîne :

$$N_{\det(Q(p))} = P_{\det(Q(p))} - Z_{\det(Q(p))} = (P_{\lambda_i} - Z_X) - (Z_{\lambda_i} - Z_X) = P_{\lambda_i} - Z_{\lambda_i} \quad (\text{III-17})$$

La stabilité est alors assurée si et seulement si (d'après la condition (III-16)) $P_{\lambda_i} = 0$. De plus, comme les fonctions $\lambda_i(p)$ sont du type « fonction de transfert équivalente », leurs zéros sont égaux à ceux obtenus en boucle ouverte. La condition nécessaire et suffisante de stabilité devient donc (de manière similaire au cas des circuits linéaire, cf. (III-6)) :

$$\text{Stabilité} \Leftrightarrow N_{\det(Q(p))} = -Z_{\text{F.T.B.O.}} \quad (\text{III-18})$$

Deux cas se posent alors :

- 1) Si toutes les variables sont observées, la boucle ouverte associée à la matrice $Q(p)$ est alors passive puisque qu'aucun sous-circuit n'est instable. Ceci impose que : $Z_{\text{F.T.B.O.}} = 0$ (cf. paragraphe III.2.1.a). La condition suffisante de stabilité devient alors (de manière similaire au cas des circuits linéaires, cf. (III-7)) :

$$\text{Stabilité} \Leftrightarrow N_{\det(Q(p))} = 0 \quad (\text{III-19})$$

Ce cas de figure est rencontré lors l'étude de la stabilité par le formalisme de l'équilibrage spectral puisque toutes les variables du système sont alors considérées.

- 2) Toutes les variables ne sont pas observées, il est alors impossible de connaître la valeur de $Z_{\text{F.T.B.O.}}$ et donc de conclure sur la stabilité connaissant $N_{\det(Q(p))}$.

Dans notre cas, l'observation de toutes les variables serait contradictoire avec l'objectif de simplicité fixé. Nous nous placerons donc, dans la suite de ce manuscrit, dans le second cas décrit où toutes les variables du système ne sont pas observées.

Afin de pouvoir statuer sur la stabilité du système dans ce cas, il est nécessaire d'introduire le calcul du déterminant de $Q(p)$ lorsque tous les éléments actifs, à l'origine des instabilités éventuelles, sont éteints ($g_m = 0$). Le circuit ainsi obtenu possède les propriétés suivantes :

- Il est stable car toutes les sources potentielles d'instabilité sont éteintes. La fonction de transfert équivalente $\det(Q_0)$ ne possède donc pas de pôles à partie réelle positive : $P_{\det(Q_0)}=0$
- Il est une boucle ouverte du circuit initial car sont éliminées les réactions de V_{gs} (ou V_{be}) vers I_{ds} (ou I_{ce}). Sa fonction de transfert $\det(Q_0)$ à un nombre de zéros à partie réelle positive égal à $Z_{F.T.B.O.}$

Ces remarques font du rapport $\det(Q)/\det(Q_0)$ une fonction de transfert équivalente ne possédant pas de zéros à partie réelle positive (compensation de $Z_{F.T.B.O.}$ entre $\det(Q)$ et $\det(Q_0)$) et dont les pôles à partie réelle positive sont les pôles instables du système (pas de compensation). La condition nécessaire et suffisante de stabilité est donc :

$$\text{Stabilité} \Leftrightarrow N_{\frac{\det(Q(p))}{\det(Q_0(p))}} = 0 \quad (\text{III-20})$$

Cette fonction de transfert équivalente est idéale pour l'étude de la stabilité car elle aboutit à une condition nécessaire et suffisante simple et sans condition, portant uniquement sur le nombre d'encerclements du point critique par le lieu de Nyquist résultant.

De même, le critère de stabilité (III-20) n'impose aucune condition sur les variables observées au travers de la matrice $Q(p)$. Cependant, quelques règles intuitives tendent à imposer l'observation de certaines variables particulières. Le paragraphe suivant apporte quelques règles pour le choix de ces variables.

c) Variables à observer

L'étude faite au paragraphe précédent permet la détermination du nombre de pôles instables du système à partir de la connaissance des pôles instables de toutes les valeurs propres $\lambda_i(p)$ de la matrice $Q(p)$ représentant ce système.

Il est évident que cette stabilité ne porte exclusivement que sur les variables d'entrée et de sortie $\vec{E}(p)$ et $\vec{S}(p)$ définies au paragraphe III.3.2 a) par la relation (III-12). Si dans le système, il existe d'autres variables qui ne se déduisent pas des variables contenues dans les vecteurs $\vec{E}(p)$ et $\vec{S}(p)$, alors aucune conclusion quant à la stabilité du système vis à vis de ces variables ne peut être déduite.

La Figure III-6 donne un exemple démonstratif d'une telle situation. Le système est modélisé par une matrice (par exemple impédance) définie dans les accès 1 et 2. Si l'impédance Z est passive, la stabilité du système ainsi modélisé est assurée alors que le système contient un système instable : un oscillateur. Cette situation se produit car aucune variable de l'oscillateur n'est observée.

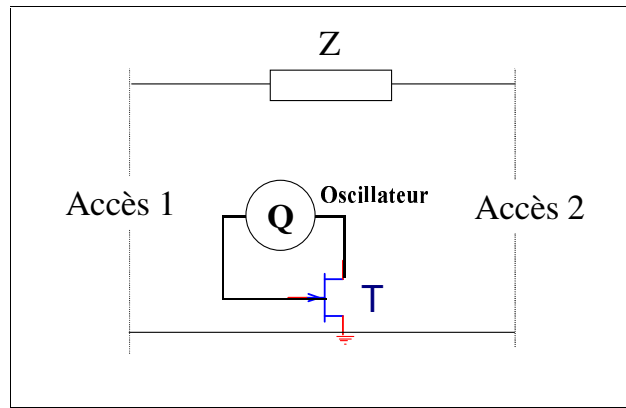


Figure III-6 : Défaut d'observabilité

Afin d'éviter ce type de défaut d'observabilité, il convient d'observer les variables *sensibles* vis à vis de la stabilité. On entend par variables *sensibles*, celles qui précèdent (ou qui suivent) les éléments actifs uniques causes d'instabilité dans les circuits électroniques. Dans le cas de la Figure III-6, il aurait été ainsi nécessaire d'ajouter la tension v_{gs} du transistor dans l'ensemble des variables d'état.

d) Conclusion

L'étude de la stabilité que nous mènerons sur les mélangeurs sera donc basée sur l'application du critère de Nyquist à la fonction de transfert équivalente suivante :

$$F.T.E. = \frac{\det(Q(p))}{\det(Q_0(p))} \quad (III-21)$$

Où $Q(p)$ sera une matrice de conversion (cf paragraphe II de ce chapitre) construite à partir de variables qui devront respecter les règles d'observabilité précédemment définies.

La stabilité étant assurée si le lieu de Nyquist de cette F.T.E. n'entoure pas le point critique et ceci sans aucune autre condition.

III.4. Application : stabilisation de la cellule de mélange

Nous allons utiliser la méthode précédemment définie pour étudier la stabilité de notre cellule originale de mélange.

Nous décrirons dans un premier temps, l'application de la méthode d'étude de la stabilité non-linéaire à cette cellule puis nous présenterons l'optimisation de l'inductance L_{stab} dont le rôle est d'assurer la stabilité de celle-ci.

*

III.4.1. Etude de la stabilité non-linéaire de mélangeur

La matrice considérée est la matrice impédance de conversion Z_{cv}^2 définie au paragraphe II.5.2.b. L'avantage de cette méthode et de ce choix de matrice est de ne nécessiter aucune mise en œuvre supplémentaire par rapport au formalisme utilisé au paragraphe II et duquel découle directement l'évaluation de Z_{cv}^2 .

Seuls sont à prévoir dans l'algorithme de calcul l'extinction de tous les éléments actifs qui se réalise aisément par l'annulation de toutes les matrices transconductances : $[G_m]=[0]$ et le calcul de la fonction de transfert équivalente définie par :

$$F.T.E. = \frac{\det(Z_{cv}^2(p))}{\det\left(Z_{cv}^2|_{[G_m]=0}(p)\right)} \quad (\text{III-22})$$

Le schéma équivalent paramétrique pour cette étude de stabilité est alors le suivant :

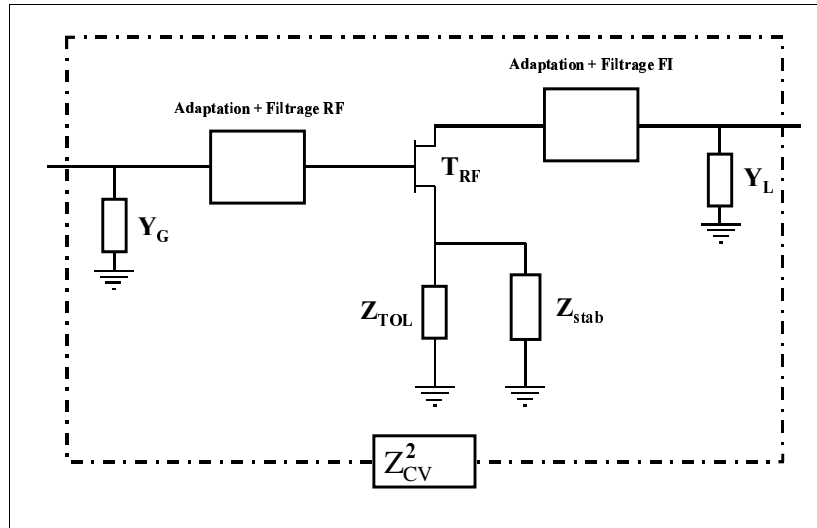


Figure III-7 : Schéma paramétrique pour l'étude de la stabilité

Dans le schéma de la Figure III-7, le transistor T_{OL} a été remplacé par son impédance de sortie. En effet, étant donnée l'extinction du signal OL (voir la construction du schéma paramétrique au paragraphe II), l'entrée de ce transistor n'est soumise à aucun signal, ce qui légitime cette modélisation.

Ce schéma permet de constater la présence d'un seul transistor qui est le siège de contre réaction potentiellement déstabilisante. Les variables d'états calculées aux deux accès du circuit suffisent donc pour assurer une observabilité suffisante du système.

En théorie, l'observation soit de l'entrée, soit de la sortie du circuit suffirait, l'observation des deux accès étant redondante. Cette redondance n'a d'effet que sur la dimension des matrices traitées qui est alors doublée. En pratique, le logiciel de traitement numérique que nous avons utilisé est suffisamment performant pour le traitement de Z_{cv}^2 . Nous n'avons donc pas cherché à éliminer cette redondance.

La fonction de transfert équivalente résultante possède deux propriétés qu'il est nécessaire d'aborder car elles simplifient grandement l'étude de Nyquist :

- $F.T.E.(-\omega) = F.T.E.^*(\omega)$: Le tracé de Nyquist pourra donc être limité aux fréquences positives, le lieu complet étant obtenu par duplication symétrique par rapport à l'axe des réels.
- $F.T.E.(\omega) \xrightarrow{\omega \rightarrow \infty} 1$: le lieu de Nyquist tend vers 1 à fréquences élevées. En effet, à très hautes fréquences le gain tous les éléments tend vers zéros, le déterminant $\det(Z_{cv}^2(i\omega))$ tend alors vers celui qui est calculé pour $[G_m]=0$. Cette propriété permet de vérifier que le balayage fréquentiel est suffisant en constatant la proximité du lieu de Nyquist à hautes fréquences du point (1,0).

Nous sommes alors en mesure de tracer le lieu de Nyquist de la F.T.E. de notre cellule de mélange, d'en déduire le nombre d'encerclement du point critique et donc de statuer sur la stabilité du circuit.

III.4.2. Optimisation de Lstab

L'objet de ce paragraphe est l'optimisation de la valeur de l'inductance L_{stab} définie sur la Figure II-9 afin d'assurer la stabilité inconditionnelle de la cellule de mélange envisagée.

Nous avons, dans un premier temps, tracé le lieu de Nyquist pour le circuit sans filtre stabilisant. La Figure III-8 représente un tel tracé pour des fréquences allant de 4GHz à 500GHz, le sens de la flèche indiquant un parcours suivant les fréquences croissantes. Notons que le lieu de Nyquist tend vers le point (1,0) à très hautes fréquences.

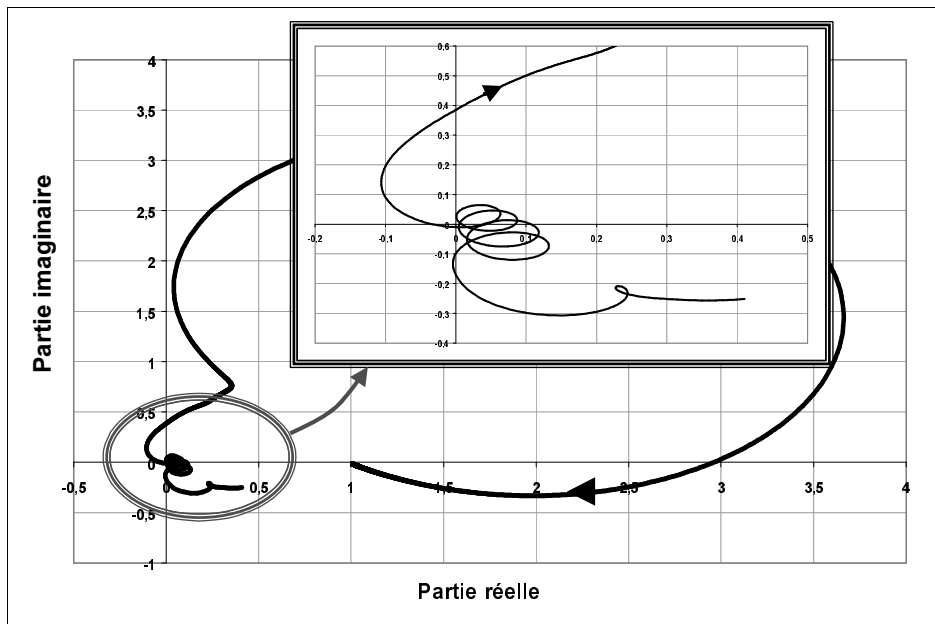


Figure III-8 : Lieu de Nyquist du circuit sans stabilisation

Le lieu de Nyquist encercle le point critique situé à l'origine et met ainsi en évidence la présence d'un fonctionnement instable de la cellule.

Déjà évoquée au chapitre 1, la justification de cette instabilité est la présence d'un réseau de contre-réaction (dû à l'impédance de sortie du transistor T_{OL}) sur la source du transistor de mélange T_{RF} . Ce réseau, à caractère capacitif, entraîne un comportement fortement instable bien connu des concepteurs d'oscillateur qui utilisent cette configuration.

Nous avons déjà évoqué dans ce chapitre ainsi qu'au premier la solution qui consiste à rajouter une impédance résonante en charge de présenter une faible impédance aux fréquences fonctionnelles du circuit : f_{RF} et f_{FI} . Rappelons que cette impédance doit garantir, certes un fonctionnement stable, mais aussi un gain de conversion suffisant, car les contres réactions sur source peuvent diminuer dramatiquement le gain des montages. C'est la raison pour laquelle nous avons fixé la valeur de la fréquence de résonance au milieu des fréquences fonctionnelles f_{RF} et f_{FI} .

L'optimisation de la valeur de l'inductance L_{stab} a été conduite afin d'éloigner autant que faire ce peut le lieu de Nyquist du point critique. La Figure III-9 présente la distance minimale D du lieu de Nyquist par rapport au point critique en fonction de la valeur de L_{stab} . La valeur de la fréquence de résonance du réseau a été maintenue constante et la plage de variation de L_{stab} a été fixée afin de garantir l'intégrabilité de cette inductance et de la capacité qui lui est associée.

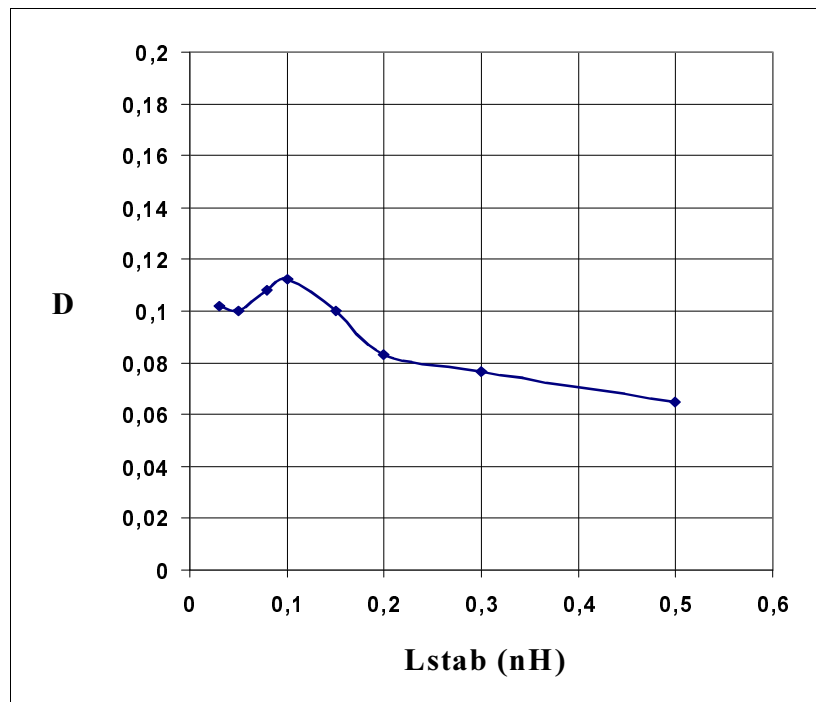


Figure III-9 : Optimisation de L_{stab}

L'optimum de la valeur de l'inductance L_{stab} se situe au alentour de 0.1nH, valeur pour laquelle l'éloignement du lieu de Nyquist du point critique est maximal prévenant ainsi tout risque d'oscillation en cas de dispersion technologique des éléments. Dans ces conditions, l'impédance globale Z_{stab} (voir la Figure III-7) est totalement définie puisque la capacité C_{stab} (voir la Figure II-9) est calculée de manière à résonner avec l'inductance L_{stab} à la fréquence $(f_{RF}+f_{FI})/2$.

La Figure III-10 superpose les lieux de Nyquist correspondant à la cellule sans stabilisation et à celle munie de l'impédance Z_{stab} optimale. Ce tracé confirme la stabilité de notre cellule originale de mélange optimisée ainsi que la nécessité de l'impédance stabilisante.

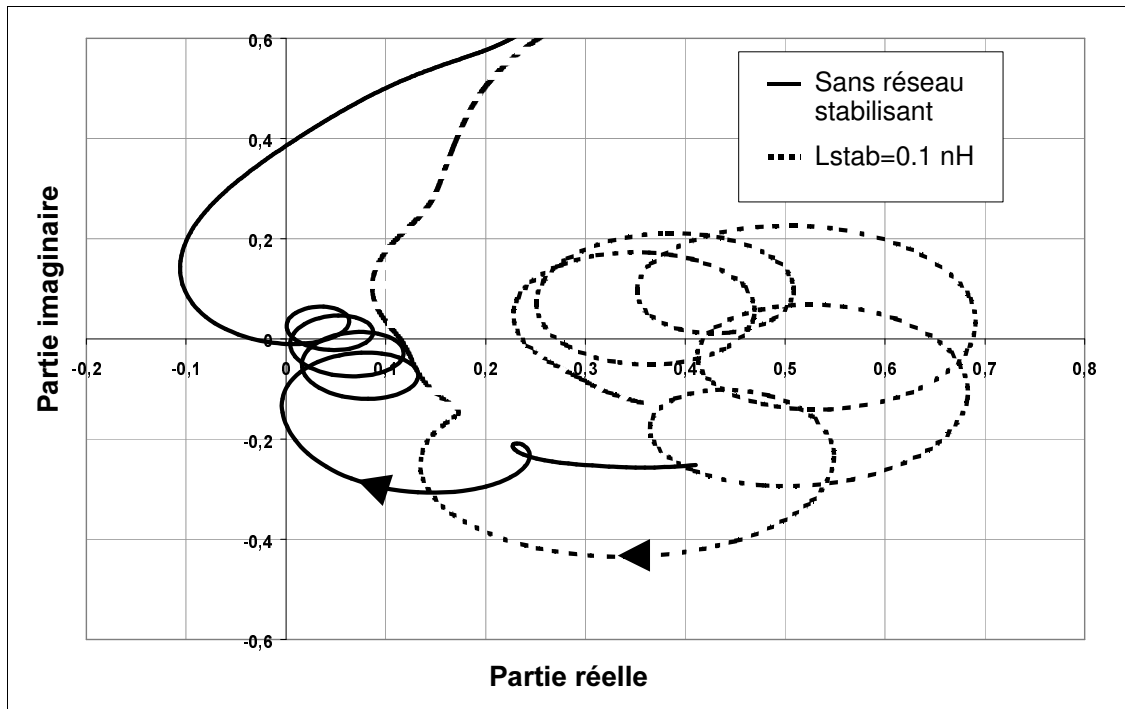


Figure III-10 : Lieu de Nyquist du circuit optimisé

Enfin, les valeurs des éléments du réseau de stabilisation étant fixées, nous avons mené une ultime vérification de la stabilité du circuit en faisant varier la puissance du signal OL. Cette variation s'est accompagnée d'une faible variation du lieu de Nyquist au voisinage du point critique, ce qui nous a permis de conclure à la stabilité du circuit pour les valeurs nominales de puissances P_{OL} .

IV. Conclusion

Nous avons dans ce chapitre présenté l'optimisation complète d'une cellule originale de mélange. Nous avons, dans un premier temps, optimisé la taille et les polarisations des transistors. Nous avons, ensuite, défini les filtres et les adaptations nécessaires et suffisantes pour assurer un fonctionnement correct à cette cellule. Nous avons de plus, grâce au formalisme des matrices de conversion optimisé les différents éléments passifs du circuit. Enfin, nous avons mené la délicate étude de la stabilité non-linéaire, qui nous a permis l'optimisation des éléments assurant la stabilité inconditionnelle de notre cellule originale de mélange.

Le développement de nombreux circuits non-linéaires originaux est souvent stoppé par leurs instabilités potentielles. Difficile à mettre en œuvre, les méthodes d'étude de la stabilité non-linéaire restent marginales et aucun simulateur électrique commercial n'intègre une telle étude. Ce défaut a entraîné l'oubli de certaines configurations pourtant performantes mais

potentiellement instables et pour lesquelles aucune vérification de stabilité à priori n'était possible. Parmi ces structures, les mélangeurs sur source en sont un exemple important. La méthode d'étude de la stabilité que nous avons présentée permet le franchissement simple de cet obstacle et a permis l'optimisation d'une cellule originale de mélange. De cette étude, ont été déduits les éléments optimaux garantissant la stabilité inconditionnelle de cette cellule.

Tous les éléments du circuit étant optimisés, nous avons été en mesure d'entreprendre l'intégration en technologie monolithique sur Arséniure de Gallium ainsi que sur Silicium de cette cellule de mélange. Ce travail fait l'objet du chapitre suivant.

REFERENCES BIBLIOGRAPHIQUES DU CHAPITRE 2

- [II.1] M. Madihian, L. Desclos, K. Maruhashi, K. Onda, M. Kuzuhara, « A K-Band CPW upconverter utilizing a source mixing concept », IEEE MTT-S Digest 1995, pp. 127-130.
- [II.2] : S.A. Maas, « The RF and microwave circuit design cookbook », Artech House, 1998.
- [II.3] : J. Reina-Tosina, C. Crespo, J.I. Alonso, F. Pèrez, « GaAs MMIC mixer based on the Gilbert cell with HEMTs biased on the subthreshold region », Microwave and optical technology letters, vol. 28, n°4, February 20 2001, pp. 241-244.
- [II.4] : S.A. Maas, « Microwave mixers », Artech House, 2nd edition, 1993.
- [II.5] : K. Kawakami, N. Uehara, K. Matsuo, « A high-gain 50-GHz-band monolithic balanced gate mixer with an external IF balun », IEEE-MTT, vol. 46, n° 6, June 1998, pp. 829-833.
- [II.6] : R. Pucel, D. Massé, R. Bera, « Performance of GaAs MESFET mixers at X band », IEEE MTT, vol. 24, n° 6, June 1976, pp. 351-360.
- [II.7] : S.A. Maas, « Design and performances of a 45-GHz HEMT mixer », IEEE MTT, vol. 34, n° 7, July 1986, pp. 799-803.
- [II.8] : G. Begmann, A. Jacob, « Conversion gain of MESFET drain mixers », Electronics letters, vol. 15, n° 18, 30th August 1979, pp. 567-568.
- [II.9] : S.A. Maas, « A GaAs MESFET balanced mixer with very low intermodulation », IEEE MTT-S Digest, 1987, pp. 895-898.
- [II.10] : F. Danneville, B. Tamen, G. Dambrine, A. Cappy « Design of ultra low noise mixer using an analytical formulae of the noise figure », IEEE MTT-S Digest, 1999, pp. 117-120.
- [II.11] : J.C. Cayrou, « Modélisation, conception et caractérisation de convertisseurs de fréquence micro-ondes à transistor à effet de champ à grille schottky sur aséniure de Gallium », Thèse de doctorat de l'Université Paul Sabatier de Toulouse, Juillet 1993.
- [II.12] : P. Penfield, « Circuit theory of periodically driven nonlinear systems », Proceedings of the IEEE, vol. 54, n°2, February 1966, pp. 266-280.

- [II.13] : M. Gayral, « Contribution à la simulation des circuits non-linéaires micro-ondes par la méthode de l'équilibrage harmonique et spectral », Thèse de l'université de Limoges, 18 décembre 1987.
- [II.14] : S.A. Maas, « Nonlinear microwave circuits », Artech House, 1988.
- [II.15] : M. Prigent, « Contribution à l'étude de la conversion de fréquence dans les circuits non-linéaires : application à la C.A.O. d'oscillateur à bruit de phase minimum », Thèse de l'université de Limoges, 23 septembre 1987.
- [II.16] : ED02AH Design Manual V2.0CD, Philips Microwave Limeil, April 1999.
- [II.17] : M. Abdellaoui, « Utilisation de la matrice de conversion pour l'étude de dispositif non linéaire à transposition de fréquences- Application au mélangeur », Thèse de l'ENSEA, 14 mars 1991.
- [II.18] : **D. Dubuc**, T. Parra et J. Graffeuil, « Original Topology of GaAs-PHEMT Mixer », 30th European Microwave Conference, September 2000.
- [II.19] R.I. Jury, « Remembering four stability theory pioneers of nineteenth century », IEEE transaction on circuits and systems I : fundamental theory and applications, vol. 43, n°10, October 1996, pp. 821-823.
- [II.20] : H. Nyquist, « Regeneration theory », Bell Systems Technical Journal 1932, vol.11, pp 126-147.
- [II.21] : H.W. Bode, « Network analysis and feedback amplifier design », D. Van Nostrand, 1945.
- [II.22] : J.M. Rollet, « Stability and power-gain invariants of linear twoports », IRE on transactions on circuit theory, March 1962, pp 29-32.
- [II.23] : A. Platzker, W. Struble, K. Hetzler, « Instabilities diagnosis and the role of K in microwave circuits », IEEE MTT-S Digest, 1993, pp1185-14188.
- [II.24] : J.M. Collantes, A. Suarez, « Period-doubling analysis and chaos detection using commercial harmonic balance simulators », IEEE MTT, vol. 49, n°4, april 2000, pp. 574-581.
- [II.25] : S. Mons, « Nouvelles méthodes d'analyse de stabilité intégrées à la CAO des circuits monolithiques micro-ondes non-linéaires », thèse de l'université de Limoges, 1999.
- [II.26] : S. Mons, J.C. Nallatamby, R. Quéré, P. Savary, J. Obregon, « A unified approach for linear and nonlinear stability analysis of microwave circuits using commercially available tools », IEEE-MTT, vol. 47, n°12, December 1999, pp. 2403-2409.

- [II.27] : A.G.J. MacFarlane, I. Postlethwaite «The generalized Nyquist stability criterion and multivariable root loci », International Journal of Control, vol.25, n°1, 1977, pp. 81-127.

Chapitre 3 : Intégration monolithique de convertisseurs de fréquence

CHAPITRE 3 PARTIE A

Mélangeur double équilibré
sur GaAs en bande Ku

Le contexte actuel des conceptions de circuits hyperfréquences tend vers l'intégration d'un maximum de fonctions au sein d'un même circuit monolithique. Ce concept de puce multifonction [III.1] est en adéquation avec la montée en fréquence, en complexité et en performance des systèmes de communication. D'une part, il élimine les interconnexions parasites de l'hybridation, limitatives de bien des caractéristiques et d'autre part, il optimise les degrés d'intégration diminuant à la fois les coûts et la masse des systèmes.

Nous avons donc envisagé l'intégration de notre cellule originale de mélange au sein d'un système de conversion de fréquence doublement équilibré constitué de deux mélangeurs simplement équilibrés (quatre mélangeurs simples) et de trois coupleurs de signaux. La figure 1 donne le synoptique du système de conversion complet.

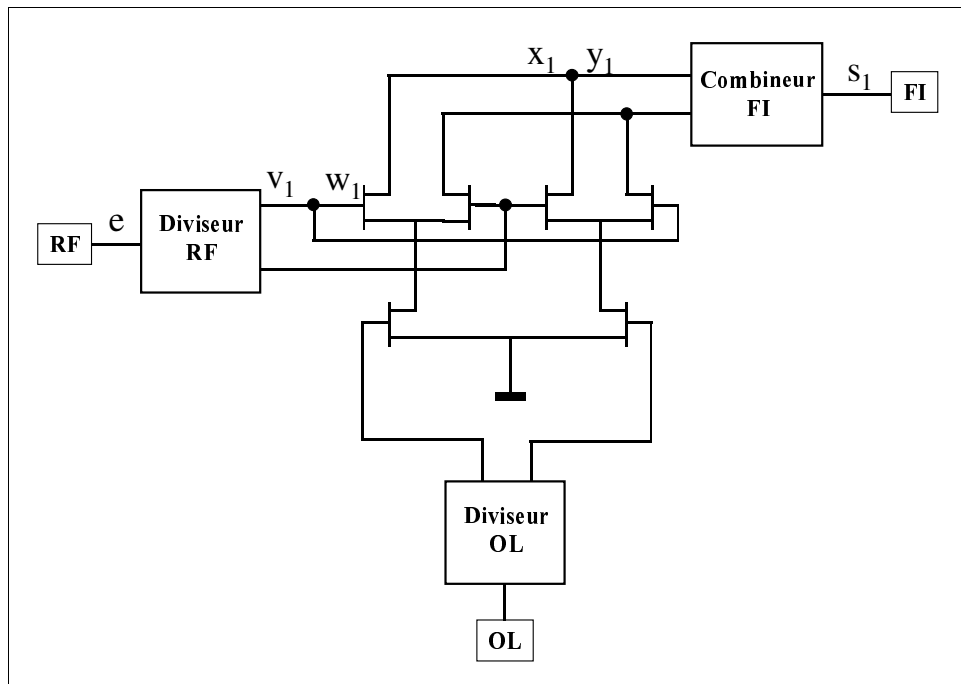


Figure 1 : Synoptique du mélangeur double équilibré

Cette partie de chapitre traite de l'intégration monolithique d'un mélangeur double équilibré dans une technologie monolithique sur Arséniure de Gallium. Nous présenterons, tout d'abord, une étude système justifiant les choix des divers éléments constituant le mélangeur global. Puis, nous détaillerons l'étude et l'optimisation des coupleurs actifs. Enfin, nous présenterons l'intégration de la fonction complète ainsi que les résultats de simulation et de caractérisation.

I. Choix structurel

Ce paragraphe a pour but, s'appuyant sur un cahier des charges précis, de guider le choix, via une étude système, des différentes cellules de base constituant le mélangeur doublement équilibré envisagé.

I.1. Cahier des charges

Pour cette première intégration dans la technologie monolithique à base de PHEMT sur Arséniure de Gallium du fondeur PML/OMMIC, nous avons considéré le cahier des charges typique d'un répéteur spatial en bande Ku [III.2]. Ci-dessous est rappelé l'essentiel des caractéristiques que doit présenter le mélangeur élément central du système répéteur.

- Fréquences :
 - Fr_f : 12.75 - 14.80 GHz
 - F_{fi} : 10.70 - 12.75 GHz
 - F_{ol} : 1.25 - 3.85 GHz
- Puissances : Pol = -20 dBm
Pr_f = -18 dBm
- Gain : > 20 dB
 - Largeur canal = 140 MHz
 - variation : 0.3 dB/250 MHz(Fr_f)
0.3 dB/dBm(Pol)
- Bruit : NF < 5 dB
- Linéarité :
 - IP₃ > 17 dBm. C/I₃ > 30 dB avec P_c = 2 dBm.
- Réjection des raies parasites :
 - Pr_f = -18 dBm et Pol = -20 dBm.
 - mFr_f +/- nF_{ol} < -46 dBc dans le canal.
< -56 dBc hors canal.
 - nF_{ol} < -70 dBc dans le canal.
< -58 dBc hors canal.
- Pertes par réflexion : entrée RF > 15 dB
entrée OL > 15 dB
sortie FI > 15 dB
- Consommation : P < 600 mW , Valim : -5v à +5v

I.2. Etude système

Cette étude est menée en considérant la chaîne de transmission, présentée en Figure I-1, et qui donne une représentation fonctionnelle du mélangeur présenté sur la figure 1 vis à vis des signaux RF et FI.

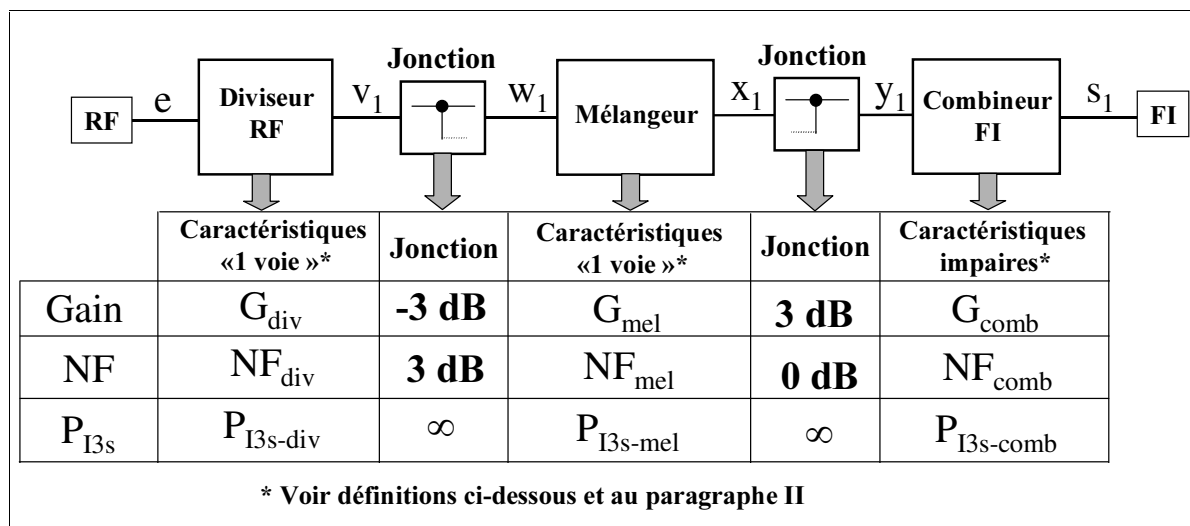


Figure I-1 : Représentation fonctionnelle

L'équivalence de ces deux synoptiques est assurée si les étages diviseur et mélangeur sont associés à leurs caractéristiques « 1 voie » pour lesquelles les définitions sont relatives à une sortie par rapport à une entrée. Le combineur devant être, quant à lui, défini par ses caractéristiques impaires pour lesquelles sont considérées la sortie par rapport à la moitié de la différence des deux entrées.

L'intérêt de l'étude système est d'orienter le choix de chacun des éléments afin que le système réponde au mieux au cahier des charges. Il apparaîtra, en effet, que les composants, optimaux lorsque considérés seuls, ne constituent pas forcément le meilleur choix lorsqu'ils sont insérés dans un système plus complet dans lequel leurs interactions influencent grandement les performances globales [III.3].

Cette affirmation trouvera sa première confirmation lors du choix du type de coupleurs. Pris séparément, malgré les pertes qu'ils présentent, les coupleurs passifs semblent mieux adaptés que leurs homologues actifs car ils assurent les meilleures caractéristiques en linéarité. Pourtant, l'étude du système complet, présentée au Tableau I-1 démontre que les coupleurs actifs amènent de bien meilleures caractéristiques globales, même en terme de linéarité.

	Diviseur RF		Mélangeur	Combineur FI		Système global	
	actif	passifs		actif	passifs	actif	passifs
Gain (dB)	10	-3	5	10	3	25	5
NF (dB)	4	3	6	6	0	5	12
P_{I3s} (dBm)	10	∞	7	20	∞	16	13
$Q^*(dB)$	16	∞	6	24	∞	36	6

* Q : coefficient usuel d'aptitude à l'intégration = $Gain(dB) - NF(dB) + P_{I3s}(dBm)$

Tableau I-1 : Tableau comparatif permettant le choix des coupleurs

Dans ce tableau, ainsi que dans tous les suivants, n'apparaîtront plus les deux jonctions passives qui seront néanmoins toujours prises en considération.

De même, l'étude du choix du type de mélangeur, présenté au Tableau I-2, montre qu'un mélangeur actif conduit à de meilleures performances du système global, même en terme de linéarité, par rapport à un mélangeur passif.

	Diviseur RF	Mélangeur		Combineur FI	Système global	
		actif	passifs		actif	passifs
Gain (dB)	10	5	-7	10	25	13
NF (dB)	4	6	11	6	5	8
P _{13s} (dBm)	10	7	6	20	16	11
Q(dB)	16	6	-12	24	36	16

Tableau I-2 : Tableau comparatif permettant le choix du mélangeur

D'une manière générale, le premier étage d'une chaîne de transmission imposera le facteur de bruit du système global tandis que le dernier élément imposera sa linéarité à condition que tous les éléments constitutifs aient un gain suffisant.

Ce constat motiva, d'une part l'étude faite sur la cellule de mélange à fort gain de conversion que nous avons présentée au second chapitre de ce manuscrit, et d'autre part la conception de coupleurs actifs que nous allons présenter au paragraphe II.

II. Etude des coupleurs

Un coupleur est un octopôle linéaire comportant deux entrées et deux sorties, comme présenté sur la Figure II-1.

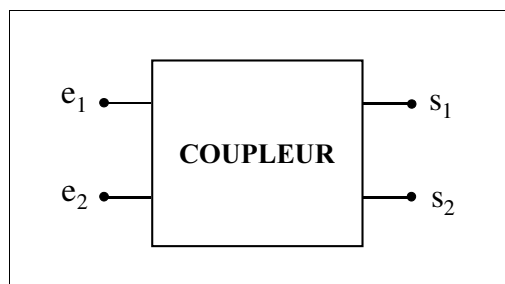


Figure II-1 : Schéma fonctionnel d'un coupleur

D'une manière générale, cet élément est modélisé, en régime petit signal, par une matrice 2x2 liant les quantités de sortie aux quantités d'entrée.

Deux autres représentations, extrêmement utiles et équivalentes, existent afin d'établir les relations d'entrée-sortie de cet octopôle :

1. La représentation par les modes pair (p) et impair (i) :

$$\begin{array}{c|c} e_1 = e_p + e_i & e_p = \frac{e_1 + e_2}{2} \\ \hline e_2 = e_p - e_i & e_i = \frac{e_1 - e_2}{2} \end{array} \quad \begin{array}{c} G_p = \frac{s_p}{e_p} \\ \longleftrightarrow \\ G_i = \frac{s_i}{e_i} \end{array} \quad \begin{array}{c|c} s_1 = s_p + s_i & s_p = \frac{s_1 + s_2}{2} \\ \hline s_2 = s_p - s_i & s_i = \frac{s_1 - s_2}{2} \end{array} \quad (\text{II-1})$$

2. La représentation par le mode différentiel (d) et le mode commun (mc) :

$$\begin{array}{c|c} e_1 = e_d + 2e_{mc} & e_d = e_1 - e_2 \\ \hline e_2 = e_d - 2e_{mc} & e_{mc} = \frac{e_1 + e_2}{2} \end{array} \quad \begin{array}{c} G_d = \frac{s_d}{e_d} \\ \longleftrightarrow \\ G_{mc} = \frac{s_{mc}}{e_{mc}} \end{array} \quad \begin{array}{c|c} s_1 = s_d + 2s_{mc} & s_d = s_1 - s_2 \\ \hline s_2 = s_d - 2s_{mc} & s_{mc} = \frac{s_1 + s_2}{2} \end{array} \quad (\text{II-2})$$

Quelle que soit la représentation, les coupleurs idéaux ont un gain en mode commun (ou un gain en mode pair) nul et peuvent donc être utilisés :

1. Soit en diviseur 180° de puissance pour lequel une seule entrée est utilisée, l'autre étant à zéro. Les sorties sont alors en opposition de phase :

$$\begin{array}{c|c} e_1 = e & \text{Diviseur} \\ e_2 = 0 & \text{180}^\circ \end{array} \quad \begin{array}{c} s_1 = \frac{G_i}{2} e = G_d e \\ s_2 = -\frac{G_i}{2} e = -G_d e \end{array} \quad (\text{II-3})$$

2. Soit en combineur 180° de puissance pour lequel les deux signaux à combiner sont appliqués aux deux entrées tandis qu'une sortie est utile et proportionnelle à la différence des deux entrées :

$$\begin{array}{c|c} e_1 & \text{Combineur} \\ e_2 & \text{180}^\circ \end{array} \quad \begin{array}{c} s_1 = \frac{G_i}{2} (e_1 - e_2) = G_d (e_1 - e_2) \\ s_2 = -\frac{G_i}{2} (e_1 - e_2) = -G_d (e_1 - e_2) \end{array} \quad (\text{II-4})$$

Quelle que soit la topologie utilisée, le gain différentiel (ou le gain du mode impair qui en est le double) correspond au gain utile du coupleur dont résulte la configuration.

Nous allons, dans un premier temps, définir l'une des caractéristiques essentielles des coupleurs pour l'utilisation dans les cellules équilibrées. Puis nous présenterons la topologie active retenue ainsi que son optimisation analytique.

II.1. Critères de sélection

Comme pour les amplificateurs, les caractéristiques principales des coupleurs de signaux sont le gain différentiel, le facteur de bruit, la linéarité, la surface occupée, ...

De plus, le taux de réjection du mode commun (TRMC, ou Common Mode Rejection Rate en anglo-saxon) est une caractéristique, spécifique aux coupleurs, qui est d'une importance capitale pour les cellules équilibrées. Il est défini par :

$$\text{TRMC(dB)} = -20 \bullet \log \left(\frac{G_{mc}}{G_d} \right) \quad (\text{II-5})$$

Cette caractéristique, qui mesure la qualité du coupleur, intervient en effet, dans les performances globales de structures équilibrées.

Un exemple est donné sur la Figure II-2 qui illustre la combinaison des signaux de sortie des deux mélangeurs dans le cas d'une structure simplement équilibré.

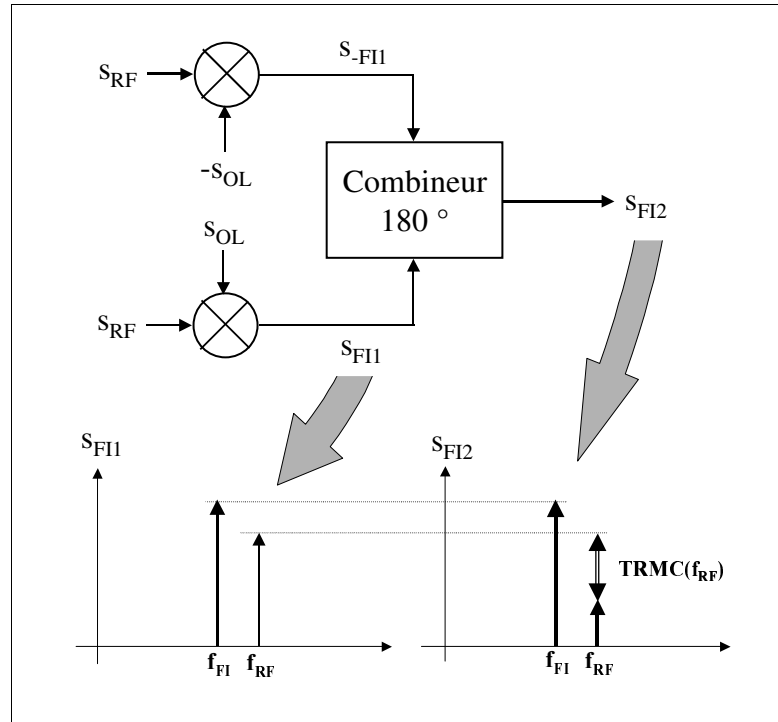


Figure II-2 : impact du TRMC sur la réjection de raie

Pour simplifier, nous n'avons considéré en sortie des mélangeurs que les raies spectrales aux fréquences \$f_{RF}\$ et \$f_{FI}\$. De plus, le gain différentiel a été choisi de -3dB afin que les résultats soient directement comparables (se reporter aux spectres de la Figure II-2).

L'interprétation du TRMC apparaît clairement en comparant les spectres du signal en sortie des mélangeurs simples et du signal de sortie de la cellule simple équilibrée présentés sur la Figure II-2. La raie parasite à la fréquence \$f_{RF}\$ est en effet atténuée de la valeur du TRMC calculée à cette fréquence lors de l'équilibrage de la structure [III.4].

Le taux de réjection du mode commun sera donc choisi le plus grand possible afin d'améliorer la réjection des raies parasites.

Souvent, sont cités les écarts de gain et/ou de phase existant entre les deux voies. Une relation existe liant le TRMC à ces caractéristiques :

$$\text{TRMC(dB)} = -10 \bullet \log \left(2 \bullet \frac{1 - \frac{2\eta}{1+\eta^2} \cos(\Delta\psi)}{1 + \frac{2\eta}{1+\eta^2} \cos(\Delta\psi)} \right), \text{ avec } \eta = 1 + \frac{\Delta G}{G} \quad (\text{II-6})$$

Où $\Delta\psi$ et ΔG sont respectivement l'écart de phase et l'écart d'amplitude des deux sorties, tandis que G représente le gain petit signal d'une des deux sorties prise pour référence sur l'entrée. La Figure II-3 permet, à partir de ces deux caractéristiques, de déduire la valeur du TRMC et donc la valeur de l'atténuation de la raie à la fréquence f_{RF} lors de l'équilibrage de la structure.

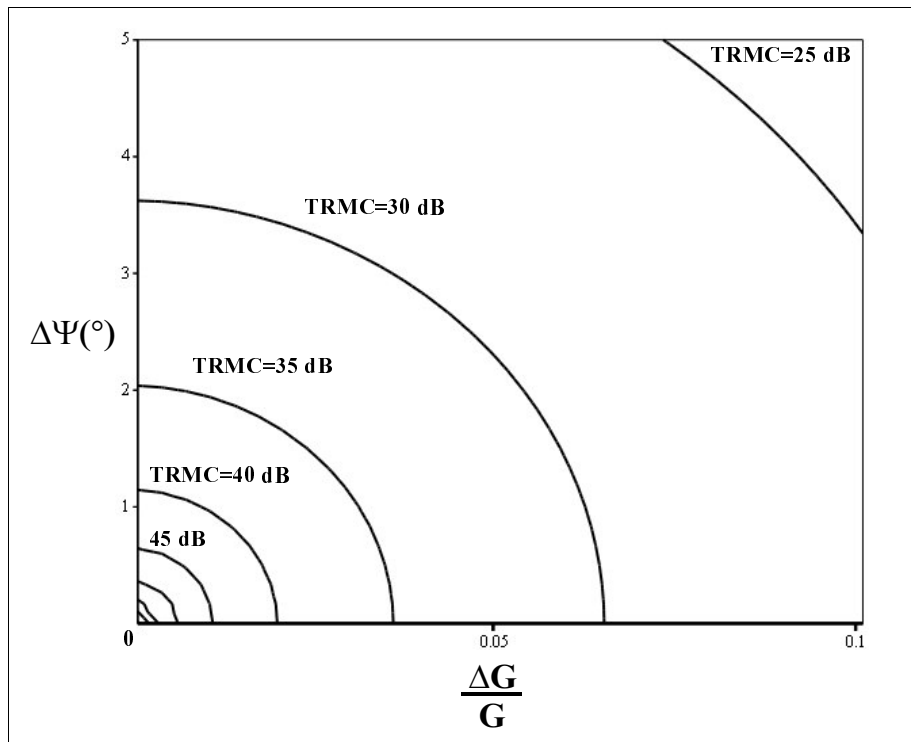


Figure II-3 : Lien entre le TRMC et les écarts de gain et de phase

Nous allons, dans le prochain paragraphe, présenter les coupleurs actifs retenus pour la réalisation de la structure double équilibrée.

II.2. Les coupleurs actifs à paire différentielle

Bien connue, la paire différentielle, présentée sur la Figure II-4, réalise la fonction de coupleur de signaux définie précédemment.

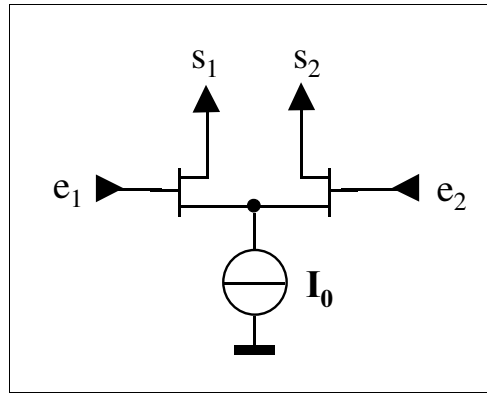


Figure II-4 : Schéma électrique d'une paire différentielle

En basse fréquence, la polarisation de la structure se fait par une source de courant I_0 connectée sur les sources (ou les émetteurs dans le cas de transistors bipolaires) des deux transistors. A ces fréquences, la source de courant et les transistors peuvent être considérés comme idéaux (l'impédance interne de la source de courant est infinie et les impédances Z_{12} et Z_{22} des transistors sont infinies), conduisant à un taux de réjection du mode commun infinie.

Dans le domaine des micro-ondes, ces hypothèses sont loin d'être vérifiées et, lorsque aucune précaution n'est prise, le taux de réjection du mode commun peut atteindre des valeurs inadmissibles.

Nous allons tout d'abord faire le bilan des éléments du coupleur à optimiser. Pour cette étude, nous traiterons le cas du diviseur de puissance. Il est bien évident que cette étude restera toute aussi valable pour les combineurs de puissance.

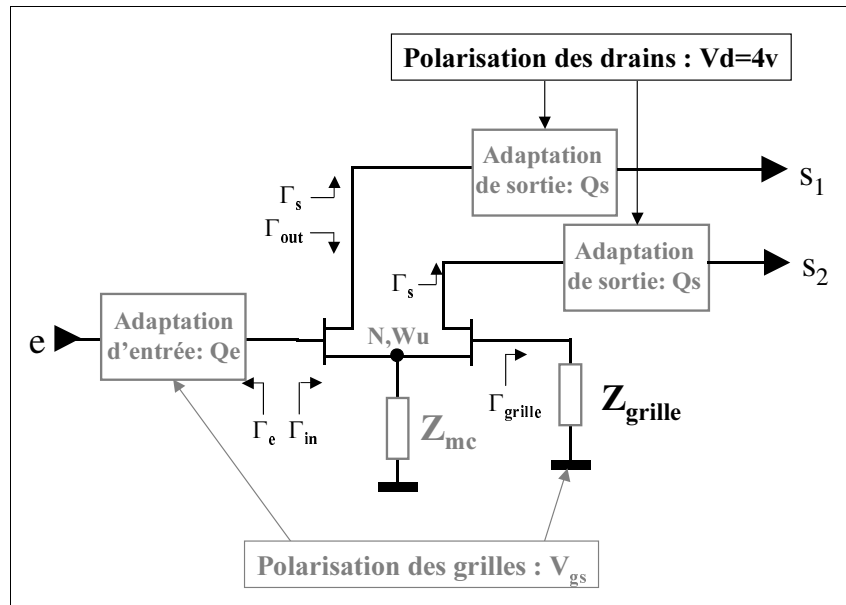


Figure II-5 : Schéma électrique du diviseur de puissance

La Figure II-5 donne une représentation du schéma électrique du diviseur de puissance dans laquelle les éléments à optimiser apparaissent grisés. Les quadripôles Q_e et Q_s , qui

correspondent à des cellules d'adaptation, doivent de même permettre les amenées de polarisation des transistors.

Le Tableau II-1 décrit l'influence des paramètres à déterminer sur les principales caractéristiques du coupleur résultant. Les coches définissent les caractéristiques que nous avons optimisées par ajustement du paramètre associé, tandis que le sigle  marque celles à ne pas dégrader lors de ces opérations.

	Gain	Bande passante	Consommation	Linéarité	Bruit	TRMC	Stabilité
V_{gs}			✓	✓	✓		
N	✓						
W_u						✓	
Z_{mc}						✓	
Z_{grille}						✓	
Q_e	✓						
Q_s	✓						

Tableau II-1 : Influences des paramètres sur les caractéristiques du coupleur

Les deux premiers éléments de ce tableau ont ainsi été déterminés par simulations électriques :

- La polarisation de grille V_{gs} a résulté d'un compromis entre le bruit, la linéarité et la consommation. Nous avons opté pour une valeur de -0.6v, qui correspond d'une part au minimum du facteur de bruit et d'autre part réalise un bon compromis entre la linéarité et la consommation de la paire différentielle.
- Le nombre de doigt de grille N a été fixé à 6 afin de maximiser le gain du circuit sans aboutir à des transistors de taille prohibitive.

Les autres éléments ont été déterminés à l'aide d'une étude analytique [III.5] de la paire différentielle menée sur un logiciel de traitement mathématique. Cette étude permet, à partir de la matrice des paramètres de dispersion du transistor $[S_T]$ et des différentes impédances du circuit, de déterminer à la fois le gain en puissance disponible du diviseur de puissance, mais aussi d'en calculer son taux de réjection du mode commun.

Comme énoncé en introduction de ce paragraphe, l'étude des performances du diviseur de puissance est considérablement simplifiée en décomposant l'ensemble constitué de la paire différentielle et de l'impédance Z_{mc} en mode pair et impair, comme présenté sur la Figure II-6 [III.6].

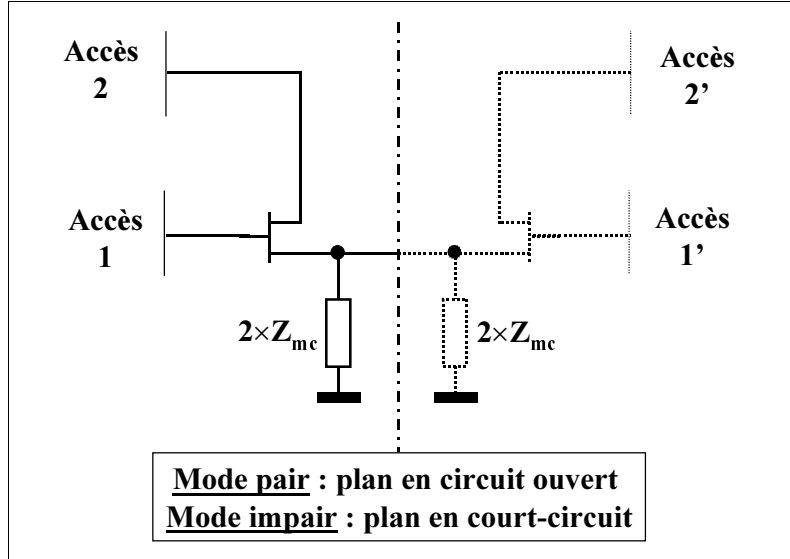


Figure II-6 : décomposition en mode pair et impair

Les matrices de dispersion des modes pair et impair (respectivement $[S_p]$ et $[S_i]$) sont facilement calculées en fonction de la matrice S du transistor seul ($[S_T]$) [III.6]:

- mode pair :

$$[S_p] = \left\{ [I_d] + \left(([I_d] - [S_T])^{-1} \times [S_T] + \frac{Z_{mc}}{Z_{ref}} \begin{bmatrix} 1 & 1 \\ 1 & 1 \end{bmatrix} \right)^{-1} \right\} \quad (\text{II-7})$$

- mode impair :

$$[S_i] = [S_T] \quad (\text{II-8})$$

Notons que la matrice $[S_p]$ est une fonction de la valeur de l'impédance Z_{mc} .

Avec ces notations, sont déduites les expressions du gain en puissance disponible et du taux de réjection du mode commun :

$$G_{\text{dmode impair}} = \frac{|b_{2\text{mode impair}}|}{|b_{\text{générateur disp.}}|}$$

$$= \left| \frac{1}{2} \times \frac{S_{21i}}{1 - \Gamma_s S_{22i}} \times \frac{1 - \Gamma_{\text{grille}} \Gamma_{\text{in-p}}}{\left[1 - \frac{\Gamma_e + \Gamma_{\text{grille}}}{2} \Gamma_{\text{in-i}} \right] \times \left[1 - \frac{\Gamma_e + \Gamma_{\text{grille}}}{2} \Gamma_{\text{in-p}} \right] - \left(\frac{\Gamma_e - \Gamma_{\text{grille}}}{2} \right)^2 \Gamma_{\text{in-i}} \times \Gamma_{\text{in-p}}} \right| \quad (\text{II-9})$$

$$\text{TRMC} = \left| \left(\frac{b_{2\text{mode commun}}}{b_{2\text{mode différentiel}}} \right)^{-1} \right| = \left| \left(\frac{1}{2} \times \frac{b_{2\text{mode pair}}}{b_{2\text{mode impair}}} \right)^{-1} \right| = \left| \left(\frac{1}{2} \times \underbrace{\frac{S_{21p}}{1-\Gamma_s S_{22p}}}_{\text{fonction de } Z_{mc}} \times \underbrace{\frac{1-\Gamma_{\text{grille}} \Gamma_{\text{in-i}}}{1-\Gamma_{\text{grille}} \Gamma_{\text{in-p}}}}_{\text{fonction de } Z_{\text{grille}}} \times \underbrace{\frac{1-\Gamma_s S_{22i}}{S_{21i}}}_{\text{fonction du transistor}} \right)^{-1} \right| \quad (\text{II-10})$$

Où Γ_s , Γ_e et Γ_{grille} sont des coefficients de réflexion définis sur la Figure II-5, tandis que $\Gamma_{\text{in-i}}$ et $\Gamma_{\text{in-p}}$ sont les coefficients de réflexion, en mode impair et pair, en entrée de la paire différentielle (Figure II-5) et définis par :

$$\Gamma_{\text{in-i}} = S_{1i} + \frac{\Gamma_s S_{12i} S_{21i}}{1 - S_{22i} \Gamma_s} \quad \text{et} \quad \Gamma_{\text{in-p}} = S_{1p} + \frac{\Gamma_s S_{12p} S_{21p}}{1 - S_{22p} \Gamma_s} \quad (\text{II-11})$$

Remarquons que ces deux coefficients de réflexion sont, bien évidemment, fonction de Γ_s .

Le taux de réjection du mode commun, qui est à maximiser, est donc constitué de trois facteurs :

- Le premier facteur n'est fonction que de Z_{mc} (au travers de $[S_p]$) et sera maximiser au paragraphe II.2.1.
- Le second facteur dépend de Z_{grille} et de Z_s (au travers de Γ_{ei} et Γ_{ep}). Nous le maximiserons au paragraphe II.2.2 en optimisant la valeur de l'impédance Z_{grille} .
- Le dernier facteur est quant à lui fonction de Z_s et des paramètres dispersion des transistors. Nous l'avons maximisé par un choix judicieux ($W_u=20\mu\text{m}$) de la largeur de grille des transistors mis en œuvre.

Enfin, le gain en puissance sera maximisé au paragraphe II.2.3 par le choix des quadripôles Q_e et Q_s (Figure II-5) réalisant l'adaptation d'impédance. Lors de l'optimisation des quadripôles, il faudra de plus veiller à maintenir le fonctionnement stable de la paire différentielle.

II.2.1. Optimisation de Z_{mc}

Dans le domaine des hyperfréquences, deux éléments vont influencer sur la valeur du taux de réjection du mode commun :

- L'impédance Z_{mc} connectée entre les sources des transistors et la masse. Cette impédance peut être intentionnelle où ramenée par les sources de polarisation (source de courant, par exemple). La Figure II-7 décrit l'évolution du TRMC

du coupleur en fonction de R_{mc} , partie réelle de Z_{mc} , lorsque la partie imaginaire de Z_{mc} est nulle.

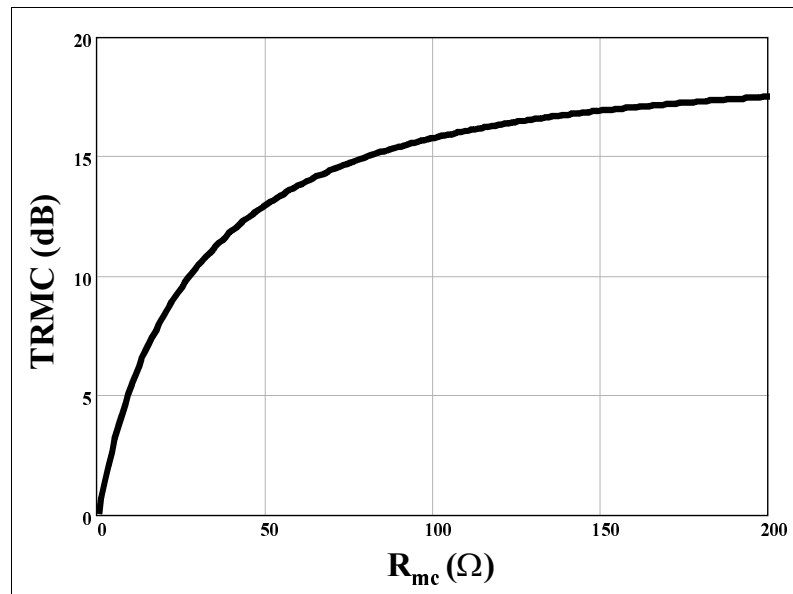


Figure II-7 : TRMC en fonction de la partie réelle l'impédance Z_{mc}

Cette courbe montre que l'impédance réelle R_{mc} influence grandement le taux de réjection du mode commun et doit être choisie la plus grande possible.

- L'impédance de sortie du transistor. Nous avons représenté à la Figure II-8 les variations du TRMC d'une paire différentielle pour laquelle les impédances de sortie des transistors étaient, artificiellement, multipliés par un coefficient k . L'impédance Z_{mc} a été prise infinie pour ce tracé.

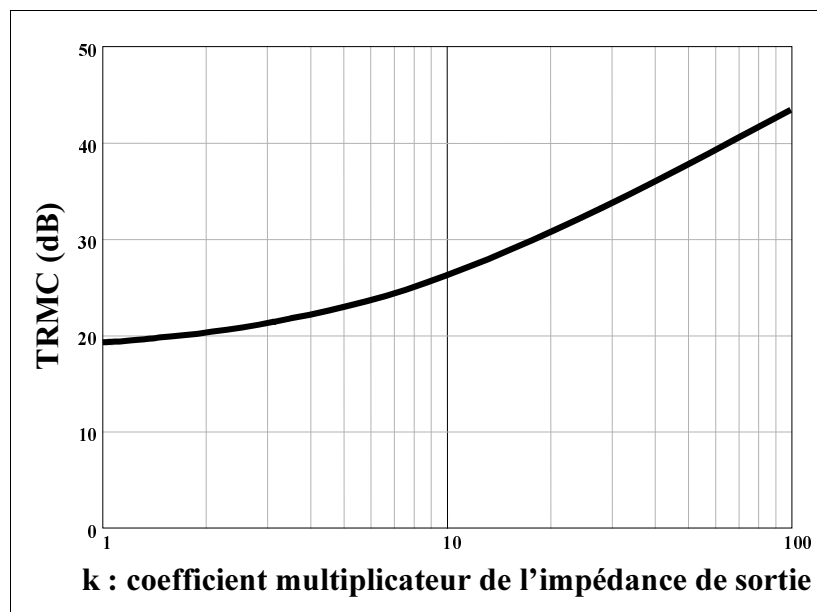


Figure II-8 : TRMC en fonction k .

Nous constatons que lorsque l'impédance de sortie du transistor augmente, tendant vers le transistor idéal, le taux de réjection du mode commun tend vers

l'infini. Ceci s'interprète facilement, puisqu'en définitive, les impédances de sortie des transistors contribuent à l'impédance globale entre leurs sources et la masse qui fixe le mode commun.

La conclusion est que le taux de rejection en mode commun est fixé par l'impédance présentée entre les sources de la paire différentielle et la masse. Cette impédance ne peut être infinie car elle est limitée par d'une part la source de courant polarisant la paire différentielle, et d'autre part les impédances de sortie des transistors. Dans notre cas, la limite supérieure de cette impédance est imposée par la conductance de sortie des transistors et majore le TRMC aux alentours de 19 dB (se reporter à la Figure II-8 pour $k=1$).

Une solution a été trouvée afin d'augmenter encore le TRMC. Cette solution découle d'une étude analytique [III.5] qui montre qu'en rendant capacitive l'impédance Z_{mc} , le TRMC pouvait être augmenté. La Figure II-9 illustre ce principe en représentant les variations du TRMC en fonction de la partie capacitive de l'impédance Z_{mc} .

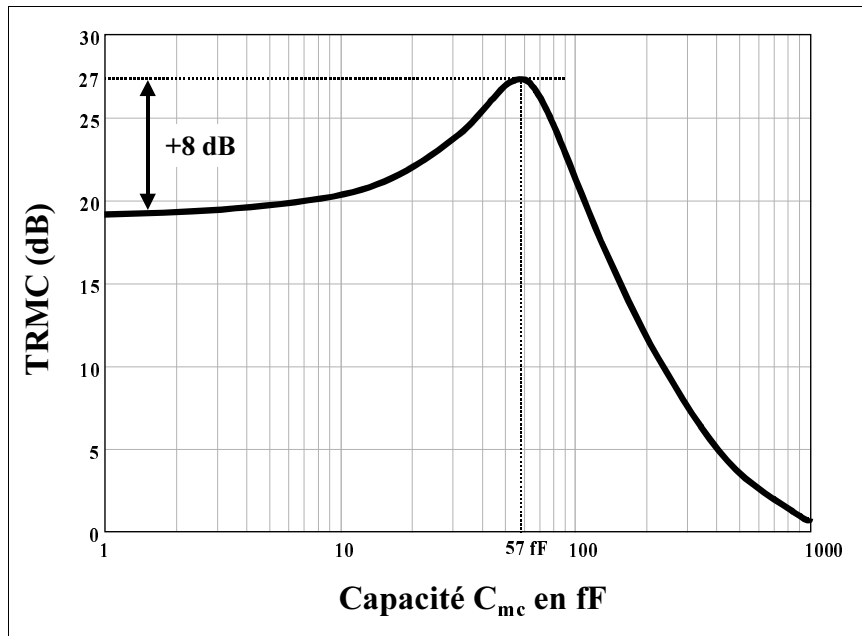


Figure II-9 : TRMC en fonction de C_{mc}

L'optimum de la valeur de cette capacité, notée C_{mc} , vaut 57fF. Cette valeur correspond à un taux de réjection du mode commun de 27dB, soit une amélioration de 8dB par rapport à une impédance Z_{mc} purement résistive.

Le premier facteur de l'expression (II-10) a ainsi été optimisé. Nous allons, dans le prochain paragraphe porter notre attention sur le second facteur de cette expression.

II.2.2. Optimisation de Z_{grille}

Cette étude vise l'optimisation de l'impédance Z_{grille} afin de maximiser le second terme du TRMC décrit par l'expression (II-10).

La Figure II-10 montre la dépendance de ce facteur en fonction de la valeur de l'inductance série de l'impédance Z_{grille} , paramétrée pour trois valeurs de sa résistance série.

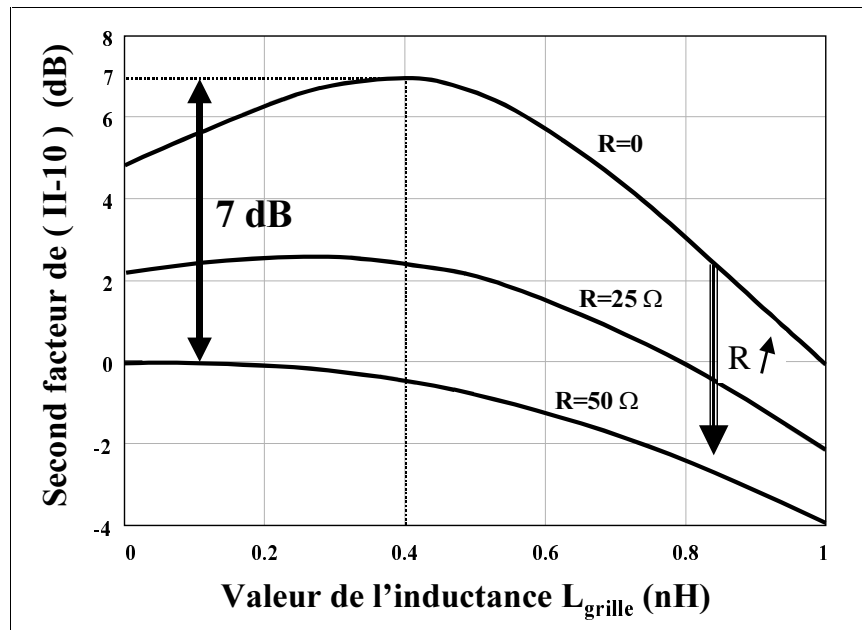


Figure II-10 : optimisation de Z_{grille}

Il apparaît donc que le taux de réjection du mode commun est maximum pour une impédance Z_{grille} purement inductive (des valeurs de réactance capacitive ont été testées et dégradent fortement le TRMC). Le TRMC est maximum pour une valeur d'inductance de grille de 0,4 nH.

Cette étude démontre que le taux de réjection du mode commun est amélioré de 7dB si la grille du transistor non-utilisée en diviseur de puissance est connectée à une inductance de valeur 0.4 nH par rapport à une charge 50Ω.

Pour conclure, les impédances, respectivement de couplage des sources des transistors de la paire différentielle Z_{mc} et de charge de grille du transistor non utilisé en diviseur Z_{grille} , ont comme nous l'avons montré une grande influence sur le TRMC : des valeurs optimales pour ces impédances se traduisent par une amélioration du TRMC de 15dB, sa valeur augmentant de 19dB à 34dB.

II.2.3. Adaptation et stabilité

Enfin, les calculs analytiques menés sur cette topologie nous ont permis de déterminer les éléments d'adaptation en entrée et en sortie afin de garantir un gain maximal sur la bande de fréquence désirée. Nous avons donc réalisé l'adaptation conjuguée de l'entrée et de la sortie maximisant la valeur du gain en puissance disponible (expression (II-9)) conduisant aux coefficients de réflexion à présenter en entrée (Γ_e) et en sortie (Γ_s) de la paire différentielle (se reporter à la Figure II-5) suivants :

$$\Gamma_e = (\Gamma_{in})^* = \left(\frac{\frac{\Gamma_{in-i} + \Gamma_{in-p}}{2} - \Gamma_{grille} \Gamma_{in-i} \Gamma_{in-p}}{1 - \Gamma_{grille} \frac{\Gamma_{in-i} + \Gamma_{in-p}}{2}} \right)^* \quad (\text{II-12})$$

$$\Gamma_s = (\Gamma_{out})^* = \left(\frac{\frac{\Gamma_{out-i} + \Gamma_{out-p}}{2} - \Gamma_s \Gamma_{out-i} \Gamma_{out-p}}{1 - \Gamma_s \frac{\Gamma_{out-i} + \Gamma_{out-p}}{2}} \right)^* \quad (\text{II-13})$$

Remarque : Les deux expressions précédentes sont des équations implicites couplées. Leur résolution est alors effectuée par optimisation des valeurs de Γ_e et Γ_s afin que les égalités (II-12) et (II-13) soient simultanément vérifiées.

Une attention particulière, comme pour la cellule de mélange, a été portée sur la stabilité de ces topologies dans lesquelles une contre réaction capacitive (C_{mc}) se retrouve sur la source des transistors. La présence de deux transistors doit nécessiter l'observation de deux variables du circuit. Il est alors judicieux, non pas d'observer les deux entrées e_1 et e_2 (se reporter aux notations sur la Figure II-1) mais de rendre compte de la stabilité des variables e_p et e_i (se reporter aux expressions (II-1)) respectivement entrée en mode pair et entrée en mode impair.

L'étude de la stabilité a donc été menée simplement en décomposant le circuit en deux sous-circuits indépendants représentatifs des modes pair et impair.

La stabilité est assurée si et seulement si les deux lieux de Nyquist de $\Gamma_{e-p} \Gamma_{in-p}$ pour le mode pair et $\Gamma_{e-i} \Gamma_{in-i}$ (en traits pleins) pour le mode impair n'entoure pas le point critique (1,0).

La Figure II-11 montre ces deux lieux de Nyquist pour la configuration optimale retenue.

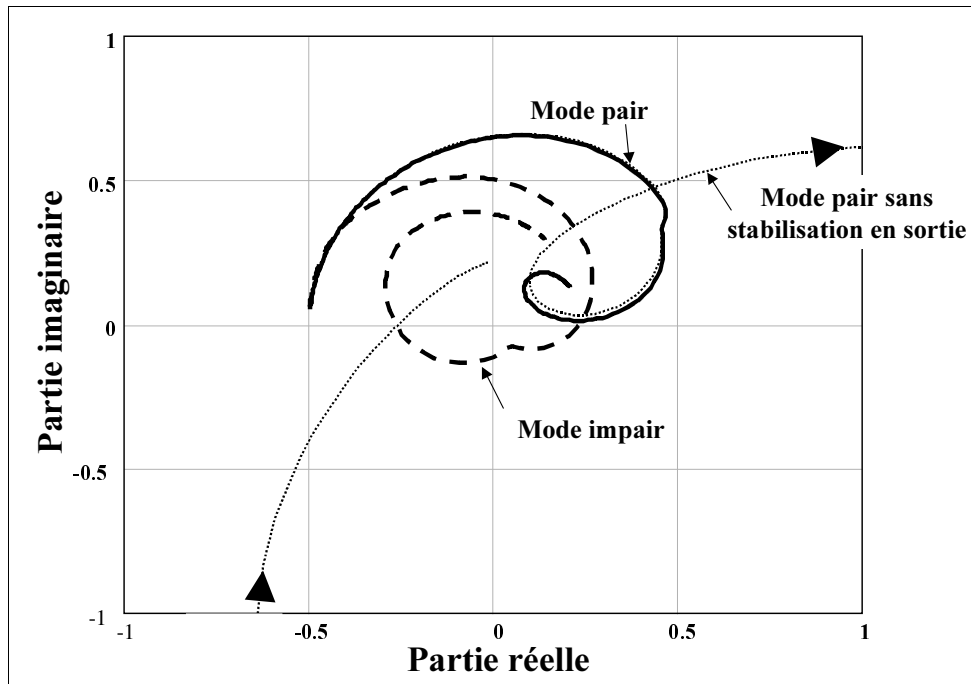


Figure II-11 : Lieux de Nyquist des modes pair et impair

Comme le montre cette figure, la stabilité est assurée pour la cellule car les lieux de Nyquist des modes pair et impair n'entourent pas le point critique (1,0). Cette stabilité a nécessité un choix correct de la topologie d'adaptation en entrée ainsi que l'ajout d'un réseau de stabilisation en sortie.

En pointillé sur la Figure II-11, le lieu de Nyquist du mode pair de la structure sans réseau stabilisant en sortie présente un encerclement du point critique (1,0) démontrant la nette instabilité de la structure et prouvant la nécessité de ce réseau.

II.2.4. Conclusion

Nous venons de présenter les différentes étapes d'optimisation à mener lors de la conception des deux diviseurs de puissance (RF et OL) ainsi que du combineur de puissance (FI) nécessaire à la réalisation d'une structure double équilibrée performante. Cette optimisation a été conduite en vue d'obtenir les valeurs du gain différentiel et du taux de réjection du mode commun les plus grandes, mais aussi afin de garantir la stabilité des coupleurs.

Nous allons, dans le prochain paragraphe, présenter l'intégration monolithique des différentes cellules conçues au sein d'un système complet de conversion de fréquence.

III. Intégration monolithique sur GaAs

Ce paragraphe décrit l'intégration monolithique du mélangeur doublement équilibré en technologie sur Arséniure de Gallium. Nous présentons tout d'abord la technologie utilisée et les règles de conception que nous avons établies. Puis, nous présenterons le circuit conçu ainsi que ses performances simulées.

III.1. La technologie PML/OMMIC

La technologie utilisée pour cette réalisation correspond à la filière ED02AH [III.7] de chez Philips Microwave Limeil (PML) récemment renommé OMMIC. Cette technologie d'intégration sur Arséniure de Gallium est basée sur des transistors à effet de champ de type PHEMT (Pseudomorphique High Electron Mobility Transistor) dont la fréquence de transition atteint 60 GHz grâce notamment à une longueur de grille de $0.2\mu\text{m}$ et au couple ternaire de l'hétérojonction GaInAs(Canal)-GaAlAs(Spacer).

Cette filière comporte des transistors à déplétion mais aussi des transistors à enrichissement (Enhancement) possédant une tension de pincement de 0.225V que nous n'avons pas utilisés.

Au niveau des interconnexions, deux niveaux métalliques sont disponibles. Le niveau supérieur (IN) est à préférer puisqu'il possède une résistance carré deux fois moins importante que le niveau inférieur (BE). Ce dernier ne sera donc utilisé que pour les croisements de lignes. Nous avons utilisé des interconnexions par lignes micro-bandes pour lesquelles le substrat est aminci à $100\mu\text{m}$ autorisant la réalisation de prises de masse par trous métallisés (via-holes). L'importance de l'inductance parasite ($L=32\text{pH}$) de ces remontées de masse à travers le substrat n'est pas à ignorer lors de l'intégration, et interdit l'utilisation d'un seul et même trou métallisé pour plusieurs prises de masse. En effet, le couplage résultant de cette configuration entraînerait une chute parfois désastreuse des performances ou encore l'apparition d'instabilité.

Deux types de capacités MIM (Métal Isolant Métal) existent afin de couvrir une large plage de valeurs. Technologiquement, la différence se situe au niveau de la distance inter-électrode et du diélectrique dominant. Qualitativement, les valeurs de capacités vont de 1fF à 50pF avec une fréquence maximale d'utilisation de 50GHz.

Nous allons traiter en détail deux points plus particuliers que sont les inductances intégrées et l'éloignement minimal entre les éléments du circuit.

III.1.1. Les inductances intégrées

Ce paragraphe traite du dimensionnement des différentes inductances que nous devons mettre en œuvre dans le circuit.

Les inductances se définissent principalement par quatre caractéristiques : leur valeur (L), leur fréquence de résonance (F_r), leur coefficient de qualité (Q) et la surface qu'elles occupent (S). Pour les inductances intégrées, l'optimisation de ces caractéristiques repose sur l'ajustement de trois paramètres : la largeur du conducteur (W), la distance entre les enroulements (G) et la longueur totale de l'enroulement (P).

Les relations entre caractéristiques et paramètres sont complexes et chaque caractéristique dépend de tous les paramètres. D'une manière générale, les valeurs de W et de G seront figées, et la valeur de P sera ajustée afin d'obtenir la valeur d'inductance désirée.

Nous allons tenter cependant de définir un choix optimal d'inductance en dictant quelques règles de conception.

1. Il est apparu, lors de simulations, que la fréquence de résonance dominante des inductances dépendait principalement de leur valeur. La courbe ci-dessous représente les variations de la fréquence de résonance en fonction de la valeur de l'inductance.

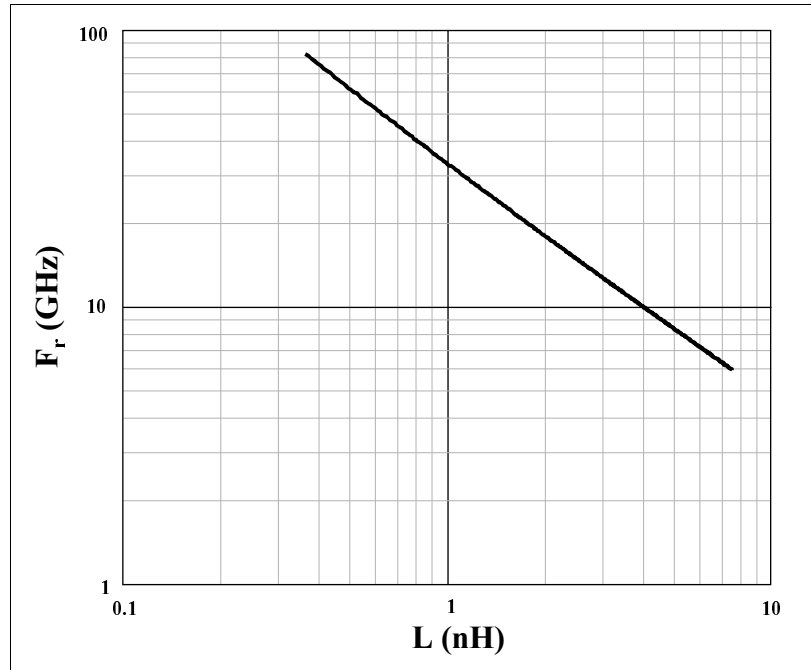


Figure III-1 : Choix de l'inductance maximale

Travaillant en bande Ku, nous avons donc limité les valeurs d'inductances à 2nH, qui correspondent à une fréquence de résonance de 18GHz. Pour les inductances devant filtrer des harmoniques du signal RF, nous nous sommes limités à des valeurs inférieures à 2nH, afin de disposer d'inductances de fréquence de résonance supérieure à 20GHz.

2. La valeur de G a été fixée à sa borne technologique inférieure, soit $5\mu\text{m}$. Cette valeur optimise à la fois la compacité de l'inductance et son coefficient de qualité.
3. Plus délicat est le choix de la valeur de W . Cette valeur sera ajustée en fonction de la valeur de l'inductance désirée. En effet, pour de forte valeur d'inductance, une trop grande valeur de W entraînerait une surface occupée prohibitive ; nous avons donc fixé, dans ce cas, $W=5\mu\text{m}$. A l'opposé, une forte valeur de W entraîne une faible valeur d'inductance minimale réalisable ; nous avons donc choisi, pour la réalisation de faibles inductances, $W=15\mu\text{m}$. Enfin, pour des valeurs d'inductance intermédiaires, un choix intermédiaire de W réalise un bon compromis entre une grande valeur de Q et une faible surface occupée. Le Tableau III-1 résume le choix de la valeur de W .

Faibles valeurs d'inductance	Valeurs intermédiaires d'inductance	Fortes valeurs d'inductance
$L < 0.4 \text{ nH}$	$0.4 \text{ nH} < L < 1.5 \text{ nH}$	$1.5 \text{ nH} < L$
$W = 15 \mu\text{m}$	$W = 10 \mu\text{m}$	$W = 5 \mu\text{m}$
$\nearrow Q$	Compromis Q/S	$\searrow S \quad \searrow Q$

Tableau III-1 : Choix de la largeur du ruban métallique des inductances

Enfin, lorsque des valeurs d'inductance ont été inférieures à la limite inférieure des inductances réalisables, des inductances à méandres ont été conçues.

III.1.2. Eloignement minimal des composants

La dernière étape de conception est le dessin des masques que requiert le fondeur pour la fabrication du circuit monolithique. Se basant sur le schéma électrique, le concepteur doit placer les différents éléments et réaliser les interconnexions à l'aide de lignes métalliques.

Le placement des composants n'est pas indifférent car, en hyperfréquences, des couplages d'origine électromagnétique existent lesquels doivent être pris en considération. Or, les simulations électriques, par les modèles qu'elles mettent en œuvre, ne permettent pas de rendre compte de ces couplages.

Deux processus de conception sont alors possibles :

1. Soit, le concepteur place les différents éléments suffisamment éloignés afin qu'aucun couplage électromagnétique n'existe entre eux. L'avantage de cette solution est de réaliser un circuit conforme aux simulations électriques. Par contre, l'éloignement des différents éléments entraîne inévitablement une surface de circuit non-optimisée.
2. Soit, le concepteur souhaite optimiser la surface du circuit monolithique en rapprochant au maximum les différents éléments entre eux. Dans ce cas, il est nécessaire d'effectuer des retro-simulations électromagnétiques afin de tenir compte des éventuels couplages inter-éléments.

Pour notre part, étant donnée la complexité du circuit envisagé, ce second processus de conception nous est apparu insurmontable pour tout simulateur électromagnétique. Nous avons donc opté pour la première méthode et tenté de définir un écartement minimal entre les différents éléments afin d'éviter tout couplage électromagnétique.

Nous avons, pour ce faire, effectué des simulations électriques sur un modèle de ligne couplée donnée par le fondeur. Le critère retenu est l'obtention d'une isolation de 30dB, suffisante pour garantir la non-interaction entre les éléments, pour des largeurs de ruban de $10 \mu\text{m}$. Notons que cette largeur représente une valeur standard pour notre conception.

La Figure III-2 définit l'espacement entre les deux lignes couplées en fonction de la longueur de couplage pour satisfaire un tel critère. Les points sur cette figure représentent les résultats de simulation et la droite (de pente 0.75) une interpolation de ces diverses simulations.

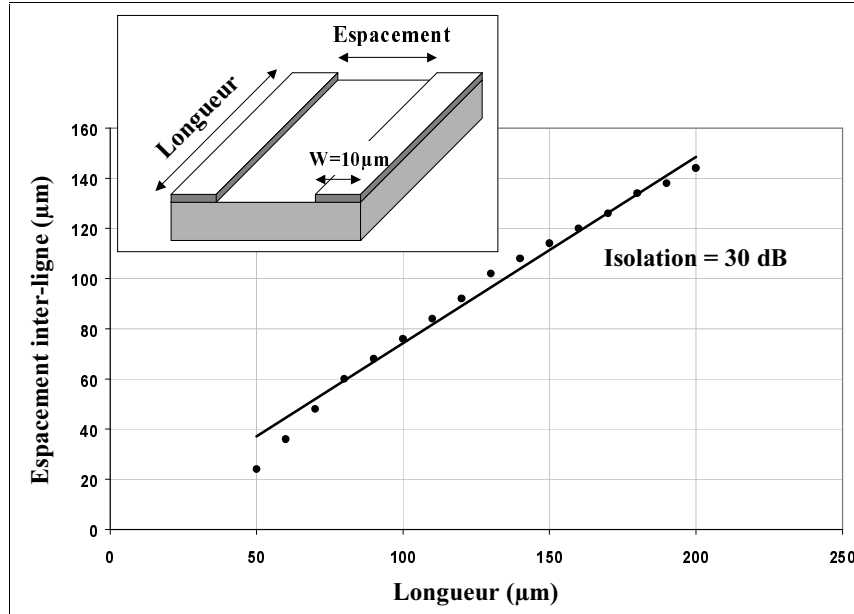


Figure III-2 : Espacement minimal fonction de la longueur de couplage pour une isolation de 30dB

La règle empirique qu'il est possible de tirer de cette figure est qu'une distance entre les différents éléments égale à 75% de leur dimension de couplage (longueur de couplage) est suffisante pour garantir qu'aucune interaction électromagnétique entre ces éléments n'existe.

III.2. Intégration MMIC

Nous avons donc effectué l'intégration monolithique du mélangeur double équilibré à l'aide de cette technologie en observant les différentes précautions citées précédemment [III.3].

Lors de cette étape finale, nous avons, autant que faire ce peut, veillé à conserver les symétries de la structure double équilibré afin de ne pas dégrader les réjections des raies parasites. De plus, nous avons tenu compte de toutes les inductances ramenées par les lignes d'interconnexion lors de la conception des cellules passives.

L'assemblage des différentes cellules n'a nécessité aucune étape de conception supplémentaire, chaque fonction ayant déjà été, lors de son optimisation théorique, optimisée en tenant compte des étages amont et aval. Seules, quelques ajustements supplémentaires ont été nécessaires afin de garantir des performances optimales, certes quelque peu dégradées lors de la prise en compte de tous les éléments parasites des divers composants mis en œuvre.

Enfin, notons que l'intégration d'une des inductances du filtre en sortie du mélangeur, a fortement dégradé le gain de la structure. En effet, la valeur importante de $L_{s0}=1.9\text{nH}$ (définie à la figure II-9 du chapitre 2) a nécessité, afin de garantir la compacité du circuit résultant, une valeur de largeur de ruban de $5\mu\text{m}$, dégradant fortement son coefficient de qualité (cf. Tableau III-1). Cette diminution du coefficient de qualité du filtre en sortie du mélangeur, pour lequel nous avons vu l'importance (résumée sur le tableau I-2 du chapitre 2), s'est accompagné d'une forte diminution du gain du mélangeur. Le Tableau III-2 résume les dégradations du gain de la cellule de mélange qui chute à 5dB.

Inductances idéales $Q=\infty$	Gc=10 dB
Inductances réelles $W=15\mu\text{m}$ $\Rightarrow$ Trop volumineux	Gc=7.5 dB
Inductances réelles $W=5\mu\text{m} \Rightarrow \searrow Q \text{ et } \searrow S$	Gc=5 dB

Tableau III-2 : Dégradation des performances du mélangeur

Une solution aurait été de fixer une borne supérieure d'inductance inférieure à 1nH afin de satisfaire aux impératifs de compacité sans dégradation des performances.

Le dessin des masques du mélangeur doublement équilibré est présenté sur la Figure III-3. Le circuit monolithique présente les caractéristiques suivantes :

- Dimensions : $2 \times 3 \text{ mm}^2$
- Epaisseur du substrat : $100\mu\text{m}$
- Métallisation arrière : oui
- Nombre de trous métallisés : 42
- Nombre de transistors : 12
- Nombre d'accès : 3
- Nombre de plots de polarisation : 10
- Consommation : 800mW

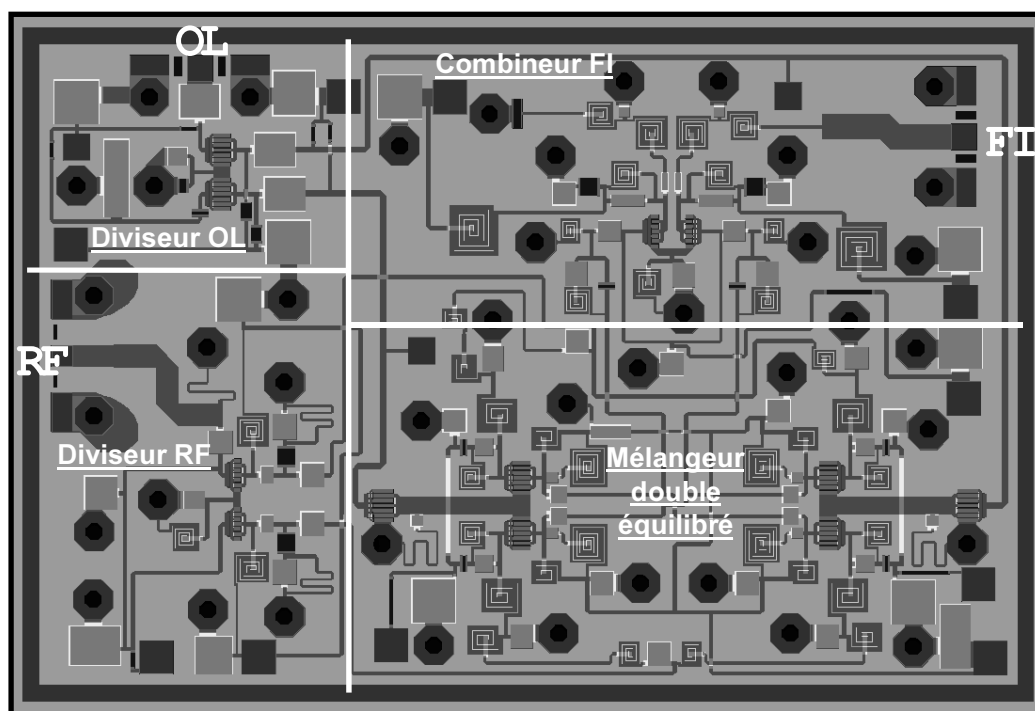


Figure III-3 : Dessin des masques du mélangeur double équilibré ($2 \times 3 \text{ mm}^2$)

Enfin, nous avons intégré séparément divers éléments de test parmi lesquels : un coupleur à 12 GHz et une cellule simple équilibrée de mélange.

III.3. Performances simulées

L'essentiel des caractéristiques simulées des différentes cellules constitutives du mélangeur double équilibré est résumé dans les trois premières colonnes du Tableau III-3. Ces valeurs ont été simulées pour une fréquence d'entrée $f_{RF}=14\text{GHz}$ et une fréquence f_{OL} de 2GHz.

	Diviseur RF	Mélangeur simple équilibré Pol=0 dBm	Combineur FI	Mélangeur double équilibré Pol=0 dBm	
				Estimation*	Simulation
Gain (dB)	9.6	5	9.2	23.8	23
NF (dB)	3.75	6	3.35	4.8	4.5
P_{13s} (dBm)	22.7	7	23.6	18	16

* Données déduites de l'application des formules de mise en cascade (voir formules II-7 et II-10 au paragraphe 2 du premier chapitre) aux trois éléments : diviseur RF, mélangeur et combineur FI.

Tableau III-3 : Synthèse des résultats

L'avant dernière colonne de ce tableau présente les caractéristiques estimées du mélangeur double équilibré déduites de la mise en cascade des trois cellules qui le constituent. La correspondance de ces valeurs avec celles présentées à la dernière colonne, déduites de

simulations globales du mélangeur, montre que l'étape ultime d'assemblage des différents blocs ne s'est accompagné d'aucune perte de performances, validant toute la méthode de conception élaborée.

Le mélangeur simple équilibré démontre des performances simulées supérieures à celle de l'état de l'art [III.8], ce qui confirme la pertinence de toute l'approche décrite au chapitre 2. De plus, dans l'analyse de ces performances, la dégradation du gain de conversion pour raison de compacité est à prendre en considération puisque sa valeur initiale était de 10 dB, ce qui, dans la gamme fréquentielle de fonctionnement, est tout à fait remarquable.

La Figure III-4 présente le gain de conversion et le facteur de bruit simulés de la structure globale, en fonction de la fréquence f_{RF} pour $f_{OL}=2$ GHz et $P_{OL}=0$ dBm.

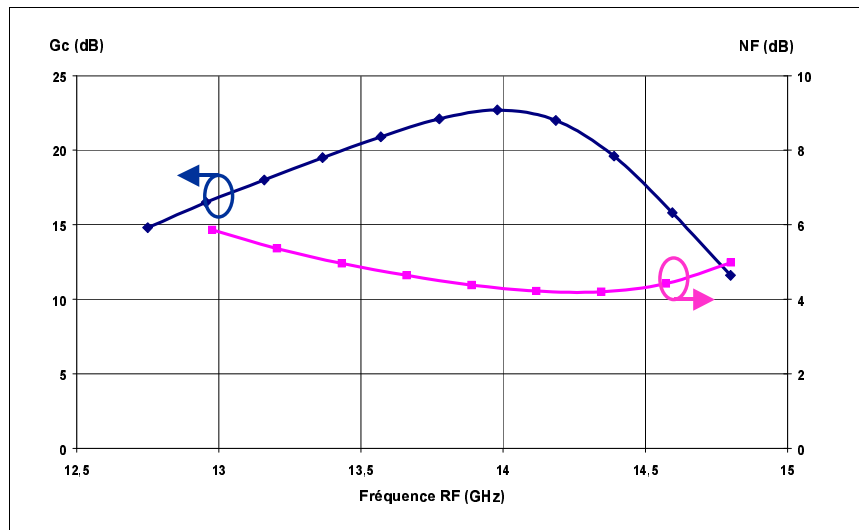


Figure III-4 : Gain et facteur de bruit en fonction de f_{RF}

Le gain de conversion est supérieur à 20 dB sur plus de 1 GHz de bande RF et reste supérieur à 10 dB sur toute la plage fréquentielle imposée par le cahier de charges. Le facteur de bruit reste, quant à lui, inférieur à 5 dB sur 1,5 GHz de bande et est inférieur à 6 dB sur toute la plage de fréquence d'entrée.

La Figure III-5 présente le gain de conversion en fonction de la fréquence f_{RF} utile pour trois valeurs de la fréquence f_{OL} correspondant à la valeur minimale (1,25 GHz), la valeur nominale (2 GHz) et la valeur maximale (3,85 GHz) données par le cahier des charges. La puissance P_{OL} appliquée étant maintenue à 0 dBm.

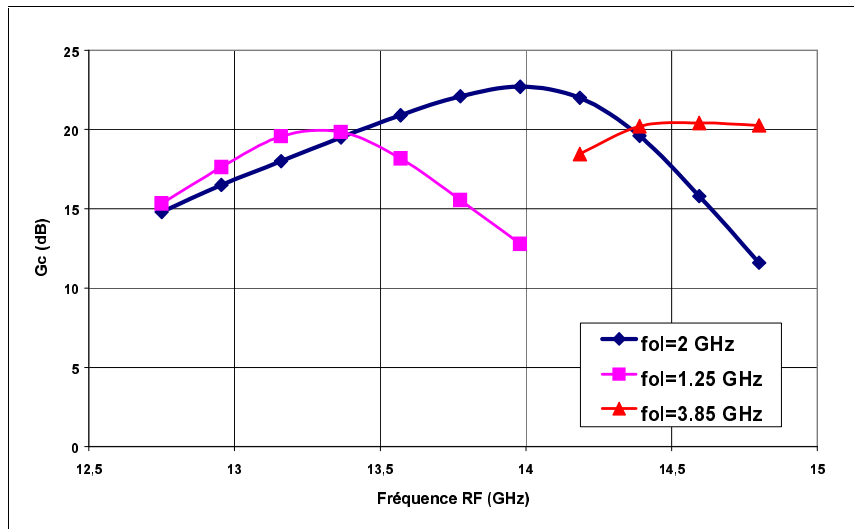


Figure III-5 : Gain de conversion en fonction de f_{RF} paramétré par la fréquence f_{OL}

Ces différentes courbes montrent que le gain de conversion n'est pas fortement dégradé lorsque la valeur de la fréquence f_{OL} n'est plus nominale, et reste supérieur à 10 dB sur toute la plage utile de la fréquence f_{RF} pour les valeurs extrémales de la fréquence f_{OL} .

Les variations de la puissance d'interception d'ordre 3 en sortie et celle du gain de conversion en fonction de la puissance OL pour les valeurs nominales des fréquences f_{RF} et f_{OL} sont présentées sur la Figure III-6. Une valeur de $P_{I3s}=17\text{dBm}$, imposée par le cahier des charges, est atteinte pour une puissance OL relativement faible de -10dBm et correspond à une valeur de gain de conversion de 21 dB.

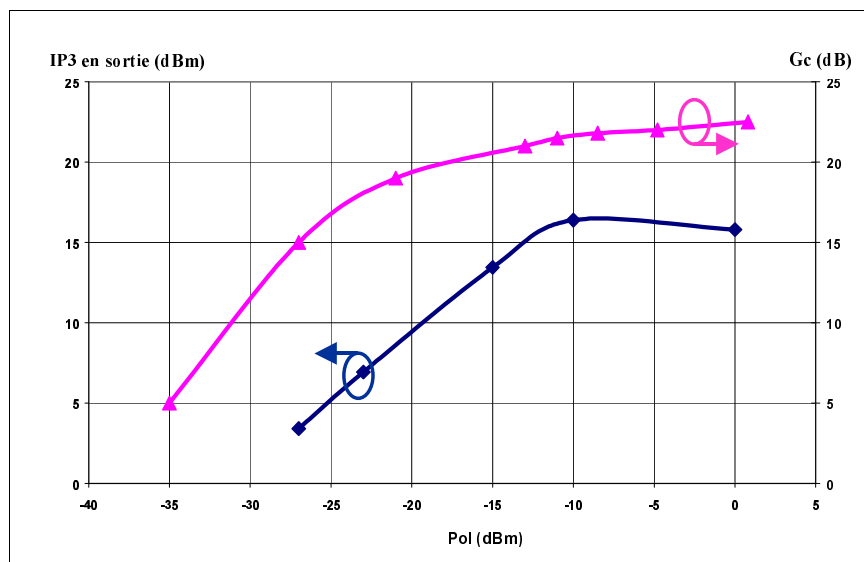


Figure III-6 : Gain de conversion et P_{I3s} en fonction de P_{OL}

Les points suivants résument l'essentiel des caractéristiques simulées du mélangeur double équilibré :

- **Gain de conversion** : **23dB** à $f_{RF}=14\text{GHz}$ et $>20\text{dB}$ sur 1GHz de bande RF pour **Pol=-10dBm**.
- **Facteur de bruit (SSB)** : **4.5dB** à 14GHz et $< 6\text{dB}$ sur la bande RF.
- **IP3 en sortie** : **17dBm** à 14GHz (11.8 et 13.3dBm pour 14.8 et 13GHz) pour $P_{OL}=0\text{dBm}$.
- **Isolations**: RF $\rightarrow$ FI = 35dBc, LO $\rightarrow$ FI = 30dBc.
- **Réjection des raies parasites** supérieure à 40dBc. (Sauf 5OL = -23dBc et RF+kOL = -20dBc)

Les performances simulées du mélangeur double équilibré ainsi intégré sont tout à fait compatibles avec les contraintes imposées par le cahier des charges. Notons que ces performances sont comparables à l'état de l'art [III.9] des systèmes répéteurs en bande Ku.

En conclusion, nous avons donc conçu, optimisé et intégré un mélangeur double équilibré constitué de trois coupleurs de signaux et de deux cellules de mélange simplement équilibrées sur une puce de dimensions $2 \times 3 \text{ mm}^2$.

Les performances simulées de cette conception sont conformes au cahier des charges imposé, à l'exception de certaines réjections de raies parasites dont le respect se réaliserait aisément par l'intégration supplémentaire de cellules de filtrage.

Nous allons, dans le paragraphe suivant, présenter les résultats de caractérisation menée sur les circuits monolithiques réalisés.

IV. Caractérisations des circuits MMIC

IV.1. Montage et caractérisation

La puce de la fonction globale et celle regroupant les éléments de test ont été conçues en vue d'une caractérisation sous pointes hyperfréquences et intègrent donc des plots coplanaires d'accès pour les signaux micro-ondes. Elles ont été fabriquées dans le cadre d'un multi-projet géré par le CMP.

Afin de garantir la stabilité mécanique des ces puces de *faibles* dimensions (quelques millimètres carrés) lors du posé des deux, trois ou quatre pointes coplanaires, et ainsi de garantir les contacts, nous avons reporté ces circuits sur un substrat hôte (tranche de silicium métallisé) de plus *grandes* dimensions (quelques centimètres carrés).

Un avantage supplémentaire de cette solution est de permettre l'adjonction de capacités MIM de découplage de forte valeur (100 pF) à proximité des plots de polarisation des circuits. Ce renforcement des capacités de découplage des alimentations évite tout risque d'oscillation aux *basses* fréquences (1MHz à quelques 100MHz).

La Figure IV-1 montre la photographie d'une réalisation, pour laquelle le découplage des alimentations a été renforcé par l'adjonction de capacités CMS de valeur 100nF.

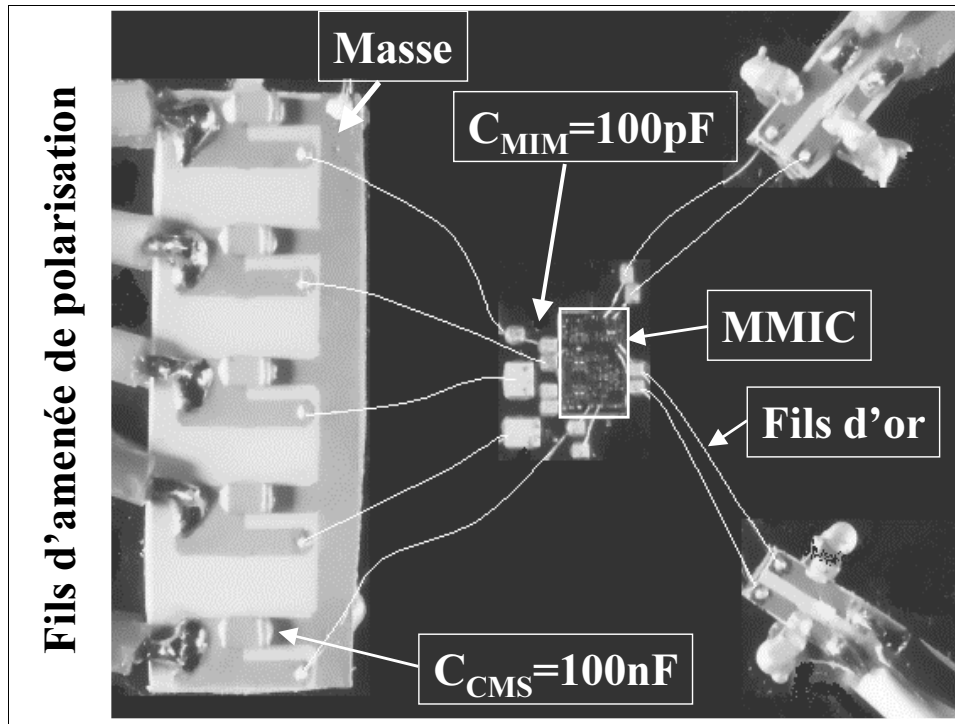


Figure IV-1 : Photographie d'un montage de puce en vue du test hyperfréquence

Enfin, lors des caractérisations, nous avons ajusté les sources de tension continues externes des différents circuits afin de nous affranchir des fortes dispersions sur les valeurs des tensions de pincement des transistors. Ces ajustements se sont effectués afin de faire correspondre les courants de polarisation expérimentaux à ceux issus des simulations assurant alors des points de polarisation optimaux.

IV.2. Cellules de base

Nous allons, dans un premier temps, présenter les caractéristiques des cellules de base que nous avons intégrées. L'objectif de ces caractérisations n'étant pas l'obtention des performances à l'état de l'art mais la validation des différents principes mis en œuvre. En effet, ces cellules, copies de celles intégrées dans le mélangeur doublement équilibré, seront testées sous des impédances de référence de 50Ω non-optimales, et ne présenteront donc pas qualitativement les performances visées dans le système complet.

IV.2.1. Mélangeur simple équilibré

Nous avons, pour cette cellule, effectué uniquement la mesure du gain de conversion. Cette caractéristique est, en effet, la seule intervenant pour l'évaluation des performances globales du système.

La Figure IV-2 montre les résultats de mesure en les comparant à ceux issus des simulations électriques. Notons tout d'abord que la valeur maximale du gain de conversion simulé n'est que de 3dB contre 5dB annoncé au Tableau III-2. Cet écart est uniquement la conséquence de la caractérisation sur des impédances de références de 50Ω non optimales pour la cellule.

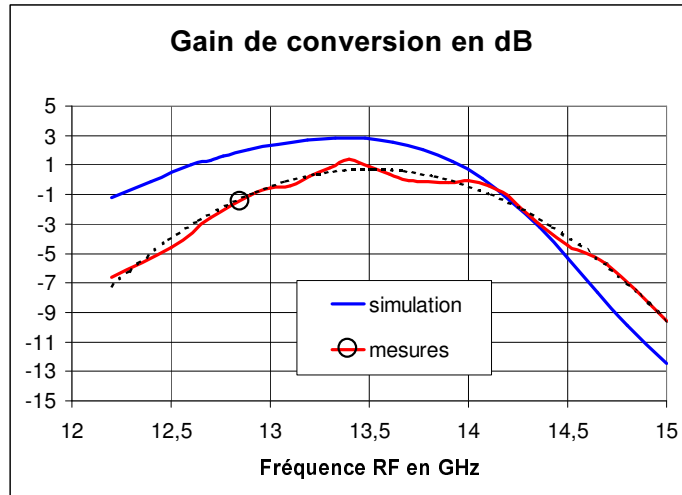


Figure IV-2 : Gain de conversion de la cellule de mélange

Un écart de 2 dB sur la valeur du gain de conversion ainsi qu'un faible décalage fréquentiel de 0.2 GHz entre simulations et mesures rendent compte des seuls effets non pris en considération lors des simulations électriques : les couplages électromagnétiques entre notamment les éléments passifs.

Rappelons que nous avons, lors des caractérisations, ajusté les sources de polarisations extérieures afin de garantir les points de polarisation optimaux. Nous avons ainsi automatiquement compensé les dispersions sur les tensions de pincement des transistors. Ce type de dispersion, très sensible dans le cas général, n'est donc pas la cause des différences observées entre simulations et expérimentations. C'est la raison pour laquelle, nous nous sommes acheminé vers une justification fondée sur des couplages électromagnétiques entre éléments contigus.

En effet, la règle empirique d'éloignement entre les différents éléments passifs énoncée au paragraphe III.1.2 n'a pu être respectée compte tenu que notre système global devait être intégré sur une puce de taille maximale $2 \times 3 \text{ mm}^2$ (imposée par le CMP, organisme réalisant la phase finale d'assemblage du réticule global). Cette forte intégration a donc entraîné des couplages électromagnétiques non négligeables entre éléments contigus.

Afin de prouver que ces couplages de proximité sont effectivement la cause des différences entre les mesures et les simulations électriques, nous avons réalisé des simulations électromagnétiques sur une des cellules passives les plus *critiques* vis à vis des performances électriques : le filtre en sortie du mélangeur, pour lequel nous avons déjà évoqué la sensibilité sur le gain de conversion au début du paragraphe III.2. Nous découvrons en Figure IV-3 le résultat de ces simulations électromagnétiques.

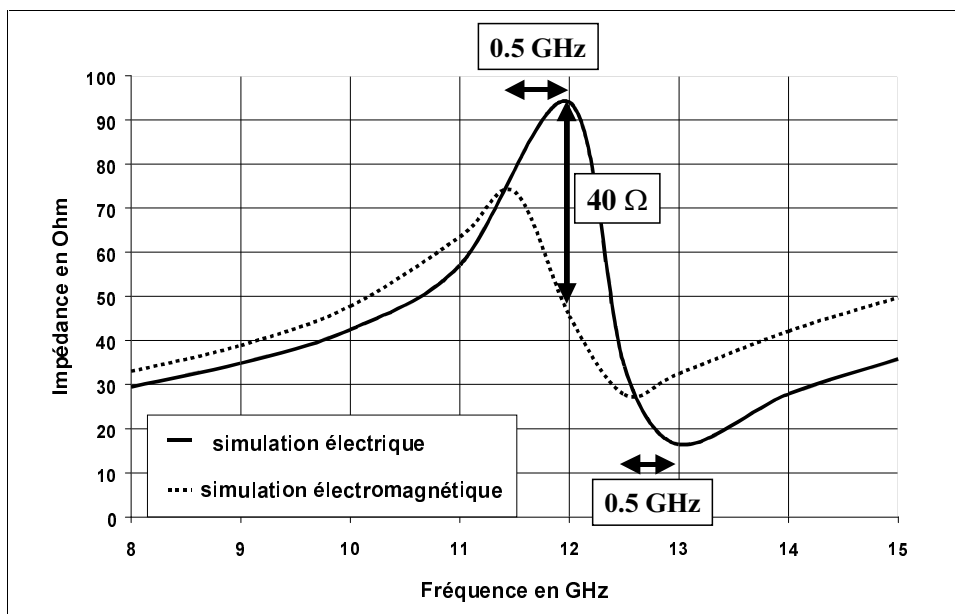


Figure IV-3 : Simulations électromagnétiques

Cette figure présente les impédances issues de la simulation électromagnétique, qui rend bien évidemment compte des couplages de proximité, et de la simulation électrique, qui n'en tient pas compte. Nous constatons, outre un faible décalage fréquentiel (4%) de 0.5 GHz entre les deux types de simulation, un décalage important (45 %) de la valeur de l'impédance de la cellule de sortie du mélangeur à la fréquence de travail f_{PI} de cet accès et dont il est manifestement fondamental de tenir compte pour l'évaluation des performances électriques globales du mélangeur.

Nous avons donc effectué une retro-simulation électrique de notre cellule de mélange en tenant compte, pour le filtre de sortie, de la simulation électromagnétique précédente. La Figure IV-4 présente le gain de conversion mesuré, simulé initialement sans tenir compte des effets de proximité et simulé en tenant compte de la simulation électromagnétique précédente (simulation corrigée).

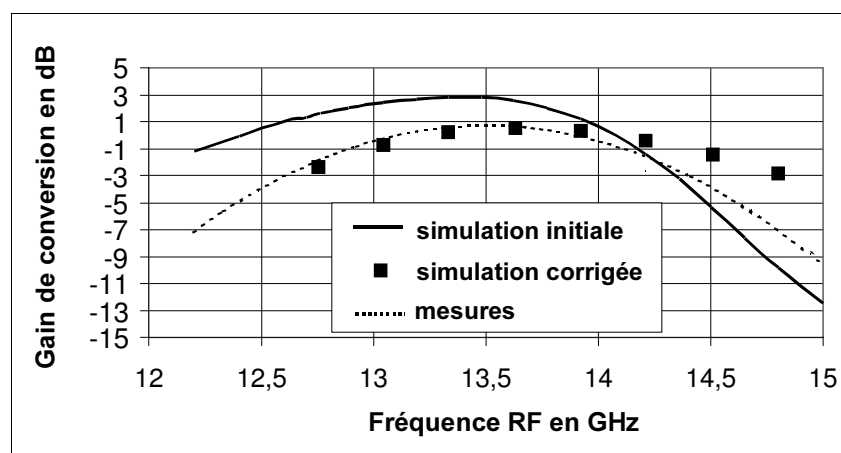


Figure IV-4 : Rétro-simulation du gain de conversion du mélangeur

Ainsi, nous constatons l'amélioration de nos simulations lors de la prise en compte des couplages électromagnétiques de notre filtre en sortie du mélangeur. Ceci démontre d'une part

la forte sensibilité de ce filtre sur les performances du mélangeur et d'autre part que l'écart initialement de 2 dB entre simulations et mesures était principalement du à la non-prise en compte des phénomènes de couplage électromagnétique entre éléments contigus, conséquences du trop fort degré d'intégration.

IV.2.2. Diviseur de puissance

Nous avons de plus caractérisé un coupleur conçu aux alentours de 12GHz et configuré en diviseur de puissance.

La Figure IV-5 présente les résultats de simulation et de caractérisation du gain différentiel et du taux de réjection du mode commun du circuit conçu. Notons toujours un écart de 2dB entre les valeurs de gain mesurées et simulées (simulation initiale). Cependant, la similitude du comportement fréquentiel de ces courbes prouve le fonctionnement correct du circuit.

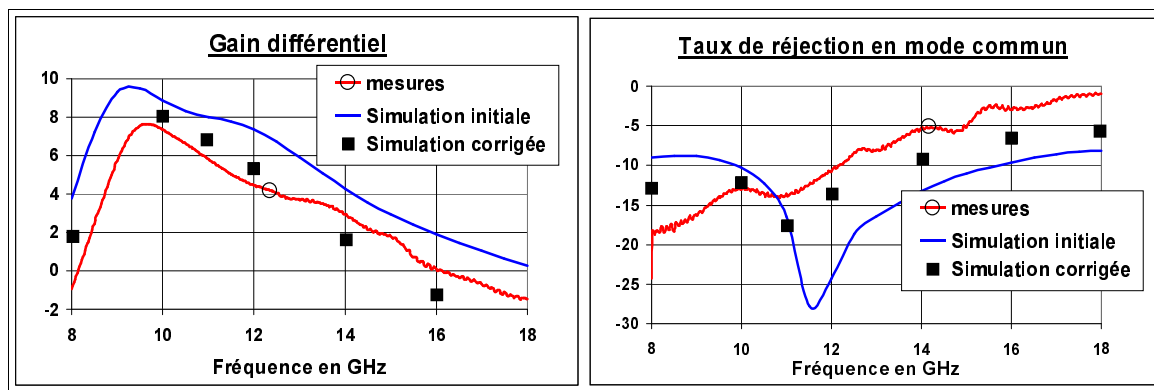


Figure IV-5 : Performances en gain et en TRMC du coupleur

Comme dans le cas du mélangeur, nous avons effectué la simulation électromagnétique d'un élément passif critique du diviseur de puissance : l'impédance Z_{MC} . La prise en compte résultante des effets de couplage des différents éléments de Z_{MC} a été ensuite introduite dans une simulation électrique du diviseur amenant au résultat de simulation corrigée présenté à la Figure IV-5. Nous aboutissons ainsi à de meilleures prédictions des performances électriques du diviseur de puissance. Cependant, les écarts résiduels entre simulations et mesures prouvent qu'il serait nécessaire de simuler électromagnétiquement d'autres parties du circuit afin d'affiner ces prédictions. Quoiqu'il en soit, ces rétro-simulations électromagnétiques et électriques démontrent, là encore, la trop forte densité d'intégration du MMIC.

La Figure IV-6 présente la mesure et la simulation du facteur de bruit ainsi que la détermination du point d'interception d'ordre 3. L'écart de 2dB sur la valeur du facteur de bruit est une conséquence de celui que nous avons observé sur le gain. Notons que la valeur simulée du facteur de bruit de 6dB est la conséquence d'une connexion de l'accès inutilisé à une résistance de 50Ω bruyante qui n'est pas dans la conception de la fonction globale.

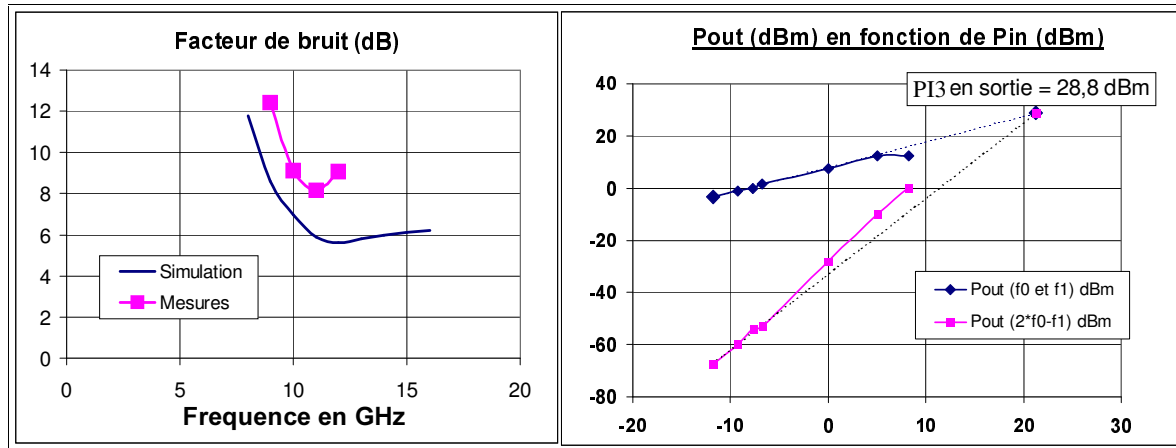


Figure IV-6 : Facteur de bruit et linéarité du coupleur

Enfin, la détermination expérimentale de P_{13s} , par la mesure des puissances appropriées du spectre en sortie du diviseur excité par deux fréquences d'entrée contiguës, démontre la bonne linéarité du circuit conçu car une valeur de $P_{13s}=28.8\text{dBm}$ est mesurée pour une fréquence d'entrée de 10GHz (25dBm en simulation).

Notons que, malgré les imperfections dues à une intégration trop compacte, les coupleurs conçus présentent des performances supérieures à celles obtenues avec des structures passives [III.10] notamment en terme de surface occupée et de gain. Les coupleurs passifs souffrent en effet du compromis insoluble équilibrage-surface, pour lequel les paires différentielles apportent une solution tout en présentant un gain supérieur à 0dB compatible avec les exigences imposées par l'étude système.

Ces caractérisations ont validé d'une part un fonctionnement correct des cellules mises en œuvre et d'autre part les principes de conception considérés.

Néanmoins, l'étude qualitative de ces résultats de mesure ainsi que la réalisation de rétro-simulations électromagnétiques ont révélé la présence de couplages électromagnétiques dus à une intégration trop dense. Ces couplages ont principalement entraîné la dégradation des valeurs de gain et de réjection du mode commun dans le cas des coupleurs, conséquences des décalages sur les valeurs des éléments passifs n'étant plus alors optimales.

Nous allons dans le paragraphe suivant présenter les résultats de simulation de la puce intégrant la fonction globale de transposition de fréquence.

IV.3. Cellule de mélange complète

Nous avons, dans un premier temps, mesuré le spectre du signal de sortie (Figure IV-7) du mélangeur doublement équilibré pour les conditions de fonctionnement optimales suivantes : $f_{RF}=13.85\text{GHz}$, $f_{OL}=2\text{GHz}$, $P_{OL}=0\text{dBm}$ et $P_{RF}=-30\text{dBm}$. Pour cette mesure, nous avons compensé les diverses pertes en puissance des câbles lors du réglage des puissances d'entrées P_{OL} et P_{RF} .

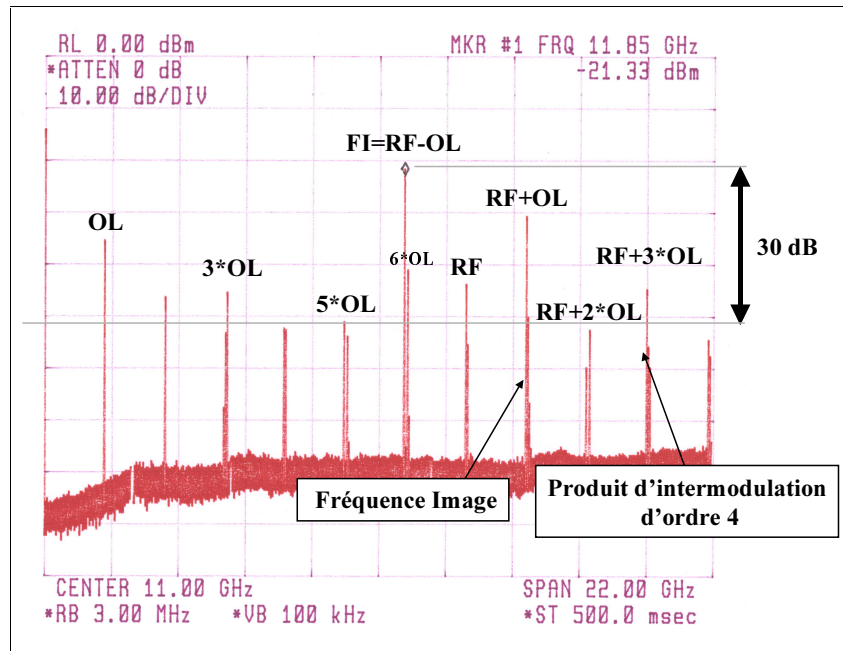


Figure IV-7 : spectre en puissance du signal de sortie

Qualitativement, ce spectre démontre le fonctionnement correct du mélangeur doublement équilibré car, d'une part la raie spectrale prépondérante correspond à la raie utile (de fréquence f_{FI}) ce qui démontre une translation de fréquence correcte, et d'autre part les autres raies spectrales, outre certaines que nous expliciterons ci-après, ont une puissance 30 dB inférieure à celle de la raie d'intérêt.

Quantitativement, l'observation de ce spectre nous amène aux remarques suivantes :

- La raie spectrale à la fréquence f_{FI} a une puissance de -21dBm démontrant un faible gain de conversion de 9dB au lieu de 23dB en simulation.
- Par rapport à la puissance P_{OL} injectée de 0dBm, les harmoniques OL du signal de sortie ont des niveaux de puissance de -35dBc à la fréquence f_{OL} , ce qui correspond à l'isolation de l'accès FI par rapport à l'accès OL, -47dBc à la fréquence $2 \times f_{OL}$, -44dBc à la fréquence $3 \times f_{OL}$ et -41dBc à la fréquence $6 \times f_{OL}$. La forte valeur de ces réjections démontre le bon fonctionnement du diviseur OL et du combineur FI à la fréquence f_{OL} et à ses harmoniques.
- La mesure de l'isolation de l'accès RF vers l'accès FI donne une valeur de 13dB, contre 35dB en simulation. Cet écart démontre un équilibrage médiocre soit du diviseur RF à la fréquence f_{RF} soit du combineur FI à la même fréquence conséquences de la dégradation du taux de réjection du mode commun précédemment observée (paragraphe IV.2.2).
- Enfin, les raies aux fréquences $f_{RF}+f_{OL}$ et $f_{RF}+3 \times f_{OL}$, non rejetées par la structure double équilibrée, ont des puissances remarquables. Si plus de réjection est nécessaire, un filtrage en sortie devra être envisagé.

Nous avons alors mesuré le gain de conversion en fonction de la puissance P_{OL} injectée (Figure IV-8) et de la fréquence f_{RF} (Figure IV-9).

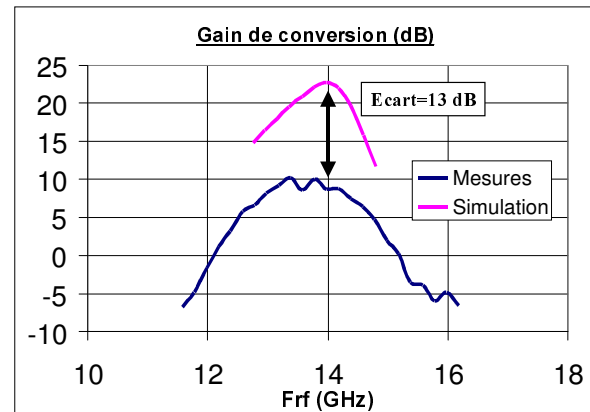
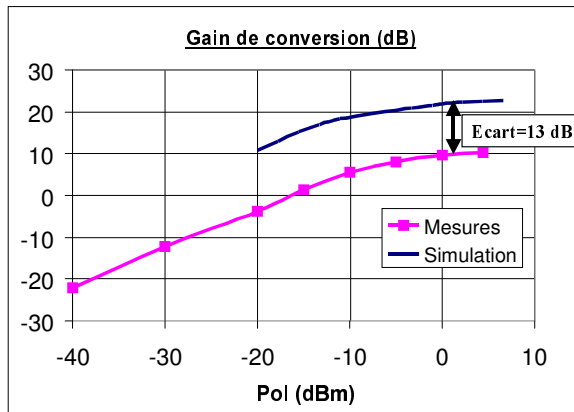


Figure IV-8 : Gain de conversion fonction de P_{OL} **Figure IV-9 : Gain de conversion fonction de f_{RF}**

La comparaison des données expérimentales et de simulation amène aux mêmes conclusions que celles tirées au paragraphe IV.2, à savoir un faible décalage fréquentiel et un décalage en gain notable, qui est alors de 13 dB. Cet écart important ne se justifie pas seulement du cumul des écarts des différents blocs mais d'une part de forts couplages électromagnétiques entre les interconnexions des différents blocs et d'autre part des interactions potentielles entre les cellules elles-mêmes.

De plus, nous avons vérifié l'absence de forte réflexion aux accès RF et FI. Les courbes de la Figure IV-10 montrent que les coefficients de réflexion mesurés aux accès RF et FI sont inférieurs à -10 dB sur les plages de fréquences utiles et ne sont donc pas la cause de la différence entre le gain simulé et mesuré.

Enfin, la Figure IV-11 montre la puissance en sortie du point d'interception d'ordre 3 en fonction de la puissance P_{OL} injectée. Malgré leur différence quantitative, ces courbes présentent la même allure, le décalage horizontal montre que la linéarité maximale est atteinte pour une puissance P_{OL} supérieure à celle prévue en simulation et le décalage vertical traduit une dégradation de la linéarité globale. Ces deux décalages ont pour cause la dégradation des différents gains des cellules constituant le mélangeur double équilibré.

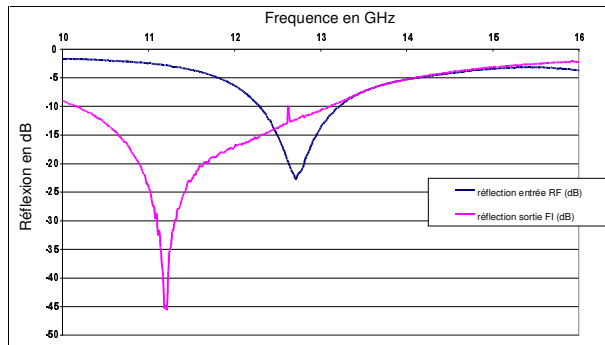


Figure IV-10 : Réflexions aux accès

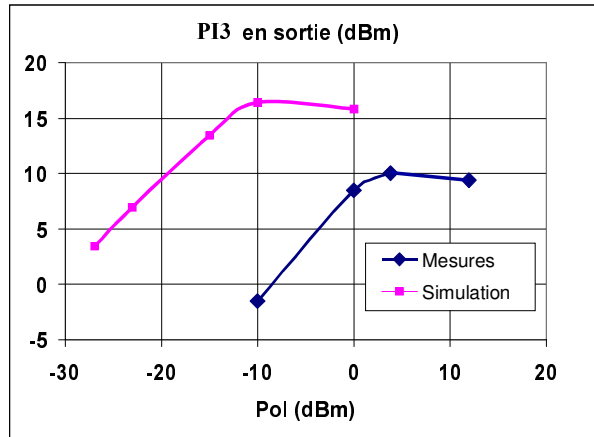


Figure IV-11 : P_{13s} fonction de P_{OL}

V. Conclusions

Nous avons, dans cette première partie du troisième chapitre, mené d'une part l'étude théorique des coupleurs actifs indispensables à l'équilibrage de structures et d'autre part l'intégration monolithique sur Arséniure de Gallium d'un mélangeur doublement équilibré.

Les diverses optimisations méthodiques mise en place lors des études théoriques des différentes cellules se sont traduites par une intégration monolithique performante puisque les caractéristiques simulées de la structure globale se sont avérées conformes au cahier des charges fixé et comparable, voir supérieure, à l'état de l'art [III.9, III.11].

Les caractérisations des structures ont, qualitativement, prouvées la validité des différents concepts mis en œuvre et quantitativement révélées des écarts avec les simulations électriques initiales. Des simulations électromagnétiques ont démontré la présence de couplages électromagnétiques parasites entre des éléments passifs du circuit trop rapprochés, ce qui s'est traduit par une différence marquée entre les performances simulées et mesurées. Ces rétro-simulations électromagnétiques et électriques ont ainsi prouvé le trop fort degré d'intégration du circuit, conséquence des contraintes de taille maximale imposée par le CMP ($2 \times 3 \text{ mm}^2$).

Les perspectives envisageables de ces travaux sont alors :

- Soit de tenir compte, ponctuellement aux endroits sensibles, des couplages par des simulations électromagnétiques. Les éléments passifs devraient être optimisés afin d'en retrouver les valeurs optimales et d'aboutir à une réalisation présentant les performances visées.
- Soit de réaliser une nouvelle intégration sur une puce de taille supérieure afin de garantir un éloignement entre éléments suffisant et de s'affranchir de simulations électromagnétiques lourdes tout en obtenant les performances escomptées.

REFERENCES BIBLIOGRAPHIQUES DU CHAPITRE 3 PARTIE A

- [III.1] : M. Soulard, J.F. Villemazet, E. Rogeaux, J.L. Cazaux, "Experience and prospective of MMIC's for space programmes". GAAS 1998, pp.592-597.
- [III.2] : D. Prieto, « Conception et caractérisation de circuits intégrés micro-ondes monolithiques (MMICs) en technologie d'interconnexions uniplanaires. Application à la conception d'un convertisseur de fréquences en bande Ku. », Thèse de doctorat de l'université Paul Sabatier-Toulouse, 1999.
- [III.3] : **D. Dubuc**, T. Parra, J. Graffeuil, "Conception d'un mélangeur actif en bande Ku" 3^{ème} Journée Micro-ondes et Electromagnétisme de Toulouse, Janvier 2000.
- [III.4] : S.A. Maas, « Nonlinear microwave circuits », Artech House, 1988.
- [III.5] : J.-L. Gautier, « Etude des amplificateurs différentiels dans la gamme micro-onde, applications », thèse de doctorat de l'université de paris sud, 1977.
- [III.6] : T. KHELIFI, « Etude, conception et applications de structures différentielles à transistors bipolaires à hétérojonction », Thèse de l'institut national polytechnique de Grenoble, 1996.
- [III.7] : ED02AH Design Manual V2.0CD, Philips Microwave Limeil, April 1999.
- [III.8] : P.S. Tsenes, G. E. Stratakis, N.K. Uzunoglu, M. Lagadas, G. Deligeorgis : « An X-band MMIC Down-Converter ». WOCSDICE 2000 , pp X.7-X.8
- [III.9] : JF Villemazet, J. Dubouloy, M. Soulard, JC Cayrou, E. Husse, B.Cogo, JL Cazaux, « New compact double balanced monolithic down-converter. Application to a single chip MMIC receiver for satellite equipment », EuMC 2001.
- [III.10] : C. Boyavalle, « Conception de récepteurs a faible bruit dans le domaine millimétrique en étudiant le bruit électrique dans les circuits non linéaires micro-ondes. », Thèse de doctorat de l'université des sciences et technologie de Lille, 1997.
- [III.11] : JF Villemazet, E. Rogeaux, D. Roques, JC Cayrou, B.Cogo, M. Soulard, JL Cazaux, « Novel compact double balanced coplanar active mixer. Application to a single chip MMIC receiver for satellite repeater », EuMC 2001.

CHAPITRE 3 PARTIE B

Mélangeur simple **sur BiCMOS-SiGe en bande Ka**

Le développement des applications grand public dans la gamme des fréquences millimétriques, telles que le LDMS et les applications multimédia, a bouleversé le marché des circuits monolithiques millimétriques. Deux mouvements concurrents se sont alors initiés :

- Les fabricants de circuits sur GaAs et InP utilisables à ces fréquences ont diminué leur coût de production en augmentant notamment la taille des *wafers* fabriqués.
- Les fabricants de circuits sur silicium, matériau indétrônable en terme de coût de production, tendent à améliorer leurs technologies afin d'en augmenter les fréquences d'utilisation.

C'est dans ce second objectif que la société ST Microelectronics développe la fabrication de transistors bipolaires à hétérojonction (TBH) Si-SiGe fonctionnant aux fréquences millimétriques et s'intégrant parfaitement dans le processus de fabrication des circuits BiCMOS alliant ainsi la grande maturité et le faible coût de cette technologie avec les hautes performances des TBH Si-SiGe.

Notre travail, pour cette technologie, a porté sur l'intégration monolithique de notre cellule de mélange en bande Ka (20-30GHz). Nous présenterons tout d'abord les éléments actifs et passifs de cette technologie en précisant leurs limites fréquentielles. Nous exposerons ensuite l'étude et l'optimisation que nous avons conduites pour la réalisation d'inductance fonctionnant aux fréquences millimétriques. Enfin, nous présenterons la conception et l'intégration de la cellule de mélange envisagée en bande Ka, les performances simulées ainsi que les premières caractérisations menées sur ces circuits.

I. La technologie BiCMOS-SiGe de ST-Microelectronics

Nous présentons dans ce paragraphe la technologie monolithique fondée sur les composants actifs à hétérojonction Si-SiGe et sur les divers éléments passifs issus de la technologie BiCMOS et quelque peu améliorés pour fonctionner aux fréquences basses micro-ondes. Puis nous présenterons la première partie de nos travaux qui a consisté à concevoir des inductances performantes, fonctionnant aux fréquences millimétriques, indispensables à toute conception monolithique dans cette gamme.

I.1. Présentation de la technologie

Nous allons, tout d'abord, présenter les divers composants passifs et actifs proposés par le fondeur au travers d'une bibliothèque constituée de modèles électriques et de dessins de masques.

I.1.1. Les composants actifs

La technologie BiCMOS Si-SiGe repose essentiellement sur les performances fréquentielles des TBH SiGe qu'elle propose. D'un concept simple, inventé par Herbert Kroemer [III.1] en 1957 et largement développée dans la littérature [III.2], les TBH ont réellement vu leur développement dans les laboratoires d'IBM dans les années 1980. Depuis, les diverses avancées technologiques ont permis l'utilisation des ces transistors dans la gamme des fréquences millimétriques conduisant les concepteurs de circuits monolithiques à s'intéresser à ces technologies pour les nouvelles applications principalement envisagées dans ces bandes fréquentielles.

Actuellement cinq types de TBH sont proposées par la technologie BICMOS6G de ST Microelectronics. Le Tableau I-1 décrit, pour le courant de polarisation de collecteur correspondant au maximum de la fréquence de transition, les caractéristiques de chaque type.

$N \times (W \times L) \mu m^2$	$I_c (mA)$	$F_{max} (GHz)$	$F_t (GHz)$	$P_{cldB} (dBm)$	$NF (dB)$	$G_{max} (dB)$
0.4×0.8	0.37	50	44	-18	5.5	6
0.4×2	1.05	62	48	-15	6.5	8
0.4×10	5.96	75	52	-6	7	10
3×(0.4×20)	31.4	94	53	9	6.8	11
5×(0.4×20)	46.3	98	53	15	6.9	11

à 30 GHz

Tableau I-1 : Classification des transistors de ST Microelectronics

Nous avons opté pour le transistor à 3 doigts de base de largeur $0.4\mu m$ et de longueur $20\mu m$ ($3 \times (0.4 \times 20)$). En effet, les trois premiers transistors de ce tableau, de tailles réduites, présentent des performances en linéarité trop médiocres pour être retenus, tandis que le dernier consomme un fort courant sans amélioration notable des caractéristiques.

Ayant choisi le transistor, nous allons maintenant présenter le catalogue des divers éléments passifs dont nous pouvons envisager l'intégration à 30GHz sachant qu'ils ne sont actuellement spécifiés par le fondeur que jusqu'à 10 GHz.

I.1.2. Les composants passifs

La conception des éléments passifs réactifs (lignes, inductances, capacités) s'appuie sur les différentes couches métalliques et d'oxyde déposées au dessus de la plaque de silicium sur laquelle sont intégrés les composants actifs.

Nous allons donc tout d'abord présenter ces différents niveaux de métallisation puis nous expliciterons en détail les capacités et inductances proposées par le fondeur.

a) Les différents niveaux métalliques

Dans les technologies sur GaAs ou InP, le semi-conducteur est fortement isolant et peut donc servir de milieu diélectrique dans lequel l'extension des lignes de champ électromagnétique ne s'accompagne pas de pertes prohibitives.

Par contre, le substrat de silicium, fortement conducteur ($15\Omega\cdot\text{cm}$ pour les substrats BiCMOS et $0.1\Omega\cdot\text{cm}$ pour les substrats CMOS), ne peut pas constituer un milieu faibles pertes pour la propagation des champs électromagnétiques. Les éléments passifs sont donc conçus sur des niveaux diélectriques déposés sur le substrat et parfois même isolés de celui-ci par des couches conductrices ou semi-conductrices fortement dopées.

La technologie de ST Microelectronics comporte cinq niveaux d'interconnexions métalliques dont l'étude qualitative des résistances carrées conduit aux règles de conception suivantes :

- Le niveau 5 (niveau supérieur) est à préférer pour les interconnexions et la réalisation de passifs car il présente la résistance carrée la plus faible (car il est d'épaisseur la plus grande).
- Les niveaux 3 et 4 servent aux croisements entre les interconnexions.
- Le niveau 2 est typiquement dédié au plan de masse pour le circuit et notamment pour les interconnexions micro-rubans.
- Le niveau 1 est à éviter pour les interconnexions car il présente une résistance carrée trop élevée.

Nous allons au prochain paragraphe présenter les performances fréquentielles des capacités proposées par ST Microelectronics.

b) Les capacités

Ce paragraphe a pour but la détermination des valeurs de capacité utilisables pour des conceptions dans le domaine millimétrique.

La Figure I-1 présente la fréquence de résonance mesurée pour trois valeurs de capacités qui ont été réalisées et testées. La zone non-grisée représente le lieu de fonctionnement nominal pour lequel la valeur de la capacité dévie seulement de 10% de sa valeur basse fréquence. Ainsi, pour une application à 30GHz, les capacités de valeur supérieure à 1.3pF ne doivent pas être utilisées.

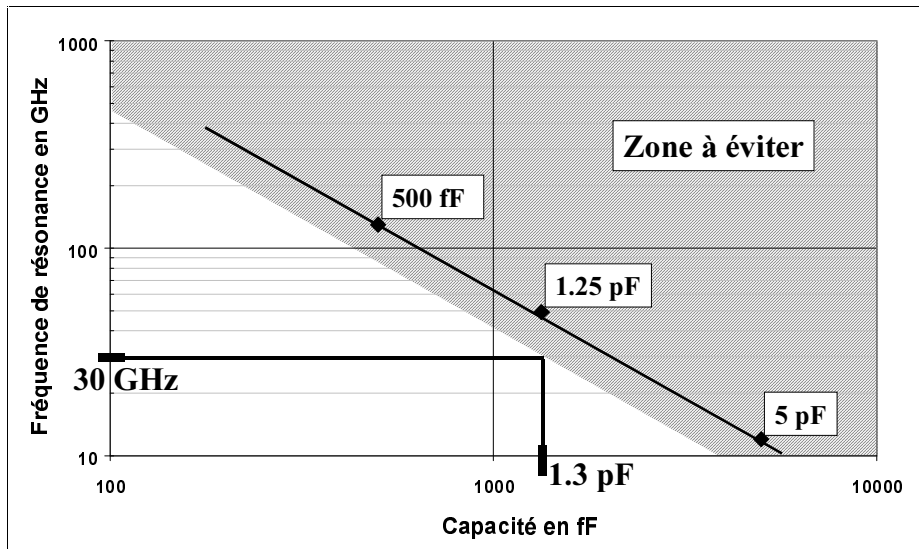


Figure I-1 : La fréquence de résonance des capacités

Les inductances sont généralement les éléments les plus critiques pour la qualité des performances d'une fonction intégrée, comme nous l'avons déjà signalé dans la partie A de ce chapitre.

Nous allons, dans le prochain paragraphe, décrire les inductances proposées par le fondeur et montrer que ces composants ne peuvent satisfaire à une application en bande millimétrique.

c) Les inductances

En hyperfréquences, deux solutions sont possibles pour réaliser des éléments inductifs suivant les valeurs désirées. Les fortes valeurs d'inductance seront réalisées par des inductances localisées intégrées de forme spirale, octogonale ou rectangulaire. Pour les faibles valeurs, des tronçons de ligne à caractère inductif de faible longueur (inférieur à environ $1/20^{\text{ième}}$ de la longueur d'onde) et de forte impédance caractéristique (la largeur de la bande doit être faible) seront utilisés pour réaliser des inductances semi-distribuées.

Nous allons évaluer les performances de trois familles d'inductances proposées par le fondeur:

- des inductances octogonales classiques,
- des inductances semi-distribuées réalisées par des tronçons de ligne micro-bande utilisant les niveaux métalliques 5 et 2,
- des inductances possédant un blindage au niveau du substrat semi-conducteur (patterned ground shield).

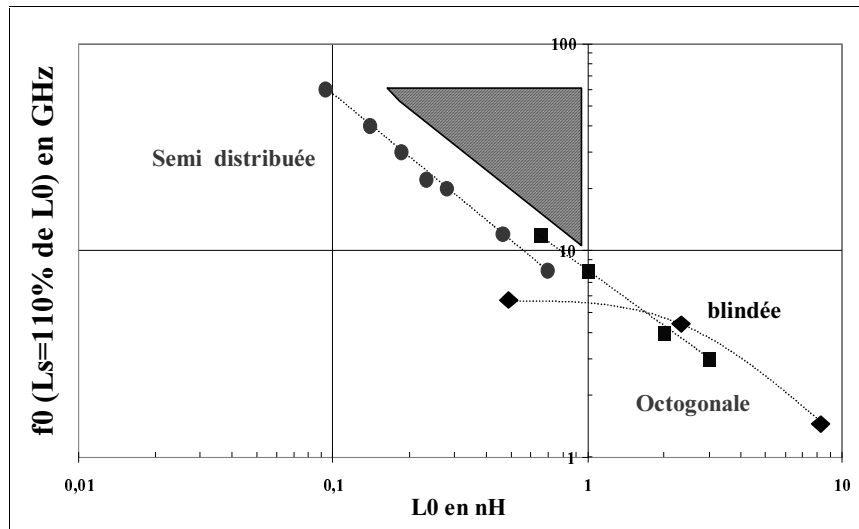


Figure I-2 : Fréquence de fonctionnement nominal des inductances

La Figure I-2 présente la fréquence pour laquelle la valeur de l'inductance dévie de 10% de sa valeur basse fréquence en fonction de cette dernière et la Figure I-3, la valeur du coefficient de qualité à cette fréquence.

Ces courbes permettent de constater que les différentes familles d'inductance couvrent des domaines d'utilisation différents. Ainsi, les fréquences d'utilisation des inductances semi-distribuées sont élevées mais leur coefficient de qualité reste faible tandis que celui des inductances blindées est nettement amélioré au prix de fréquences d'utilisation inférieures.

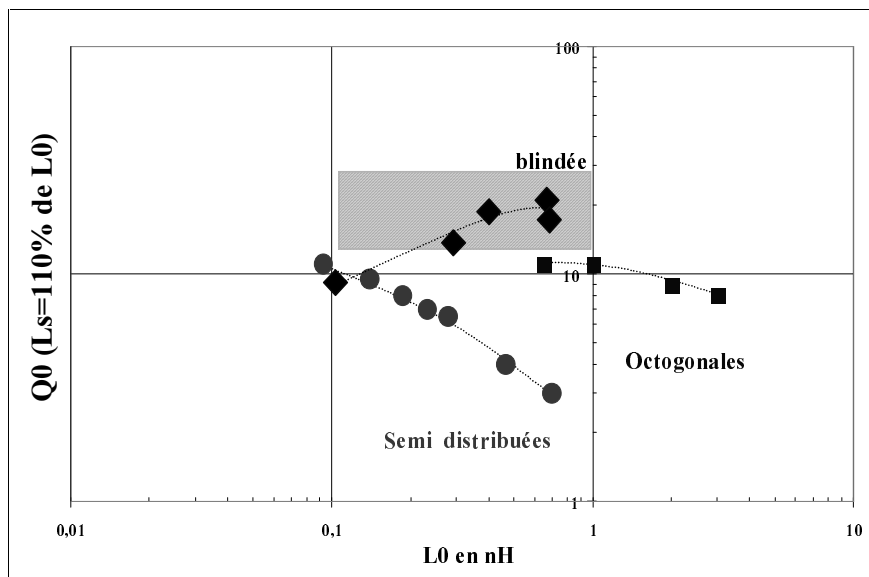


Figure I-3 : Coefficient de qualité au fonctionnement nominal des inductances

Pour une fréquence de fonctionnement de 30GHz, seules les inductances semi-distribuées sont donc utilisables et limitent les valeurs de réalisation à 0,2nH. De plus, la valeur de leur coefficient de qualité reste inférieure à 10, ce qui est très insuffisant.

En conclusion, les inductances mises à disposition par le fondeur ne conviennent pas aux conceptions millimétriques qui nécessitent des inductances dont les caractéristiques

devraient idéalement se situer dans les zones grisées des figures I-2 et I-3. Nous avons donc entrepris la conception d'inductances pouvant convenir aux fréquences millimétriques.

I.2. Inductance dans la gamme millimétrique

Les conceptions monolithiques dans le domaine millimétrique (30 à 100GHz) se heurtent simultanément à trois limites dues aux éléments passifs :

- Les éléments distribués restent de taille prohibitive pour une intégration compacte en dessous de 50GHz.
- Les éléments semi-distribués ne permettent pas d'accéder à de fortes valeurs d'inductance et de coefficient de qualité.
- Les inductances localisées *classiques* présentent des performances fréquentielles seulement acceptables dans le domaine centimétrique.

Nous avons donc mené une étude visant la conception d'une inductance localisée fonctionnant dans le domaine millimétrique qui puisse ainsi répondre à notre besoin pour l'intégration monolithique d'un mélangeur 30GHz / 20GHz.

I.2.1. Le dimensionnement

Le dimensionnement d'une inductance dépend fortement du domaine fréquentiel visé. Ainsi, suivant le domaine, les conclusions peuvent être opposées.

La Figure I-4 présente la définition des paramètres géométriques à optimiser pour la réalisation d'inductance carrée dont les caractéristiques sont sa valeur L , sa fréquence de résonance (F_{res}) et son coefficient de qualité à 30 GHz Q .

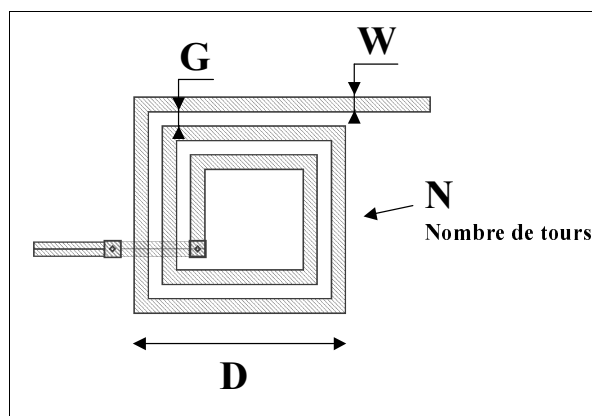


Figure I-4 : Paramètres géométriques d'une inductance carrée

L'utilisation d'un modèle électrique équivalent analytique d'inductance [III.3] nous a permis de déterminer la sensibilité des caractéristiques d'une inductance en fonction de ses paramètres géométriques. Nous n'avons pas cherché à préciser finement ces valeurs mais plutôt à les définir qualitativement.

Le Tableau I-2 résume ces investigations effectuées pour une inductance dont les valeurs initiales des paramètres géométriques ont été choisies de façon empirique afin de se rapprocher au mieux des contraintes fixées: $N=2.5$, $W=10\mu\text{m}$, $G=10\mu\text{m}$ et $D=100\mu\text{m}$. Les performances estimées de cette inductance sont une fréquence de résonance de 60 GHz et une valeur maximale du coefficient de qualité de 12 à 15 GHz.

	L	Fres	Q(30 GHz)
N	500 pH/tours	- 4 GHz/tours	-10/tours
W	- 92 pH/ μm	- 8 GHz/ μm	-2.5 μm^{-1}
G	- 63 pH/ μm	3 GHz/ μm	0.3 μm^{-1}
D	10 pH/ μm	- 0.58 GHz/ μm	-0.07 μm^{-1}

Tableau I-2 : Sensibilités des paramètres sur les caractéristiques d'une inductance

Ce tableau permet un certain nombre de remarques :

- La valeur de l'inductance sera ajustée grossièrement par le nombre de tours de l'enroulement étant donnée la forte influence de ce paramètre sur cette caractéristique.
- La valeur du nombre de tours est à minimiser car elle influe fortement sur les valeurs de la fréquence de résonance et du coefficient de qualité.
- La valeur de D n'a que peu d'influence sur les valeurs de la fréquence de résonance et du coefficient de qualité et servira donc pour ajuster finement la valeur de l'inductance.
- La valeur de W sera à minimiser afin d'obtenir une inductance à fort coefficient de qualité et à haute fréquence de résonance. Notons que cette minimisation de W est d'autant plus souhaitable qu'elle s'accompagne d'une augmentation bénéfique de la valeur de l'inductance.
- Enfin, la valeur de G n'a d'influence importante que sur les valeurs de l'inductance et de sa fréquence de résonance. Sa valeur découlera alors du compromis entre l'obtention d'une forte valeur de l'inductance et d'une forte valeur du coefficient de qualité.

Le respect de ces points, nous a amené à choisir les paramètres géométriques suivants : $W=5\mu\text{m}$, $G=5\mu\text{m}$, $D=75\mu\text{m}$ et $N=2.5$ tours conduisant ainsi à une inductance de valeur estimée à $L=0.7\text{nH}$ [III.4] (en tenant compte des lignes d'accès à l'inductance). A partir du modèle électrique équivalent, nous avons de plus estimé la fréquence de résonance à 80GHz et le coefficient de qualité à 15 pour une fréquence de 30GHz, ce qui constitue une amélioration notable par rapport aux performances initiales. Nous avons néanmoins cherché à améliorer ce coefficient de qualité ce qui passe par la réduction des pertes.

1.2.2. Les différents types de pertes

Trois types de pertes entrent en ligne de compte dans les inductances intégrées :

1. Les pertes métalliques qui traduisent l'effet joule dans le conducteur réalisant l'inductance.
2. Les pertes diélectriques, résultantes des pertes par conduction et par relaxation dues au champ électrique dans les diélectriques à pertes [III.3].
3. Les pertes par courant de Foucault. Elles sont dues au champ magnétique enlaçant un substrat conducteur et générant des courants tourbillonnants [III.3].

Le premier type de perte est incontournable et la seule solution pour le minimiser est d'augmenter la conductivité du matériau conducteur ($\sigma_{Al}=2.7 \times 10^7$ S/m $\rightarrow$ $\sigma_{Ag}=4.1 \times 10^7$ S/m $\rightarrow$ $\sigma_{Cu}=5.6 \times 10^7$ S/m). Le coefficient de qualité résultant est proportionnel à la fréquence puis à sa racine carrée lors de l'apparition de l'effet de peau dans le conducteur.

Les deux autres types de pertes dépendent à la fois de la conductivité du substrat et de la fréquence des ondes électromagnétiques. La Figure I-5 présente les quatre modes apparaissant dans les structures passives des technologies BiCMOS, pour lesquelles le ruban métallique est séparé du silicium par une couche d'oxyde, en fonction de la fréquence [III.5].

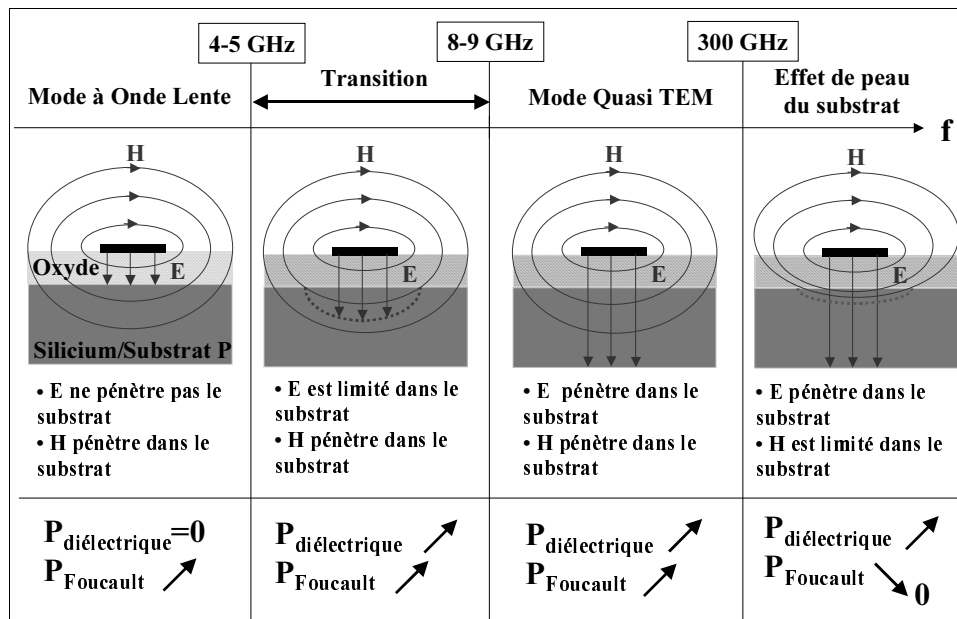


Figure I-5 : Les différents modes d'un substrat MIS

Cette figure montre qu'aux fréquences millimétriques, les deux types de pertes, diélectriques et de Foucault, apparaissent dans le substrat BiCMOS et que la propagation s'effectue alors dans le mode quasi-TEM. La Figure I-6 présente alors le coefficient de qualité des différents type de pertes en fonction de la fréquence, ainsi que le coefficient de qualité total.

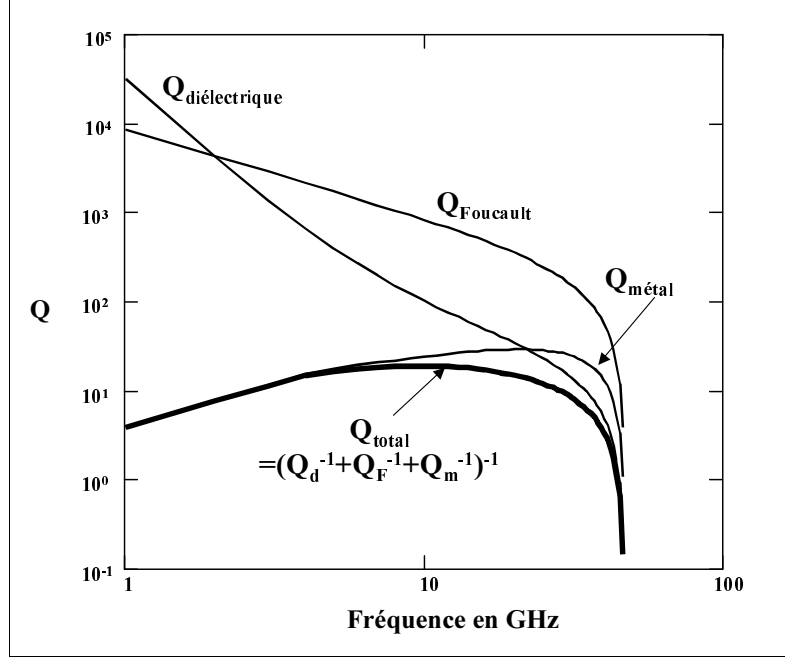


Figure I-6 : Facteurs de qualité liés aux différents types de pertes

Cette figure permet de relever, en fonction de la fréquence, le type de pertes qui sera à l'origine de la limitation des performances. Pour des fréquences jusqu'à 8-9GHz, ce sont les pertes métalliques qui dominent, puis de 10 à 30GHz, les pertes diélectriques apportent aussi leur contribution et enfin, au-delà de 30GHz, la coexistence des trois types de pertes se traduit par une chute brusque du coefficient de qualité.

Sur la base de ces constats, nous avons mis en œuvre des solutions en vue d'améliorer le coefficient de qualité total de l'inductance aux fréquences millimétriques.

Ce travail fait l'objet du prochain paragraphe.

I.2.3. Réalisation d'une inductance millimétrique

Une solution afin de limiter les pertes dans le substrat est de rajouter une couche semi-conductrice à la surface du substrat BiCMOS [III. 6] présentant les caractéristiques suivantes :

- Sa conductivité (σ) doit être élevée afin d'obtenir un mode à onde lente aux fréquences de travail. Dans ces conditions, le champ électrique ne pénétrera pas cette couche et par voie de conséquence ne pénétrera pas non plus le substrat silicium, annulant théoriquement ainsi les pertes diélectriques. Cette condition s'exprime par [III.5] :

$$\sigma \gg \frac{t}{h_{ox}} 2\pi \epsilon_{ox} f_{MAX} \quad (I-1)$$

où σ et t sont respectivement la conductivité et l'épaisseur de la couche semi-conductrice rajoutée ; h_{ox} et ϵ_{ox} sont respectivement l'épaisseur d'oxyde entre le métal et le substrat et sa permittivité; et enfin, f_{MAX} est la fréquence maximale d'utilisation de l'inductance conçue.

- Son épaisseur (t) doit être suffisamment faible afin que l'effet de peau, limitant la pénétration du champ magnétique, n'apparaisse pas aux fréquences d'intérêt [III.5] :

$$t \ll \sqrt{\frac{1}{\sigma_{f_{MAX}} \pi \mu_0}} \quad (I-2)$$

où μ_0 est la perméabilité du vide.

Cette dernière condition est primordiale pour la réalisation d'inductance de bonnes performances. En effet une couche, ne permettant pas la pénétration du champ magnétique, placée au voisinage d'une inductance intégrée, aura des effets désastreux. D'une part elle diminuera fortement la valeur de l'inductance car les lignes de champ magnétique seraient alors fortement perturbées, et d'autre part, elle générera des courants magnétiques diminuant fortement le coefficient de qualité. C'est la raison qui justifie la non-utilisation du métal 1 ou 2 en dessous de l'inductance qui respecte certes la condition (I-1) mais certainement pas celle en (I-2).

La couche semi-conductrice ainsi employée est la couche Népi, disponible dans les filières BiCMOS, ayant une épaisseur de 0.6 μ m et une résistivité de 0.3 Ω .cm. Cette couche s'intercale entre le substrat et l'oxyde remplissant parfaitement le rôle demandé. Ses caractéristiques satisfaisant aux conditions (I-1) et (I-2) jusqu'à 500GHz.

Des simulations électromagnétiques effectuées avec le logiciel Ansoft-HFSS, présentées à la Figure I-7, ont montré qualitativement que le champ électrique ne pénétrait pas cette couche Népi tandis que le champ magnétique n'était pas perturbé par cette dernière.

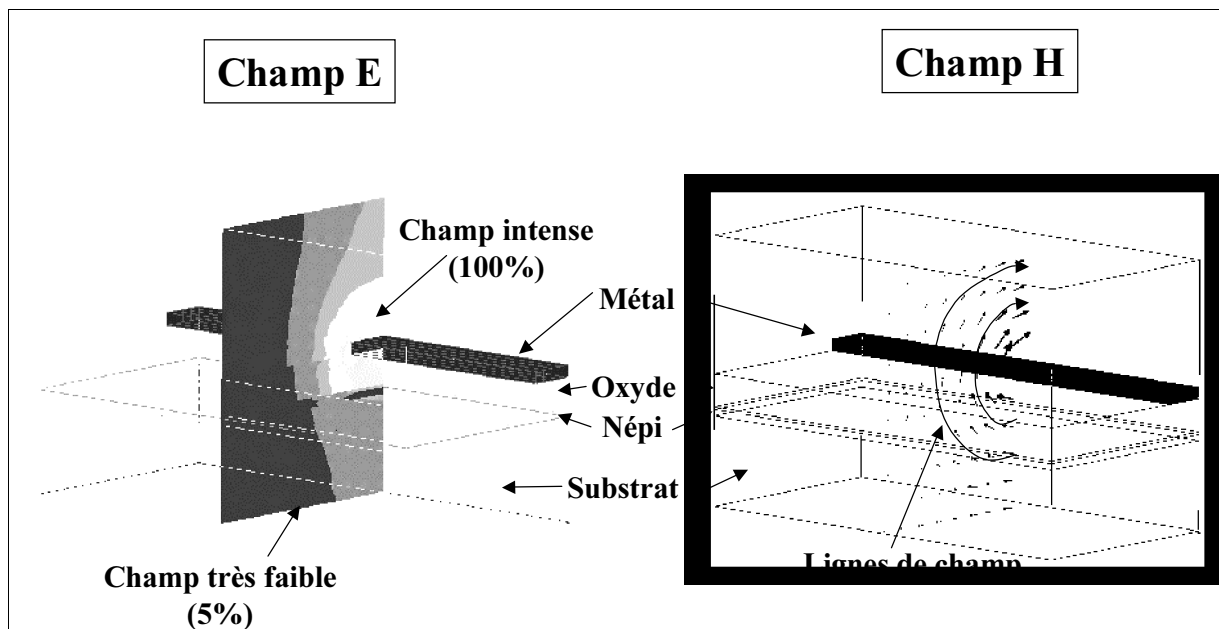


Figure I-7 : Simulation électromagnétique

L'inductance ainsi réalisée fonctionne, aux fréquences millimétriques, dans le mode à onde lente, dans lequel les pertes diélectriques sont nulles, ou fortement négligeables. Par rapport au substrat massif, le gain est appréciable dans la gamme des fréquences comprises entre 10GHz et 30GHz où les pertes diélectriques prédominent.

Nous avons donc réalisé l'inductance géométriquement définie au paragraphe I.2.1 avec et sans la couche Népi. La Figure I-8 compare les mesures aux simulations, effectuées à partir du modèle électrique équivalent, pour les deux configurations de substrat.

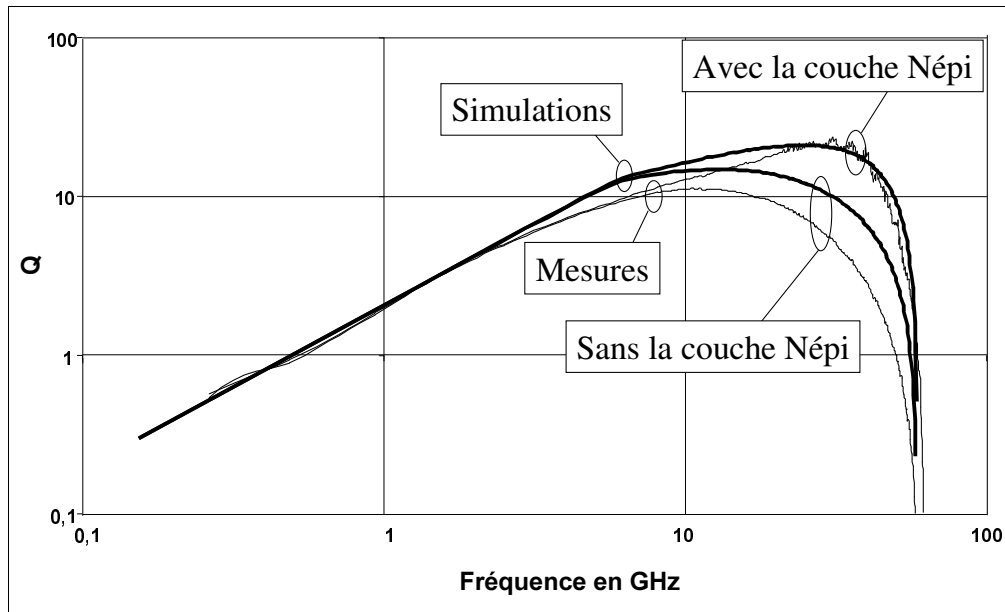
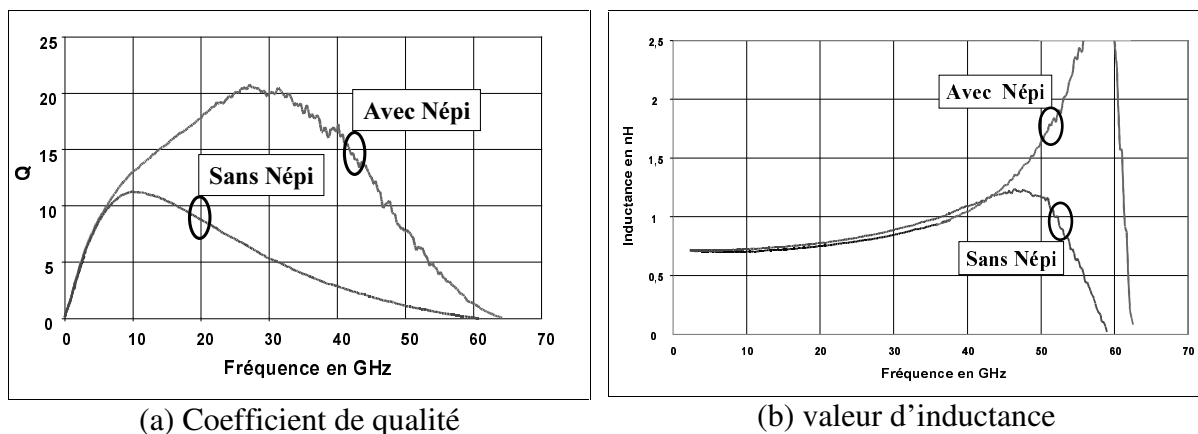


Figure I-8 : Coefficient de qualité mesuré avec et sans la couche Népi

Les résultats présentés sur cette figure valident le modèle électrique équivalent mis en œuvre pour les différentes optimisations et permettent de vérifier toute la théorie concernant la minimisation des pertes dans les inductances. Qualitativement, ces mesures prouvent, en effet, que l'ajout de la couche Népi supprime les pertes diélectriques et que le coefficient de qualité suit jusqu'à 30 GHz une loi en racine carrée de la fréquence imposée par les pertes métalliques. Au-delà de 30GHz, les pertes par courant de Foucault, inévitables, font chuter fortement le coefficient de qualité.

Les validations des différentes théories étant faites, nous allons maintenant résumer l'essentiel des caractéristiques de l'inductance réalisée. Les graphiques de la Figure I-9 permettent de constater le gain appréciable apporté par la présence de la couche Népi, notamment aux fréquences millimétriques qui nous intéressent.



(a) Coefficient de qualité

(b) valeur d'inductance

Figure I-9 : Résultats de caractérisation

L'inductance, optimisée pour fonctionner aux fréquences millimétriques, présente une fréquence de résonance de 62 GHz (cf. Figure I-9 (b)), rendant son utilisation largement possible dans la gamme des fréquences d'intérêts 20-30GHz. Le gain apporté par la couche Népi s'observe sur le coefficient de qualité présenté sur la Figure I-9 (a). L'inductance optimisée présente un coefficient de qualité de 20 à 30GHz contre 6 sans la couche Népi soit une amélioration de 230%.

Afin de positionner notre travail, nous avons comparé les caractéristiques de l'inductance que nous venons de présenter par rapport à l'état de l'art des inductances sur Silicium et sur Arséniure de Gallium.

Sur la Figure I-10, nous avons tout d'abord comparé la fréquence de résonance de notre inductance avec certains résultats faisant référence dans le domaine et concernant des inductances sur substrat BiCMOS [III.7,III.8] et sur GaAs [III.9].

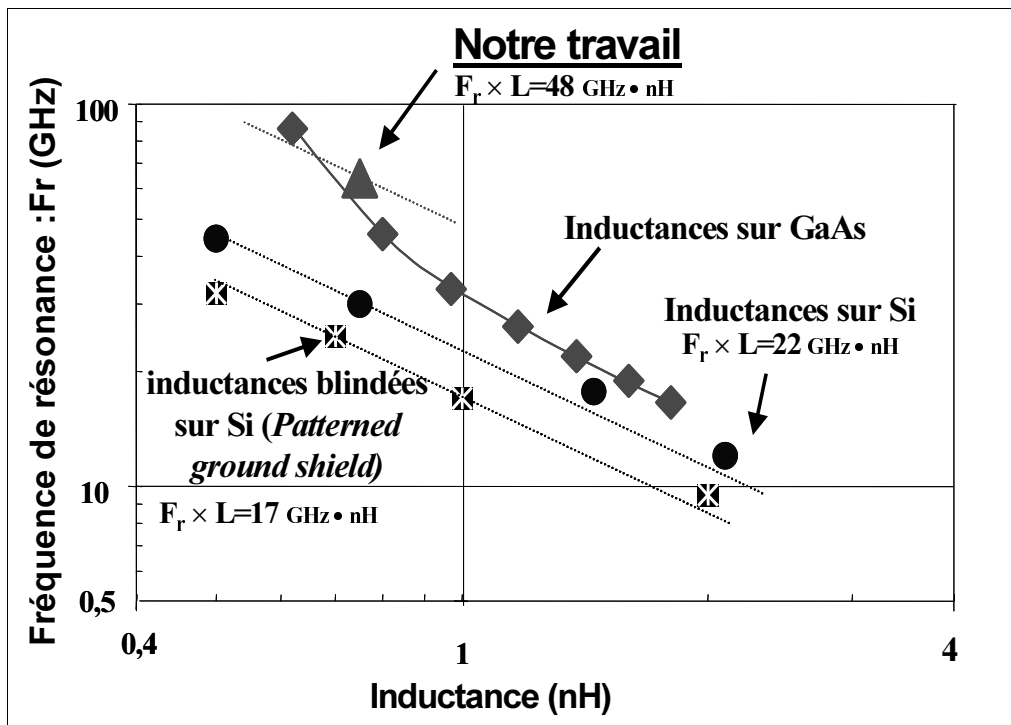


Figure I-10 : Fréquence de résonance de notre inductance

La fréquence de résonance étant inversement proportionnelle à la valeur de l'inductance réalisée, nous avons introduit le paramètre normalisé $F_{\text{res}} \times L$ afin de rendre possible une comparaison. Notre inductance présente une valeur du paramètre $F_{\text{res}} \times L$ de 48 GHz×nH contre 17 à 22 pour les inductances classiques sur silicium soit une amélioration d'environ 150%. Notons que l'inductance conçue surpasse même celles de la technologie sur GaAs.

Enfin, la Figure I-11 présente le coefficient de qualité maximum des inductances précédemment citées [III.7, III.8, III.9]. Nous constatons que le coefficient de qualité de notre inductance est bien supérieur à celui d'autres réalisations. Seules les inductances blindées sur silicium atteignent de telles valeurs de coefficient de qualité mais au détriment de leur

fréquence de résonance qui sont inférieures à 30GHz (voir la Figure I-10) rendant ces inductances inutilisables pour des réalisations dans la gamme millimétrique.

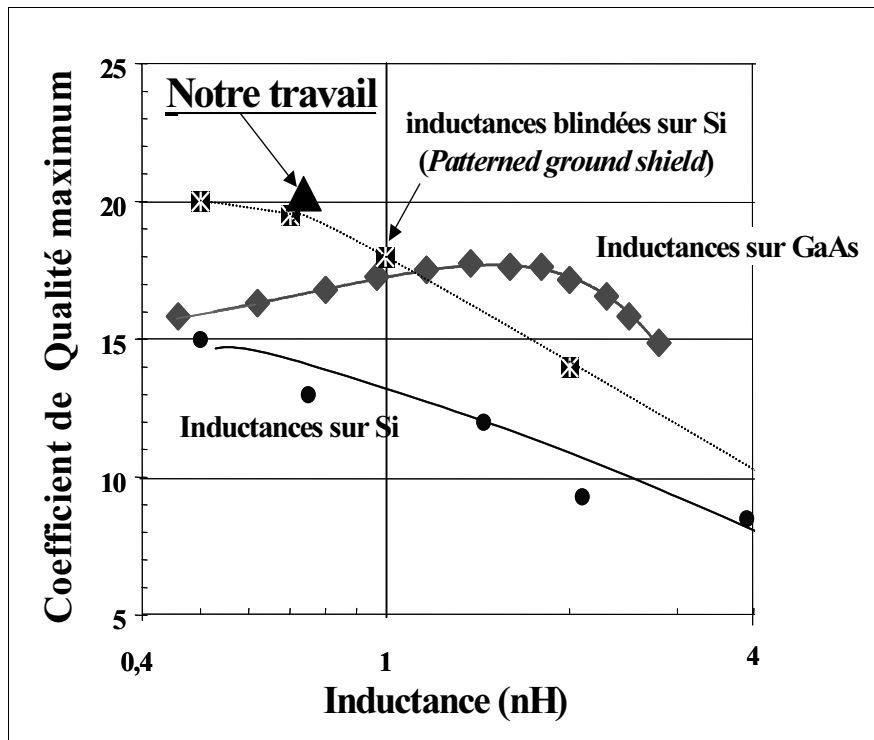


Figure I-11 : Coefficient de qualité de notre inductance

L'inductance ainsi optimisée présente des performances à l'état de l'art [III.3, III.7, III.8] tant au niveau de sa fréquence de résonance qui vaut 62GHz que sur le plan de son coefficient de qualité qui atteint une valeur de 20 à 30GHz. Notons que ces performances n'étaient jusqu'alors accessibles que par les techniques de micro-usinage [III.10].

I.3. Conclusion

Nous avons, dans ce premier paragraphe, présenté une inductance parfaitement adaptée à la conception monolithique en bande Ka. Ses performances mesurées dépassent largement celles des inductances publiées et se situent ainsi à l'état de l'art des réalisations d'inductances à ces fréquences [III. 6], notamment en terme de coefficient de qualité [III.11] que seules les inductances micro-usinées dépassent.

Disposant ainsi de composants actifs et passifs fonctionnant aux fréquences millimétriques, nous avons pu envisager l'intégration monolithique de la cellule de mélange 30GHz / 20GHz. Ce travail est décrit au prochain paragraphe.

II. Intégration de la cellule de mélange

Le but de cette conception étant l'évaluation de la technologie SiGe pour les applications en bande Ka, nous n'avons pas retenu un cahier des charges précis. Nous avons ainsi envisagé la réalisation de notre cellule de mélange simple avec les fréquences suivantes : $f_{RF}=30\text{GHz}$, $f_{OL}=10\text{GHz}$ et $f_{FI}=20\text{GHz}$ correspondant aux applications LMDS par satellite.

Nous présentons tout d'abord le circuit mélangeur conçu ainsi que ses performances simulées. Puis, nous présenterons la caractérisation complète du circuit monolithique.

II.1. Le schéma électrique

Compte tenu des doutes sur la validité des modèles des composants pour la bande de fréquences visée (20/30GHz), nous avons, pour cette première conception, limité au maximum la complexité du circuit et ainsi envisagé l'intégration d'une cellule simple de mélange ne nécessitant aucun coupleur de signaux.

De plus, l'étude de l'inductance présentée au paragraphe précédent fut menée parallèlement à la conception de ce mélangeur. Aussi, ce dernier n'intègre pas cette inductance mais utilise uniquement des inductances semi-distribuées mettant en œuvre la couche Népi afin de leur assurer le meilleur coefficient de qualité possible.

Le schéma électrique, présenté sur la Figure II-1, reprend les concepts largement développés dans la première partie de ce chapitre ainsi que dans le second chapitre. Le mélangeur est ainsi constitué des éléments suivants.

- Des cellules d'adaptation en PI constituées de deux inductances semi-distribuées et d'une capacité, pour les accès RF et OL.
- Un filtre résonant à la fréquence f_{RF} et présentant une faible impédance sur la base du transistor T_{RF} à la fréquence f_{OL} (pour un pompage optimal) et à la fréquence f_{FI} (limitation des contre-réactions).
- Un autre filtre résonant à la fréquence f_{FI} sur le collecteur de T_{RF} et présentant une faible impédance à la fréquence f_{OL} (pour un pompage optimal) et à la fréquence f_{RF} (maximisation de l'isolation de l'accès RF vers l'accès FI).
- Deux accès de polarisation en tension munis de capacités de découplage (C_d) et une polarisation en courant pour le transistor T_{OL} .
- Un filtre stabilisant sur l'émetteur du transistor T_{OL} qui doit présenter une faible impédance aux fréquences f_{RF} et f_{FI} afin, d'une part, de maximiser le gain de conversion et d'autre part d'assurer la stabilité du mélangeur.

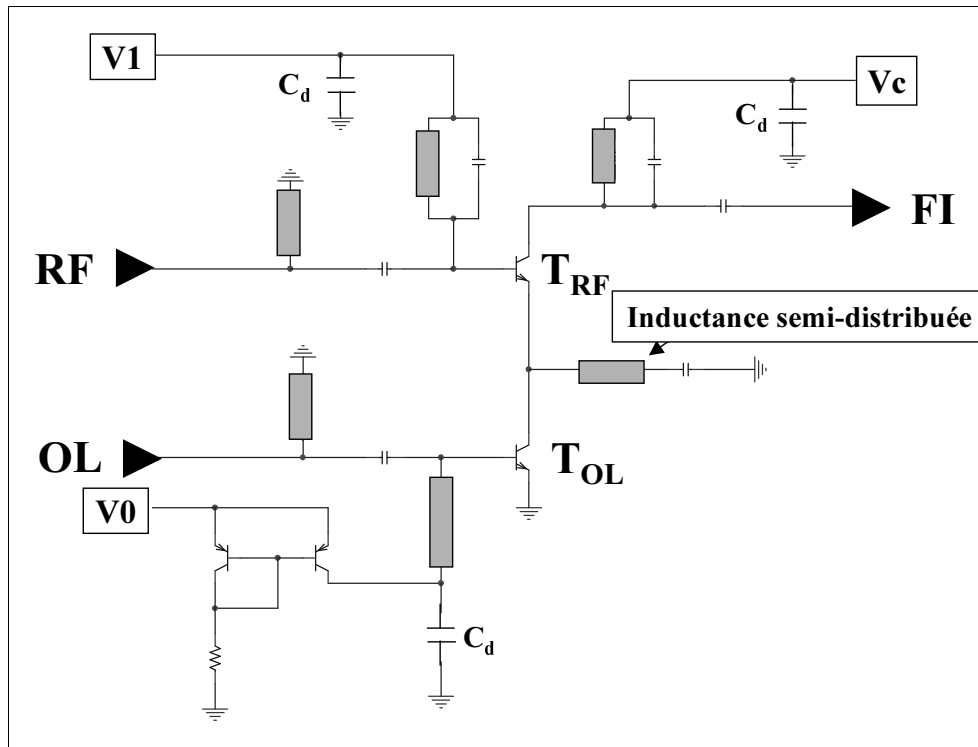


Figure II-1 : Schéma électrique du mélangeur 30GHz/20GHz

L'optimisation de tous les éléments passifs, suivant la méthodologie décrite au second chapitre, nous a conduit à l'obtention d'une valeur du gain de conversion maximale de 7dB atteinte pour une puissance P_{OL} de -15dBm .

L'intégration monolithique a été conduite en tentant de minimiser, autant que faire se peut, tous les parasites introduits par cette étape et liés d'une part à tous les éléments d'interconnexions pour lesquels il n'existe pas de modèles électriques fournis par le fondeur (coudes, croisements entre deux lignes, ...) et d'autre part aux nombreux couplages, induits par la compacité de cette technologie, notamment au niveau du substrat.

La Figure II-2 montre une photographie du MMIC dont les dimensions de la partie utile (sans les plots de polarisation et les accès coplanaires hyperfréquences) sont de $300 \times 450 \mu\text{m}^2$.

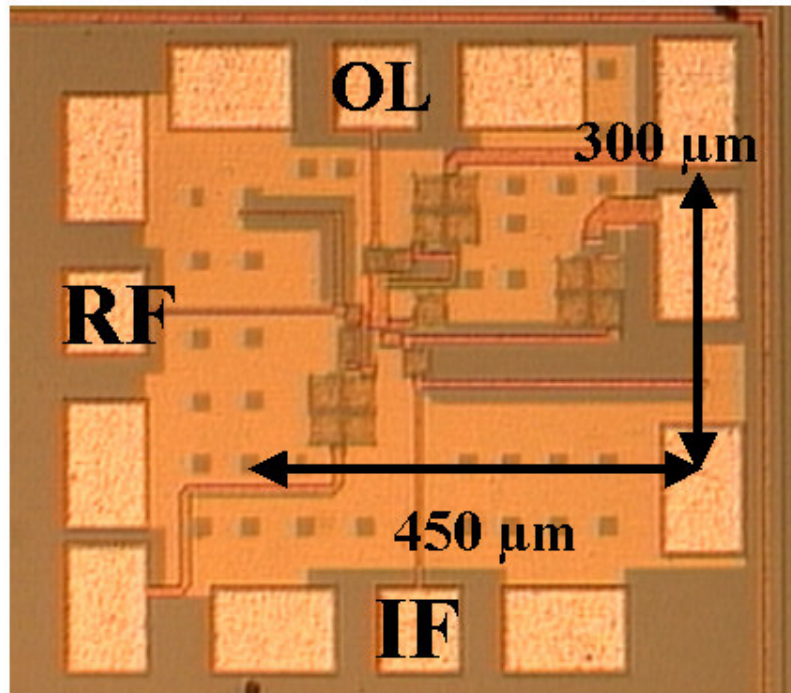


Figure II-2 : Photographie du mélangeur

La conception monolithique étant effectuée, nous allons présenter dans le prochain paragraphe la caractérisation complète du circuit conçu.

II.2. Les performances mesurées

Comme pour les circuits sur GaAs, la puce à tester a été reportée sur un substrat hôte afin de garantir sa stabilité mécanique lors des caractérisations tout en facilitant les amenées de polarisation. Nous avons de plus porté une attention particulière au découplage des alimentations, en rajoutant des capacités MIM (100pF) au plus proche des plots du circuit.

La caractérisation hyperfréquence du circuit fut effectuée sous pointes coplanaires et les points de polarisation optimaux ont été ajustés par les alimentations continues extérieures.

Nous présentons tout d'abord les caractérisations effectuées sur ce circuit, puis nous en discuterons les résultats par rapport à l'état de l'art.

II.2.1. Caractérisation du circuit

Nous avons tout d'abord procédé à la mesure des coefficients de réflexion aux différents accès hyperfréquences, présentés à la Figure II-3, afin d'en vérifier les bandes passantes.

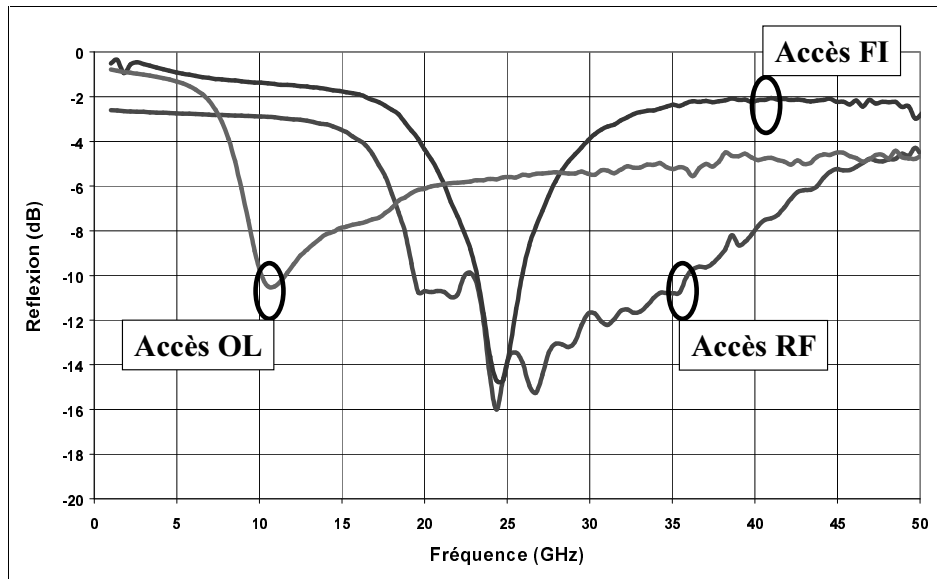


Figure II-3 : Mesure des coefficients de réflexion aux trois accès

Notons que, pour les accès OL et RF, les coefficients de réflexion sont inférieurs à 10dB pour les fréquences de fonctionnement nominales. Par contre, pour l'accès FI, un décalage fréquentiel d'environ 5GHz est observé puisque la valeur minimale du coefficient de réflexion est atteinte pour une fréquence voisine de 25GHz au lieu de 20GHz désiré.

Ce décalage est imputable à la présence d'un circuit résonnant en sortie (voir le schéma Figure II-1) pour lequel les variations des éléments le constituant sont beaucoup plus sensibles sur ses performances fréquentielles que dans le cas des réseaux d'adaptation classiques présents aux autres accès. Comme pour les conceptions sur AsGa (voir paragraphe IV.2.1 du chapitre 3 partie A), ces variations sont dues à :

- La présence, importante à 30GHz, de phénomènes de couplage non pris en compte dans les simulations électriques et essentiellement dus à la forte densité d'intégration ainsi qu'au couplage par le substrat.
- D'une manière générale, tous les éléments non-considerés lors des simulations électriques : croisements d'interconnexions, passages d'un niveau métallique à un autre par une matrice de via, ...

Nous avons ensuite mesuré le gain de conversion de notre circuit en fonction de la puissance P_{OL} (Figure II-4).

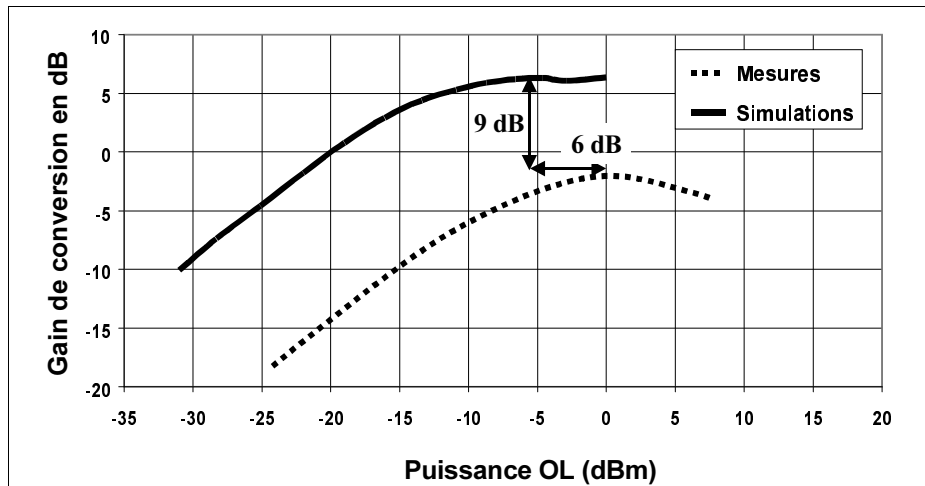


Figure II-4 : Mesure du gain de conversion

Une valeur maximale du gain de -2 dB est atteinte pour une puissance P_{OL} de 0 dBm . Les écarts importants entre les simulations et les mesures, de 9 dB sur la valeur du gain et de 6 dB sur celle de la puissance P_{OL} , s'expliquent par :

- Le coefficient de réflexion en sortie de -4 dB pour une fréquence $f_{FI}=20\text{ GHz}$, qui entraîne des pertes par réflexion de 2.2 dB .
- La dégradation des performances du filtrage en sortie, notamment le faible filtrage de la raie à la fréquence f_{RF} en sortie, qui entraîne une forte perturbation des différents phénomènes de conversion de fréquence, dégradant indirectement le gain de conversion (voir le chapitre 2).

Notons que la valeur du gain de conversion de -2 dB correspond tout à fait à l'état de l'art des circuits silicium en bande Ka. En effet, une valeur de 2 dB , obtenue à l'aide d'une cellule de Gilbert ne fonctionnant qu'à 20 GHz avec des transistors de performances identiques aux nôtres, fut récemment publiée et fait référence dans ce domaine [III.12].

Enfin, pour une puissance OL de 0 dBm , ont été effectuées les mesures de la puissance d'interception d'ordre 3 en sortie ainsi que du facteur de bruit. Le Tableau II-1 résume l'essentiel des performances électriques mesurées de notre mélangeur, performances qui vont être discutées et évaluées par rapport à l'état de l'art des circuits sur silicium opérant à 30 GHz dans le prochain paragraphe.

Performances pour $f_{RF}=30\text{ GHz}$ et $f_{IF}=20\text{ GHz}$	
P_{LO}	0 dBm
G_c	-2 dB
NF	20 dB
OIP3	1 dB
Consommation statique	25 mW
Surface du MMIC	$300 \times 4500\text{ }\mu\text{m}^2$

Tableau II-1 : Résumé des performances électriques

II.2.2. Discussion des caractérisations

La première discussion porte sur la linéarité de notre mélangeur. En effet, la mauvaise linéarité des circuits à base de transistors bipolaires (par rapport à leurs homologues à effet de champ) est un souci majeur lors de la conception de circuit. L'évaluation de cette linéarité se fait par la puissance d'interception d'ordre 3 en sortie, mais cette puissance est mal adaptée à l'évaluation de la linéarité de plusieurs mélangeurs étant donné qu'elle dépend de la puissance OL pour laquelle elle est évaluée. Aussi, un coefficient efficace permettant une évaluation correcte est Q_3 défini par : $Q_3(\text{dB}) = \text{OIP}_3(\text{dBm}) - P_{\text{OL}}(\text{dBm})$.

De plus, comme un compromis existe entre forte linéarité et faible puissance consommée, la Figure II-5 présente les valeurs Q_3 en fonction de la puissance consommée de quelques mélangeurs intégrés récemment publiés [III.12, III.13, III.14]. Comme les mélangeurs mentionnés correspondent à des cellules double équilibrées, 6dB doivent être ajoutés au Q_3 de notre cellule simple et sa puissance consommée doit être quadruplée, ceci afin de permettre la comparaison.

Nous voyons ainsi que notre topologie réalise le meilleur compromis par rapport à l'état de l'art entre une bonne linéarité et une faible consommation.

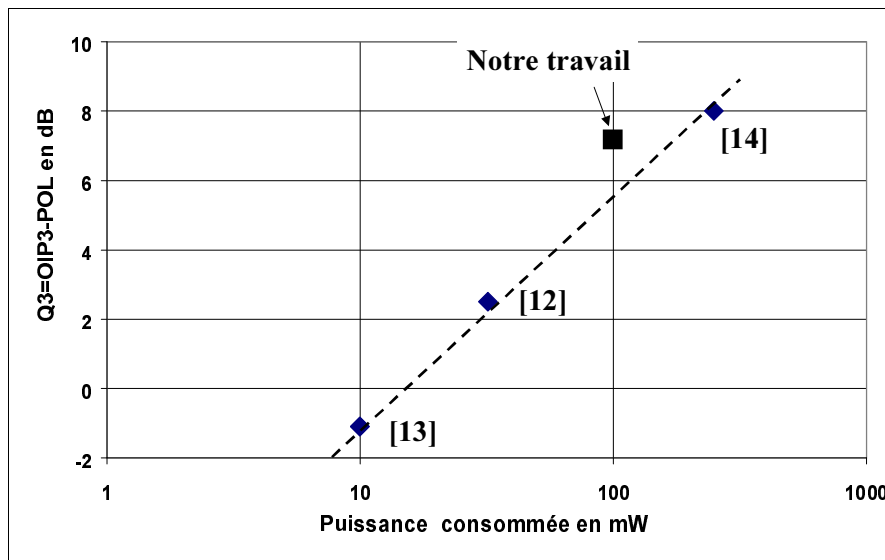


Figure II-5 : Linéarité de mélangeurs

De plus, la plupart des optimisations de mélangeurs se font en vue de maximiser une performance électrique mais toujours au détriment d'autres. Notre conception présente, outre une bonne linéarité, de bonnes performances globales en terme de gain de conversion, de bruit et de puissance OL. Afin d'évaluer globalement les performances de notre mélangeur, nous avons introduit le facteur de mérite suivant : $Q_T(\text{dB}) = G_c(\text{dB}) - \text{NF}(\text{dB}) + Q_3(\text{dB}) - P_{\text{OL}}(\text{dBm})$ qui tient compte du gain de conversion, du bruit, de la linéarité et de la puissance OL nécessaire. Pour notre topologie ce coefficient atteint -21 dB, ce qui donnerait donc pour une cellule double équilibrée -15dB. Ce résultat est à l'état de l'art des mélangeurs silicium qui se situe à $Q_T = -14\text{dB}$ [III.13].

Notons que la valeur relativement élevée du facteur de bruit mesuré, qui est due aux fréquences de fonctionnement f_{RF} et f_{FI} élevées, est comparable à celles présentées par des cellules de Gilbert qu'elles soient constituées d'HBT sur silicium [III.12, III.13] ou sur AsGa [III.]. De plus, cette caractéristique n'est pas forcément critique pour le facteur de bruit global d'un système car ce dernier peut être simplement amélioré en ajoutant un préamplificateur RF [III.15].

Ainsi, les caractérisations se situent à l'état de l'art des réalisations sur silicium [III.16], ce qui démontre tout l'intérêt de notre topologie de mélange et valide sa méthodologie de conception. Bien qu'un effort de modélisation reste à fournir (notamment concernant les interconnexions et les couplages électromagnétiques dus à la forte densité d'intégration et au substrat) pour permettre la conception de circuits monolithiques encore plus performants, ces résultats remarquables prouvent les fortes potentialités de la filière BiCMOS SiGe de ST Microelectronics dans la gamme des fréquences millimétriques.

III. Conclusion

Le premier paragraphe de cette partie a concerné la conception d'une inductance fonctionnant aux fréquences millimétriques. L'optimisation des dimensions géométriques de cette inductance ainsi que le choix des moyens technologiques à mettre en œuvre pour pallier la faible résistivité du silicium ont abouti à la réalisation d'une inductance à l'état de l'art [III.6], de valeur 0.75nH et présentant une fréquence de résonance de 62 GHz ainsi qu'une valeur de 20 du coefficients de qualité à 30 GHz.

Nous avons ensuite présenté la conception d'un mélangeur simple fonctionnant en bande Ka, reprenant les techniques génériques de conception développées au chapitre 2. Les caractérisations de ce circuit ont démontré un fonctionnement performant validant notre topologie originale de mélangeur ainsi que la démarche de son optimisation.

Pour un fonctionnement aux fréquences millimétriques, notre mélangeur présente des performances globales mesurées se situant au niveau de l'état de l'art des circuits sur silicium [III.16], notamment en terme de linéarité puisque notre topologie assure le meilleur compromis linéarité-consommation par rapport aux cellules de Gilbert classiquement utilisées.

Cependant, des efforts restent à fournir pour développer les éléments passifs de performances adaptées aux fréquences millimétriques ainsi que pour parfaire et valider les modèles disponibles pour ces bandes fréquentielles.

REFERENCES BIBLIOGRAPHIQUES DU CHAPITRE 3 PARTIE B

- [III.1] : H. Kroemer, « Theory of a wide-gap emitter for transistors », Proceedings of the IRE, 1957, pp. 1535-1537.
- [III.2] : J.-F. Luy, P. Russer, « Silicon-based millimeter-wave devices », Springer-verlag.
- [III.3] : John R. Long, « integrating Passive Components and RF/MMIC Design », BCTM 2000 Short course.
- [III.4] : <http://www-smirc.stanford.edu/spiralcalc.html>
- [III.5] : H. Hasegawa, M. Furukawa and H. Yanai « Properties of microstrip line on Si-SiO₂ system », IEEE MTT-19, n°11, Nov. 1971, pp869-881.
- [III.6] : **D. Dubuc** , É. Tournier , I. Telliez , T. Parra, C. Boulanger and J. Graffeuil, “ High Quality Factor and High Self-Resonant Frequency Monolithic Inductor for Millimeter-Wave Si-based IC's”, IEEE MTT-S , Seattle USA, 2-7 June 2002, pp 193-196.
- [III.7] : Joachim N. Burghartz “Status and Trends of Silicon RF Technology”, Technical Digest ESSDERC, Sept. 13-15, 1999, Leuven, Belgium.
- [III.8] Joachim N. Burghartz “Integrated Multilayer RF Passives in Silicon Technology”, Topical Meeting on Silicon monolithic integrated circuits in RF systems, 17-18 sept 1998, Ann Arbor, Michigan pp 141-147.
- [III.9] : ED02AH Design Manual V200CD, Philips Microwave Limeuil (OMMIC), April 1999.
- [III.10] : C-Y Chi, G.M. Rebeiz, « Planar microwave and millimeter-wave lumped and coupled-line filters using micro-machinig techniques », IEEE MTT-43, n°4 april 1995, pp.730-738.
- [III.11] : S.P. Voinigescu, D. Marchesan, M.A. Copeland, « A family of monolithic inductor-varactor SiGe-HBT VCOs for 20GHz to 30GHz LMDS fiber-optic receiver applications », IEEE RFIC Symposium 2000, pp.173-176.
- [III.12] : K.B. Schad, H. Schumacher, A. Schüppen, « Low-power active mixer for Ku-Band application using SiGe HBT MMIC technology», IEEE RFIC Symposium 2000, pp.263-266.
- [III.13] : M. Wurzer, “30 GHz active mixer in a Si/SiGe Bipolar Technology“, Asia Pacific Microwave conference- Sydney Australia, December 2000, pp 780-783.

- [III.14] : J. Gleen, M. Case, D. Harame, B. Meyerson, R. Poisson, “ 12-GHz Gilbert Mixers using A Manufacturable Si/SiGe Epitaxial-Base Bipolar Technology”, BTCM 1995, pp 186-189.
- [III.15] : S. Hackl, J. Böck, M. Wurzer, A.L. Scholtz, “40 GHz Monolithic Integrated Mixer in SiGe Bipolar Technology“, IEEE MTT-S, June 2002, pp 1241-1244.
- [III.16] : **D. Dubuc**, É. Tournier, T. Parra, I. Telliez, C.Boulanger, R. Plana and J. Graffeuil, ”high performances 30 GHz active mixer in a SiGe bipolar technology“, en soumission à *electronics letters*.

Conclusion générale

Les travaux présentés dans ce mémoire s'inscrivent dans le développement incessant des communications hyperfréquences qui nécessite l'augmentation des performances, des fréquences de fonctionnement et de la complexité des systèmes mais aussi impose des contraintes de coût, de masse, de volume et de rapidité de réalisation.

Ces impératifs se sont traduits depuis peu par la généralisation des circuits monolithiques conjuguée au développement des méthodes de conception. La nécessité de proposer rapidement des applications performantes est en effet devenue vitale par exemple pour les systèmes de communication satellitaires afin qu'ils restent compétitifs avec le segment terrestre. Il est donc nécessaire à la fois de développer les technologies et les topologies des circuits monolithiques, afin de satisfaire aux besoins de performances, mais aussi de formaliser les différents moyens de conception dans le but de pouvoir répondre encore plus rapidement aux nécessités imposées par le marché.

Nos travaux de recherche se sont donc articulés d'une part autour de la conception de circuits mélangeurs en technologie monolithique et d'autre part, plus fondamentalement, à la définition de méthodes génériques de conception de ces circuits.

Au cours du premier chapitre, nous avons présenté, outre les définitions relatives aux systèmes de conversion de fréquence, la classification des différentes cellules de base composant le mélangeur. Nous avons ainsi pu dégager les principes de fonctionnement, les avantages et les inconvénients de chaque structure permettant ainsi leur choix judicieux en fonction des contraintes imposées.

S'appuyant sur cette classification, nous avons développé, au second chapitre, la conception d'une cellule originale de mélange présentant d'excellentes performances mais aussi de bonnes perspectives d'intégrabilité en technologie monolithique. Une partie importante de nos travaux a porté sur l'établissement d'une méthode de conception générique permettant l'optimisation du circuit de manière structurée. Cette méthode a, de plus, rendu possible l'étude délicate de la stabilité non-linéaire pour laquelle l'absence d'outils commerciaux pénalise le développement de circuits micro-ondes de fortes potentialités.

Enfin, le troisième chapitre a porté sur la description de l'intégration monolithique de la cellule de mélange conçue précédemment. Cette intégration, fondée sur un cahier des charges de répéteurs spatiaux en bandes Ku et Ka a été réalisée pour deux technologies : la première développée autour de transistors à haute mobilité électronique (PHEMT) sur Arséniure de Gallium, la seconde basée sur des transistors bipolaires à hétérojonction Silicium/Silicium-Germanium. Une part importante des travaux a également porté sur l'optimisation d'inductance sur substrat silicium fonctionnant dans la gamme des fréquences millimétriques. Cette étude a ainsi abouti à la réalisation d'une inductance de performance à l'état de l'art des réalisations du même type. Enfin, la cellule de mélange silicium, dont les performances sont aussi à l'état de l'art des circuits fonctionnant dans cette gamme de fréquences, a validé d'une part les différents concepts mis en œuvre et d'autre part les fortes potentialités de la technologie employée.

Ces travaux de recherche, notamment sur le silicium dans la bande des fréquences 20 / 30GHz, s'inscrivent dans l'effort actuel des développements des applications millimétriques pour lesquelles les perspectives sont très vastes car devenant pluridisciplinaires.

En effet, l'intégration augmentant, les concepteurs doivent maîtriser d'une part la connaissance des éléments actifs semi-conducteurs qui deviennent extrêmement performants et, d'autre part, la conception de circuits par les simulations électriques. Ils doivent de plus maîtriser la modélisation électromagnétique indispensable étant donnée l'augmentation

fréquentielle conjuguée à la forte densité d'intégration de la technologie silicium. Ces différents aspects de conception doivent être mis en œuvre de façon complémentaire, ouvrant ainsi des perspectives de recherche dans le domaine des circuits silicium pour la gamme des fréquences millimétriques, tant au niveau topologie de circuit qu'au niveau modélisation et optimisation des éléments passifs et des différentes interactions entre les éléments.

De plus, le développement des micro-systèmes mécaniques pour les applications micro-ondes va, à moyen terme, inévitablement se conjuguer aux efforts portés sur les technologies silicium. Il ne sera alors plus suffisant de restreindre la conception aux études électriques et électromagnétiques, déjà très complexes, mais il faudra aussi envisager la prise en compte des aspects mécaniques voire thermiques. Ces problématiques importantes doivent être abordées afin de profiter au maximum des potentialités offertes par les diverses technologies, chacune apportant sa contribution à l'augmentation des performances globales des systèmes.

Ainsi, aux fréquences micro-ondes et millimétriques, l'intégration monolithique sur silicium de circuits actifs et passifs performants avec des fonctions électro-mécaniques pourra tirer partie des avantages apportés par le silicium (bat coût, maturité de production, possibilité d'intégration des fonctions numériques,...), de ceux apportés par les transistors bipolaires SiGe, notamment en terme de bruit basse fréquence et de gain, mais aussi des potentialités, quotidiennement améliorées, des structures électromécaniques notamment en terme de pertes d'insertion et d'isolation. A ce niveau d'intégration, d'autres aspects interviendront pour les circuits actifs tels que les protections contre les décharges électrostatiques (ESD) indispensables à leur développement industriel. Tous ces domaines sont très riches en perspectives d'avenir.

“ Contribution à la conception de convertisseurs de fréquence. Intégration en technologie Arséniure de Gallium et Silicium Germanium.”

Résumé des travaux :

Ces dernières années ont vu le fort développement des communications spatiales et avec lui le changement des contraintes de conception de ces systèmes. En effet, la multiplicité des applications hyperfréquences ainsi que leur ouverture au domaine grand public ont entraîné l'augmentation des densités d'intégration et des performances des systèmes, mais aussi la nécessité de prendre en considération leur fiabilité et leurs coûts de production. Ainsi, nos travaux portent sur l'étude et l'intégration de convertisseurs de fréquence micro-ondes répondant à ces impératifs.

Dans un premier temps, nous abordons les définitions relatives aux mélangeurs ainsi que les principales caractéristiques nécessaires à leurs évaluations. Nous présentons ensuite les différentes topologies de mélangeurs existantes en citant, pour chacune, ses avantages et ses inconvénients.

La seconde partie de nos travaux de recherche est dédiée à la définition d'une cellule originale de mélange. Une méthode analytique de conception, basée sur le formalisme des matrices de conversion, nous a permis d'une part de définir des règles génériques de conception de mélangeur et d'autre part d'optimiser notre cellule de mélange suivant des critères fixés. Enfin, une partie importante de ces investigations est dédiée à l'étude de la stabilité non-linéaire des mélangeurs, point sensible de la topologie retenue.

La dernière partie de ce mémoire est consacrée à l'intégration monolithique et à la caractérisation du mélangeur performant envisagé pour deux technologies concurrentes : l'une basée sur des transistors à effet de champ à haute mobilité électronique sur Arséniure de Gallium et l'autre fondée sur des transistors bipolaires à hétérojonction Silicium/Silicium-Germanium. Pour cette dernière technologie, un effort a été porté sur la conception de circuits passifs fonctionnant aux fréquences millimétriques. La caractérisation des circuits a démontré le bien fondé des études présentées précédemment ainsi que l'aptitude de la technologie Silicium-Germanium pour l'intégration monolithique de systèmes aux fréquences millimétriques.

Mots-clefs : Micro-ondes, Hyperfréquences, fréquences millimétriques, MMIC, mélangeur, mélange, convertisseur de fréquence, Silicium Germanium, Arséniure de Gallium, matrice de conversion, stabilité.

“ Conception of frequency converters. Monolithic integration with PHEMT Gallium Arsenide and HBT Silicon Germanium technologies ”

Abstract :

The current great development of spatial communication comes with the change of conception constraints of these systems. Indeed, the multiplicity of new spatial systems as well as public applications must always raise integration density and system performances but also have to satisfy reliability and production cost needs. Thus, our work deals with the study and the integration of microwave frequency converter which satisfies these imperatives.

We first present mixers characteristics in order to classify the reported topologies with respect to their drawbacks and advantages.

The second part of this work is dedicated to the conception of an original mixer's topology. Based on an analytical method of conception, using conversion matrix formalism, we first define generic rules of mixer conception and on the other hand we optimise the mixer performances. Finally, a great part of our investigation deals with the study of the non-linear stability of mixers, key point of our topology.

The last part of this manuscript addresses the monolithic integration and the characterization of our mixer with two concurrent technologies : one based on high electron mobility field effect transistor on Gallium Arsenide and the other one based on heterojunction bipolar transistor Silicon/Silicon-Germanium. For the last technology, an effort was made on the conception of passive circuits acting at millimeter wave frequencies. The characterization of the resulting circuits demonstrates the validity of the theoretical part and the aptitude of Silicon-Germanium technology for millimetre-wave monolithic integration.

Key words : Microwave, millimeter-wave frequency, MMIC, mixer, frequency converter, Silicon Germanium, Gallium Arsenide, conversion matrix, stability.