

Intégration en technologie CMOS de circuits de contrôle pour antennes reconfigurables

par

Elouan MILLET

MÉMOIRE PRÉSENTÉ À L'ÉCOLE DE TECHNOLOGIE SUPÉRIEURE
COMME EXIGENCE PARTIELLE À L'OBTENTION DE LA MAÎTRISE
AVEC MÉMOIRE EN GÉNIE ÉLECTRIQUE
M. Sc. A.

MONTRÉAL, LE 19 FÉVRIER 2026

ÉCOLE DE TECHNOLOGIE SUPÉRIEURE
UNIVERSITÉ DU QUÉBEC



Elouan Millet, 2026



Cette licence Creative Commons signifie qu'il est permis de diffuser, d'imprimer ou de sauvegarder sur un autre support une partie ou la totalité de cette oeuvre à condition de mentionner l'auteur, que ces utilisations soient faites à des fins non commerciales et que le contenu de l'oeuvre n'ait pas été modifié.

PRÉSENTATION DU JURY

CE MÉMOIRE A ÉTÉ ÉVALUÉ

PAR UN JURY COMPOSÉ DE:

M. Nicolas Constantin, directeur de mémoire
Département de génie électrique à l'École de Technologie Supérieure

M. Gheorghe Marcel Gabrea, président du jury
Département de génie électrique à l'École de Technologie Supérieure

M. Richard Al Hadi, membre du jury
Département de génie électrique à l'École de Technologie Supérieure

IL A FAIT L'OBJET D'UNE SOUTENANCE DEVANT JURY ET PUBLIC

LE 4 FÉVRIER 2026

À L'ÉCOLE DE TECHNOLOGIE SUPÉRIEURE

REMERCIEMENTS

J'aimerais tout d'abord remercier CMC Microsystems sans lequel ce projet ne se serait pas matérialisé. Cet organisme permet aux universités canadiennes de réaliser des microsystèmes en mettant à leur disposition des logiciels, des équipements, des formations et du soutien technique. Dans le cadre de notre projet, CMC Microsystems a notamment servi d'intermédiaire entre nous et le fabricant de circuits intégrés TSMC.

Je souhaite également remercier Mathieu Gratuze et Amin Pourvali, du LaCIME, pour leur soutien technique lors de la phase de conception du circuit intégré, ainsi que pour leurs précieux conseils pour mieux réussir le projet. Je n'oublie pas non plus l'aide que m'ont apportée M. Normand Gravel et M. Rigoberto Avelar lors de l'assemblage des circuits. Ensuite, en plus de m'avoir donné des conseils lors de la conception de mes circuits, M. Pascal Burasa de Polytechnique Montréal m'a permis d'effectuer des mesures au laboratoire de Poly-Grames.

Je suis également reconnaissant envers M. Nicolas Constantin, professeur au département de génie électrique à l'École de technologie supérieure, qui m'a donné l'opportunité de travailler sur ce projet et qui m'a permis de m'initier dans le domaine de la microélectronique analogique.

Enfin, je souhaite remercier mes amis, mon frère, ma mère et mon père pour leurs mots d'encouragement tout au long du projet.

Intégration en technologie CMOS de circuits de contrôle pour antennes reconfigurables

Elouan MILLET

RÉSUMÉ

Ce mémoire présente la conception et l'essai d'un circuit intégré en technologie CMOS de 65 nm, capable de fonctionner simultanément comme détecteur de puissance et mélangeur de fréquence.

Les futurs réseaux 5G et 6G intégreront un grand nombre de nœuds satellitaires, nécessitant des antennes simples et capables de modeler leur faisceau. Les antennes réseau reconfigurables avec des diodes de déphasage sont idéales pour cette application. Notre circuit contribue à l'amélioration du contrôle de la polarisation des diodes à capacité variable, qui ont tendance à engendrer de la distorsion, en mesurant la puissance et la distorsion du signal déphasé.

Suite à une analyse théorique, nous avons conçu un circuit intégré simple de détecteur-mélangeur en nous basant sur un circuit de détecteur de puissance existant. Nous avons également conçu un circuit imprimé dédié pour tester la puce.

Lors de ces mesures, nous avons pu confirmer la double fonctionnalité du circuit. Le circuit bénéficie d'une plage dynamique de 27,7 dB en tant que détecteur et de 42 dB en tant que mélangeur, ainsi que d'un taux de réjection des produits d'intermodulation néfastes de plus de 30 dB sur une grande partie de la plage dynamique. Il peut recevoir des signaux allant jusqu'à 10,5 GHz ; et le signal à sa sortie peut aller jusqu'à 1 GHz. Le circuit actif n'occupe qu'une surface de $1\,507,7\,\mu\text{m}^2$; et il ne consomme que 3,3 mW maximum au repos. Même s'il existe d'autres détecteurs quadratiques et mélangeurs avec de meilleures performances, les résultats nous encouragent à approfondir les études théoriques sur la non-linéarité des transistors, ainsi qu'à améliorer le circuit présenté.

Mots-clés: détecteur, mélangeur, CMOS, quadratique, non-linéarité

CMOS Integration of a reconfigurable antenna array control circuit

Elouan MILLET

ABSTRACT

This thesis covers the design and the trial of a 65-nm CMOS circuit that functions both as a power detector and a mixer, thanks to the square-law properties of its transistors.

Future 5G and 6G networks will integrate a great number of satellite nodes, requiring simple antennas with beam-forming capabilities. Reconfigurable reflectarray antennas, using diode-based phase shifters, are well-suited for such applications. This work contributes to the enhancement of these antennas. The goal of our circuit is to allow smarter control of the biasing of varicap diodes, which are notorious for their tendency to distort signals, by measuring the power and the distortion of the signal.

Following a theoretical analysis, we designed a simple and easily integrable square-law detector-mixer. The circuit is based on the design of an existing power envelope detector. Then, we conceived a printed circuit board designed to host the chip for measurements.

Through experimental validation, we confirmed the circuit's dual functionality. The circuit boasts a dynamic range of 27.7 dB as a detector and 42 dB as a mixer. It also demonstrates an intermodulation rejection ratio exceeding 30 dB across most of the dynamic range. The circuit can process signals of up to 10.5 GHz, with an estimated output bandwidth of 1 GHz. This very simple circuit covers an area of $1,507.7 \mu\text{m}^2$ and consumes at most 3.3 mW of power. Although other detectors and mixers have proven to have better performance figures, these results are encouraging and motivate further exploration of square-law circuits and improvement upon our design.

Keywords: detector, mixer, CMOS, square-law, nonlinearity

TABLE DES MATIÈRES

	Page
INTRODUCTION	1
0.0.1 Motivation de la recherche	1
0.0.2 Objectifs de recherche	2
0.1 Contribution à la recherche	4
0.2 Plan du mémoire	4
CHAPITRE 1 REVUE DE LITTÉRATURE	5
1.1 Le besoin d’antennes reconfigurables pour les réseaux 5G et 6G	5
1.2 Les antennes reconfigurables	6
1.2.1 Différents types d’antennes reconfigurables	7
1.2.1.1 Les RTA et les RRA	7
1.2.1.2 Les RIS	8
1.2.2 Les applications pour les RRA, RTA et RIS	10
1.2.3 Les déphaseurs	11
1.2.3.1 Les technologies de déphaseurs	11
1.2.3.2 Les diodes à capacité variable	12
1.3 Les détecteurs	13
1.3.1 Les détecteurs d’enveloppe de tension	13
1.3.2 Détecteurs de puissance	15
1.3.2.1 Les détecteurs de puissance moyenne	16
1.3.2.2 Les détecteurs d’enveloppe de puissance	17
1.3.3 Les caractéristiques d’un détecteur d’enveloppe de puissance	18
1.3.3.1 La responsivité	18
1.3.3.2 La plage dynamique	19
1.3.3.3 La borne inférieure de la plage dynamique	20
1.3.3.4 Le point de compression	21
1.3.3.5 La plage de fréquence	21
1.3.3.6 La bande passante de sortie (ou la rapidité)	21
1.3.3.7 La consommation de puissance	22
1.3.3.8 La technologie	22
1.3.3.9 La taille	23
1.3.4 Les applications pour les détecteurs de puissance	23
1.3.4.1 La démodulation d’amplitude	23
1.3.4.2 Built-in test	23
1.3.4.3 Le contrôle de niveau automatique	24
1.3.4.4 Autres applications	24
1.3.5 Les travaux précédents de détecteurs de puissance quadratiques	25
1.3.5.1 Distorsion dans un détecteur	25
1.4 Les mélangeurs de fréquence	26
1.4.1 Le principe d’un mélangeur de fréquence	26

1.4.2	Les types de mélangeurs de fréquence	27
1.4.3	Les caractéristiques d'un mélangeur	27
1.4.3.1	Le gain de conversion	27
1.4.3.2	L'alimentation	28
1.4.3.3	La bande passante	28
1.4.3.4	La figure de bruit	29
1.4.3.5	Le point de compression	29
1.4.3.6	Le point d'interception d'ordre trois	29
1.4.3.7	L'isolation entre les ports	30
1.4.4	Les topologies de mélangeurs communes	30
1.4.4.1	Les mélangeurs quadratiques à base de diodes	31
1.4.4.2	Les mélangeurs à double équilibrage en pont de diodes	31
1.4.4.3	La cellule de Gilbert	33
1.4.4.4	Les mélangeurs quadratiques à base d'un transistors actif	34
1.4.5	Les précédents travaux de mélangeurs utilisant la non-linéarité d'un transistor	36
1.5	Un circuit détecteur-mélangeur	37
1.6	Conclusion de chapitre	37
CHAPITRE 2 ANALYSE ET THÉORIE SUR LES CIRCUITS QUADRATIQUES		39
2.1	La non-linéarité d'un transistor	39
2.1.1	Le modèle «simple» d'un MOSFET en saturation	39
2.1.2	Les autres phénomènes non linéaire	43
2.1.2.1	La modulation de longueur de canal	43
2.1.2.2	L'effet de substrat	44
2.2	Le fonctionnement d'un détecteur de puissance quadratique	44
2.2.1	La détection d'un signal à amplitude constante	44
2.2.2	La détection d'un signal modulé en amplitude	46
2.3	Un mélangeur de fréquences grâce à la caractéristique quadratique	46
2.3.1	Rappel du fonctionnement d'un mélangeur	47
2.3.2	La mise au carré pour la multiplication	48
2.3.3	Les contraintes pour un mélangeur quadratique	49
2.3.3.1	Le chevauchement du signal IF et de l'enveloppe de puissance à la sortie	50
2.3.3.2	Le chevauchement du signal IF et RF ou LO rémanent	50
2.4	La fonction de transfert non linéaire du circuit	51
2.4.1	La responsivité du détecteur	51
2.4.2	Le gain de conversion d'un mélangeur	52
2.5	La contribution de la responsivité à l'augmentation de la plage dynamique	52
2.6	Produits d'intermodulation nuisibles	53
2.6.1	Les ordres des produits d'intermodulation néfastes	53
2.6.1.1	Pour les détecteurs	54
2.6.1.2	Pour les mélangeurs	54

2.6.2	Les conséquences de l'intermodulation d'ordre quatre pour les détecteurs	55
2.6.2.1	Le produit d'intermodulation d'ordre quatre pour un détecteur	55
2.6.2.2	La compression de l'enveloppe de puissance	56
2.6.2.3	Le point d'interception	57
2.6.2.4	La relation entre le point de compression et le point d'interception	57
2.6.3	Les produits d'intermodulations d'ordre quatre pour les mélangeurs	58
2.6.3.1	Un produit d'intermodulation à IF	58
2.6.3.2	La compression du signal transposé	59
2.6.3.3	Le point d'interception	60
2.6.3.4	La relation entre le point de compression et le point d'interception	60
2.7	Conclusion du chapitre	60
CHAPITRE 3 CONCEPTION DU CIRCUIT		63
3.1	La conception du circuit intégré	63
3.1.1	La technologie de fabrication utilisée	63
3.1.2	L'analyse du fonctionnement du circuit intégré	63
3.1.2.1	Une explication qualitative du circuit	65
3.1.2.2	Les hypothèses pour l'analyse	66
3.1.2.3	Le gain de tension quadratique idéal	67
3.1.2.4	Le gain quadratique de tension en considérant la modulation du longueur de canal	68
3.1.2.5	L'impédance d'entrée	69
3.1.2.6	L'impédance de sortie	70
3.1.3	Dimensionnement des composants	70
3.1.3.1	Dimensionnement des transistors NMOS	70
3.1.3.2	Dimensionnement des transistors PMOS	72
3.1.3.3	Dimensionnement des composants passifs	73
3.1.4	La vérification de l'analyse par simulation	74
3.1.5	Les limites de notre analyse	75
3.1.6	La protection contre les décharges électrostatiques	76
3.1.7	La technique de montage de la puce	77
3.2	La conception du circuit imprimé	78
3.2.1	L'adaptation d'impédance pour la partie principale	78
3.2.1.1	Extraction des paramètres S	78
3.2.1.2	Les lignes de transmission	80
3.2.1.3	Le circuit d'adaptation pour les entrées	80
3.2.1.4	Le circuit d'adaptation pour la sortie	81
3.3	Conclusion sur la conception des circuits	82
CHAPITRE 4 SIMULATIONS ET MESURES		85
4.1	Les tests en courant continu	85
4.1.1	La consommation de courant	85

4.1.2	La tension à la sortie au repos	86
4.2	L'adaptation d'impédance de la partie principale	86
4.3	Les tests en mode détecteur	89
4.3.1	La détection d'un signal à amplitude constante	89
4.3.2	La détection d'un signal à deux fréquences	93
4.3.3	Les simulations transitoires et les mesures à l'oscilloscope	96
4.3.3.1	Les mesures à l'oscilloscope de l'enveloppe de puissance d'un signal à deux fréquences	98
4.3.3.2	La simulation et la mesure de la réactivité du détecteur	99
4.4	Simulations de la réjection de la porteuse	102
4.4.1	La réjection de la porteuse en fonction du déphasage	103
4.4.2	La réjection de la porteuse en fonction de la symétrie du circuit intégré ..	104
4.5	Les tests en mode mélangeur	105
4.5.1	Les mesures sous pointes	105
4.5.2	Les mesures des spectres sur la partie principale	109
4.5.3	La linéarité du mélangeur	111
4.5.4	Les mesures à l'oscilloscope sur la partie principale	113
4.6	Conclusion sur la performance du circuit	114
CONCLUSION ET RECOMMANDATIONS		119
ANNEXE I	CALCULS THÉORIQUES	123
ANNEXE II	LE CIRCUIT INTÉGRÉ	149
ANNEXE III	LE CIRCUIT IMPRIMÉ	159
BIBLIOGRAPHIE		165

LISTE DES TABLEAUX

	Page
Tableau 1.1	Tableau de synthèse des détecteurs de puissance de la revue de littérature 25
Tableau 1.2	Tableau de synthèse des mélangeurs de la revue de littérature 36
Tableau 4.1	Résultats des mesures sur le détecteur avec un signal à amplitude constante à l'entrée 93
Tableau 4.2	Points de compression, points d'interception pour un détecteur et leurs rapport selon V_p 95
Tableau 4.3	La bande passante simulée en fonction de la tension de polarisation .. 100
Tableau 4.4	Performance du mélangeur avec $P_{LO} = 5$ dBm 112
Tableau 4.5	Comparaison des performances de notre circuits et d'autres détecteurs de puissance en technologie CMOS 116
Tableau 4.6	Comparaison des performances de notre circuits et d'autres mélangeurs en technologie CMOS 117

LISTE DES FIGURES

	Page
Figure 1.1	Comparaison entre une RRA et une RTA 8
Figure 1.2	Usage d'une RIS pour améliorer la couverture malgré la présence d'obstacles 9
Figure 1.3	Un signal modulé en amplitude et son enveloppe 14
Figure 1.4	Schéma d'un redresseur comme détecteur d'enveloppe 14
Figure 1.5	Schéma d'un détecteur de puissance réalisée grâce à une diode 16
Figure 1.6	Schéma d'un mélangeur doublement symétrique de pont de diodes 32
Figure 1.7	Schéma d'un exemple de cellule de Gilbert réalisé avec des MOSFET .. 33
Figure 1.8	Mélangeurs quadratiques à base d'un MOSFET <i>gate-pumped</i> et <i>source-</i> <i>pumped</i> 35
Figure 1.9	Exemple d'usage du détecteur-mélangeur pour améliorer la performance d'une RRA 38
Figure 2.1	La relation courant-tension d'un transistor, ainsi que des approximations locales 42
Figure 2.2	Spectres de l'entrée et de la sortie d'un détecteur quadratique idéal qui reçoit un signal sinusoïdal 45
Figure 2.3	Spectre de l'entrée et de la sortie d'un détecteur quadratique idéal qui reçoit un signal à deux fréquences 47
Figure 2.4	Spectre de l'entrée et de la sortie d'un mélangeur de fréquences quadratique 49
Figure 2.5	Les produits d'intermodulation à la sortie d'un détecteur et d'un mélangeur en théorie 61
Figure 3.1	Schéma détaillé du circuit intégré 64
Figure 3.2	Schéma simplifié du circuit intégré ICSTSEM1 64
Figure 3.3	Transistor NMOS non linéaires du circuit intégré 66

Figure 3.4	Modèle de transistor en petits signaux utilisé	67
Figure 3.5	Simulation DC des dérivées de la transconductance en fonction de la tension et la longueur de la grille pour un NMOS	71
Figure 3.6	Simulation en courant continu des dérivées de la transconductance d'un NMOS en fonction de la polarisation et de la largeur de la grille ...	72
Figure 3.7	Comparaison des gains en tension quadratique selon la simulation et les analyses	75
Figure 3.8	Étalement spectral causé par un miroir de courant non linéaire	76
Figure 3.9	Schéma des diodes de protection reliées aux pattes de soudure de la puce	77
Figure 3.10	Simulation de paramètres S de réflexion des entrées et de la sortie de la puce	79
Figure 3.11	Vue en coupe d'une ligne microruban	80
Figure 3.12	Principe de l'adaptation d'impédance avec un ligne microruban à largeurs alternées	81
Figure 3.13	Illustration de la technique d'adaptation d'impédance à largeurs de ligne alternant sur une abaque de Smith	82
Figure 4.1	Consommation de courant du circuit intégré au repos en fonction de la tension de polarisation	86
Figure 4.2	Tension à la sortie du circuit intégré au repos en fonction de la tension de polarisation	87
Figure 4.3	Mesure des coefficients de réflexion aux entrées de la partie principale du circuit imprimé	88
Figure 4.4	Mesure du coefficient de réflexion à la sortie de la partie principale du circuit imprimé	88
Figure 4.5	Coefficients de transmission entre les ports RF et LO	89
Figure 4.6	Banc d'essai du détecteur avec un signal sinusoïdal	90
Figure 4.7	Responsivités du détecteur simulée et mesurée	91
Figure 4.8	V_s mesurée en fonction de la puissance du signal à amplitude constante à l'entrée	92

Figure 4.9	$V_s - V_r$ en fonction de la puissance du signal à amplitude constante à l'entrée 92
Figure 4.10	Taux d'adaptation d'impédance aux entrées de la partie principale du circuit imprimé 94
Figure 4.11	Puissances mesurées à $2f_m$ et $4f_m$ à la sortie du détecteur en fonction de P_{RF} et de V_p 94
Figure 4.12	Le taux de réjection du produit d'intermodulation d'ordre quatre en fonction de P_{RF} 96
Figure 4.13	$P_{s@2f_m}$ et $P_{s@4f_m}$ en fonction de la tension de polarisation V_p 97
Figure 4.14	Banc d'essai du détecteur branché à l'oscilloscope 97
Figure 4.15	Comparaison à l'oscilloscope du signal modulant et de la sortie du détecteur 98
Figure 4.16	Simulation transitoire de la réactivité du détecteur 100
Figure 4.17	Réactivité simulée du détecteur en fonction de la tension de polarisation 101
Figure 4.18	Mesure à l'oscilloscope de la réactivité du détecteur 102
Figure 4.19	Simulation de la réjection de la porteuse en fonction de l'erreur de phase ϕ 104
Figure 4.20	Simulation de Monte Carlo de la réjection de la porteuse 105
Figure 4.21	Schéma du banc de test pour les mesures sous pointes 107
Figure 4.22	Photographie du circuit imprimé sous les sondes de mesure 107
Figure 4.23	Spectres de la sortie du mélangeur avec f_{RF} et f_{LO} qui varient 108
Figure 4.24	Spectres de la sortie du mélangeur avec f_{IF} et f_{LO} qui varient 108
Figure 4.25	Banc d'essai du mélangeur sur la partie principale du circuit imprimé .. 109
Figure 4.26	Spectre à la sortie du mélangeur avec les signaux RF et LO rémanents 110
Figure 4.27	Spectres à la sortie du mélangeur illustrant les basses fréquences 110
Figure 4.28	Puissance mesurée à la sortie IF en fonction de P_{RF} 111

Figure 4.29	Mesures des produits d'intermodulation d'ordre deux et quatre du mélangeur112
Figure 4.30	Banc de test pour les mesures à l'oscilloscope du mélangeur113
Figure 4.31	Signal IF à la sortie du mélangeur mesuré à l'oscilloscope comparé au signal modulant l'entrée114
Figure 4.32	Signal IF à la sortie du mélangeur mesuré à l'oscilloscope comparé au signal WCDMA modulant l'entrée114

LISTE DES ABRÉVIATIONS, SIGLES ET ACRONYMES

1P9M	1-Polysilicon 9-Metal
5G	Cinquième génération
6G	Sixième génération
ADS	Advanced Design System
AM	Amplitude modulation
ASK	Amplitude-shift keying
BiCMOS	Bipolar Complementary Metal-Oxide-Semiconductor
BIST	Built-in self-test
BIT	Built-in test
CMOS	Complementary Metal Oxide Semiconductor
DES	Décharge électrostatique
DNW	Deep N-well
ENIG	Electroless Nickel Immersion Gold
ETS	École de Technologie Supérieure
ESD	Electrostatic Discharge
FR-4	Flame Retardant four
GND	Ground
HBT	Heterojunction Bipolar Transistor
HEMT	High-Electron-Mobility Transistor
IEEE	Institute of Electrical and Electronics Engineers
IF	Intermediate frequency
IGFET	Insulated Gate Field-Effect Transistors
I/Q	In-phase/Quadrature
LaCIME	Laboratoire de Communications et d'Intégration de la Microélectronique
LNA	Low noise amplifier

LO	Local oscillator
MEMS	Microelectromechanical Systems
MMIC	Monolithic Microwave Integrated Circuit
MOSFET	Metal-Oxide-Semiconductor Field-Effect Transistor
NMOS	N-type Metal–Oxide–Semiconductor
NTN	Non-terrestrial network
OOK	On-off keying
PIN	Positif-intrinsèque-négatif
PMOS	P-type Metal–Oxide–Semiconductor
RF	Radiofréquence
RIS	Reconfigurable Intelligent Surface
RMS	Root mean squared
RRA	Reconfigurable reflectarray antenna
RTA	Reconfigurable transmitarray antenna
PD	Plage dynamique
SMA	SubMiniature version A
SNR	Signal-to-Noise Ratio
STAR-RIS	Simultaneously Transmitting And Reflecting Reconfigurable Intelligent Surface
STIN	Space-Terrestrial Integrated Network
TSMC	Taiwan Semiconductor Manufacturing Company
VCO	Voltage-controlled Oscillator
VDD	Tension d'alimentation
WCDMA	Wideband Code Division Multiple Access

LISTE DES SYMBOLES ET UNITÉS DE MESURE

α_i	Coefficient d'ordre i de la série de Taylor d'une fonction de transfert en tension non linéaire
ΔV	Résolution de tension continue d'un outil de mesure
ϵ_r	permittivité relative
λ	Constante modélisant l'effet de modulation de la longueur du canal d'un MOSFET
μ	Mobilité des porteurs de charge
μA	Microampère
μm	Micromètre
τ	Constante de temps de la réactivité d'un circuit RC
ϕ	Phase ou erreur de phase
ω	Pulsation
Ω	Ohm
A	Ampère
C_{gd}	Capacité grille-drain d'un MOSFET
C_{gs}	Capacité grille-source d'un MOSFET
C_{ox}	Capacité d'oxyde de grille d'un MOSFET
dB	Décibel
dBm	Décibel référencé à 1 milliwatt
dBV	Décibel référencé à 1 volt
f	Fréquence
f_c	Fréquence porteuse d'un signal RF
f_{IF}	Fréquence centrale d'un signal IF
f_{LO}	Fréquence fondamentale d'un signal LO
f_m	Fréquence de modulation
G_{cm}	Gain de conversion d'un mélangeur de fréquences

g_{dsi}	Coefficient d'ordre i de la série de Taylor de la dérivée I_{ds} par rapport à V_{ds}
GHz	Gigahertz
g_{mi}	Coefficient d'ordre i de la série de Taylor de la dérivée I_{ds} par rapport à V_{gs}
I_{ds}	Courant de drain d'un MOSFET
$IP_{1\text{ dB,d}}$	Puissance du signal à l'entrée provoquant 1 dB de compression à la sortie du détecteur
$IP_{1\text{ dB,m}}$	Puissance du signal à l'entrée provoquant 1 dB de compression à la sortie du mélangeur
$IP_{IP,d}$	Puissance du signal à l'entrée correspondant au point d'interception du détecteur
$IP_{IP,m}$	Puissance du signal à l'entrée correspondant au point d'interception du mélangeur
L	Longueur de la grille
mA	Milliampère
MHz	Mégahertz
mV	Millivolt
mW	Milliwatt
nF	Nanofarad
nH	Nanohenry
P_{IF}	Puissance d'un signal IF
P_{LO}	Puissance d'un signal LO
P_{RF}	Puissance d'un signal RF
R_d	Responsivité d'un détecteur de puissance
r_o	Résistance modélisant l'effet de modulation de la longueur du canal d'un MOSFET pour de petits signaux
S_{ij}	Coefficient de répartition des ondes de puissance entre les ports i et j
V	Volt

$V_{1\text{ dB}}$	Amplitude du signal d'entrée causant 1 dB de compression à la sortie
V_{ds}	Tension entre le drain et la source d'un transistor FET
v_e	Tension du signal à l'entrée en modèle des petits signaux
V_{gs}	Tension entre la grille et la source d'un transistors FET
$V_{IP,d}$	Amplitude du signal d'entrée correspondant au point d'interception du détecteur
$V_{IP,m}$	Amplitude du signal d'entrée correspondant au point d'interception du mélangeur
V_{LO}	Amplitude du signal LO
V_{ov}	Différence entre V_{gs} et V_{th}
V_p	Tension de polarisation
V_r	Tension continue à la sortie au repose
V_{RF}	Amplitude du signal RF
$V_{RF, \min}$	Plus petite amplitude du signal d'entrée mesurable
v_s	Tension de la sortie en modèle des petits signaux
V_s	Tension continue à la sortie
V_{sb}	Tension entre la source et le substrat
V_{th}	Tension de seuil des transistors FET
W	Watt ou largeur de la grille d'un MOSFET

INTRODUCTION

0.0.1 Motivation de la recherche

Les futurs réseaux 5G et 6G intégreront dans les réseaux de télécommunication terrestres existants des nœuds aériens et spatiaux, notamment sous forme de drones et de satellites, dans le but d'assurer une couverture globale des réseaux sans fil (Chen *et al.* (2023)). Les réseaux 5G dépendent du déploiement massif des stations de base équipées d'antennes réseau à commande de phase (appelés *phased array antennas* en anglais) capables de modelage de faisceau. Pour les réseaux du futur, des antennes plus simples et abordables deviendront nécessaires pour démocratiser davantage les moyens de télécommunication avancés (Guo & Ziolkowski (2022)). Des antennes alternatives également capables de modelage de faisceau, suscitent l'intérêt des chercheurs depuis plusieurs années. Dans cette catégorie, nous retrouvons notamment les *reconfigurable reflectarray antennas* (RRA) et les *reconfigurable transmitarray antennas* (RTA). Ces nouvelles antennes peuvent notamment être intéressantes pour des applications satellitaires du fait de leur compacité, de leur économie d'énergie et de leur fiabilité selon Fu *et al.* (2022). En plus de ces antennes, nous connaissons l'émergence des surfaces intelligentes reconfigurables, appelés *reconfigurable intelligent surfaces* (RIS) en anglais, qui ont de nombreux cas d'usage potentiels à l'avenir (Chen *et al.* (2023)).

Les RRA, RTA et RIS comportent souvent des déphaseurs au niveau des cellules d'antennes pour permettre le modelage de faisceau. Les diodes sont des composants privilégiés là où il existe un besoin de fiabilité et de rapidité, comme dans le milieu extraterrestre. En particulier, nous nous intéressons aux diodes à capacité variable (appelées *varicap diodes* en anglais) qui présentent plusieurs avantages par rapport aux diodes *positive-intrinsic-negative* (PIN). Le défaut majeur des diodes à capacité variable demeure néanmoins leur non-linéarité inhérente, qui engendre de la distorsion (Zhu *et al.* (2024); Fischer, Steinmetz & Adel (2025)). Une façon de remédier au problème serait d'adapter la polarisation de ces diodes en fonction de la puissance du signal et

de la distorsion mesurée. Or, sur une RRA ou sur une RTA, les cellules d'antennes se trouvent à différentes distances et angles par rapport à l'antenne source. La puissance de l'onde incidente est différente pour chaque cellule. De plus, chaque cellule d'antenne a besoin d'un déphasage différent de celui de ses voisines. Par conséquent, la correction de la polarisation des diodes est différente pour chacune des nombreuses cellules de l'antenne. Un détecteur simple et compact, pour chaque cellule d'antenne, qui mesure rapidement la puissance du signal reçu, permettrait un ajustement de la polarisation du déphaseur adapté à la cellule. Une autre façon d'adapter la polarisation des déphaseurs serait de mesurer le niveau de distorsion présent dans le signal déphasé afin d'apporter une meilleure correction. La mesure du niveau de distorsion est possible grâce aux outils de traitement du signal numérique. Cependant, le signal dans l'antenne est porté à des fréquences trop élevées pour un échantillonnage direct par un processeur numérique. Il faut donc transposer le signal vers une fréquence plus basse grâce à un mélangeur.

Les détecteurs quadratiques (*square-law*) à base de transistors en CMOS, comme celui dans le mémoire de Bourgault (2023), présentent une solution simple et éprouvée. D'autre part, il existe également des mélangeurs à base de transistors, utilisant également une réponse quadratique, qui répondent au critère de simplicité (Shahriar Rashid & Rashid (2010); Bao, Jacobsson, Aspemyr, Carchon & Sun (2006)). Des liens implicites entre ces circuits ont déjà été mentionnés dans la littérature par le passé (Bigdeli, Burasa & Wu (2025)). Cependant, nous n'avons trouvé aucun ouvrage qui n'a tenté d'intégrer les deux circuits en un seul. Donc, une des principales motivations de notre travail était de démontrer la faisabilité d'un circuit à double-fonctionnalité simple.

0.0.2 Objectifs de recherche

L'objectif initial de ce projet était de développer un circuit détecteur de puissance basé sur celui de Bourgault (2023) avec des performances améliorées. Cependant, ayant réalisé plus tard qu'il était également possible d'en faire un mélangeur, nous avons décidé de ne réaliser qu'un seul

circuit capable d'accomplir les deux tâches simultanément. Nous avons ainsi posé les objectifs suivants :

- Imaginer un circuit à double fonctionnalité de détecteur d'enveloppe de puissance et de mélangeur de fréquences.
- Le concevoir en technologie CMOS de 65 nm de TSMC pour le miniaturiser et faciliter une éventuelle intégration dans un système de radiofréquence (RF) plus complexe.
- Rendre le circuit adapté pour une utilisation sur une antenne :
 - En minimisant le courant consommé par le circuit intégré. En effet, dans une matrice d'antenne avec un très grand nombre de cellules, il faut autant de détecteurs-mélangeurs que de cellules. Il faut donc minimiser l'impact des circuits intégrés sur l'efficacité énergétique de l'antenne.
 - Il fallait que le circuit puisse non seulement fonctionner avec des signaux à l'entrée de 11 GHz, une fréquence utilisée par des satellites de communication (Government of Canada (2022)), mais aussi avec des signaux RF de fréquences variées pour offrir de la flexibilité.
 - Le mélangeur devait transposer le signal RF vers des fréquences intermédiaires suffisamment basses pour un traitement du signal numérique. La fréquence intermédiaire devait également pouvoir varier selon les besoins.
 - La bande passante à la sortie du circuit intégré devait être assez importante pour que le détecteur puisse suivre l'enveloppe des signaux à haut débit et élargir la plage des fréquences intermédiaires possibles.
- Le mélangeur de fréquence devait introduire le moins de distorsion possible lors de la transposition pour ne pas fausser la mesure du niveau de distorsion engendré par les diodes à capacité variable.

0.1 Contribution à la recherche

Ce mémoire contribue à l'enrichissement du savoir de différentes façons :

- Nous démontrons, par la théorie, les simulations et les expériences, la faisabilité d'un circuit spécifiquement conçu pour accomplir les deux fonctions de détection de puissance et de mélange.
- Nous enrichissons nos connaissances sur les détecteurs et les mélangeurs, ainsi que sur le lien entre ces circuits. Nous avons réalisé cela en :
 - étudiant les propriétés d'un circuit quadratique ;
 - analysant l'apparition des produits d'intermodulation et leur impact sur la performance du circuit ;
 - proposant une méthode pour analyser mathématiquement un circuit non linéaire.

0.2 Plan du mémoire

Suite à une revue de littérature sur les antennes reconfigurables, les détecteurs et les mélangeurs, nous analyserons la théorie nécessaire pour comprendre les mécanismes derrière les détecteurs et les mélangeurs. Ensuite, nous étudierons le cheminement de la conception du circuit du détecteur-mélangeur. Puis, nous mesurerons la performance du circuit conçu par simulation et expérimentalement. Enfin, avant de conclure, nous exposerons les pistes possibles d'amélioration du circuit.

CHAPITRE 1

REVUE DE LITTÉRATURE

1.1 Le besoin d'antennes reconfigurables pour les réseaux 5G et 6G

En 2024, 2,25 milliards d'individus à travers le monde avaient accès à un réseau 5G selon Nguyen (2025). L'Union européenne a revendiqué une couverture domestique de 94,3% dans un de ses rapports en 2025 (European 5G Observatory (2025)). Tandis que la couverture des réseaux 5G continue de s'étendre, les agences gouvernementales de la Chine, des États-Unis, du Japon, de la Corée du Sud et de l'Union européenne collaborent avec des industriels et des centres de recherche pour anticiper les réseaux 6G (Chen *et al.* (2023); Doczkat, Etemad & Gosain (2025)). Maintenant que la 5G a atteint une maturité technologique, on prévoit le déploiement de la 6G d'ici les années 2030, selon European Commission (2025). Selon Doczkat *et al.* (2025), l'un des objectifs de la 6G est l'ubiquité de la couverture. Dans ce but, on cherche non seulement à intégrer des *non-terrestrial networks* (NTN) dans les réseaux terrestres pour mieux desservir les régions isolées, mais également à repenser la couverture dans des zones urbaines pour établir des connexions fiables et efficaces.

Les réseaux 5G d'aujourd'hui dépendent des antennes réseau à commande de phase. C'est un ensemble de plus petites antennes arrangées de façon matricielle (Ansys (2025)). Chaque cellule d'antenne est reliée à une chaîne de signaux radiofréquence (RF) qui modifie l'amplitude et/ou la phase du signal de façon plus ou moins indépendante, selon l'architecture du système (Guo & Ziolkowski (2022)). Modifier l'amplitude et/ou la phase au niveau de chaque élément permet de modeler la forme du faisceau et l'orientation du faisceau principal sans aucun mouvement mécanique de l'antenne. Cette capacité est au cœur des évolutions des réseaux de télécommunication 5G et 6G. En effet, le modelage de faisceau sert notamment à concentrer la puissance dans la direction de l'utilisateur, ce qui permet d'économiser de l'énergie. De plus, concentrer le rayonnement permet l'usage de la même bande de fréquences par différentes stations de base sans qu'il y ait d'interférences (Chen *et al.* (2023); Guo & Ziolkowski (2022)).

De plus, rappelons-nous qu'une directivité plus importante de l'antenne donne un niveau de SNR plus élevé, ce qui permet, par conséquent, un débit d'information supérieur (Pozar (2012)).

Les satellites de communication joueront un rôle encore plus important dans les réseaux de demain. Afin d'atteindre une couverture planétaire totale, des satellites géostationnaires, ainsi que ceux en orbite terrestre basse, seront désormais intégrés dans ce que l'on appelle parfois le *Space-Terrestrial Integrated Network* (STIN), composé de nœuds terrestres et extraterrestres de natures variées (Yao, Wang, Wang, Lu & Liu (2018); Peng *et al.* (2025); Chen *et al.* (2023)). Dans les STIN, nous trouverons également des avions pilotés et autonomes. La capacité de modelage de faisceau est une capacité indispensable pour l'intégration des plateformes aérospatiales en mouvement constant : cela permet d'éviter des interférences avec d'autres usagers et de rendre la liaison plus fiable lorsque l'on communique avec une plateforme aérienne ou céleste constamment en mouvement (Chen *et al.* (2023)). Remodeler le faisceau est très utile lorsque l'antenne doit s'adapter à un environnement changeant. De plus, l'absence de mouvement mécanique contribue à la rapidité et à la fiabilité du système. Cependant, les antennes réseau réflecteur à commande de phase peuvent être très complexes et coûteuses, d'après Tang, Xu, Yang & Li (2021). Si les futurs réseaux doivent couvrir de vastes surfaces en dehors des centres urbains, et si nous souhaitons intégrer dans le réseau des véhicules autonomes en énergie, alors il faut trouver des solutions plus simples, économes et abordables (Guo & Ziolkowski (2022)).

Que ce soit dans l'espace, dans l'air ou sur terre, les antennes réseau capables de modelage de faisceau seront des composants intégrants du futur réseau 6G.

1.2 Les antennes reconfigurables

Les antennes traditionnelles, telles que les antennes paraboliques, doivent être réorientées mécaniquement pour changer l'orientation de leur faisceau. De plus, elles ne peuvent pas modifier la forme de leur rayonnement. Modeller le faisceau et réorienter le rayonnement sans aucun mouvement mécanique et de façon réversible ne sont pas des capacités propres aux antennes réseau à commande de phase. Il existe des solutions alternatives plus abordables.

1.2.1 Différents types d'antennes reconfigurables

Dans cette revue de la littérature, nous nous intéresserons en particulier aux *reconfigurable reflectarray antennas* (RRA), aux *reconfigurable transmitarray antennas* (RTA)¹ et aux *reconfigurable intelligent surfaces* (RIS). Ce dernier n'est techniquement pas une antenne, mais il demeure très pertinent. Tous ces systèmes utilisent une matrice de cellules d'antenne qui contrôle chacune la phase et/ou l'amplitude du signal.

1.2.1.1 Les RTA et les RRA

Dans une RTA, une antenne-source fixe rayonne d'un côté d'un réseau de cellules. Ces cellules se trouvent sur une surface plane à plusieurs couches. Dans certains cas, on peut même envisager l'illumination des cellules par plusieurs sources à la fois (Di Palma *et al.* (2016)). L'onde traverse les cellules de l'antenne, qui modifient les propriétés de l'onde. Fu *et al.* (2022); Hum & Perruisseau-Carrier (2014)

Les RRA sont très similaires aux RTA. Cependant, ici, l'onde provenant de l'antenne-source ne traverse pas la matrice de cellules, mais elle est réfléchiée. Fu *et al.* (2022) souligne une similitude entre les RRA et les antennes paraboliques dans leur principe de fonctionnement. Une RRA a tout-de-même l'avantage de pouvoir modeler le faisceau.

Comparons les RRA et les RTA. Par rapport aux RTA, le désavantage des RRA est le blocage de certaines directions par l'antenne-source. De plus, selon Di Palma *et al.* (2016), le placement de l'antenne source derrière les cellules, dans une RTA, est avantageux pour l'embarquement dans un véhicule avec une forte contrainte dimensionnelle. Selon Berthiaume (2022), les antennes RTA ont l'avantage par rapport aux RRA de mieux amplifier le signal reçu avant de le retransmettre. L'amplification compense entièrement la perte d'insertion de l'antenne RTA. Sur les antennes RRA, la possibilité d'amplification est plus limitée, parce que le couplage entre l'entrée et la sortie de l'amplificateur peut le rendre instable. C'est d'ailleurs ce qui a été réalisé par Pan,

¹ Différents acronymes existent dans la littérature, ce qui peut porter à confusion. Nous reprenons ici ceux de Berthiaume (2022) en ce qui concerne les RRA et les RTA

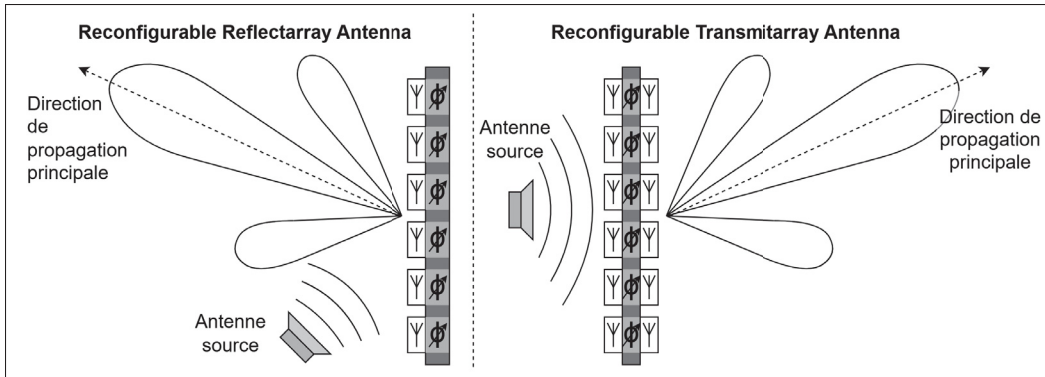


Figure 1.1 Comparaison entre une RRA et une RTA

Huang, Ma & Luo (2014) : une RTA avec des amplificateurs qui suivent les déphaseurs. L'usage des amplificateurs permet notamment de compenser les pertes d'insertion des déphaseurs. Enfin, Fu *et al.* (2022) indique que la bande passante des RRA et des RTA est intrinsèquement limitée par la dimension des cellules.

Dans une antenne réseau à commande de phase, chaque cellule d'antenne est liée à une chaîne RF, où l'on peut trouver des lignes de transmission, des filtres, des amplificateurs et des mélangeurs, en plus d'un déphaseur (Guo & Ziolkowski (2022)). En comparaison, dans une RRA ou une RTA, la liaison entre la source du signal et les cellules d'antenne se fait sans fil. Par conséquent, elles sont plus simples et moins coûteuses, car elles ont tout simplement moins de composants. Cette liaison sans fil permet également aux RRA et aux RTA de minimiser les pertes entre la source et les cellules d'antenne (Pan *et al.* (2014); Fu *et al.* (2022); Hum & Perruisseau-Carrier (2014)). Dans une antenne réseau à commande de phase, les différentes cellules ont besoin de se synchroniser, alors que dans les RRA et RTA, l'antenne-source sert de référence (Brunet, Ayling & Hajimiri (2024)). De plus, on évite les pertes entre la source et les cellules, selon Fu *et al.* (2022); Hum & Perruisseau-Carrier (2014).

1.2.1.2 Les RIS

Les RIS sont des surfaces matricielles capables de modifier intelligemment la propagation des ondes électromagnétiques. Chacune de ses cellules a la capacité de contrôler la réflexion, la

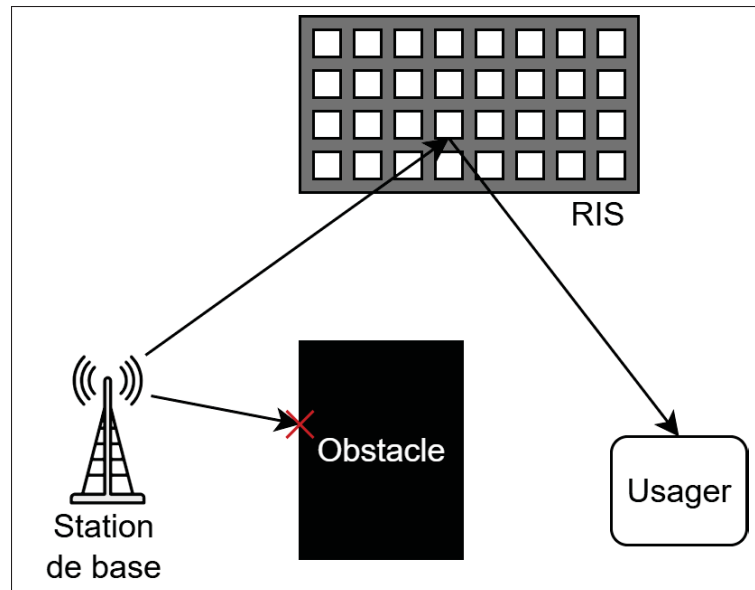


Figure 1.2 Usage d'une RIS pour améliorer la couverture malgré la présence d'obstacles

réfraction, voire la polarisation des ondes grâce à un contrôle électronique. Les RIS se présentent sous une forme passive, où les cellules ne font que déphaser le signal, et une forme active, où les cellules peuvent également amplifier pour compenser les pertes. Les RIS conventionnelles agissent comme un miroir qui réfléchit les ondes. Cependant, les STAR-RIS sont à la fois capables de réfléchir et de réfracter les ondes. Un circuit numérique, tel qu'un FPGA, contrôle ses nombreuses cellules. (Toka *et al.* (2024); Fu *et al.* (2022)) Si les RIS sont différentes des RRA, c'est parce que elles ne sont pas techniquement une antenne . Les RIS ne font que modifier les ondes incidentes provenant de sources externes.

Comme illustré par la figure 1.2, les RIS servent principalement à améliorer la couverture du réseau sans fil en réfléchissant intelligemment une onde provenant d'une station de base vers une direction qui était à l'origine inatteignable. Ils permettent d'étendre la couverture des réseaux 5G sans installer de nouvelles stations de base onéreuses, complexes et gourmandes en énergie (Zhu *et al.* (2024)). Leur usage est notamment envisagé pour les milieux urbains remplis d'obstacles où les ondes millimétriques peuvent facilement être bloquées (Toka *et al.* (2024); Doczkat *et al.* (2025)). En effet, une RIS, n'amplifiant pas de signal, consomme relativement peu d'énergie

(Fu *et al.* (2022)). Elles peuvent donc accomplir la même fonctionnalité que des relais, mais de façon quasi-passive, car elles ne sont pas la source d'un signal (Chen *et al.* (2023)).

Selon Fu *et al.* (2022), les RA, RTA et RIS sont faciles à fabriquer et peu onéreuses, car elles utilisent des technologies de circuits imprimés. De plus, une antenne plate est plus facile à ranger et à déployer là où l'espace est limité, comme sur un satellite. C'est notamment une raison d'utiliser une RA à la place d'une antenne parabolique (Huang & Encinar (2008)).

1.2.2 Les applications pour les RRA, RTA et RIS

Les satellites en orbite terrestre basse sont généralement des systèmes très compacts et peu coûteux, avec un budget énergétique limité, selon Toka *et al.* (2024). Aussi la simplicité, le bas coût, les faibles pertes, la compacité, la capacité à remodeler les faisceaux et la flexibilité font des antennes RA, RTA et RIS d'excellents candidats pour être embarquées sur ces satellites. Berthiaume, Laurin & Constantin (2021) Par exemple, Toka *et al.* (2024) propose d'embarquer les RIS sur les satellites en orbite terrestre basse pour ajuster la couverture du réseau sans fil au-dessus des régions moins bien desservies. Les ondes réfléchies sur les RIS d'un satellite peuvent contourner des obstacles terrestres ou la courbure de la planète. Ou encore, Brunet *et al.* (2024) proposent de s'en servir pour permettre le transfert de puissance entre les satellites géostationnaires et la Terre.

Les RRA et RTA peuvent être utilisées pour les radars, car la capacité d'orienter électroniquement, et non pas mécaniquement, le faisceau permet de suivre une cible avec une plus grande vitesse angulaire, selon Di Palma *et al.* (2016). Il existe des systèmes radars à balayage électronique réalisés grâce à des RRA (Romanofsky, Mueller & Chandrasekar (2009) et Huang *et al.* (2019)) et des RTA (Di Palma *et al.* (2016)). Pour cette application aussi, on préfère ces solutions aux antennes réseau à commande de phase en raison de leur prix et de leur simplicité. En effet, Romanofsky *et al.* (2009) affirment que leur solution, qui est une RRA envisagée pour être embarquée sur un satellite météorologique, est dix fois moins chère qu'une antenne réseau à commande de phase, et qu'elle est cinq fois moins énergivore.

1.2.3 Les déphaseurs

Le déphasage de l'onde incidente par chaque cellule permet aux antennes réseau de réorienter et de modeler leur faisceau. Les caractéristiques du déphaseur déterminent donc la performance globale de l'antenne. Étant donné que le but de cette recherche est de compenser les défauts des diodes à capacité variable, une des nombreuses technologies actuellement disponibles pour le déphasage, il est important de comprendre pour quelles raisons on souhaite choisir cette technologie en premier lieu.

1.2.3.1 Les technologies de déphaseurs

Pour le contrôle du déphasage, on peut se servir des diodes PIN, des diodes à capacité variable, des éléments ferromagnétiques, des MEMS ou encore des matériaux ajustables. Pour le bon fonctionnement de l'antenne, il est important que ces éléments offrent une grande plage de déphasage. De plus, il faut une faible perte d'insertion causée par ces éléments. Le changement de phase doit se faire linéairement avec la fréquence pour assurer une bonne bande passante. Enfin, ils doivent être fiables, compacts et économes en énergie. (Hum & Perruisseau-Carrier (2014); Berthiaume (2022))

Les diodes sont des technologies fiables et bien établies. Ces composants électroniques permettent la modification de faisceau très rapide, beaucoup plus rapide que les solutions mécaniques ou le déphasage à base de cristaux liquides, selon Fischer *et al.* (2025). L'avantage des diodes à capacité variable par rapport aux diodes PIN est leur capacité à faire varier la phase de façon continue. Les antennes utilisant des diodes PIN doivent augmenter le nombre de diodes PIN par cellule afin d'augmenter la résolution de contrôle et de diminuer la dégradation du gain, tandis que Fischer *et al.* (2025) et Huang *et al.* (2019) ont démontré que l'on peut réaliser un circuit de déphasage avec une seule diode à capacité variable. Par ailleurs, les diodes à capacité variable ne consomment pas de courant continu, car elles sont polarisées en inverse dans les circuits de déphaseur. Cependant, elles sont non linéaires par nature, causant ainsi la distorsion du signal de haute puissance (Hum & Perruisseau-Carrier (2014) Berthiaume (2022)). L'usage

d'un moyen de déphasage n'empêche pas l'emploi d'autres moyens. Par exemple, Tang *et al.* (2021) utilisent une combinaison de diodes à capacité variable et de diodes PIN pour obtenir une plage de déphasage de 360° en combinant les avantages des deux types de diodes.

1.2.3.2 Les diodes à capacité variable

Les diodes à capacité variable sont spécifiquement conçues pour être utilisées comme réactances variables. En contrôlant la tension de polarisation aux bornes de la diode, on modifie l'épaisseur de la zone de déplétion de la diode. C'est ainsi que l'on fait varier la capacité de jonction (Norwood & Shatz (1968)). Depuis la seconde moitié du 20^{ème} siècle, ces diodes sont utilisées dans les circuits RF pour le réglage de circuits d'adaptation, des oscillateurs contrôlés par tension (VCO) ou des filtres ajustables, d'après Norwood & Shatz (1968), de Dieuleveult & Romain (2008) et Buisman *et al.* (2005). En ce qui concerne les antennes à réseau, elles ont fait leurs preuves dans des antennes RRA, RTA et RIS (Zhu *et al.* (2024); Tang *et al.* (2021); Huang *et al.* (2019)).

La relation entre la capacité et la tension d'une diode à capacité variable est la suivante :

$$C(V) = \left(\frac{K}{\Phi + V} \right)^n$$

(tirée de Meyer & Stephens (1975) et de Buisman *et al.* (2005).) $C(V)$ est la capacité de la diode, K est une constante de capacité, Φ est le potentiel intrinsèque, V est la tension aux bornes de la diode et n est l'exposant décrivant empiriquement la diode. n dépend notamment de la distribution des impuretés dans la diode, selon Norwood & Shatz (1968).

Meyer & Stephens (1975) indique qu'en raison de la relation non linéaire entre la capacité de jonction des diodes et la tension à leurs bornes, ces diodes engendrent de la distorsion. La distorsion est d'autant plus apparente que le signal est puissant. Pour remédier à cela, Buisman *et al.* (2005) et Berthiaume *et al.* (2021) ont proposé des configurations utilisant plusieurs diodes pour réduire la distorsion. Ces efforts réduisent fortement le taux de distorsion, mais ils ne constituent pas des solutions absolues en soi.

1.3 Les détecteurs

Un détecteur est un circuit qui permet de détecter la présence d'un signal et de mesurer la puissance ou l'amplitude de la tension. Ce qu'il mesure exactement varie selon les besoins et les applications.

Il existe de nombreux circuits portant l'appellation de « détecteur » utilisant différentes technologies, avec des applications très variées. Par conséquent, nous nous concentrerons surtout sur les détecteurs dont le fonctionnement se base sur des semi-conducteurs. De plus, nous nous intéresserons seulement aux détecteurs conçus pour des applications RF. Malgré leur grande variété, nous pouvons affirmer que le fonctionnement des détecteurs repose sur la non-linéarité d'un composant. Par conséquent, contrairement à un amplificateur dont les signaux d'entrée et de sortie se situent sur la même bande de fréquence, le signal sortant d'un détecteur occupe la bande de base, une bande de fréquence plus basse que celle du signal modulé entrant.

1.3.1 Les détecteurs d'enveloppe de tension

Les détecteurs d'enveloppe de tension (aussi appelés détecteurs de crête ou d'amplitude) mesurent une ligne imaginaire reliant les crêtes d'un signal, comme celle illustrée sur la figure 1.3. Sur cette figure, l'amplitude de la porteuse à haute fréquence est modulée par un signal de plus basse fréquence. L'enveloppe suit la variation du signal modulant. Il faut retenir que ce genre de détecteur a une réponse proportionnelle à l'amplitude du signal à l'entrée.

La topologie du détecteur d'enveloppe de tension la plus connue et la plus simple est composée d'une diode suivie d'un filtre passe-bas : le redresseur, représenté sur la figure 1.4. La diode, l'élément non linéaire, ne laisse passer le courant que dans un seul sens, puis le condensateur maintient la tension jusqu'à la prochaine crête. Le montage basique n'utilise qu'une diode pour redresser le signal que sur un demi-cycle du signal. Utiliser plus de diodes permet de redresser le signal également pendant le cycle négatif (Bourgault (2023)). Les redresseurs ne peuvent opérer qu'avec des signaux d'amplitude importante, car il faut dépasser la tension de seuil de la diode pour qu'un courant puisse le traverser ; on dit que ce genre de détecteur utilise des

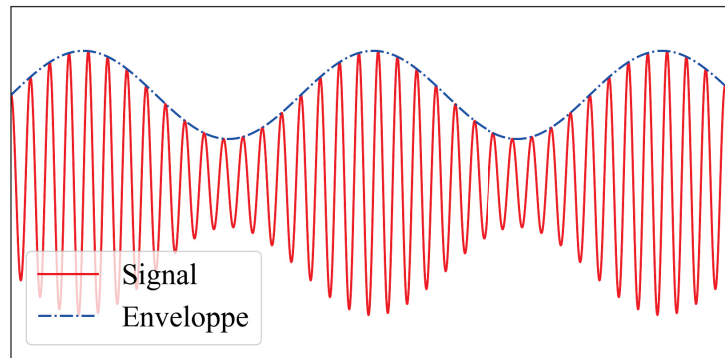


Figure 1.3 Un signal modulé en amplitude et son enveloppe

principes « grands signaux ». Par conséquent, les redresseurs ne sont pas les circuits les plus sensibles (Meyer (1995)). Puisque l'on travaille avec des signaux de grande amplitude, on peut supposer que la relation courant-tension de la diode (en réalité exponentielle) est quasi-idéale : on peut supposer que la diode est un court-circuit pour une tension supérieure à celle du seuil. Les détecteurs d'enveloppe de tension à base de diodes ont l'avantage d'être simples, car ils ne nécessitent même pas de polarisation. De plus, les redresseurs à diodes sont connus pour pouvoir opérer à de hautes fréquences, selon Meyer (1995).

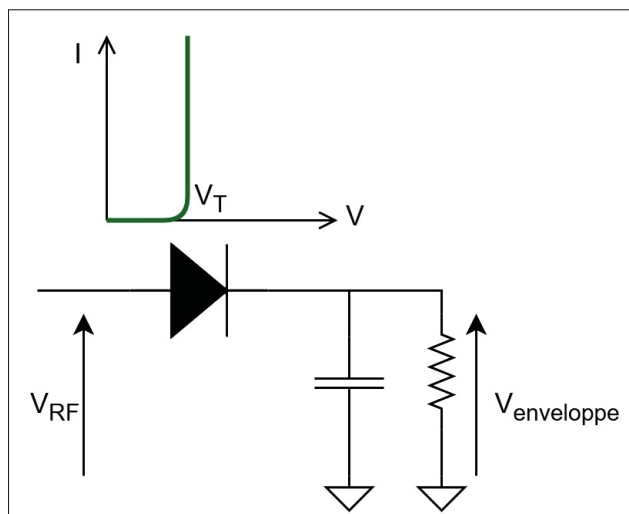


Figure 1.4 Schéma d'un redresseur comme détecteur d'enveloppe

Il existe évidemment d'autres manières plus avancées, afin de réaliser un détecteur de crête. Su, Xia, Geng & Liu (2019) utilisent des MOSFET et leur relation courant-tension quadratique pour réaliser un détecteur de crête. Cha *et al.* (2009) et Xia & Boumaiza (2015) réalisent également cela grâce à des redresseurs en MOSFET. Quant à Bourgault (2023), il réalise également un détecteur dont la sortie suit la variation d'amplitude du signal à l'entrée à base de MOSFET. Cependant, contrairement aux détecteurs précédemment mentionnés, l'amplitude de la sortie de son détecteur n'est pas proportionnelle à celle de l'entrée ; son circuit est un détecteur de puissance.

1.3.2 Détecteurs de puissance

Avec les détecteurs de puissance, la tension à la sortie est proportionnelle au carré de celle à l'entrée. Cela peut nous rappeler la relation entre la tension et la puissance.

Les détecteurs de puissance peuvent être réalisés en utilisant les propriétés thermiques d'un matériau ou les propriétés électriques des diodes ou des transistors. La détection thermique, bien qu'elle soit très utilisée pour des instruments de laboratoire, est difficile à mettre en œuvre dans un circuit intégré (Zhou & Chia (2008); Zhang, Eisenstadt, Fox & Yin (2006)). Ensuite, les détecteurs de puissance à base de diodes peuvent également être très simples ; ils peuvent fonctionner à très haute fréquence et ils peuvent être extrêmement sensibles. D'ailleurs, la topologie des détecteurs de puissance à base de diode la plus simple est très similaire à celle d'un redresseur. Saeed *et al.* (2021) mentionne que les détecteurs à base de diodes sont avantageux en termes d'efficacité énergétique par rapport aux circuits de transistors. Cependant, les détecteurs de puissance à base de diode peuvent être plus sensibles que les redresseurs, car les premiers utilisent les principes dits « petits-signaux ». Pour leur capacité à détecter des signaux faibles, ils ont été utilisés pour la radioastronomie dans le passé (Reid, Gardner & Stelzried (1975); Frater (1964)). Plus récemment, Zirath & He (2010) et Saeed *et al.* (2021) ont développé des détecteurs de puissance respectivement à base d'une diode Schottky et d'une diode en graphène. Leurs circuits pouvaient respectivement opérer jusqu'à 60 GHz et 70 GHz. Cependant, malgré sa simplicité, selon Zhang *et al.* (2006) et Ivanov, Lavrov & Matveev (2015), les circuits à base

de diodes sont plutôt limités en plage dynamique : au-delà d'une certaine puissance, la relation quadratique entre l'entrée et la sortie n'est plus valide et ce détecteur devient un redresseur. En outre, la performance des détecteurs qui utilisent des diodes est à la merci du procédé de fabrication et des caractéristiques exactes du composant (Zhang *et al.* (2006); Zhou & Chia (2008)). Nous nous intéresserons par la suite principalement aux détecteurs de puissance à base de transistors.

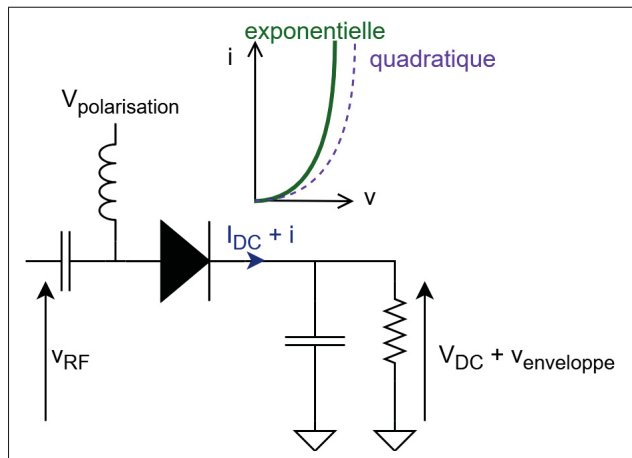


Figure 1.5 Schéma d'un détecteur de puissance réalisée grâce à une diode

Tous les détecteurs à base de diodes et de transistors utilisent la *square-law* : la relation (approximativement) quadratique entre le courant et la tension. Cependant, les détecteurs de puissance peuvent exister sous différentes formes. Nous pouvons les regrouper dans deux grandes familles : celle des détecteurs d'enveloppe de puissance et celle des détecteurs de puissance moyenne.

1.3.2.1 Les détecteurs de puissance moyenne

Le détecteur de puissance moyenne, aussi connu sous le nom *root mean square (RMS) power detector* en anglais, mesure la puissance moyenne au lieu de la puissance instantanée d'un signal. Sa réponse reste identique, peu importe la forme du signal, c'est-à-dire la fréquence ou la modulation du signal (Zhang *et al.* (2006)). Selon Zhou & Chia (2008) et Ravanne, Then,

Su & Hijazin (2021), cette caractéristique permet au détecteur de puissance moyenne de mesurer précisément la puissance d'un signal dont l'amplitude varie beaucoup.

Son fonctionnement devient plus clair lorsque en revoyant la formule de la puissance moyenne d'un signal :

$$\langle P(t) \rangle = \frac{1}{T} \int_{t_0}^{t_0+T} \frac{V^2(t)}{R} dt \quad (1.1)$$

Les détecteurs de puissance moyenne obtiennent la moyenne d'un signal, après l'avoir mis au carré, grâce à un filtre passe-bas ; c'est notamment ainsi que fonctionnent ceux de Ravanne *et al.* (2021), de Zhang *et al.* (2006) et de Yang, Uchida, Liu & Yoshimasu (2012), et c'est ainsi que fonctionne l'étage d'entrée du circuit de Zhou & Chia (2008). La période sur laquelle la moyenne est faite est largement supérieure à la période du signal porteur et du signal modulant ; c'est pourquoi la tension à la sortie paraît constante. La vitesse à laquelle ces détecteurs réagissent au changement d'amplitude du signal entrant dépend du filtrage passe-bas à la sortie. Cependant, elle n'importe point pour les détecteurs de puissance moyenne. Par ailleurs, on reconnaît un détecteur de puissance moyenne si la tension continue à la sortie s'accroît proportionnellement à la puissance du signal à l'entrée ; c'est-à-dire que la tension à la sortie augmente d'une décade lorsque la puissance à l'entrée augmente d'une décade. Zhou & Chia (2008) ont fait le choix particulier d'ajouter un générateur de courant exponentiel, ce qui fait que la tension à la sortie de leur circuit est proportionnelle à la puissance à l'entrée en décibels.

Dans notre cas d'usage, nous souhaitons ajuster la polarisation des diodes à capacité variable dès que le signal augmente en amplitude. Le temps de réaction d'un détecteur de puissance moyenne est beaucoup trop important pour notre cas. Par conséquent, ce type de détecteurs n'est pas d'intérêt pour nous.

1.3.2.2 Les détecteurs d'enveloppe de puissance

Les détecteurs d'enveloppe de puissance ont également une réponse quadratique, comme les détecteurs de puissance moyenne. Cependant, comme les détecteurs d'enveloppe, leurs sorties

suivent la modulation d'amplitude du signal d'entrée. Le détecteur de Bourgault (2023), celui de Serhan, Lauga-Larroze & Fournier (2015a), celui de Saeed *et al.* (2021) et le circuit actif proposé par Zirath & He (2010) tombent clairement dans cette catégorie, car ils démontrent une capacité à réagir rapidement aux changements d'amplitude du signal. Ils possèdent donc une bande passante de sortie importante. Ces circuits sont très similaires aux détecteurs vus précédemment. Cependant, la fréquence de coupure du filtre passe-bas à leurs sorties est fixée plus haut que celle des détecteurs de puissance moyenne. (Dans le cas de Bourgault (2023), il ne se trouve même pas de condensateur de filtrage à la sortie ; le filtrage est réalisé par les capacités parasites.) Pour apprendre sur la réponse d'un détecteur d'enveloppe de puissance, nous pouvons de nouveau mesurer la tension continue à la sortie en fonction de la puissance d'un signal à amplitude constante à l'entrée (Serhan *et al.* (2015a)), comme pour les détecteurs de puissance moyenne.

1.3.3 Les caractéristiques d'un détecteur d'enveloppe de puissance

Plusieurs caractéristiques, qui reviennent souvent dans la littérature, permettent d'évaluer la performance d'un détecteur de puissance.

1.3.3.1 La responsivité

Sur les détecteurs de puissance (moyenne ou d'enveloppe), la variation de la tension de sortie est proportionnelle au carré de la tension d'entrée, c'est-à-dire à la puissance du signal d'entrée. La responsivité (*responsivity* en anglais) quantifie la capacité du circuit à traduire une puissance en une tension. Manifestement, l'appellation même de cette notion n'est pas universelle, car certains ouvrages s'y réfèrent comme *conversion gain* ou comme *sensitivity* (Serhan *et al.* (2015a); Ivanov *et al.* (2015)). Cependant, le terme *sensitivity* prête à confusion, car Bourgault (2023) et Li, Yi, Boon & Yang (2019) utilisent ce terme pour désigner la plus petite puissance détectable. D'autre part, nous réserverons le terme « gain de conversion » pour les mélangeurs. Par conséquent, nous privilégierons le terme « responsivité » par la suite en nous appuyant sur la définition de l'Office québécois de la langue française (2025).

Outre la terminologie, cette grandeur est généralement exprimée en V/W (ou dans une unité qui en découle). Dans un grand nombre de travaux concernant les détecteurs, les auteurs illustrent le lien entre la tension de sortie et l'amplitude ou la puissance du signal d'entrée sur des graphes en échelle logarithmique. En effet, cette représentation met en évidence la réponse quadratique d'un détecteur de puissance : la variation d'une décade de la tension de sortie pour une décade de variation de la puissance du signal à l'entrée.

Par conséquent, nous proposons de définir, pour la suite de ce mémoire, la responsivité d'un détecteur R_d ainsi :

$$R_d = \frac{|V_s - V_r|}{P_{RF}} \quad (1.2)$$

où V_s est la tension continue à la sortie en volts, V_r est la tension à la sortie au repos en volts et P_{RF} est la puissance du signal RF à l'entrée en watts. La tension au repos est le potentiel à la sortie lorsque le circuit ne reçoit aucun signal à l'entrée. La partie 2.4.1 explique plus en détail cette définition du gain. Bien que cette formule ne soit pas apparue dans notre recherche bibliographique, elle nous permet tout de même de caractériser et de comparer les détecteurs de puissance que nous avons trouvés dans la littérature, grâce aux graphes inclus dans les ouvrages.

Par ailleurs, nous tenions à remarquer que certains circuits, tels que ceux de Zhou & Chia (2008); Ravanne *et al.* (2021); Zhang *et al.* (2006), utilisent des topologies quasi-symétriques, où l'entrée est asymétrique et la sortie est différentielle. Dans le cas de ces circuits, sans signal à l'entrée, la tension différentielle à la sortie est nulle. Donc, la tension différentielle mesurée à leur sortie correspond à $|V_s - V_r|$.

1.3.3.2 La plage dynamique

La plage dynamique (PD) est la largeur de l'intervalle de puissance d'entrée sur lequel le détecteur de puissance démontre avoir la réponse quadratique désirée. En ce qui concerne les détecteurs de puissance, on ne mentionne que la PD de la puissance du signal à l'entrée, et non pas celle de la tension de sortie. La plage dynamique se mesure habituellement en décibels.

La définition exacte semble dépendre de l'ouvrage. Certains le définissent comme la plage sur laquelle la responsivité varie de moins de 1 dB par rapport à sa valeur nominale (Li *et al.* (2019); Xia & Boumaiza (2015)). D'autres, tels que Qayyum & Negra (2017) et Ivanov *et al.* (2015), la définissent comme la différence entre le point de compression de 1 dB et la plus faible puissance mesurée. Enfin, nombreux sont ceux qui ne la définissent pas (Serhan *et al.* (2015a); Zhou & Chia (2008); Zhang *et al.* (2006); Serhan, Lauga-Larroze & Fournier (2015b); Saeed *et al.* (2021); Bigdeli *et al.* (2025)).

La PD peut tout de même être étendue grâce à des techniques telles que la rétroaction ; c'est ce qui a été réalisé par Ivanov *et al.* (2015).

1.3.3.3 La borne inférieure de la plage dynamique

Dans plusieurs ouvrages sur les détecteurs de puissance, nous pouvons retrouver les termes *minimum detectable power*, *minimum detectable input power* ou *minimum detectable signal power*. Cette grandeur décrit la plus faible puissance que le détecteur peut mesurer de manière fiable. Elle est généralement mesurée en dBm. Si elle est importante, c'est parce qu'elle correspond souvent à la borne inférieure de la plage dynamique. Nous remarquons que, dans les travaux, cette notion est traitée de façons variées. D'abord, certains, tels que Xu, Yang, Wang & Yoshimasu (2014), Valdes-Garcia, Venkatasubramanian, Silva-Martinez & Sanchez-Sinencio (2008) et Yang *et al.* (2012) omettent une définition. Ensuite, Serhan *et al.* (2015a) définissent cette grandeur comme « la puissance d'entrée pour laquelle le gain de conversion est égal à la moitié de sa valeur maximale ». (En effet, le « gain de conversion » de leur circuit s'accroît puis décroît avec la puissance d'entrée qui augmente.) Or, dans un autre ouvrage, les mêmes auteurs définissent la *minimum detectable power* comme la puissance d'entrée donnant un ratio signal sur bruit (SNR) de 10 dB (Serhan *et al.* (2015b)). Quant à Ivanov *et al.* (2015), ils mesurent la puissance minimale du signal d'entrée à deux fréquences qui donne un SNR de 8 dB avec une bande passante vidéo de 1 kHz sur l'analyseur de spectre. La définition utilisée varie d'un auteur à l'autre et semble dépendre du cas d'usage particulier du détecteur en question.

Dans notre cas, nous mesurons dans la partie 4.3.1 la tension continue de sortie de notre détecteur V_s avec un multimètre en faisant varier la puissance du signal à l'entrée P_{RF} . Lors de cette mesure, nous considérons que la limite inférieure de notre plage dynamique est là où la variation de la tension de sortie $|V_s - V_r|$ est tellement faible que la lecture au multimètre cesse d'être fiable. En effet, en dessous de ce niveau de puissance, la fluctuation aléatoire de la tension, causée par le bruit, est supérieure à la variation de tension, que l'on souhaite mesurer. Étant donné que le but de notre détecteur n'est pas de mesurer de faibles puissances, cette définition nous suffit et nous permet de garder un banc d'essai simple.

1.3.3.4 Le point de compression

Au-delà d'un certain niveau de puissance, la tension de sortie ne croît généralement plus de manière quadratique, et la responsivité diminue. Au point de compression à 1 dB, la tension de sortie est inférieure de 1 dB à la réponse quadratique idéale. (Voir la figure 2.5.) Ce point est souvent considéré comme la limite supérieure de la plage dynamique. Pour les détecteurs, on mesure la puissance d'entrée (exprimée en dBm) correspondant au point de compression. Par la suite, nous noterons cette puissance $IP_{1\text{ dB,d}}$.

1.3.3.5 La plage de fréquence

La plage de fréquence (ou la bande passante de l'entrée), exprimée en hertzs, désigne l'intervalle de fréquence du signal d'entrée sur lequel un circuit de détecteur a démontré être fonctionnel avec une responsivité plutôt constante. Cette plage peut être limitée par les composants qui servent à adapter l'impédance et par les éléments parasites.

1.3.3.6 La bande passante de sortie (ou la rapidité)

En ce qui concerne les détecteurs d'enveloppe de puissance, nous nous intéressons également à la bande passante de sortie, aussi appelée « réactivité du circuit ». Elle caractérise la capacité du détecteur à réagir aux changements rapides d'amplitude du signal d'entrée. D'un point de

vue fréquentiel, si le signal d'entrée RF occupe une certaine largeur de bande, alors le signal de l'enveloppe de puissance à la sortie est tout aussi large, mais il se trouve en bande de base. La bande passante de sortie représente la largeur de bande maximale que le détecteur peut tolérer avant que les éléments capacitifs ne filtrent une partie du signal. La bande passante peut être exprimée en hertzs. Elle peut également être exprimée en bande passante fractionnelle, qui correspond au rapport de la bande passante sur la fréquence de la porteuse du signal (Bourgault (2023)). Ou encore, on peut parler du temps de réaction du circuit à un changement brusque d'amplitude de signal (Serhan *et al.* (2015a)). La bande passante fractionnelle de sortie dans les ouvrages précédant Bourgault (2023) est faible, selon lui.

1.3.3.7 La consommation de puissance

Pour de nombreuses applications, la consommation de puissance est un critère important. Dans les publications, on cite la puissance ou le courant consommé par le circuit au repos.

1.3.3.8 La technologie

De nombreuses technologies de semi-conducteurs peuvent être employées pour réaliser un détecteur.

La technologie de fabrication de circuits intégrés la plus répandue au monde est sans doute le CMOS. Son avantage majeur est son coût de fabrication et sa facilité d'intégration dans des systèmes sur puce, appelés *system on chip* en anglais. Cependant, purement en termes de performance, Serhan *et al.* (2015b) indiquent que les détecteurs CMOS ont tendance à être limités en PD et en fréquence d'opération par rapport aux détecteurs en technologie bipolaire. Qayyum & Negra (2017) affirme carrément une infériorité de performance des détecteurs CMOS par rapport à leurs cousins bipolaires ou à haute mobilité électronique (HEMT). Selon Ravanne *et al.* (2021), les détecteurs en BiCMOS présentent notamment un avantage en termes de bruit, tandis que les technologies HEMT offrent l'avantage d'une haute fréquence de coupure. Zirath & He (2010) développent des circuits de détecteur à base d'une diode Schottky et

d'un transistor en technologie HEMT pour des fréquences allant jusqu'à 90 GHz. Ce genre de technologie est adapté à des opérations à très haute fréquence, dans la gamme des ondes millimétriques.

1.3.3.9 La taille

Dans des circuits intégrés, les détecteurs font partie d'un ensemble de circuits plus large. Par conséquent, il est important de minimiser la surface occupée par le circuit. Dans les publications, on peut inclure ou non les composants périphériques, tels que des condensateurs de découplage, lorsqu'on parle de la surface occupée. La surface s'exprime souvent en μm^2 pour les circuits intégrés.

1.3.4 Les applications pour les détecteurs de puissance

Il existe de nombreuses applications pour les détecteurs de puissance. Nous n'examinerons que le petit nombre d'applications qui revient le plus souvent dans les ouvrages.

1.3.4.1 La démodulation d'amplitude

Un détecteur d'enveloppe peut servir à la démodulation d'un signal purement modulé en amplitude, tel qu'un signal modulé en AM, OOK ou ASK. En effet, un détecteur d'enveloppe élimine la porteuse du signal et ne garde que l'amplitude du modulant. C'est d'ailleurs sur un signal OOK que Zirath & He (2010) testent leur circuit.

1.3.4.2 Built-in test

La tendance actuelle à intégrer des systèmes RF entiers dans une seule puce pose des défis pour l'essai de ces puces lors de la fabrication (Serhan *et al.* (2015b)). L'essai des puces constitue un obstacle majeur à la réduction des coûts de fabrication des circuits intégrés, selon Zhang, Gharpurey & Abraham (2007). Pour simplifier la vérification du fonctionnement des parties des puces, il est possible d'intégrer dans la même puce une partie dédiée à l'essai. C'est ce qu'on

appelle *built-in test* (BIT) ou *built-in-self-test* (BIST) en anglais. Les détecteurs intégrés sont un moyen d'essai simple et peu coûteux. L'ajout d'un petit circuit de détection n'engendre pas de coût additionnel lors de la fabrication. De plus, cela ne nécessite pas l'utilisation d'instruments RF précis et sophistiqués, car le signal de sortie d'un tel détecteur est en bande de base (Serhan *et al.* (2015b)). En outre, comme démontré par Zhang *et al.* (2007), il est possible d'utiliser un tel détecteur pour mesurer des caractéristiques précises d'un circuit. Dans leur cas, un mélangeur était testé grâce à un détecteur d'enveloppe. Le gain de conversion et le taux de distorsion du mélangeur furent mesurés grâce au détecteur.

1.3.4.3 Le contrôle de niveau automatique

Les détecteurs de puissance sont utiles pour l'ajustement des circuits en fonction du niveau de puissance du signal. Cela peut être, par exemple, pour ajuster le gain d'un amplificateur (Ozeren *et al.* (2016)). Selon Yang *et al.* (2012), les amplificateurs de puissance, qui sont énergivores, peuvent être contrôlés grâce à une chaîne de rétroaction pour réduire la consommation de puissance. Grâce au niveau de puissance à la sortie de l'amplificateur mesuré par le détecteur, on peut ajuster la polarisation de l'amplificateur (Bourgault (2023)). Dans ce genre d'application, la plus petite puissance mesurable par le détecteur est un critère secondaire, selon Bourgault (2023) et Serhan *et al.* (2015a).

1.3.4.4 Autres applications

Li *et al.* (2019) ont conçu un détecteur de puissance dont le but est de mesurer la puissance transmise lors d'un rechargement sans fil. Ils justifient le besoin d'un détecteur à grande plage dynamique et avec un haut point de compression par le fait que la puissance transmise peut beaucoup varier lors d'un rechargement sans fil. Ils n'ont pas opté pour un détecteur de puissance quadratique avec un transistor, parce qu'il devait pouvoir détecter des puissances importantes. Placer un atténuateur en amont du détecteur n'aurait pas satisfait leurs besoins non plus.

1.3.5 Les travaux précédents de détecteurs de puissance quadratiques

Dans le but de comparer quantitativement différents détecteurs de puissance de la littérature, nous avons regroupé dans le tableau 1.1 plusieurs travaux utilisant des topologies et des technologies différentes. Le tableau comprend à la fois les détecteurs de puissance moyenne et les détecteurs d'enveloppe de puissance. (Puisque certains ouvrages ne citaient pas toutes les caractéristiques, nous avons donc dû déduire les responsivités et les plages dynamiques à partir des graphiques présentés dans ces travaux.)

Tableau 1.1 Tableau de synthèse des détecteurs de puissance de la revue de littérature

Ouvrages	Technologie	R_d [V/mW]	PD [dB]	Fréquence
Xu <i>et al.</i> (2014)	CMOS 65 nm	0,25*	25 à 30 *	100 MHz à 44 GHz
Li <i>et al.</i> (2019)	CMOS 65 nm	$\sim 5.10^{-3}$ *	34	4 à 6 GHz
Qayyum & Negra (2017)	CMOS 130 nm	7*	24	7 à 77 GHz
Serhan <i>et al.</i> (2015a)	BiCMOS 55 nm	0,32*	30	50 à 66 GHz
Ivanov <i>et al.</i> (2015)	HBT discrets	1	55	4 GHz
Zhang <i>et al.</i> (2006)	IBM 7WL bipolaire	1	40	60 GHz
Ravanne <i>et al.</i> (2021)	MOSFET discrets	0,8*	60	2 GHz
Yang <i>et al.</i> (2012)	CMOS 130 nm	1,5*	10*	100 MHz à 40 GHz
Zirath & He (2010) ^A	MMIC HEMT 150 nm	0,5	17,5*	60 GHz
Zirath & He (2010) ^B	MMIC HEMT 150 nm	2,8	55*	10 à 60 GHz

* : donnée déduite à partir des graphes

A : circuit à base de diode de la première partie

B : circuit à base de transistor de la seconde partie

1.3.5.1 Distorsion dans un détecteur

Bourgault (2023) mesure, dans son mémoire, différents produits d'intermodulation générés par son détecteur d'enveloppe de puissance, ce qui soulève la question de la distorsion concernant les détecteurs. Cependant, notre recherche bibliographique révèle que peu d'articles traitant des détecteurs de puissance évoquent la distorsion. Bien que plusieurs auteurs mesurent le point de compression, celui-ci n'est qu'une facette de la non-linéarité. À propos des amplificateurs et des mélangeurs, les chercheurs et les fabricants de circuits citent le point d'interception

comme critère de distorsion. En revanche, à notre connaissance, aucun critère équivalent n'est systématiquement utilisé pour évaluer les détecteurs. Nous pouvons donc constater un manque d'information dans la littérature scientifique au sujet de la cause, des conséquences et des solutions liées à la distorsion dans les détecteurs.

1.4 Les mélangeurs de fréquence

1.4.1 Le principe d'un mélangeur de fréquence

Les mélangeurs de fréquences sont des composants essentiels pour les systèmes de communication sans fil. Ils sont omniprésents à la fois dans les chaînes de transmission et de réception. La principale utilité d'un mélangeur en télécommunication est la transposition de fréquence : c'est-à-dire le changement de la fréquence porteuse d'un signal. Le mélangeur fait cela en multipliant le signal par un signal d'oscillation appelé LO, comme *local oscillator* en anglais. (Parfois, le signal LO est connu sous le nom de *pump signal* (Ludwig & Bretchko (2000))) Les trois ports d'un mélangeur sont désignés RF, LO et IF. Le signal de radiofréquence (RF) est de haute fréquence et il porte l'information. Le signal LO est une oscillation sinusoïdale ou carrée. Enfin, le signal de fréquence intermédiaire (IF) est à plus basse fréquence que le signal RF, mais il contient les mêmes informations. Par la multiplication de deux signaux, un mélange transpose simultanément un signal vers le haut et vers le bas (cf. partie 2.3.1). Selon le cas d'usage, on ne garde qu'un seul des deux produits. Pour le mélangeur d'un transmetteur, les entrées sont les ports IF et LO ; le mélangeur transpose le signal IF vers le haut afin d'obtenir le signal RF. Dans un récepteur, on multiplie les signaux RF et LO pour transposer le premier vers le bas. Cela donne le signal IF. (de Dieuleveult & Romain (2008)) La multiplication se fait soit grâce à la non-linéarité des composants, soit grâce à une variance du circuit dans le temps (telle une commutation). Les composants non linéaires sont en pratique des semi-conducteurs, tels que des diodes ou des transistors (Lee (2003)).

1.4.2 Les types de mélangeurs de fréquence

Il existe différents cas d'usage pour les mélangeurs, ce qui exige différents fonctionnements et critères de la part de ces derniers. Tout d'abord, dans un récepteur à fréquence intermédiaire nulle, aussi appelé *homodyne*, on transpose un signal de la bande RF directement vers la bande de base. La fréquence LO vaut donc exactement la fréquence porteuse RF (de Dieuleveult & Romain (2008)).

Dans un récepteur à double changement de fréquence, un premier mélangeur transpose le signal RF vers une fréquence IF non nulle, avant qu'un second mélangeur ne transpose le signal vers la bande de base (de Dieuleveult & Romain (2008)).

Dans un mélangeur sous-harmonique, on mélange le signal IF non pas avec le signal LO, mais avec une harmonique de ce dernier. Par conséquent, on peut obtenir une fréquence RF plus élevée à la sortie. Par exemple, en utilisant la seconde harmonique du signal LO, on obtient une fréquence à la sortie de $f_{RF} = 2f_{LO} \pm f_{IF}$. Il est plus simple d'utiliser ce genre de mélangeur que de générer un signal d'oscillation LO fiable à très haute fréquence, notamment à plusieurs centaines de gigahertz, voire des térahertz. (Bao *et al.* (2006); Liu *et al.* (2022))

1.4.3 Les caractéristiques d'un mélangeur

Des figures de mérite propres aux mélangeurs existent, bien que certaines soient similaires, voire identiques, à celles des détecteurs.

1.4.3.1 Le gain de conversion

Dans un mélangeur, malgré l'usage de la non linéarité, la relation entre les signaux IF et RF est censée être linéaire. On s'intéresse donc au rapport des puissances à l'entrée et à la sortie. On parle de gain de conversion, souvent dans un système actif, lorsque la puissance à la sortie est supérieure à celle à l'entrée, tandis que l'on parle de perte de conversion dans les mélangeurs passifs, où la puissance à la sortie est inférieure à celle à l'entrée. La formule du gain de

conversion est la puissance à la sortie divisée par celle à l'entrée ; la formule de la perte de conversion en est l'inverse. Ces grandeurs sont exprimées en décibels. (Ludwig & Bretchko (2000) de Dieuleveult & Romain (2008))

Par exemple, pour un mélangeur d'un récepteur, le gain de conversion est le rapport de la puissance du signal IF sur celle du signal RF.

$$G_{cm} = \frac{P_{IF}}{P_{RF}} \quad (1.3)$$

En pratique, le gain de conversion dépend de la puissance du signal LO. Pour plusieurs topologies de mélangeurs, comme celle de Gilbert, il existe un niveau minimum de LO nécessaire pour assurer le fonctionnement nominal du mélangeur.

1.4.3.2 L'alimentation

Les mélangeurs peuvent être soit passifs, soit actifs. Les mélangeurs passifs n'ont besoin que de la puissance des signaux pour effectuer le mélange, tandis que les mélangeurs actifs nécessitent une alimentation externe pour fonctionner de manière optimale. Les mélangeurs passifs peuvent être réalisés également grâce à des diodes ou des transistors. En effet, il existe des mélangeurs à base de transistors qui fonctionnent sans courant d'alimentation (Bao *et al.* (2006)), et il existe des circuits à base de diodes qui nécessitent une polarisation. Les mélangeurs actifs peuvent avoir un meilleur gain de conversion qu'un mélangeur passif. Cependant, le prix de ce gain de conversion supérieur est la consommation énergétique. Autre que la méthode d'alimentation, dans les ouvrages, on mentionne parfois la consommation de puissance du mélangeur.

1.4.3.3 La bande passante

Souvent, dans les ouvrages, les auteurs mentionnent la plage de fréquence RF sur laquelle le circuit maintient un gain de conversion constant, pour une fréquence IF fixe. Parfois, on mentionne également la plage de fréquence IF sur laquelle le circuit maintient un gain constant pour une fréquence LO fixe.

Avec le gain (ou la perte) de conversion, la bande passante est une caractéristique immanquablement citée dans une publication.

1.4.3.4 La figure de bruit

Le facteur de bruit est le rapport du ratio signal-bruit (SNR) de l'entrée sur le ratio SNR de la sortie. La figure de bruit est le facteur de bruit exprimé en décibels. Ils témoignent de la dégradation entre l'entrée et la sortie d'un circuit de l'immunité du signal contre le bruit. Dans le cas d'un mélangeur, contrairement à un amplificateur, le bruit provenant de différentes bandes de fréquences peut être transposé vers la même bande de fréquence à la sortie (Fong & Meyer (1999)). Le filtrage en amont est donc important pour réduire ces sources de bruit additionnelles. La figure de bruit pour un mélangeur dans un récepteur est très importante, car ce dernier peut se situer juste en aval du LNA, second dans la chaîne de réception. Elle est donc très souvent citée également dans les publications.

1.4.3.5 Le point de compression

Comme sur un détecteur, la réponse d'un mélangeur cesse d'être linéaire au delà d'une certaine puissance de signal. Le point de compression de 1 dB référé à l'entrée ($IP_{1\text{ dB,m}}$) est le niveau de puissance du signal à l'entrée où le gain de conversion diminue 1 dB (Pozar (2012)). Tandis que la plage dynamique ne semble que rarement mentionnée pour les mélangeurs, le point de compression est souvent cité comme critère de performance. $IP_{1\text{ dB,m}}$ est exprimé en dBm.

1.4.3.6 Le point d'interception d'ordre trois

Pour caractériser la non linéarité d'un mélangeur, comme pour un amplificateur, nous pouvons injecter un signal à deux fréquences à l'entrée du mélangeur, puis mesurer la puissance des produits d'intermodulation à la sortie. En plus du signal à deux fréquences transposées, nous retrouvons des produits d'intermodulation à des fréquences adjacentes. La puissance des produits d'intermodulation les plus proches du signal désiré augmente de 30 dB pour chaque

décade d'augmentation de la puissance du signal à l'entrée. (cf. partie 2.6.3.1.) Sur un graphe logarithmique, on peut alors tracer une prolongation de la courbe de l'amplitude du signal de sortie désiré et de celle des produits d'intermodulation d'ordre trois. Les prolongations des courbes se croisent en un point. La puissance d'entrée correspondant à ce point d'intersection, bien que ce dernier ne soit qu'imaginaire, est un facteur de mérite pour la linéarité du mélangeur. La puissance du signal d'entrée correspondant à la position de ce point, exprimée en dBm, peut être notée $IP_{IP,m}$. (Pozar (2012); Fong & Meyer (1999))

1.4.3.7 L'isolation entre les ports

L'isolation caractérise la quantité de puissance de signal qui peut fuir d'un port à un autre. On peut définir l'isolation entre les ports RF et LO, entre les ports RF et IF et entre les ports LO et IF. Elle est mesurée en décibels, et plus sa valeur est importante, mieux le circuit est isolé.

L'isolation entre les ports est importante pour minimiser les produits d'intermodulation néfastes (de Dieuleveult & Romain (2008)). Dans un récepteur, une mauvaise isolation aux entrées du mélangeur, entre les ports RF et LO, peut faire en sorte que le signal LO remonte la chaîne de réception et qu'il soit rayonné par l'antenne (Ludwig & Bretchko (2000); Fong & Meyer (1999)). La fuite des signaux des entrées vers la sortie peut être assez facilement remédiée grâce à un filtre. Cependant, si les niveaux de fuite sont importants, alors les signaux fuités peuvent engendrer une compression imprévue (Fong & Meyer (1999)).

1.4.4 Les topologies de mélangeurs communes

Il existe également plusieurs topologies de mélangeurs qui sont plutôt communes et qui reviennent souvent dans la littérature. Ceci n'est pas une revue exhaustive des topologies de mélangeurs, mais nous mettrons en avant celles qui paraissent les plus pertinentes pour notre projet.

1.4.4.1 Les mélangeurs quadratiques à base de diodes

Un mélangeur peut être aussi simple qu'une seule diode. Les signaux à mélanger sont additionnés avant l'entrée. Puis, parmi les nombreux produits d'intermodulation créés par la non linéarité de la diode à la sortie, on peut retrouver le produit des signaux. Plus précisément, on dit que la diode suit une loi quadratique (*square-law*), comme dans un détecteur de puissance, puisque, sur une certaine plage dynamique, la relation courant-tension exponentielle de la diode peut être approximée par une relation quadratique. Un filtre passe-bande peut être utilisé pour ne garder que le signal désiré. La diode peut être polarisée à une certaine tension continue pour la faire fonctionner dans une région précise (Pozar (2012)). Si l'on souhaite éliminer des produits d'intermodulations et des harmoniques superflus, alors il existe des configurations de diodes, telles que la configuration antiparallèle utilisée par Liu, Niu, Wang & Yu (2024). Pour une haute fréquence de coupure, on se sert spécifiquement d'une diode Schottky. De plus, la basse tension de seuil des diodes Schottky permet de diminuer la puissance LO minimale nécessaire pour faire fonctionner le circuit (Hsiao, Meng, Wang & Huang (2014); Liu *et al.* (2024)). Les mélangeurs ayant pour base une diode Schottky peuvent fonctionner jusqu'à des centaines de gigahertz, voire quelques térahertz, pour des systèmes à très haut débit (Liu *et al.* (2022, 2024)). Ce genre de mélangeur est le plus simple et rudimentaire qu'il puisse exister. Généralement, sa performance ne satisfait pas les besoins des systèmes de communication modernes, selon Lee (2003). Par exemple, ce mélangeur n'a aucun gain, mais une perte de conversion. De plus, l'isolation entre les ports est mauvaise. Cependant, à de très hautes fréquences, là où d'autres topologies plus complexes échouent, il peut s'agir de la seule solution fonctionnelle.

1.4.4.2 Les mélangeurs à double équilibrage en pont de diodes

Les mélangeurs à double équilibrage en pont de diodes (*double-balanced diode ring mixers*), comme celui illustré sur la figure (1.6), sont constitués d'un pont de diodes relié à des transformateurs. On dit que cette topologie est à double équilibrage, car l'usage des quatre diodes et des deux *baluns* apportent deux degrés de symétrie. Si l'on utilise quatre diodes, c'est pour éliminer les modes indésirables, donc pour minimiser les distorsions du signal à la sortie.

Ils ont l'avantage d'offrir une bonne isolation entre les ports. Cette topologie a également une empreinte supérieure à celle vue précédemment, parce qu'il nécessite deux *baluns*.

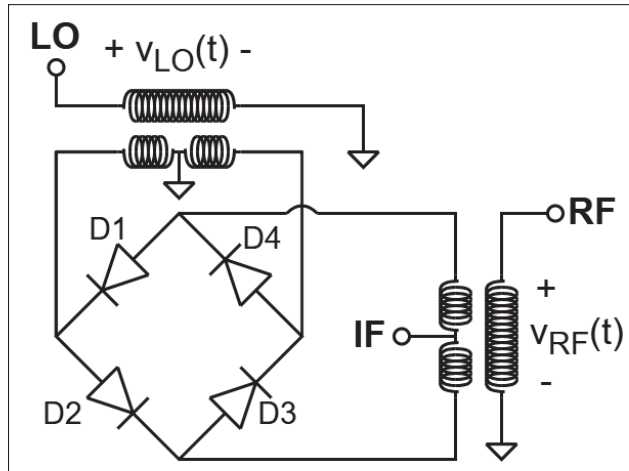


Figure 1.6 Schéma d'un mélangeur doublement symétrique de pont de diodes

Au port LO, un signal assez puissant pour activer des diodes est injecté. La nécessité d'un signal LO d'amplitude importante est d'ailleurs le plus grand défaut de ce type de mélangeur, selon de Dieuleveult & Romain (2008). Cependant, on peut utiliser des diodes Schottky pour une tension d'activation plus basse (Hsiao *et al.* (2014)). Supposons que la tension asymétrique LO injectée soit positive à un instant donné, alors les diodes D2 et D3 deviennent conductrices, tandis que D1 et D4 bloquent le courant. Pendant le cycle négatif, D1 et D4 deviennent conductrices, tandis que D2 et D3 bloquent le courant. Le circuit équivalent change en fonction du signe de la tension de LO. Par conséquent, on peut approximer la tension IF à la sortie comme le produit entre le signal RF et le signe de LO :

$$v_{IF}(t) = \text{signe}(v_{LO}(t))v_{RF}(t) \quad (1.4)$$

Ce type de mélangeur peut être réalisé autrement qu'avec des diodes et des transformateurs passifs. Dans l'ouvrage de Michaelson, Johansen, Tamborg & Zhurbenko (2015), le mélangeur a été réalisé grâce à un procédé BiCMOS. Au lieu des diodes, des transistors bipolaires reliés en

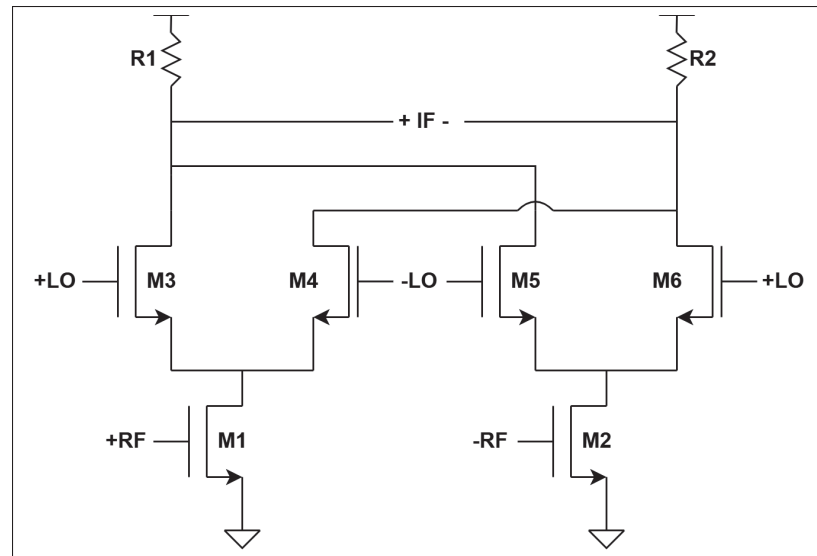


Figure 1.7 Schéma d'un exemple de cellule de Gilbert réalisé avec des MOSFET

mode diode ont été utilisés. De plus, les auteurs de cet ouvrage utilisent un *balun* actif pour le port RF et un *balun* passif pour le port LO.

1.4.4.3 La cellule de Gilbert

La cellule de Gilbert fait partie des topologies de mélangeur les plus reconnaissables et les plus répandues selon Shahriar Rashid & Rashid (2010). La cellule de Gilbert peut être décomposée en trois parties : la partie de transconductance, la partie commutative et les charges. Le signal RF est injecté en mode différentiel au niveau des grilles de transistors dans la partie de transconductance (M1 et M2). Les tensions sont converties en courant. Le signal RF module les courants tirés par les transistors. Reliés aux drains de ces derniers, se trouve la partie commutative (M3, M4, M5 et M6). Les transistors de cette partie reçoivent le signal LO qui les fait commuter. Ensuite, les charges (R1 et R2), reliées aux drains des transistors en commutation, transforment le courant en tension pour la sortie.

Cette topologie est également dite active doublement symétrique, car les entrées et la sortie sont symétriques. Ces symétries permettent d'annuler les signaux RF et LO à la sortie, selon

Ludwig & Bretchko (2000). Cette topologie est reconnue pour bénéficier d'un bon taux d'isolation entre les ports (Ludwig & Bretchko (2000); Bekkaoui (2017)). De plus, en tant que mélangeur actif, il peut avoir un gain de conversion supérieur à 0 dB. Un désavantage de cette topologie est la consommation de puissance, selon Shahriar Rashid & Rashid (2010). En effet, par rapport aux autres topologies vues précédemment, la cellule de Gilbert comporte le plus grand nombre de transistors actifs. Toujours selon eux, cette topologie est également plus susceptible au bruit par rapport aux mélangeurs quadratiques à un seul transistor à cause du plus grand nombre de composants.

1.4.4.4 Les mélangeurs quadratiques à base d'un transistors actif

Il est également possible de se servir de la non linéarité d'un transistor bipolaire ou à effet de champ pour accomplir le mélange de fréquences. Comme pour les mélangeurs à une diode, ce type de mélangeur est également considéré comme *square-law*. L'avantage d'un tel mélangeur, qui est actif, est la possibilité d'obtenir un gain de conversion supérieur à 0 dB, ce qui est impossible avec un mélangeur passif. Par exemple, Shahriar Rashid & Rashid (2010) obtiennent un gain de conversion allant jusqu'à 11,4 dB. Pour obtenir un tel gain de conversion, selon Pozar (2012), on peut polariser un transistor à effet de champ en dessous de sa tension de seuil V_{th} pour exploiter au mieux la non linéarité de sa transconductance.

Très souvent, on utilise une topologie de source commune (ou émetteur commun), avec les signaux à mélanger se trouvant au niveau de la grille (ou à la base). Pour cette raison, on appelle ce circuit *gate-pumped transistor mixer* en anglais. Dans le courant du drain (ou du collecteur), on peut retrouver le produit des signaux parmi d'autres produits d'intermodulation. Cependant, selon Ludwig & Bretchko (2000), le problème avec cette topologie très simple est la difficulté d'à la fois adapter les impédances des ports d'entrée reliés à la grille et d'assurer une bonne isolation entre ces ports. Shahriar Rashid & Rashid (2010) ont tenté de résoudre le manque d'isolation en utilisant des condensateurs et des bobines formant des filtres passe-bande reliés à la grille. Le filtre, relié à un des ports de l'entrée, isole ce dernier de l'autre port d'entrée aussi relié à la grille.

Pour améliorer l'isolation entre les ports d'entrée, il est possible d'injecter, d'une part, le signal à transposer par la grille (ou la base) et d'injecter, d'autre part, le signal LO par la source (ou l'émetteur) du transistor. Cette topologie est donc appelée *source-pumped mixer* en anglais. La tension grille-source (ou base-émetteur) est modulée par la différence entre les signaux. Pour ne pas court-circuiter le signal LO à la masse, cette topologie ajoute une charge ou une source de courant reliée à la source (ou l'émetteur) du transistor.

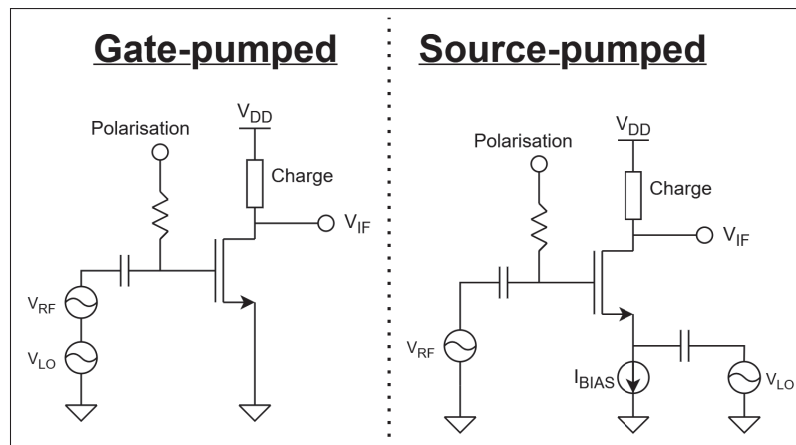


Figure 1.8 Mélangeurs quadratiques à base d'un MOSFET *gate-pumped* et *source-pumped*

Pour diminuer la consommation de puissance, les auteurs de Bao *et al.* (2006) ont dérivé la topologie de mélangeur *source-pumped* pour la rendre passive. S'ils ont choisi la topologie *source-pumped* au lieu de *gate-pumped*, c'est également pour améliorer le gain de conversion.

Nous pouvons constater que les topologies *gate-pumped* et *source-pumped*, qui n'utilisent qu'un seul transistor, ne bénéficient pas de la symétrie comme les mélangeurs à double équilibrage en pont de diodes ou les cellules de Gilbert. Par conséquent, sans filtrage, des produits d'interférences inintéressants, voire néfastes, sont présents à leurs sorties. Comme pour les mélangeurs quadratiques à base de diode, les mélangeurs quadratiques à base de transistors souffrent également d'une plage dynamique limitée. Au-delà d'un certain niveau de puissance, la caractéristique quadratique laisse place à d'autres types de non linéarité. Cependant, Bigdeli *et al.* (2025) ont tout-de-même réussi à étendre cette plage dynamique grâce à de la prédistorsion.

Tableau 1.2 Tableau de synthèse des mélangeurs de la revue de littérature

Ouvrage	Technologie	G_{cm} [dB]	IP_{1dB} [dBm]	IP3 [dBm]	Puissance consommée [mW]	Fréquence RF
Hsiao <i>et al.</i> (2014)	CMOS 180 nm	6	-10	-2	32,8	60 GHz
Bekkaoui (2017)	CMOS 65 nm	10	-2,54	6	2,17	1,9 GHz
Liu <i>et al.</i> (2022)	MMIC	-7,5	?	?	3*	292 à 356 GHz
Liu <i>et al.</i> (2024)	MMIC	-9,4	2,6	?	19,95*	311 à 349 GHz
Tsai <i>et al.</i> (2007a)	CMOS 90 nm	3	-4,5 à 2,1	?	93	25 à 75 GHz
Michaelsen <i>et al.</i> (2015)	SiGe BiCMOS	4	-11	13	31,62*	10,5 GHz
Shahriar Rashid & Rashid (2010)	CMOS 90 nm	11,4	-10,8	0	9,25	17,3 à 21,8 GHz
Nouri, Nezhad-Ahamdi, Mirabbasi & Safavi-Naeini (2010)	CMOS 130 nm	1,7	?	14,2	6 + 1,4*	60 GHz
Bao <i>et al.</i> (2006)	CMOS 90 nm	-11 à -8	?	3 à 7	9,33 à 18,2*	9 à 31 GHz

* : puissance du signal LO

Toutefois, l'avantage de ces mélangeurs demeure leur simplicité et le fait qu'ils ne nécessitent pas un signal LO de haute puissance pour fonctionner, contrairement aux topologies qui fonctionnent sur la base de la commutation. Selon Bigdeli *et al.* (2025), l'usage de la *square-law* permet d'économiser massivement de l'énergie sur un système massif, tel qu'une antenne réseau, où l'on peut trouver plusieurs mélangeurs par cellule d'antenne. Par exemple, sur une antenne réseau embarquée à bord d'un satellite en orbite basse, avec un budget énergétique minimal, l'usage de mélangeurs quadratiques paraît plus approprié que celui des cellules de Gilbert.

1.4.5 Les précédents travaux de mélangeurs utilisant la non-linéarité d'un transistor

Des caractéristiques de différents mélangeurs mentionnés dans cette revue de littérature se trouvent regroupées dans le tableau 1.2. La comparaison des mélangeurs s'avère plus difficile que celle des détecteurs, car ils n'opèrent pas sur les mêmes fréquences et ne les transposent pas tous dans le même sens. De plus, les mélangeurs existent sous formes passive et active, et ils utilisent différents principes. Bref, il existe une plus grande variété de mélangeurs que de détecteurs de puissance quadratiques. Ce qui complexifie la tâche est encore le manque de certaines données dans des ouvrages.

1.5 Un circuit détecteur-mélangeur

Les mélangeurs quadratiques et les détecteurs quadratiques ont souvent des topologies similaires. Toutefois, ils ont des applications différentes. Après tout, un mélangeur est un circuit qui multiplie deux signaux distincts, tandis qu'un détecteur quadratique multiplie en quelque sorte un signal par lui-même. C'est pour cela que certains ouvrages, tels que Frater (1964) et Bigdeli *et al.* (2025), confondent les termes « mélangeurs » (*multiplier*) et « détecteur quadratique » (*square-law detector*). Frater (1964) parle de réaliser la multiplication de signaux grâce à la mise au carré de leur somme et de leur différence. Ils réalisent un détecteur quadratique avec un multiplicateur dont les entrées sont reliées.

1.6 Conclusion de chapitre

Dans ce chapitre, nous avons vu le besoin croissant d'antennes et de surfaces reconfigurables moins coûteuses que les antennes réseau à commande de phase. Si ces nouvelles antennes sont en demande, c'est qu'elles permettent le développement des réseaux 6G nécessitant notamment des plates-formes satellitaires en masse. Les RRA, RTA et RIS peuvent utiliser des déphaseurs à base de diodes à capacité variable. Ces dernières ont l'avantage d'être fiables et de pouvoir contrôler la phase de façon continue. Or, elles admettent le défaut d'engendrer une distorsion du signal. Par conséquent, nous envisageons l'usage d'un détecteur de puissance et d'un mélangeur de fréquence pour appliquer une technique de prédistorsion sur la polarisation des diodes, tel que représenté dans la figure 1.9. Le détecteur peut mesurer le niveau de puissance de l'échantillon du signal déphasé, car la distorsion engendrée par une diode est directement liée à la puissance du signal. Le mélangeur permet également de mesurer la distorsion engendrée par le déphaseur en transposant l'échantillon du signal vers le bas. Ce signal transposé peut être ensuite analysé par un circuit de traitement du signal numérique.

De nombreuses variantes de détecteurs et de mélangeurs existent dans la littérature. Parmi elles, se trouvent les détecteurs d'enveloppe de puissance et les mélangeurs utilisant le principe de la *square-law*. Ces circuits s'avèrent relativement simples, peuvent être facilement intégrés dans

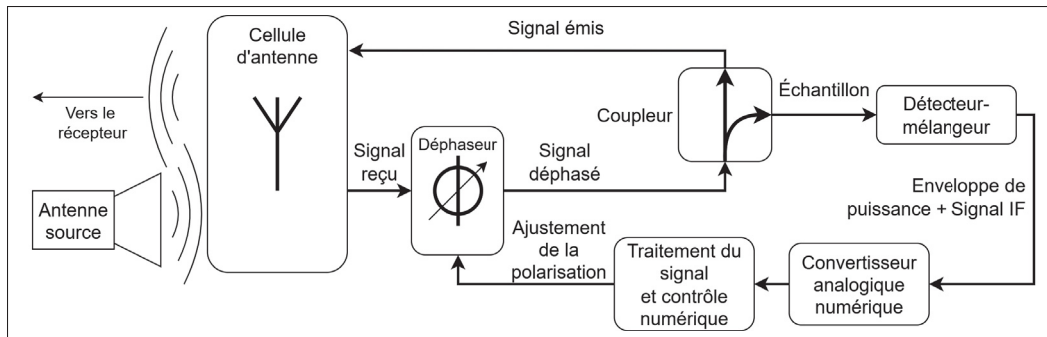


Figure 1.9 Exemple d'usage du détecteur-mélangeur pour améliorer la performance d'une RRA

une puce en CMOS et consomment peu d'énergie. Nous avons remarqué qu'ils partagent assez de similarités pour que nous concevions un circuit de détecteur-mélangeur qui accomplie les deux tâches à la fois. A notre connaissance, aucun article ne présente un seul circuit qui est à la fois optimisé pour fonctionner en tant que détecteur et en tant que mélangeur pour améliorer la performance d'une antenne reconfigurable.

CHAPITRE 2

ANALYSE ET THÉORIE SUR LES CIRCUITS QUADRATIQUES

Dans ce chapitre, nous étudierons la théorie du fonctionnement d'un détecteur ou d'un mélangeur quadratique, parce que les détails de leurs mécanismes ne sont pas extensivement couverts dans la littérature. Pour la conception du circuit, nous devons comprendre les mécanismes derrière le mélange de fréquences et la détection de l'enveloppe de puissance, ainsi que retrouver les origines des produits d'intermodulation néfastes.

2.1 La non-linéarité d'un transistor

2.1.1 Le modèle «simple» d'un MOSFET en saturation

Soit un MOSFET en mode saturation. Si l'on n'ignore pas l'effet de modulation de la longueur du canal, la relation entre sa tension de grille et son courant de drain est la suivante :

$$I_{ds} = \frac{1}{2} \mu C_{ox} \frac{W}{L} (V_{gs} - V_{th})^2 (1 + \lambda V_{ds}) \quad (2.1)$$

(tiré de Razavi (2013) p. 289.) I_{ds} est le courant allant du drain à la source, μ la mobilité des porteurs de charge, C_{ox} la capacité de grille, W la largeur de la grille, L la longueur de la grille, V_{gs} la différence de potentiel entre la grille et la source, V_{th} la tension de seuil pour la saturation, λ une grandeur caractérisant l'effet de modulation de longueur de canal, et V_{ds} la tension drain-source. Cette relation cesse d'être valide en dehors de la région de saturation et pour les transistors dits à « canal court » (Gray, Hurst, Lewis & Meyer (2009)). Dans cet ouvrage, nous nous contenterons (dans un premier temps) de l'expression ci-dessus. Nous n'utiliserons pas de modèle plus sophistiqué, tel que celui proposé par Tsividis, Suyama & Vavelidis (1995) ou le *Berkeley Short-channel IGFET Model* (BSIM Group (2025)).

Posons pour la suite i_{ds} et v_{gs} , qui sont respectivement une petite variation du courant de drain et une petite variation de la tension grille-source. Ces grandeurs sont dites « petits signaux », car elles n'affectent pas la polarisation. (Pour la suite, de façon conventionnelle, toutes les grandeurs

petits signaux seront notées en minuscule, contrairement aux tensions et aux courants continus en majuscule.) Examinons la relation entre les deux grandeurs. Si l'on stimule un transistor à la grille avec v_{gs} , alors, en plus du courant continu I_{ds} , qui est engendré par la tension continue V_{gs} , nous obtenons i_{ds} :

$$I_{ds} + i_{ds} = \frac{1}{2} C_{ox} \mu \frac{W}{L} [(V_{gs} + v_{gs}) - V_{th}]^2 (1 + \lambda V_{ds}) \quad (2.2)$$

$$I_{ds} + i_{ds} = \frac{1}{2} C_{ox} \mu \frac{W}{L} [V_{gs}^2 + V_{th}^2 - 2V_{gs}V_{th} + 2v_{gs}(V_{gs} - V_{th}) + v_{gs}^2] (1 + \lambda V_{ds}) \quad (2.3)$$

i_{ds} dépend donc à la fois de v_{gs} et v_{gs}^2 . Le terme $2v_{gs}(V_{gs} - V_{th})$ devient prépondérant devant le terme v_{gs}^2 lorsque la tension de polarisation V_{gs} augmente et, au contraire, v_{gs}^2 est prépondérant lorsque $V_{gs} \approx V_{th}$. Donc, rien qu'avec ce modèle simple, nous pouvons démontrer qu'un MOSFET a une fonction de transfert quadratique lorsque la tension de polarisation est proche de la tension de seuil V_{th} .

Comme pour beaucoup de relations non linéaires, nous pouvons modéliser la relation courant-tension des petits signaux grâce à une série de Taylor. Pour une polarisation donnée, si nous supposons qu'il n'y a pas de variation de la tension de drain ($v_{ds} = 0$), alors nous avons :

$$i_{ds}|_{v_{ds}=0} = \sum_{i=1}^{\infty} \frac{1}{i!} \frac{\partial^i I_{ds}}{\partial V_{gs}^i} v_{gs}^i \quad (2.4)$$

Puis, en s'inspirant de la notation de Qayyum & Negra (2017) et Ou & Farahmand (2012), nous pouvons poser :

$$\forall i \in \mathbb{N}^*, g_{mi} = \frac{1}{i!} \frac{\partial^i I_{ds}}{\partial V_{gs}^i} \quad (2.5)$$

avec g_{mi} ayant comme unité $A.V^{-i}$. Cela donne :

$$i_{ds}|_{v_{ds}=0} = g_{m1}v_{gs} + g_{m2}v_{gs}^2 + g_{m3}v_{gs}^3 + \dots \quad (2.6)$$

Chaque terme de transconductance g_{mi} dépend de la polarisation, des caractéristiques du transistor, ainsi que des facteurs tels que la température. De plus, dans cette modélisation qui

est dite « statique », nous supposons qu'il n'y a pas d'effet mémoire. Pour nous, la valeur du courant à un instant t ne dépend pas de la valeur du courant à un moment antérieur. Nous ne nous mêlons pas des séries de Volterra (Narayanan (1967)). Enfin, dans ce chapitre, nous ne nous préoccupons pas de la modulation de la longueur de canal non linéaire. Si nous devons prendre en compte ces effets, i_{ds} dépendrait des tensions drain-source et substrat-source. Nous aurions eu une série de Taylor multidimensionnelle, comme dans Ou & Farahmand (2012). Nous nous contenterons d'une série de Taylor unidimensionnelle.

Pour les petits signaux, le plus souvent, en se basant toujours sur la relation (2.1), on pose la transconductance g_{m1} (plus classiquement notée g_m) comme la dérivée de I_{ds} par rapport à V_{gs}

$$g_{m1} = \frac{\partial I_{ds}}{\partial V_{gs}} = \mu C_{ox} \frac{W}{L} (V_{gs} - V_{th})(1 + \lambda V_{ds}) \quad (2.7)$$

Grâce à cette grandeur, on peut approximer qu'une petite variation de courant de drain i_{ds} est proportionnelle à une petite variation de tension de grille v_{gs} .

$$i_{ds}|_{v_{ds}=0} \approx g_{m1} v_{gs} \quad (2.8)$$

Cependant, cette approximation n'est valide que si (2.1) est valide, si les variations de courant et de tension sont petites et si $2v_{gs}(V_{gs} - V_{th}) \ll v_{gs}^2$. Surtout, une relation linéaire n'est pas intéressante pour notre application. Si nous ne nous basions que sur l'équation (2.3), nous pourrions trouver la valeur de g_{m2} :

$$g_{m2} = \mu C_{ox} \frac{W}{L} (1 + \lambda V_{ds}) \quad (2.9)$$

Selon cette expression, la transconductance quadratique serait indépendante de V_{gs} . Donc, selon (2.1), la série de Taylor ne comporterait que deux termes : g_{m1} et g_{m2} .

Cela s'avère faux, et le modèle ci-dessus est trop simple pour décrire la réalité. En effet, plus de termes sont nécessaires dans le développement en série pour décrire les phénomènes de non-linéarité de la réalité. Sur la figure 2.1, nous avons tracé le courant de drain I_{ds} simulé

d'un transistor NMOS de 75 nm de longueur et de 1,56 μm de largeur en fonction de la tension V_{gs} . Nous avons également tracé des approximations linéaires et quadratiques autour de $I_{ds}(V_{gs} = 400mV)$. L'approximation du premier ordre ne fonctionne que sur une plage très limitée. En effet, près de la tension V_{th} , le terme $2v_{gs}(V_{gs} - V_{th})$ de la relation (2.3) n'est pas prépondérant devant v_{gs}^2 , donc le transistor n'a pas une réponse très linéaire. L'approximation du second ordre imite mieux la courbe réelle sur une plus grande plage de valeurs que l'approximation du premier ordre. Nous pouvons visuellement observer que, près de la tension de seuil, le MOSFET a une réponse plutôt quadratique. C'est d'ailleurs la raison pour laquelle autant de circuits de détecteurs et de mélangeurs quadratiques à base de MOSFET opèrent près de la tension de seuil. Cependant, même l'approximation du second ordre laisse place à l'erreur. Nous pouvons donc déduire que la série de Taylor (2.4), découlant de (2.1), a plus de deux termes. En effet, la relation courant-tension d'un MOSFET n'est pas purement quadratique, d'où l'existence des modèles informatiques plus précis pour les simulations.

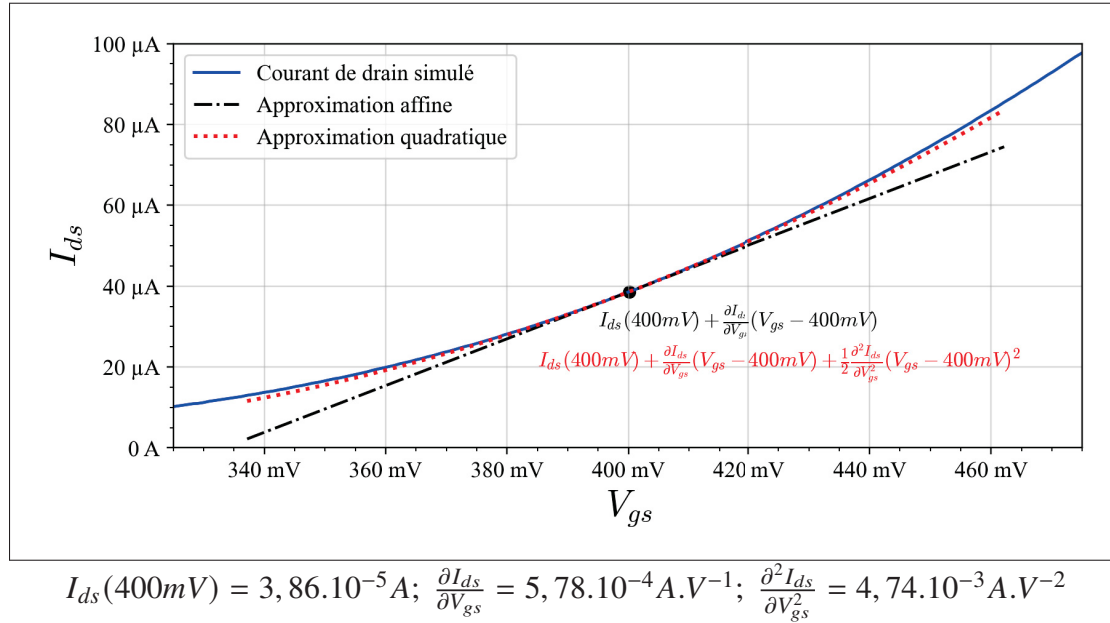


Figure 2.1 La relation courant-tension d'un transistor, ainsi que des approximations locales

Dans cette partie, nous avons pu démontrer qu'un MOSFET a un comportement approximativement quadratique, même en modèle des petits signaux, surtout près de la tension de seuil. Cependant,

sa réponse n'est pas purement quadratique non plus. Des termes d'ordres supérieurs sont également nécessaires pour décrire la non linéarité d'un MOSFET. Dans le but de concevoir un détecteur et un mélangeur, nous devons faire en sorte que le comportement du transistor soit le plus quadratique possible.

2.1.2 Les autres phénomènes non linéaire

Investiguons brièvement d'autres phénomènes non linéaires au sein d'un MOSFET.

2.1.2.1 La modulation de longueur de canal

Si l'on ne néglige pas l'effet de modulation de longueur de canal, alors, selon 2.1, nous pouvons poser

$$g_{ds} = \frac{\partial I_{ds}}{\partial V_{ds}} = \frac{1}{2} \mu C_{ox} \frac{W}{L} (V_{gs} - V_{th})^2 \lambda \quad (2.10)$$

Par conséquent, l'approximation de premier ordre bidimensionnelle du courant dans le modèle des petits signaux est :

$$i_{ds} \approx g_{m1} v_{gs} + g_{ds} v_{ds} \quad (2.11)$$

La relation entre i_{ds} et v_{ds} semble linéaire. Et puisque, selon le modèle de (2.1), λ est une constante, on pose plus communément r_o telle que :

$$r_o = \frac{1}{g_{ds}} = \frac{2}{\lambda \mu C_{ox} (W/L) (V_{gs} - V_{th})^2} \quad (2.12)$$

Or, en réalité, la relation entre i_{ds} et v_{ds} n'est pas purement linéaire : les simulations le confirment. Il serait donc plus réaliste de modéliser l'effet de modulation de longueur de canal par une série de Taylor :

$$i_{ds}|_{v_{gs}=0} = \sum_{i=1}^{\infty} g_{dsi} v_{ds}^i \quad (2.13)$$

Rien qu'en prenant en compte la non-linéarité de l'effet de modulation de la longueur du canal, nous pouvons complexifier le modèle des petits signaux du transistor grâce à une série de Taylor

bidimensionnelle, comme dans l'ouvrage de Ou & Farahmand (2012) :

$$i_{ds} = \sum_{i=1}^{\infty} g_{mi} v_{gs}^i + \sum_{j=1}^{\infty} g_{dsj} v_{ds}^j + \underbrace{\sum_{i=1}^{\infty} \sum_{j=1}^{\infty} \frac{1}{i!j!} \frac{\partial^{i+j} I_{ds}}{\partial V_{gs}^i \partial V_{ds}^j}}_{\text{Termes mixtes}} v_{gs}^i v_{ds}^j \quad (2.14)$$

Même à seulement deux dimensions, l'apparition des termes mixtes complexifie beaucoup l'analyse de circuits. Pour la suite, pour une question de simplicité, nous ne prendrons pas en compte ces termes. De plus, nous supposons que l'effet de modulation de longueur de canal est linéaire ($i_{ds}|_{v_{gs}=0} \approx v_{ds}/r_o$). En effet, nous nous intéresserons surtout aux termes $g_{mi} v_{gs}^i$.

2.1.2.2 L'effet de substrat

Si le substrat et la source d'un transistor ne se trouvent pas au même potentiel, alors la tension de seuil V_{th} varie. En effet, V_{th} augmente (de façon non linéaire) avec la tension entre la source et le substrat V_{sb} (Razavi (2013)). Donc, la variation v_{sb} a un impact sur le courant i_{ds} . Cependant, une fois de plus, nous négligerons pour la suite cet effet.

2.2 Le fonctionnement d'un détecteur de puissance quadratique

Étudions de façon analytique le fonctionnement d'un détecteur quadratique. Nous pouvons y procéder en élevant au carré des signaux simples.

2.2.1 La détection d'un signal à amplitude constante

Rappelons-nous que la puissance instantanée d'un signal électrique est proportionnelle au carré de sa tension (ou de son courant). Notre type de détecteur indique la puissance grâce à la loi quadratique des transistors. Soit $v_{RF}(t) = V_{RF} \cos(2\pi ft + \phi)$ la tension d'un signal sinusoïdal à l'entrée de notre détecteur, dont nous voulons mesurer la puissance. V_{RF} est son amplitude, ϕ

est sa phase et f est sa fréquence. En mettant cette tension au carré, nous obtenons

$$v_{RF}^2(t) = V_{RF}^2 \left[\frac{1 + \cos(4\pi f t + 2\phi)}{2} \right] \quad (2.15)$$

Comme également représenté dans la figure 2.2, le signal à la sortie est composé de deux termes : l'un à 0 Hz et l'autre à deux fois la fréquence à l'entrée. L'amplitude des deux est proportionnelle au carré de l'amplitude à l'entrée, soit à la puissance à l'entrée. Cependant, on filtre généralement le terme à haute fréquence.

À la sortie d'un détecteur, cette tension continue, proportionnelle au carré de l'amplitude du signal, s'ajoute (ou se soustrait) à V_r , la tension continue à la sortie au repos (*i.e.* la tension à la sortie sans signal). Par conséquent, si nous posons V_s la tension continue à la sortie du détecteur, alors $|V_s - V_r|$ est proportionnel à V_{RF}^2 .

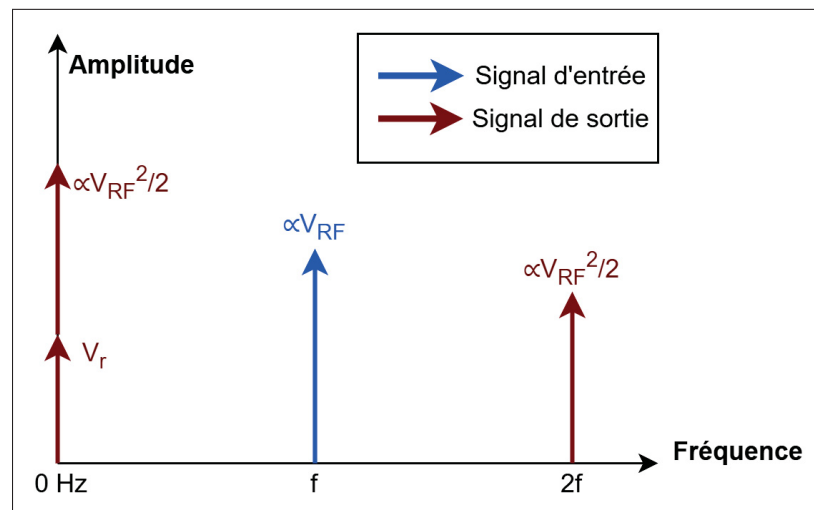


Figure 2.2 Spectres de l'entrée et de la sortie d'un détecteur quadratique idéal qui reçoit un signal sinusoïdal

Mesurer le changement de la tension continue à la sortie du détecteur de puissance en fonction de la puissance d'un signal sinusoïdal est une façon simple et fiable de tester et de caractériser le circuit. Nous avons vu, lors de la revue de littérature, que cela permet notamment de mesurer la responsivité.

2.2.2 La détection d'un signal modulé en amplitude

Nous ne nous intéressons pas qu'aux signaux dont l'amplitude est constante. Et nous n'avons pas l'intention de concevoir un détecteur de puissance moyenne. Il faut donc étudier la réponse d'un détecteur d'enveloppe de puissance aux signaux avec une modulation d'amplitude. Posons $v_{RF}(t) = V_{RF} \cos(2\pi f_c t) \cos(2\pi f_m t)$. f_c est la fréquence porteuse et f_m est la fréquence du modulant. Le signal n'est composé que de deux fréquences ; c'est le signal modulé en amplitude le plus simple que l'on puisse imaginer. Posons également $f_1 = f_c + f_m$ et $f_2 = f_c - f_m$. Donc $v_{RF}(t) = \frac{V_{RF}}{2} (\cos(2\pi f_1 t) + \cos(2\pi f_2 t))$. Supposons également que $f_c \gg f_m$.

Mettons au carré cette tension :

$$v_{RF}^2 = \frac{V_{RF}^2}{4} \left(1 + \underbrace{\cos((\omega_1 - \omega_2)t)}_{\cos(2\omega_m t)} + \underbrace{\cos((\omega_1 + \omega_2)t)}_{\cos(2\omega_c t)} + \frac{\cos(2\omega_1 t) + \cos(2\omega_2 t)}{2} \right) \quad (2.16)$$

(Voir la partie 1.2.1 de l'annexe I pour le développement entier.) Le carré du signal contient plusieurs termes à différentes fréquences, comme nous l'avons représenté sur la figure 2.3. Celui à $2f_m$ correspond à l'enveloppe de la puissance du signal modulant. Celui à $2f_c$ donne l'information sur la puissance de la porteuse : elle ne nous intéresse pas. Les termes à $2f_1$ et $2f_2$ ne nous sont pas utiles non plus. Étant donné que $f_c \gg f_m$, nous pouvons facilement filtrer les termes aux fréquences $2f_c$, $2f_1$ et $2f_2$. Nous retrouvons également le terme à 0 Hz. Ce dernier donne toujours la valeur moyenne de la puissance.

2.3 Un mélangeur de fréquences grâce à la caractéristique quadratique

Maintenant que nous avons compris le principe de fonctionnement d'un détecteur de puissance quadratique, nous pouvons établir le lien entre ce dernier et le mélangeur quadratique.

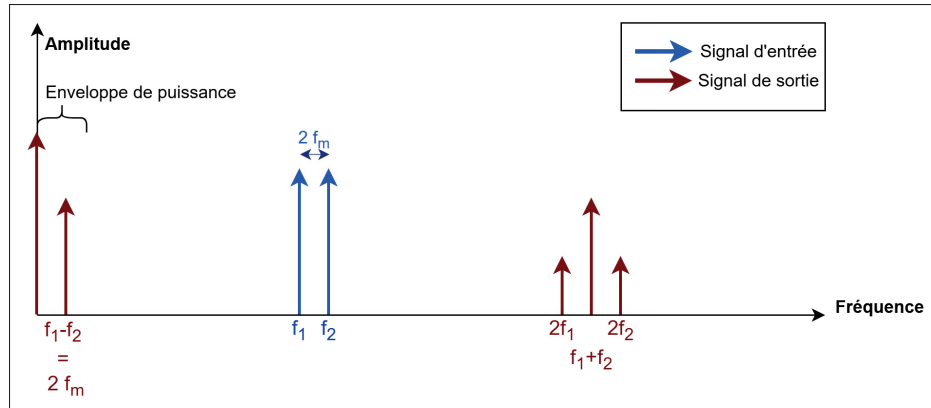


Figure 2.3 Spectre de l'entrée et de la sortie d'un détecteur quadratique idéal qui reçoit un signal à deux fréquences

2.3.1 Rappel du fonctionnement d'un mélangeur

Rappelons-nous dans cette partie comment la multiplication de signaux permet la transposition de fréquence. Soit $v_{RF}(t)$ toujours le signal RF avec la même expression que dans la partie précédente. Soit $v_{LO}(t) = V_{LO} \cos(\omega_{LO}t)$ une tension d'oscillation. f_{LO} détermine à quelle fréquence on transpose le signal RF. Supposons que $f_{LO} < f_c$. Posons f_{IF} telle que $f_c = f_{LO} + f_{IF}$.

$$v_{RF}v_{LO} = \frac{V_{RF}V_{LO}}{4} [\cos((2\omega_c + \omega_m + \omega_{IF})t) + \cos((\omega_{IF} - \omega_m)t) + \cos((2\omega_c - \omega_m + \omega_{IF})t) + \cos((\omega_{IF} + \omega_m)t)] \quad (2.17)$$

La multiplication transpose le signal RF à la fois vers le haut et vers le bas. Nous nous intéressons ici au mélangeur pour la réception de signaux, donc à transposer vers le bas. Appliquons donc un filtrage passe-bas au signal ci-dessus :

$$v_{RF}v_{LO} = \frac{V_{RF}V_{LO}}{4} [\cos((\omega_{IF} - \omega_m)t) + \cos((\omega_{IF} + \omega_m)t) + \dots] \quad (2.18)$$

Cette multiplication nous a permis de réduire la fréquence qui porte le signal RF de f_c à f_{IF} . Si nous souhaitons ramener le signal RF en bande de base, il faudrait que $f_{LO} = f_c$, ce qui équivaut à $f_{IF} = 0$.

Nous remarquons un facteur 1/2 dans l'équation ci-dessus. Il apparaît à cause de la multiplication, qui transpose une moitié du signal vers les hautes fréquences et une autre partie vers les basses fréquences. La puissance du signal RF est en quelque sorte divisée par deux à la sortie IF.

2.3.2 La mise au carré pour la multiplication

Mais si la multiplication de signaux permet effectivement de transposer des signaux, comment peut-on réaliser une telle opération grâce à un circuit quadratique ? On peut multiplier deux signaux en les additionnant avant d'injecter la somme dans notre circuit quadratique. C'est le principe de l'identité mathématique $(a + b)^2 = a^2 + 2ab + b^2$, où a et b sont des réels. On s'intéresse ici à n'extraire que le terme $2ab$.

Si l'on met au carré la somme des signaux RF et LO :

$$\begin{aligned}
 (v_{RF}(t) + v_{LO}(t))^2 &= \frac{V_{RF}^2}{4} \underbrace{[1 + \cos(2\omega_m t)]}_{\text{Enveloppe}} \\
 &+ \frac{V_{LO}^2}{2} + \underbrace{\frac{V_{RF}V_{LO}}{2} [\cos((\omega_{IF} + \omega_m)t) + \cos((\omega_{IF} - \omega_m)t)]}_{\text{Signal transposé}} + \dots
 \end{aligned} \tag{2.19}$$

(Voir la partie 1.2.2 de l'annexe I pour le développement.) Nous obtenons le carré de RF, le carré de LO et leur produit. Nous obtenons des termes à 0 Hz, un terme à $2f_m$ (qui correspond à l'enveloppe du signal modulant), le signal porté à la fréquence f_{IF} et encore d'autres termes à haute fréquence que nous filtrons. Ces produits d'intermodulation sont représentés sur la figure 2.4. On peut appliquer un filtrage passe-bande autour de f_{IF} , si l'on souhaite éliminer le signal en bande de base.

Nous constatons que nous obtenons toujours l'enveloppe de puissance du signal RF. En effet, le circuit ne cesse pas de fonctionner comme un détecteur de puissance simplement parce que l'on ajoute un signal LO. Le circuit fonctionne à la fois comme détecteur et mélangeur en

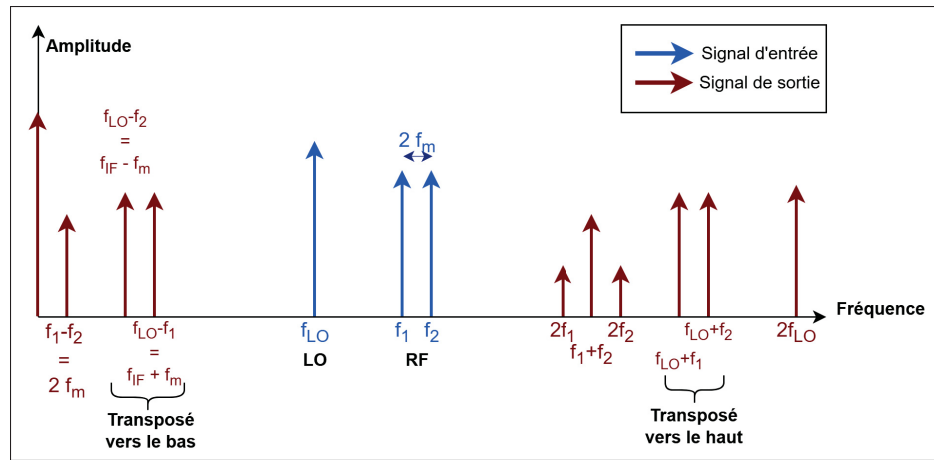


Figure 2.4 Spectre de l'entrée et de la sortie d'un mélangeur de fréquences quadratique

permanence. Ceci présente une contrainte incontournable sur la fréquence IF que nous pouvons choisir, comme nous le verrons ultérieurement. Il est tout de même essentiel de souligner que l'amplitude du signal transposé reste proportionnelle à V_{RF} , tandis que l'enveloppe de puissance est proportionnelle à V_{RF}^2 . En d'autres termes, le mécanisme de mélange est quadratique, mais la réponse du mélangeur reste linéaire.

Dans les sections ci-dessus, nous avons démontré en quoi la mise au carré d'un signal est théoriquement utile pour concevoir un détecteur de puissance ou un mélangeur de fréquence. Comme mentionné dans la partie 2.1, en pratique, nous utiliserons le terme quadratique g_{m2} de la transconductance pour accomplir la détection et le mélange.

2.3.3 Les contraintes pour un mélangeur quadratique

Nous venons de voir que la mise au carré du signal produit différents termes à différentes fréquences. Cela pose des contraintes sur l'opération des mélangeurs quadratiques.

2.3.3.1 Le chevauchement du signal IF et de l'enveloppe de puissance à la sortie

Tout d'abord, à cause de la présence inévitable de l'enveloppe du signal modulant en sortie, ce type de mélangeur de fréquence ne peut pas être utilisé dans une structure de récepteur *homodyne*. Si $f_{LO} = f_c$, alors le produit $v_{RF}v_{LO}$ se trouverait en bande de base. Or, l'enveloppe du signal, qui provient du terme v_{RF}^2 , y serait aussi. Il y aurait donc une interférence. Par conséquent, ce mélangeur ne peut être utilisé que dans une structure hétérodyne.

La coexistence de l'enveloppe du signal et du signal IF impose une contrainte sur f_{IF} et sur la largeur de bande du signal RF pour éviter tout chevauchement. Étudions les conditions qui permettent d'éviter un chevauchement. Procédons toujours avec un signal RF à deux fréquences espacées de $2f_m$.

Si $f_{LO} < f_c$, alors :

$$\begin{aligned} [2f_m < f_{IF} - f_m] &\iff [2f_m < (f_c - f_{LO}) - f_m] \\ &\iff [f_{LO} < f_c - 3f_m] \end{aligned} \quad (2.20)$$

Si $f_{LO} > f_c$, alors la condition devient $f_{LO} > f_c + 3f_m$. Dans les deux cas, la fréquence du signal LO doit donc maintenir un certain écart par rapport à celle du signal RF.

2.3.3.2 Le chevauchement du signal IF et RF ou LO rémanent

À la sortie du mélangeur, peut demeurer un signal RF ou LO à faible niveau à cause d'une isolation imparfaite. Pour éviter le chevauchement entre le signal IF en sortie et les signaux d'entrée rémanents, il faut satisfaire les conditions suivantes :

1 - Si $f_{LO} < f_c$, on cherche à éviter le chevauchement du signal IF et du signal LO rémanent.

$$\begin{aligned} [f_{IF} + f_m < f_{LO}] &\iff [f_c - f_{LO} + f_m < f_{LO}] \\ &\iff [(f_c - f_m)/2 < f_{LO}] \end{aligned} \quad (2.21)$$

2 - Si $f_{LO} > f_c$, on cherche à éviter le chevauchement du signal IF et du signal RF rémanent.

$$\begin{aligned} [f_{IF} + f_m < f_c - f_m] &\iff [f_{LO} - f_c + f_m < f_c - f_m] \\ &\iff [f_{LO} < 2f_c - 2f_m] \end{aligned} \quad (2.22)$$

2.4 La fonction de transfert non linéaire du circuit

La fonction de transfert d'un détecteur ou d'un mélangeur est non linéaire. Alors, selon une modélisation statique, où l'on ignore les effets de mémoire, la relation entre la tension du signal à l'entrée $v_e(t)$ et celle à la sortie $v_s(t)$ peut être exprimée ainsi :

$$V_s + v_s(t) = V_r + \sum_{i=1}^{\infty} \alpha_i v_e^i(t) = V_r + \alpha_1 v_e(t) + \alpha_2 v_e^2(t) + \alpha_3 v_e^3(t) + \alpha_4 v_e^4(t) + \dots \quad (2.23)$$

(Tiré de (10.37) de Pozar (2012) p.512.) Rappelons que V_s est la tension continue à la sortie et que V_r est la tension continue à la sortie si le circuit est au repos. Cette dernière ne dépend que de la polarisation. Si $v_e = 0$, alors $V_s = V_r$. Puisque nous concevons un circuit quadratique, nous pouvons supposer, dans un premier temps, que cette fonction se simplifie ainsi :

$$V_s + v_s(t) \approx V_r + \alpha_2 v_e^2(t) \quad (2.24)$$

Nous pourrions appeler α_2 le gain quadratique de tension par la suite.

2.4.1 La responsivité du détecteur

Si le détecteur reçoit un signal $v_{RF} = V_{RF} \cos(\omega t)$ à son entrée, alors la tension à sa sortie devient :

$$V_s + v_s(t) \approx V_r + \alpha_2 V_{RF}^2 \frac{(1 + \cos(2\omega t))}{2} \quad (2.25)$$

En filtrant le terme à haute fréquence, nous obtenons :

$$V_s - V_r \approx \frac{\alpha_2 V_{RF}^2}{2} \quad (2.26)$$

Nous pouvons désormais lier la responsivité défini précédemment et le coefficient quadratique de la fonction de transfert non linéaire α_2 .

$$R_d = \frac{|V_s - V_r|}{P_{RF}} = \frac{\alpha_2 V_{RF}^2 / 2}{V_{RF}^2 / (2R)} = \alpha_2 R \quad (2.27)$$

Généralement, dans un système RF, $R = 50 \, \Omega$.

2.4.2 Le gain de conversion d'un mélangeur

Soient $v_{RF} = V_{RF} \cos(\omega_c t)$ et $v_{LO} = V_{LO} \cos(\omega_{LO} t)$ à l'entrée d'un mélangeur quadratique qui transpose vers le bas. La tension du signal IF à la sortie est :

$$v_{IF}(t) = \alpha_2 (v_{RF} + v_{LO})^2|_{f=f_{IF}} = \alpha_2 V_{RF} V_{LO} \cos((\omega_c - \omega_{LO})t) \quad (2.28)$$

Le gain de conversion du mélangeur, qui était d'abord défini par (1.3), est alors :

$$G_{cm} = \frac{P_{IF}}{P_{RF}} = \frac{(\alpha_2 V_{RF} V_{LO})^2 / (2R)}{V_{RF}^2 / (2R)} = \alpha_2^2 V_{LO}^2 \quad (2.29)$$

Idéalement, dans un tel circuit, le gain de conversion doit être proportionnel à la puissance du signal LO.

2.5 La contribution de la responsivité à l'augmentation de la plage dynamique

Soit un détecteur stimulé à l'entrée par le signal $v_{RF}(t) = V_{RF} \cos(\omega t)$. La variation de la tension de sortie, qui est proportionnelle à la puissance du signal d'entrée, peut être mesurée avec un voltmètre ayant une précision limitée. Posons ΔV la résolution de mesure de la tension avec notre outil. Si la puissance du signal est suffisamment faible, la variation de tension $|V_s - V_r|$ ne

peut pas être détectée par notre instrument. Posons donc $V_{RF,min}$ l'amplitude de tension du signal V_{RF} minimale pour faire varier la tension de sortie de ΔV par rapport à V_r . En d'autres termes,

$$\exists V_{RF,min} \in \mathbb{R}^+, \Delta V = |V_s(V_{RF} = V_{RF,min}) - V_r| \quad (2.30)$$

Si nous injectons cette définition dans la formule (2.26), nous obtenons :

$$\Delta V = \frac{\alpha_2 V_{RF,min}^2}{2} \quad (2.31)$$

$$V_{RF,min} = \sqrt{2\Delta V / \alpha_2} \quad (2.32)$$

Pour une résolution de mesure donnée, il faut augmenter α_2 afin de maximiser la plage dynamique. D'ailleurs, Qayyum & Negra (2017) partagent également l'avis qu'une bonne responsivité permet de mesurer des signaux de faible puissance.

2.6 Produits d'intermodulation nuisibles

Nous avons vu que, grâce à la non-linéarité, un transistor qui génère des produits d'intermodulation peut être utilisé comme mélangeur ou détecteur. Dans notre cas, nous tirons avantage des produits d'intermodulation d'ordre deux. Idéalement, il nous faut donc une relation purement quadratique entre l'entrée et la sortie d'un système. Or, nous avons vu dans la partie 2.1.1 qu'un MOSFET n'admet pas une relation purement quadratique entre le courant de drain et la tension de grille. Alors, nous pouvons nous demander quel serait l'impact d'une fonction de transfert qui n'est pas parfaitement quadratique.

2.6.1 Les ordres des produits d'intermodulation néfastes

Dans le cas des amplificateurs, des systèmes censés être linéaires, les produits d'intermodulation les plus nuisibles sont les produits d'intermodulation d'ordre 3 (Pozar (2012)). Or, notre détecteur et notre mélangeur ne fonctionnent pas du tout comme un amplificateur. Notamment, l'entrée

et la sortie de ces circuits n'occupent pas du tout la même bande de fréquence. Donc, il faut vérifier quels produits d'intermodulation sont néfastes pour notre détecteur et notre mélangeur.

2.6.1.1 Pour les détecteurs

La détection de puissance se fait grâce au produit d'intermodulation d'ordre deux. Nous pouvons prouver que, de façon générale, seuls les produits d'intermodulation des ordres pairs se trouvent dans la bande de base. (cf. la partie 1.1.1 et le tableau I-1 de l'annexe I.) Plus précisément, si un détecteur reçoit un signal à deux fréquences ($f_c \pm f_m$), les produits d'intermodulation d'ordre pair se trouvent à un multiple pair de f_m . L'intermodulation des ordres impairs génère des produits à des fréquences très supérieures à la bande de base. Donc, elle ne mérite pas notre attention. L'intermodulation d'ordre quatre est la plus néfaste, car son produit se trouve à $4f_m$. C'est très proche de l'enveloppe de puissance du signal modulant à $2f_m$. De plus, en dessous des niveaux très importants, l'intermodulation d'ordre quatre est prépondérante devant les intermodulations d'ordre six, huit, dix, etc. Il faut retenir que la conséquence concrète de l'intermodulation d'ordre quatre (et celles d'ordres supérieurs) sur un signal plus complexe est l'étalement de spectre et la déformation de l'enveloppe de la puissance du signal.

2.6.1.2 Pour les mélangeurs

Comme pour le détecteur de puissance, le mélange de fréquences désiré est un produit d'intermodulation d'ordre deux. Nous avons également démontré que les produits d'intermodulation autour de f_{IF} sont d'ordre pair seulement. (cf. la partie 1.1.2 et le tableau I-2 de l'annexe I.) Comme pour le détecteur, les produits d'intermodulation d'ordre quatre sont les plus néfastes, car ils se trouvent à $f_{IF} \pm 3f_m$, près des fréquences du mélange désiré à $f_{IF} \pm f_m$.

2.6.2 Les conséquences de l'intermodulation d'ordre quatre pour les détecteurs

2.6.2.1 Le produit d'intermodulation d'ordre quatre pour un détecteur

Nous pouvons évaluer les effets du produit d'intermodulation d'ordre quatre sur un détecteur d'enveloppe de puissance en évaluant un signal à deux fréquences élevé à la puissance quatre.

D'après la partie 1.2.1 de l'annexe I,

$$v_{RF}^4(t) = \frac{V_{RF}^4}{16} \left[9/4 + 3 \cos(2\omega_m t) + \frac{3}{4} \cos(4\omega_m t) + \dots \right] \quad (2.33)$$

(Les « ... » représentent tous les termes qui se trouvent en dehors de la bande de base.) Nous pouvons remarquer que l'élévation à la puissance quatre génère des produits d'intermodulation à 0 Hz, à $2f_m$ et à $4f_m$. Rappelons-nous que $2f_m$ est la fréquence de l'enveloppe de puissance du signal modulant. Le produit d'intermodulation d'ordre quatre est néfaste pour plusieurs raisons :

1. Il génère un terme à $4f_m$ qui est trop proche de $2f_m$ pour être facilement filtré.
2. Il génère un terme à $2f_m$ proportionnel à v_{RF}^4 , ce qui modifie la réponse censée être quadratique de l'enveloppe de puissance. Cela a pour conséquence la compression.
3. Il génère un terme à 0 Hz proportionnel à v_{RF}^4 . Cela affecte donc également la lecture de la puissance moyenne. Même si nous avons conçu un détecteur de puissance moyenne, où la distorsion de la forme du signal n'aurait pas été un problème, l'intermodulation d'ordre quatre aurait été nuisible.

Pour de petites amplitudes du signal, v_{RF}^4 est négligeable par rapport à v_{RF}^2 , donc on peut ignorer les trois effets mentionnés ci-dessus. Cependant, v_{RF}^4 s'accroît plus vite que v_{RF}^2 . Donc, pour de grandes amplitudes, les effets néfastes ne peuvent plus être ignorés. D'ailleurs, nous pourrions également démontrer que le signal élevé à n'importe quelle puissance paire (4, 6, 8, etc.) génère des produits d'intermodulation à $2f_m$. Cependant, les produits d'intermodulation d'ordre quatre sont les plus préoccupants.

2.6.2.2 La compression de l'enveloppe de puissance

L'intermodulation d'ordre quatre génère un terme à la fréquence $2f_m$. Cela a pour conséquence la compression de l'enveloppe de puissance du signal.¹

Idéalement, le rapport entre l'amplitude de l'enveloppe de puissance du signal modulé située à $2f_m$ et celle du signal à l'entrée serait une constante proportionnelle à la responsivité défini par (1.2).

$$\frac{\|v_s|_{f=2f_m}\|}{\|v_{RF}\|^2} \approx \alpha_2 \propto R_d \quad (2.34)$$

Or, selon l'équation (2.33), il faut aussi prendre en compte le terme à $2f_m$ de $\alpha_4 v_{RF}^4$. Ce rapport devient alors :

$$\frac{\|v_s|_{f=2f_m}\|}{\|v_{RF}\|^2} \approx \frac{\alpha_2(V_{RF}^2/4) + \alpha_4(3V_{RF}^4/16)}{(V_{RF}^2)/4} = \alpha_2 + \frac{3}{4}\alpha_4 V_{RF}^2 \quad (2.35)$$

La responsivité du détecteur de puissance n'est donc plus une constante, mais il varie avec l'amplitude du signal d'entrée. Cela explique le phénomène de compression. Ensuite, pour modéliser la compression, supposons que les signes de α_2 et de α_4 soient de signes opposés. En effet, s'ils étaient du même signe, alors la responsivité s'accroîtrait sans cesse avec l'amplitude de l'entrée.

Nous pouvons désormais tenter de prédire le point de compression à 1 dB. Posons $V_{1\text{ dB,d}}$ l'amplitude de la tension RF en entrée, telle que la puissance de l'enveloppe de puissance soit 1 dB inférieure à ce qu'elle serait si le circuit était parfaitement quadratique. Pour $V_{RF} = V_{1\text{ dB,d}}$, la puissance à la sortie à $2f_m$ est 1 dB plus basse que si la responsivité était idéale. Cela donne :

$$V_{1\text{ dB,d}} = 2\sqrt{\frac{\alpha_2}{3|\alpha_4|} \left(1 - 10^{-\frac{1}{20}}\right)} \quad (2.36)$$

¹ Nous n'avons pas pris en compte les effets des intermodulations d'ordres supérieurs à 4 pour simplifier les calculs, car l'intermodulation d'ordre quatre demeure prépondérante devant celles d'ordres supérieurs jusqu'à un certain point.

(Voir la partie 2.2 de l'annexe I pour la démonstration.) Pour maximiser la plage dynamique utilisable du circuit, c'est-à-dire maximiser $V_{1\text{ dB,d}}$, il faut maximiser α_2/α_4 .

La même chose s'applique à la compression de 1 dB de la tension continue $|V_s - V_r|$ lorsqu'on stimule un détecteur avec un signal sinusoïdal à une seule fréquence. (Voir l'équation (??) la section 2.1 de l'annexe I.) La mesure du point de compression du détecteur peut donc se faire de deux manières. Une première façon serait d'injecter un signal à amplitude constante et de mesurer la croissance de la tension continue à la sortie en fonction de la puissance du signal. Une seconde méthode serait d'injecter un signal à deux fréquences, puis de mesurer la croissance de l'amplitude du terme à $2f_m$ en fonction de la puissance du signal. De plus, ces deux techniques de mesure sont censées donner les mêmes résultats selon nos calculs.

2.6.2.3 Le point d'interception

Le produit d'intermodulation d'ordre quatre à $4f_m$ est proportionnel à V_{RF}^4 , et notre enveloppe de puissance à $2f_m$ est proportionnelle à V_{RF}^2 . Tandis qu'à basse puissance, le produit d'ordre quatre est négligeable par rapport au produit d'ordre deux, il existe, en théorie, un point où les deux produits d'intermodulation ont la même puissance.

En théorie, il existe une amplitude de tension $V_{IP,d}$ du signal d'entrée telle que les deux puissances soient égales :

$$V_{IP,d} = 4 \sqrt{\frac{\alpha_2}{3|\alpha_4|}} \quad (2.37)$$

(Voir la partie 3.1.1 de l'annexe I pour la démonstration.) Encore une fois, il est de notre intérêt de maximiser $|\alpha_2/\alpha_4|$ pour minimiser la distorsion de l'enveloppe de puissance à la sortie.

2.6.2.4 La relation entre le point de compression et le point d'interception

On remarque d'ailleurs que

$$\frac{V_{IP,d}}{V_{1\text{ dB,d}}} = \frac{2}{\sqrt{1 - 10^{-\frac{1}{20}}}} \approx 6,064 = 15,66 \text{ dB} \quad (2.38)$$

(Voir la partie 3.1.2 de l'annexe I pour la démonstration.) La compression est un phénomène qui se produit bien avant qu'on atteigne le point d'interception. Par conséquent, comme avec les amplificateurs et les mélangeurs, le point d'interception est surtout théorique, car la compression du signal est inévitable.

2.6.3 Les produits d'intermodulations d'ordre quatre pour les mélangeurs

Les produits d'intermodulation d'ordre quatre ont un effet similaire mais différent pour les mélangeurs quadratiques que pour les détecteurs de puissance.

2.6.3.1 Un produit d'intermodulation à IF

Soient toujours un signal RF à deux fréquences et un signal LO. Dans la partie 1.2.2 de l'annexe II, nous avons calculé leur somme élevée à la puissance quatre :

$$\begin{aligned} (v_{RF} + v_{LO})^4 = \dots + \frac{V_{RF}^3 V_{LO}}{16} (9 \cos((\omega_{IF} \pm \omega_m)t) + 3 \cos((\omega_{IF} \pm 3\omega_m)t) \\ + \frac{3V_{RF}V_{LO}^3}{4} \cos((\omega_{IF} \pm \omega_m)t) + \dots \end{aligned} \quad (2.39)$$

(Lire les signes $\pm$ ci-dessus comme « plus **et** moins », plutôt que « plus **ou** moins ».) Nous nous apercevons que, comme dans le cas du détecteur, l'intermodulation d'ordre quatre génère des produits à des fréquences voisines de celles des produits désirés. Dans ce cas-ci, des produits, proportionnels à $V_{RF}^3 V_{LO}$, sont générés à $f_{IF} \pm 3f_m$. Ils se trouvent trop proches de $f_{IF} \pm f_m$ pour un filtrage facile. De plus, tout comme dans le cas du détecteur, l'intermodulation d'ordre quatre génère des produits aux mêmes fréquences que l'intermodulation d'ordre deux, mais à $f_{IF} \pm f_m$ cette fois-ci. Ces produits sont proportionnels à $V_{RF}^3 V_{LO}$ ou à $V_{RF} V_{LO}^3$. Étant donné que V_{LO} reste constant lors de l'opération du mélangeur, le terme le plus néfaste est celui proportionnel à $V_{RF}^3 V_{LO}$. Une fois de plus, le gain de conversion n'est pas constant à cause de l'intermodulation d'ordre quatre.

Pour le mélangeur également, nous pourrions démontrer que les intermodulations d'ordres pairs (4, 6, 8, ...) génèrent des produits à $f_{IF} \pm f_m$. Et nous avons déjà démontré que les intermodulations d'ordre pair ($2i$) génèrent des produits à des fréquences voisines de $f_{IF} \pm (2i + 1)f_m$.

2.6.3.2 La compression du signal transposé

Les mélangeurs de fréquence sont également soumis au phénomène de compression, qui est causé par les intermodulations des ordres pairs supérieurs à deux. Si l'on considère désormais les produits d'intermodulation d'ordre quatre, alors, comme pour les détecteurs, la formule du gain de conversion devient dépendante de l'amplitude du signal à l'entrée.

$$\begin{aligned} \frac{\|v_o|_{f_{IF} \pm f_m}\|}{\|v_{RF}\| \|v_{LO}\|} &\approx \frac{\frac{\alpha_2}{2} V_{RF} V_{LO} + \frac{3}{4} \alpha_4 V_{RF} V_{LO}^3 + \frac{9}{16} \alpha_4 V_{RF}^3 V_{LO}}{\frac{V_{RF}}{2} V_{LO}} \\ &= \alpha_2 + \alpha_4 \left[\frac{3V_{LO}^2}{2} + \frac{9V_{RF}^2}{8} \right] \end{aligned} \quad (2.40)$$

La compression peut à la fois dépendre de l'amplitude de v_{RF} et de celle de v_{LO} . Plus les amplitudes V_{RF} et V_{LO} sont importantes, plus le gain sera faible, puisque nous supposons toujours que $\alpha_2 > 0 > \alpha_4$.

Au point de compression de 1 dB, la puissance à la sortie est 1 dB inférieure à ce qu'elle serait si le système était parfaitement quadratique. Dans le cas du mélangeur, le point de compression $V_{1 \text{ dB,m}}$ a la formule suivante :

$$V_{1 \text{ dB,m}} = \frac{4}{3} \sqrt{\frac{\alpha_2}{2|\alpha_4|} (1 - 10^{\frac{-1}{20}}) - \frac{3V_{LO}^2}{2}} \quad (2.41)$$

(Voir la partie 2.3 de l'annexe I pour la démonstration.) La formule ci-dessus est très semblable à celle du point de compression du détecteur, l'équation (2.36), sauf pour $-2V_{LO}^2$. Pour le mélangeur, plus V_{LO} est important, plus l'amplitude V_{RF} nécessaire pour causer de la compression est faible.

2.6.3.3 Le point d'interception

Au point d'interception, les produits d'intermodulation d'ordre quatre néfastes à $f_{IF} \pm 3f_m$ ont la même amplitude que le signal transposé à $f_{IF} \pm f_m$. Il existe une amplitude du signal RF V_{IP} telle que $P_{f_{IF} \pm f_m} |_{\alpha_i=0, i>2}(V_{RF} = V_{IP,m}) = P_{f_{IF} \pm 3f_m} |_{\alpha_i=0, i>4}(V_{RF} = V_{IP,m})$.

$$V_{IP,m} = 4 \sqrt{\frac{\alpha_2}{6|\alpha_4|}} \quad (2.42)$$

(Voir la partie 3.2.1 de l'annexe I pour la démonstration.) Une fois de plus, comme dans l'équation (2.37), le point d'interception dépend du ratio $\alpha_2/|\alpha_4|$. Et la formule ressemble beaucoup à celle du point d'interception du détecteur.

2.6.3.4 La relation entre le point de compression et le point d'interception

Une fois de plus, nous pouvons établir une relation entre le point de compression à 1 dB et le point d'interception :

$$V_{IP,m} = \sqrt{\frac{3V_{1\text{ dB},m}^2 + 8V_{LO}^2}{1 - 10^{\frac{-1}{20}}}} \quad (2.43)$$

(Voir la partie 3.2.2 de l'annexe I pour la démonstration.) Pour conclure, la figure 2.5 résume cette sous-partie concernant les effets des produits d'intermodulation d'ordre quatre sur les détecteurs et les mélangeurs quadratiques. L'intermodulation d'ordre quatre génère des produits nuisibles qui s'accroissent plus vite que les signaux désirés. De plus, elle cause la compression de ces derniers. La compression a lieu avant que les produits d'intermodulation nuisibles n'atteignent la même puissance que les produits désirés.

2.7 Conclusion du chapitre

Dans ce chapitre, nous avons vu qu'un transistor MOSFET polarisé près de la tension de seuil V_{th} a un comportement plutôt, mais pas entièrement, quadratique. Puis, nous avons pu comprendre comment une réponse quadratique permet la détection de la puissance et le mélange

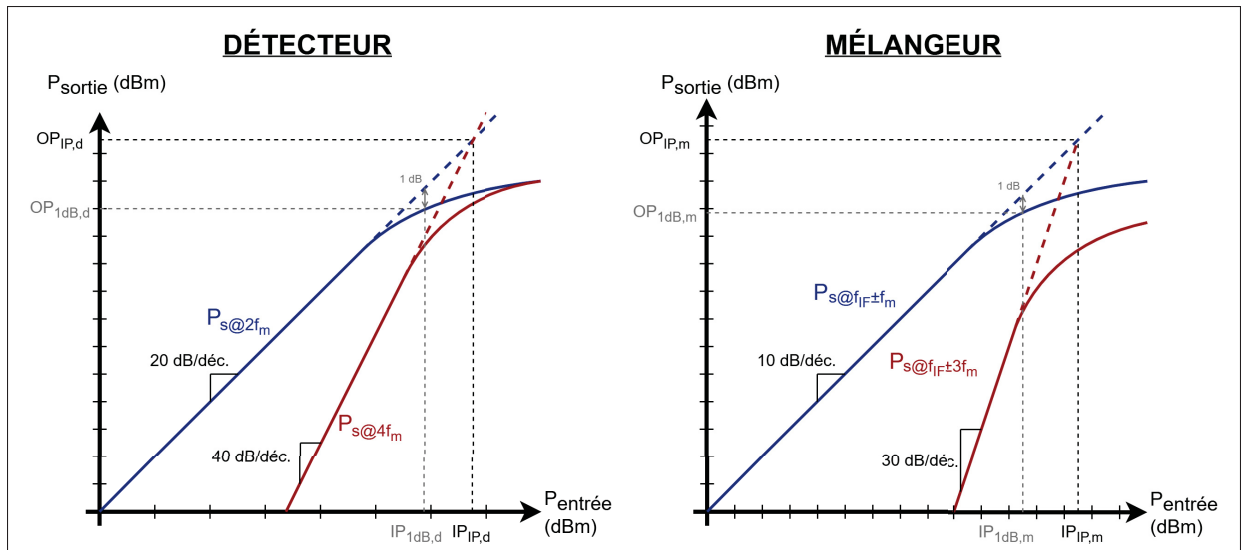


Figure 2.5 Les produits d'intermodulation à la sortie d'un détecteur et d'un mélangeur en théorie

de fréquences. Un circuit quadratique peut d'ailleurs accomplir les deux tâches simultanément, car c'est exactement le même mécanisme en question. Cependant, un MOSFET admet un comportement non linéaire d'ordre supérieur à deux. Or, l'intermodulation d'ordres pairs supérieurs à deux engendre des effets nuisibles pour nos circuits. Cela se manifeste sous la forme de compression et d'apparition de produits à des fréquences voisines de celles des signaux désirés.

CHAPITRE 3

CONCEPTION DU CIRCUIT

3.1 La conception du circuit intégré

Pour concevoir un circuit de détecteur-mélangeur quadratique, nous nous sommes basés sur la topologie utilisée par Bourgault (2023) pour son détecteur d'enveloppe de puissance. Si nous avons repris sa topologie, c'était pour sa simplicité, son bon taux de réjection de la porteuse et sa bonne bande passante de sortie (simulée).

Cette section couvre la phase de conception du circuit intégré du détecteur-mélangeur nommé ICSTSEM1. Nous mènerons notamment une étude approfondie du fonctionnement du circuit intégré. Nous calculerons analytiquement le gain en tension quadratique du circuit. Ensuite, nous expliquerons comment nous avons dimensionné chaque composant grâce à l'outil de simulation. Enfin, nous vérifierons la validité de la formule analytique. Par ailleurs, l'annexe II contient des informations supplémentaires sur le circuit intégré.

3.1.1 La technologie de fabrication utilisée

Pour réaliser notre circuit intégré, nous avons utilisé la technologie CMOS de 65 nm de TSMC. Elle permet l'intégration des transistors à effet de champ, dont la longueur de grille a une longueur minimale de 65 nm. Les neuf couches métalliques de la technologie permettent de relier les différents composants passifs et actifs de manière flexible. Des descriptions détaillées des différentes couches, des transistors et des diodes sont fournies dans la section 2 de l'annexe II.

3.1.2 L'analyse du fonctionnement du circuit intégré

Cette partie traite de l'analyse du circuit intégré du détecteur-mélangeur. Nous avons effectué une analyse en petits signaux en nous inspirant de Razavi (2013). Une telle analyse permet de comprendre comment différents paramètres des composants affectent la performance du

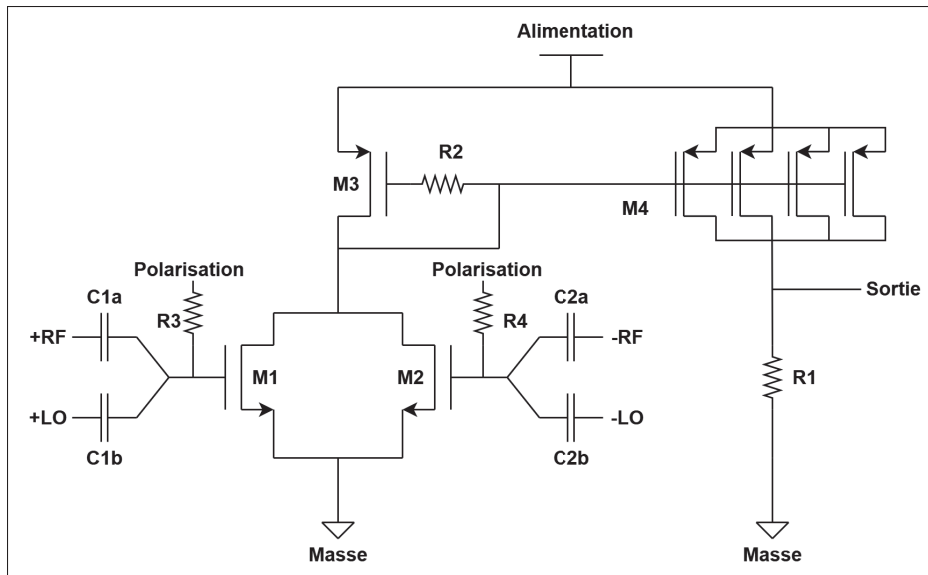


Figure 3.1 Schéma détaillé du circuit intégré

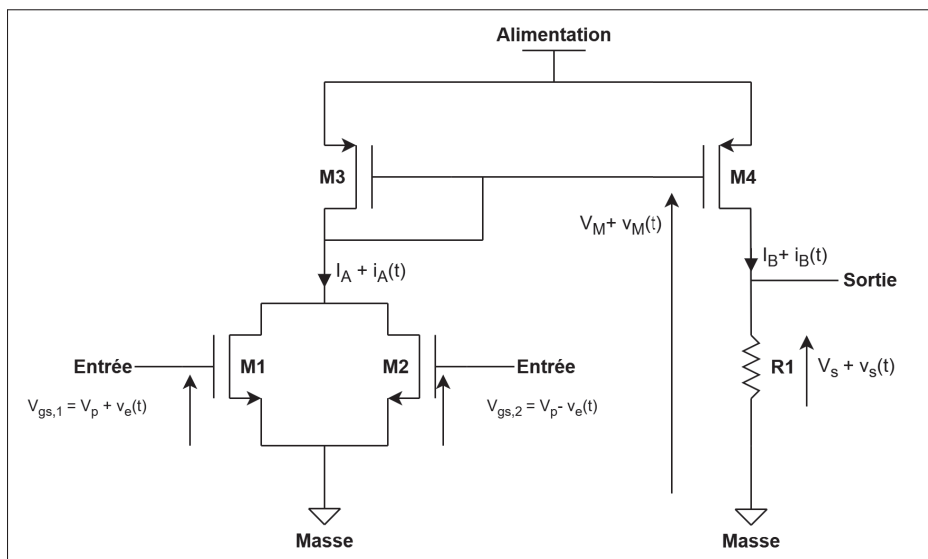


Figure 3.2 Schéma simplifié du circuit intégré ICSTSEM1

circuit. L'analyse en petits signaux classique suppose que les composants agissent linéairement. Cependant, dans notre cas, nous avons modifié la méthode de l'analyse pour prendre en compte la non linéarité des transistors.

La figure 3.1 est un schéma détaillé du circuit intégré ICSTSEM1 du détecteur-mélangeur ; tous les composants essentiels y sont représentés. Cependant, pour l'analyse du circuit, nous nous servons du schéma de la figure 3.2, qui ne retient que l'essentiel de la topologie. Ce dernier fait abstraction des condensateurs de couplage pour l'entrée, des résistances de polarisation et de la résistance R2 qui ne sert qu'à augmenter la bande passante. M4 est en réalité quatre PMOS reliés en parallèle. Cependant, il est équivalent à un transistor quatre fois plus large que M3.

3.1.2.1 Une explication qualitative du circuit

Tous les transistors sont censés opérer en région de saturation lorsque la tension de polarisation V_p est suffisamment élevée. Comme illustré sur la figure 3.2, une tension de polarisation aux grilles de M1 et de M2 les fait tirer un courant continu I_A . Des tensions de signal $+v_e(t)$ et $-v_e(t)$ s'ajoutent respectivement aux grilles des transistors M1 et M2. Ces tensions provoquent la circulation du courant alternatif $i_A(t)$. Puisque les entrées sont symétriques et que les deux drains sont reliés, les parties impaires des courants des drains s'annulent, tandis que les parties paires s'additionnent. Le courant de drain de M3 est donc principalement le carré du signal $v_e(t)$. M3 et M4 forment un miroir de courant. M4 donc réplique le courant $I_A + i_A(t)$. La résistance R1 permet de convertir ce courant en tension $V_s + v_s(t)$.

Il aurait été tout à fait possible de n'utiliser qu'un seul NMOS, opérant en mode commun, au lieu d'une entrée symétrique, et d'utiliser un filtre passe-bas pour supprimer les composants de $i_A(t)$ se trouvant à la fréquence porteuse et au-delà. Par exemple, c'est ce qui est fait dans le circuit de Serhan *et al.* (2015a). Cependant, l'avantage d'utiliser la symétrie est qu'elle n'a pas de fréquence de coupure. Si nous avions placé un condensateur, il aurait imposé une limite fixe de fréquence à $i_A(t)$, alors que la symétrie permet de supprimer la porteuse (et toute la partie impaire du signal) peu importe la fréquence de la porteuse. De plus, la symétrie donne au circuit une meilleure immunité contre le bruit selon Zhang *et al.* (2007). Cependant, le désavantage de l'entrée symétrique est la nécessité de rendre le signal symétrique.

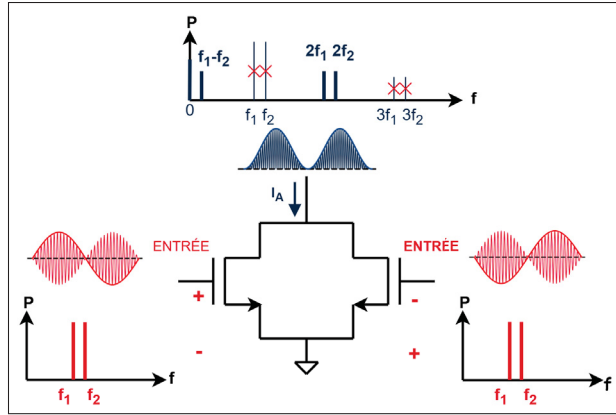


Figure 3.3 Transistor NMOS non linéaires du circuit intégré

3.1.2.2 Les hypothèses pour l'analyse

Pour permettre l'analyse, nous avons posé un certain nombre d'hypothèses simplificatrices :

1. Nous supposons que tous les transistors restent en saturation.
2. Les tensions de polarisation de M1 et de M2 sont identiques, mais des signaux parfaitement symétriques sont appliqués. ($V_{gs,1} = V_p + v_e(t)$ et $V_{gs,2} = V_p - v_e(t)$)
3. Supposons que, comme exprimé dans l'équation (2.4), la relation entre le courant de drain et la tension de grille s'exprime comme une série de Taylor.
4. Supposons également que $\forall i \in \mathbb{N}^*, \forall v_{gs} \in \mathbb{R}, g_{mi} v_{gs}^i \gg g_{m(i+1)} v_{gs}^{i+1}$. Nous pouvons justifier cela en supposant que la tension v_{gs} soit suffisamment petite.
5. Supposons que les transistors M1 et M2 soient parfaitement identiques. $\forall i \in \mathbb{N}^*, g_{mi,1} = g_{mi,2}$. Cette hypothèse perd sa validité si M1 et M2 sont asymétriques.
6. Rien n'est branché à la sortie : il y a un circuit ouvert.
7. Négligeons les capacités parasites.
8. Les tensions d'alimentation et de polarisation sont constantes.
9. La modulation de longueur de canal n'est pas négligée. Comme représenté sur la figure 3.4, nous posons une résistance r_o d'une valeur finie en parallèle avec la source de courant.
10. Nous négligerons l'effet du substrat.

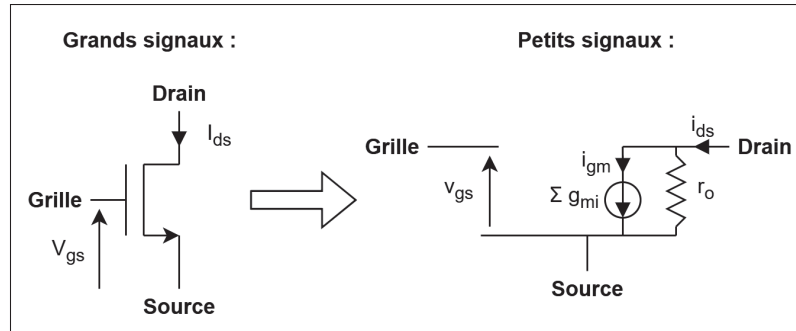


Figure 3.4 Modèle de transistor en petits signaux utilisé

3.1.2.3 Le gain de tension quadratique idéal

Avant d'entamer des calculs précis du gain, nous pouvons d'abord développer une formule simple pour comprendre l'essence du circuit. Seulement pour cette sous-partie, ignorons l'effet de modulation de la longueur du canal (*i.e.* $r_o \rightarrow \infty$). Alors, nous avons pu calculer un gain de tension quadratique idéal :

$$\alpha_2 = \frac{v_s}{v_e^2} = \frac{2R_1 g_{m1,4} g_{m2,1}}{g_{m1,3}} = 2R_1 K g_{m2,1} \quad (3.1)$$

(Voir la partie 4.2 de l'annexe I pour les détails)

Le gain est assez naturellement proportionnel à la transconductance quadratique de M1 et de M2, $g_{m2,1} = g_{m2,2}$. Cette dernière est, en quelque sorte, le cœur de ce circuit qui dépend des caractéristiques physiques et de la polarisation des transistors M1 et M2. On peut remarquer que le gain idéal est également proportionnel au ratio $g_{m1,4}/g_{m1,3} = I_B/I_A = W_4/W_3 = K$. Idéalement, on pourrait augmenter le gain en élargissant M4 ou en utilisant plusieurs transistors en parallèle avec M4.

3.1.2.4 Le gain quadratique de tension en considérant la modulation du longueur de canal

Nous pouvons ensuite calculer le gain en tension quadratique du circuit en prenant en compte l'effet de modulation du canal. Selon la partie 4.3 de l'annexe I, le gain en tension quadratique du circuit vaut donc :

$$\alpha_2 = \frac{v_s}{v_e^2} = \frac{2 (r_{o4} || R_1) (\frac{r_{o1}}{2} || r_{o3}) g_{m1,4} g_{m2,1}}{1 + (\frac{r_{o1}}{2} || r_{o3}) g_{m1,3}} \quad (3.2)$$

Nous pouvons déduire qu'à cause de l'effet de modulation de la longueur du canal, qui se manifeste sous la forme de résistances r_o , le gain est réduit par rapport au cas idéal. Or, cet effet n'est définitivement pas négligeable lorsque la longueur de la grille des transistors tend vers 65 nm.

Nous souhaitons ensuite exprimer le gain en tension quadratique en fonction des dimensions des transistors. $g_{m1,4}$ et $g_{m1,3}$ sont des grandeurs dont la formule est bel et bien connue chez les concepteurs de circuits. En effet, selon le modèle de l'équation (2.1), $g_{m1} = \frac{2I_{ds}}{V_{ov}}^1$ (où $V_{ov} = V_{gs} - V_{th}$). Les résistances r_o peuvent également être exprimées en fonction des dimensions des composants. $r_o = (I_{ds}\lambda)^{-1}$. Bien que λ soit une grandeur difficile à contrôler pour un concepteur de circuits, on peut toujours régler I_{ds} . Nous avons donc presque entièrement pu exprimer les (3.2) en fonction des dimensions et de la tension de polarisation de M1 et de M2 :

$$\alpha_2 = \frac{4 \left(\frac{R_1}{1 + R_1 \lambda_4 K \mu_n C_{ox} \frac{W_1}{L_1} V_{ov,1}^2} \right) \frac{K g_{m2,1}}{(\lambda_1 + \lambda_3) V_{ov,1}} \sqrt{\frac{\mu_p \frac{W_3}{L_3}}{2 \mu_n \frac{W_1}{L_1}}}}{1 + \frac{2}{(\lambda_1 + \lambda_3) V_{ov,1}} \sqrt{\frac{\mu_p \frac{W_3}{L_3}}{2 \mu_n \frac{W_1}{L_1}}}} \quad (3.3)$$

(c.f. la partie 4.3 de l'annexe I pour les calculs.) Le gain est toujours proportionnel à $g_{m2,1}$. Selon (2.9), cette grandeur ne dépend pas de la polarisation. Cependant, les simulations montrent qu'elle dépend effectivement de W_1 , L_1 et de $V_{ov,1}$. Quant aux autres grandeurs facilement

¹ Nous sommes conscient que ce modèle n'est pas précis. Cependant, il nous permet d'obtenir une première formule approximative du gain en fonction des dimensions.

ajustables, telles que K , R_1 , $V_{ov,1}$, W_1 , L_1 , W_3 et L_3 , nous pouvons présumer qu'il existe des valeurs optimales pour chacune de ces grandeurs afin de maximiser le gain, sans pour autant compromettre d'autres critères de performance du circuit.

D'un côté, la formule (3.3) nous permet de comprendre que l'effet de modulation de la longueur du canal contrarie l'augmentation du gain. D'un autre côté, cette formule est limitée en utilité et en précision tant que l'on utilise le modèle de MOSFET simpliste comme base, et tant que l'on n'a pas de formule précise pour $g_{m2,1}$.

3.1.2.5 L'impédance d'entrée

L'impédance d'entrée dépend des résistances de polarisation R_3 et R_4 , des condensateurs de découplage C_1 et C_2 et des capacités des grilles de M_1 et de M_2 . Si nous évaluons l'impédance perçue au niveau de la grille de M_1 , alors l'impédance d'entrée est la résistance de polarisation R_3 en parallèle avec la capacité de la grille de M_1 . Si M_1 est en saturation, alors la capacité grille-source $C_{gs,1} = \frac{2}{3}W_1L_1C_{ox}$ et la capacité drain-source $C_{gd,1} = 0$ (Gray *et al.* (2009) p.51). Par conséquent, l'impédance perçue au niveau de la grille devient $R_3 || (1/[jC_{gs,1}\omega])$. Si le circuit du détecteur-mélangeur était directement relié à la sortie d'un autre circuit (tel qu'un amplificateur) présent sur la même puce, alors nous souhaiterions que R_3 soit le plus grand possible et que la capacité de M_1 soit la plus petite possible pour ne pas affecter le fonctionnement du circuit qui les précède. La même chose s'applique à M_2 et R_4 .

Nous pourrions penser que minimiser les dimensions de M_1 et de M_2 suffit pour maximiser la fréquence d'opération du circuit de détecteur-mélangeur. Cependant, il faut également prendre en compte les capacités des diodes de protection, reliées près des pastilles de contact, et les effets parasites des carrés de remplissage. En effet, les diodes et celles des pastilles de connexion sont très grandes par rapport aux transistors M_1 et M_2 . Donc, les simulations nous ont montré que ces éléments, et non pas M_1 et M_2 , sont les facteurs limitant la fréquence de coupure au niveau de l'entrée.

3.1.2.6 L'impédance de sortie

Si l'on néglige les capacités parasites, l'impédance de sortie est la résistance R_1 en parallèle avec l'impédance du drain de M4. Elle vaut donc $R_1 || r_{o4}$. Afin d'adapter l'impédance de la sortie à 50Ω , il faut ajuster R_1 ou r_{o4} . Cependant, cela peut compromettre le gain en tension α_2 . De plus, la valeur de R_1 et les dimensions de M4 doivent être également choisies de telle façon que M4 reste en saturation. En effet, il faut que $R_1 I_B$ soit suffisamment bas, de tel que $|V_{ds,4}| > |V_{ov,4}|$. Donc, le dimensionnement de R_1 et de M4 est un compromis entre le gain, l'adaptation à 50Ω et la polarisation de M4. Si l'on prend en compte les capacités parasites, alors les capacités de jonction des diodes de protection réduisent la fréquence de coupure à la sortie du circuit.

3.1.3 Dimensionnement des composants

3.1.3.1 Dimensionnement des transistors NMOS

Le fonctionnement du circuit entier dépend des NMOS et de leur comportement quadratique. Aussi les NMOS furent les premiers composants à être dimensionnés lors de la conception du circuit intégré. Étant donné que l'on cherche à maximiser le gain quadratique α_2 et à minimiser les produits d'intermodulation néfastes, il a fallu dimensionner le circuit de telle sorte que $g_{m2,1}$ et $g_{m2,2}$ soient maximisées et que $g_{m4,1}$ et $g_{m4,2}$ soient minimisées.

Nous avons réalisé les simulations sur le logiciel de conception Cadence Virtuoso. Lors des simulations à courant continu, nous avons simulé les caractéristiques d'un NMOS isolé avec une tension de drain V_{ds} fixe. Le logiciel donnait la transconductance g_{m1} du transistor en fonction de V_{gs} . Nous avons donc dérivé g_{m1} pour obtenir g_{m2} et g_{m4} . D'abord, nous avons voulu étudier l'effet de la longueur de la grille L sur g_{m2} et g_{m4} pour une tension V_{ds} et une largeur fixes. Sur la figure 3.5, nous avons illustré g_{m2} , g_{m4} et $|g_{m4}/g_{m2}|$ en fonction de V_{gs} pour différentes longueurs de grille. En premier lieu, nous constatons que g_{m2} varie avec la tension de la grille. Nous en déduisons que la formule (2.9), qui affirme que g_{m2} est constante, est fausse. Nous constatons que g_{m2} diminue avec V_{gs} qui augmente. Le maximum de g_{m2} se situe vers la tension

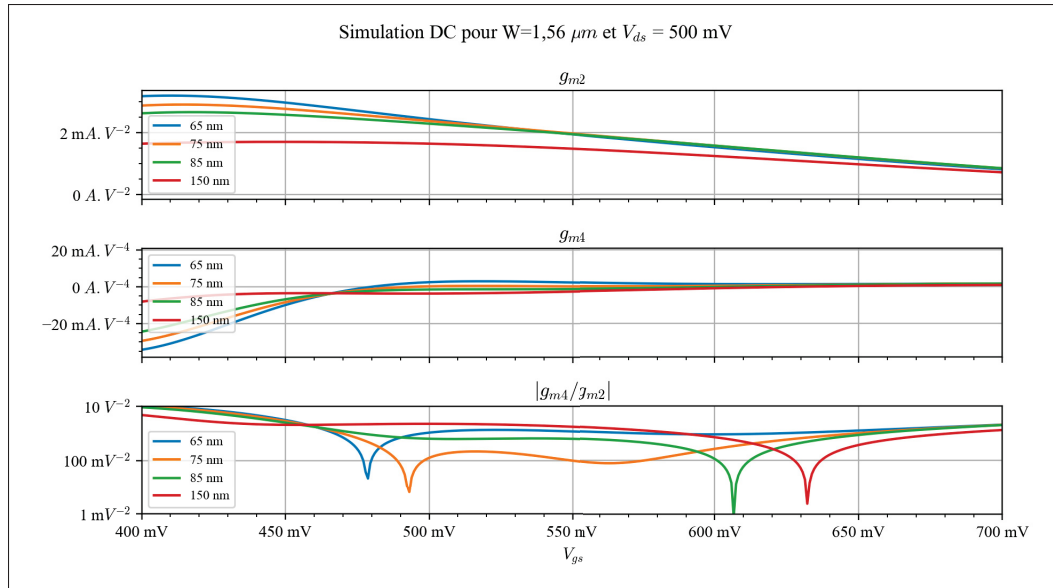


Figure 3.5 Simulation DC des dérivées de la transconductance en fonction de la tension et la longueur de la grille pour un NMOS

de seuil $V_{th} \approx 400$ mV. D'après les simulations, s'il ne s'agissait que de maximiser le gain quadratique, on choisirait la longueur minimale permise par la technologie de fabrication : 65 nm. Quant à g_{m4} , l'allure de ses courbes varie avec la longueur. Celle où $L = 75$ nm semble la plus prometteuse, car g_{m4} semble tendre vers 0 sur la plus grande plage de tension. En effet, pour $L = 75$ nm, la courbe de $|g_{m4}/g_{m2}|$ est relativement basse sur une grande plage de tension. Au final, une longueur de grille de 75 nm a été choisie pour les transistors NMOS.

Ensuite, pour évaluer l'effet de la largeur de la grille, nous avons simulé la transconductance et ses dérivées en fonction de la largeur W pour une longueur L et une tension V_{ds} fixes. La figure 3.6 illustre la relation entre les dérivées de la transconductance et la largeur. Contrairement à la longueur, la largeur ne semble pas avoir d'effet sur l'allure des courbes ; elle ne semble affecter que l'amplitude. D'un côté, augmenter la largeur de M1 et de M2 augmente $g_{m2,1}$ et $g_{m2,2}$. D'un autre côté, augmenter la largeur diminue r_{o1} et r_{o2} . un compromis doit être atteint pour équilibrer les deux grandeurs. Au final, une largeur de $1,56 \mu\text{m}$ a été choisie par optimisation numérique. Les critères d'optimisation étaient notamment le gain quadratique et la distorsion.

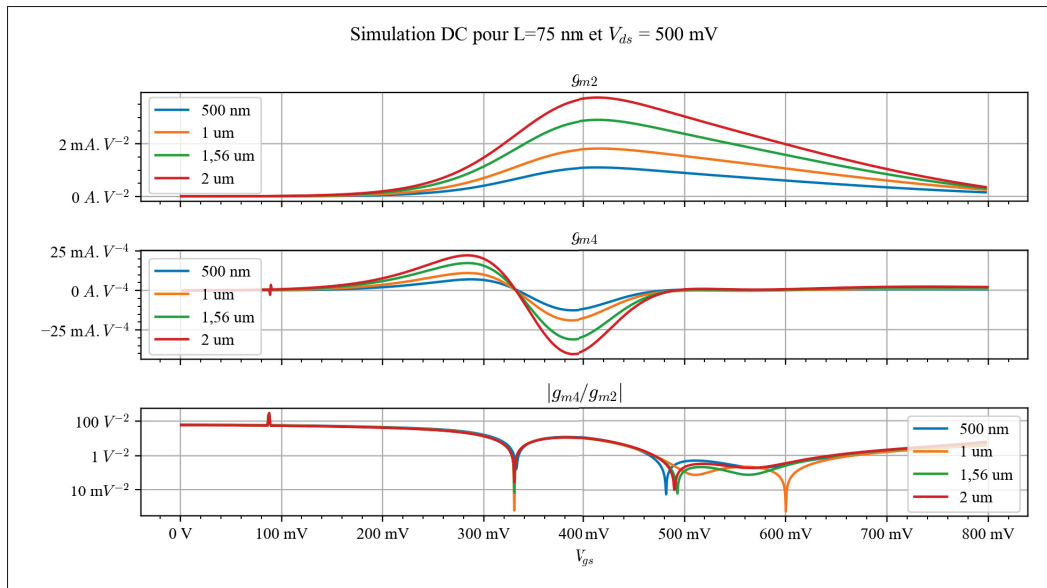


Figure 3.6 Simulation en courant continu des dérivées de la transconductance d'un NMOS en fonction de la polarisation et de la largeur de la grille

3.1.3.2 Dimensionnement des transistors PMOS

Les transistors PMOS du circuit constituent le miroir de courant. Ils ont dû être dimensionnés de telle sorte qu'ils ajoutent le moins de distorsion possible tout en maximisant le gain quadratique et la bande passante. Or, plusieurs considérations entraînent en conflit les unes avec les autres en ce qui concerne le choix des dimensions de M3 et de M4 :

1. Tout d'abord, nous souhaitons que tous les transistors fonctionnent en saturation. M3 devait avoir assez de transconductance pour pousser le courant de drain que M1 et M2 tirent, sans affecter ces derniers. En effet, la grille de M3 est reliée aux drains de M1 et de M2. Si M3 manque de transconductance, alors $|V_{gs,3}|$ doit augmenter à tel point que M1 et M2 soient mis hors de la saturation (*i.e.* $V_{ds,1} < V_{ov,1}$). Donc W_3/L_3 et W_4/L_4 doivent être suffisamment larges.
2. Ensuite, en ce qui concerne la distorsion, contrairement à M1 et M2, nous cherchions à ce que M3 et M4 agissent le plus linéairement possible. Donc, M3 et M4 devaient opérer dans

la région de saturation, loin de la tension de seuil V_{th} . Pour cela, il fallait que M3 et M4 aient des ratios W/L assez petits vis-à-vis de leurs courants de drain.

3. En plus de cela, comme indiqué dans la formule (3.2), il fallait maximiser r_{o3} et r_{o4} pour augmenter le gain. Or, pour maximiser ces valeurs, il fallait également que W_3/L_3 et W_4/L_4 soient suffisamment faibles.
4. Pour autant, la variation de la tension de drain de M3 $V_{ds,3}$ devait être restreinte, parce que la modulation de la longueur du canal n'est pas un phénomène linéaire (Ou & Farahmand (2012)). Puisque ce phénomène est également non linéaire, une trop grande variation $v_{ds,3}$ peut aussi engendrer des produits d'intermodulation importants. Pour limiter l'amplitude de $v_{ds,3}$, il fallait aussi que W_3/L_3 soit agrandie.
5. Enfin, en ce qui concerne la bande passante, il faut savoir que plus les transistors sont grands, plus il y a de capacités parasites qui réduisent la fréquence de coupure du circuit. Pour minimiser les capacités parasites, il fallait minimiser les surfaces des grilles W_3L_3 et W_4L_4 . La dimension la plus critique était notamment celle de M4, qui était en réalité K copies du transistor M3 en parallèle. K devait être suffisamment grand pour maximiser le gain quadratique, mais pas excessif pour minimiser les capacités parasites.

Ainsi, en raison de ces nombreuses considérations, le choix final de W_3/L_3 et de W_4/L_4 fut le résultat d'un compromis réalisé par l'outil d'optimisation du logiciel de CAO.

3.1.3.3 Dimensionnement des composants passifs

Comme pour les transistors, les composants passifs ont été dimensionnés à l'aide des simulations du logiciel de CAO. Les condensateurs de blocage de courant continu C1a, C1b, C2a et C2b empêchent tout courant continu provenant des ports $\pm RF$ et $\pm LO$ d'affecter la tension de polarisation. Ils sont dimensionnés suffisamment grands pour présenter une impédance très faible à des fréquences de plusieurs gigahertz. Les résistances de polarisation R3 et R4 ont été dimensionnées avec une grande résistivité pour ne pas affecter le reste du circuit. La résistance de R1, quant à elle, a été choisie suffisamment importante pour augmenter le gain en tension, mais suffisamment basse pour ne pas mettre M4 hors de la région de saturation. De plus, la

valeur de R_1 a été choisie pour être plutôt basse pour faciliter l'adaptation de l'impédance de la sortie à $50\ \Omega$. Enfin, la résistance R_2 , qui se situe au niveau du miroir de courant, sert à augmenter la bande passante du circuit selon une technique proposée par Voo & Toumazou (1995) et reprise par Bourgault (2023). À basse fréquence, cette résistance n'a aucun effet sur la polarisation. À haute fréquence, où les impédances des grilles diminuent, la résistance dissuade le courant de passer par la grille de M_3 . La valeur de R_2 a été choisie de telle sorte que la bande passante soit maximisée.

3.1.4 La vérification de l'analyse par simulation

Après avoir finalisé la conception du circuit intégré, nous avons utilisé les outils de simulation afin de vérifier notre analyse du gain en tension quadratique. Lors d'une simulation, nous avons appliqué au circuit ICSTSEM1, relié à aucune charge, un signal à deux fréquences ($f_c \pm f_m$) d'une amplitude de 100 mV. Puis, nous avons mesuré l'amplitude de la tension à la sortie de l'enveloppe du signal modulant à $2f_m$. Cette simulation a été répétée pour différentes tensions de polarisation. Ce gain simulé est illustré dans la figure 3.7. Pour obtenir des valeurs numériques du gain selon les formules de l'analyse, nous avons d'abord obtenu les caractéristiques des transistors (tels que r_o , g_m , ...) en fonction de la tension de polarisation grâce à une simulation en courant continu. Puis nous avons injecté ces valeurs dans les formules (3.1) et (3.2), donnant les courbes du gain calculé idéal et du gain calculé non idéal sur la figure 3.7.

Selon la simulation du circuit ICSTSEM1, notre formule du gain idéal (3.1) surestime largement le gain pour toute tension de polarisation. Cela montre que l'effet de modulation de la longueur du canal réduit considérablement le gain. La formule du gain non idéal (3.2) surestime également le gain pour une tension de polarisation inférieure à 500 mV. En effet, la formule découle de l'hypothèse selon laquelle tous les transistors sont en saturation. Or, d'après les simulations, tous les transistors sont en saturation si la tension de polarisation est supérieure à 425 mV. Nous pouvons donc conclure que la formule du gain en tension non idéale établie précédemment est assez fidèle à la simulation si tous les transistors sont saturés.

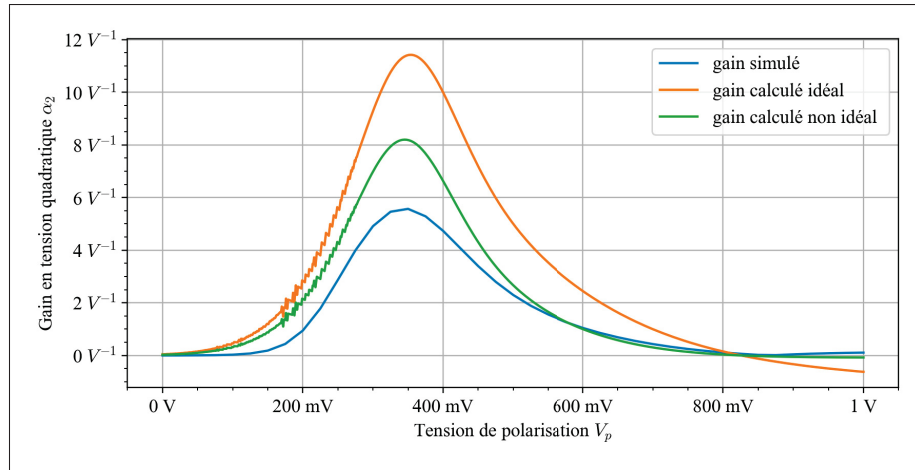


Figure 3.7 Comparaison des gains en tension quadratique selon la simulation et les analyses

3.1.5 Les limites de notre analyse

Les simulations démontrent bel et bien les limites de nos analyses précédentes en termes de précision. Cela peut être expliqué par le fait que nous nous sommes basés sur un modèle très simpliste du MOSFET en saturation (cf. équation 2.1). Ce modèle n'est valide que dans la région de saturation. En plus de cela, bien que M4 soit formé de quatre copies exactes de M3 en parallèle, le rapport $K = g_{m1,4}/g_{m1,3}$ varie avec la tension de polarisation du circuit, selon les simulations.

Non seulement notre analyse précédente ne permet pas d'estimer le gain en tension précisément pour toutes les tensions de polarisation, mais elle néglige plusieurs sources de produits d'intermodulation nuisibles. D'abord, nous avons ignoré les coefficients $g_{m(2i),1}$ et $g_{m(2i),2}$ pour $i > 1$, qui sont les causes principales des produits d'intermodulation d'ordres pairs, dont l'ordre quatre. Ensuite, nous avons considéré les transistors M3 et M4 comme linéaires. Or, ces transistors admettent aussi une (légère) non-linéarité. Contrairement aux transistors M1 et M2, qui bénéficient de la symétrie pour la suppression d'une partie des produits d'intermodulation, M3 et M4 peuvent générer des produits d'intermodulation d'ordres pairs et impairs. Les intermodulations d'ordres pairs au sein de M3 et de M4 peuvent également transposer des

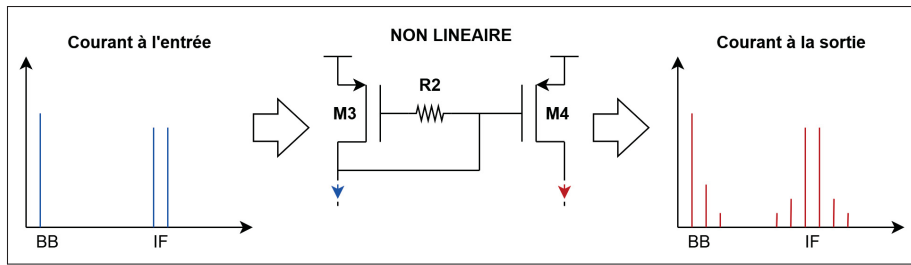


Figure 3.8 Étalement spectral causé par un miroir de courant non linéaire

signaux de haute fréquence vers la bande de base et la bande IF. De plus, les intermodulations d'ordres impairs peuvent causer l'étalement spectral des signaux en bande de base et en bande IF. Enfin, M3 et M4 peuvent générer des harmoniques du signal de bande de base. Par ailleurs, dans l'analyse, nous avons supposé que l'effet de modulation de la longueur du canal était linéaire avec r_o . Or, cela n'est pas le cas. Si nous souhaitions être très fidèles à la réalité, il aurait fallu décrire la transconductance comme une série de Taylor bidimensionnelle, comme nous l'avons mentionné dans la partie 2.1.2.1.

3.1.6 La protection contre les décharges électrostatiques

Les circuits intégrés sont très sensibles aux décharges électrostatiques (DES) (ou *electrostatic discharge* (ESD) en anglais), causées par les machines lors de la fabrication ou par des humains lors de la manipulation (Weste & Harris (2011)). Par conséquent, lors de la conception du circuit intégré, une attention particulière a été portée à la prévention des dégâts causés par ce phénomène. Il existe différentes méthodes pour protéger les circuits vitaux de la puce contre ces décharges. Une technique populaire parmi plusieurs qui engendre moins de capacité parasite est l'usage de diodes reliées en parallèle à la masse et à l'alimentation de la puce (Richier *et al.* (2000)). Toutes les pastilles de connexion de la puce ont une protection contre les DES comme représenté sur la figure 3.9. Près des pastilles de connexion, sont placées des diodes reliées à la masse ou à l'alimentation. Les diodes sont reliées de telle façon qu'en fonctionnement normal, aucun courant ne passe par elles. Mais, si à cause d'une DES le potentiel de la pastille montait au-dessus de celui de l'alimentation ou en dessous de celui de la masse, les diodes

feraient bifurquer le courant. Ainsi, ce courant potentiellement dévastateur n'atteint pas le circuit principal. La structure physique des diodes est illustrée sur la figure II-6 de l'annexe II. Pour le bon fonctionnement de la protection, les diodes doivent pouvoir supporter un courant important et offrir un chemin très peu résistif - moins résistif que le chemin menant vers le circuit fragile. Pour cela, il faut augmenter la taille des diodes. Or, cela augmente la capacité de jonction. Les capacités de jonction contribuent à la baisse de la fréquence de coupure des entrées et des sorties du circuit intégré (Weste & Harris (2011)). Sur notre circuit, les pastilles de connexion pour les courants continus (VDD, GND et BIAS) ont des diodes de protection trois fois plus larges que celles des signaux ($\pm RF$, $\pm LO$, OUT), parce que les premières ne sont pas affectées par les capacités parasites.

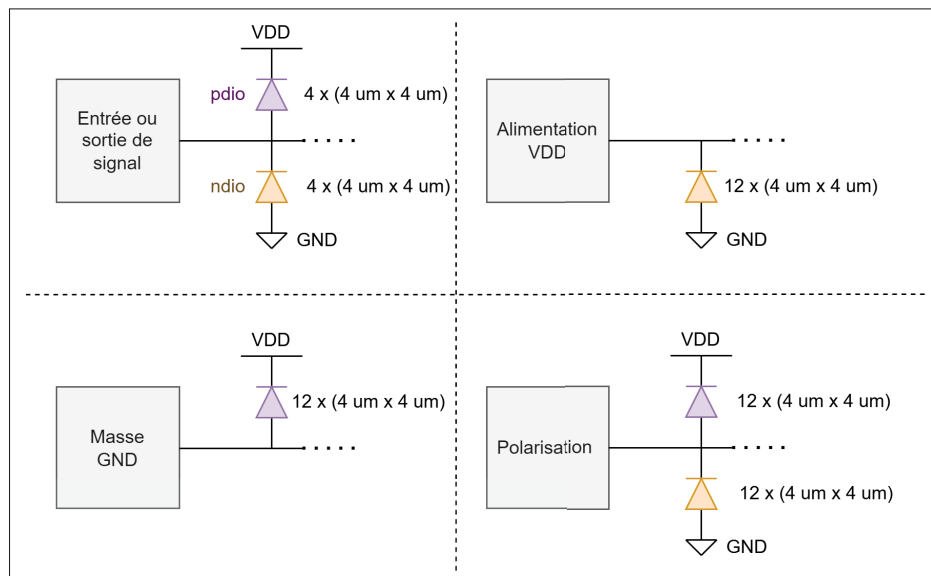


Figure 3.9 Schéma des diodes de protection reliées aux pattes de soudure de la puce

3.1.7 La technique de montage de la puce

Pour le montage de la puce sur le circuit imprimé, deux techniques étaient possibles : le câblage par fil (*wire-bond* en anglais) et la technique de la puce retournée (*flip-chip* en anglais). La conception de la puce et du circuit imprimé dépendait de la méthode choisie, car ces deux techniques permettent différents nombres et positionnements de pastilles de connexion sur la

puce. La partie 3 de l'annexe II compare ces deux techniques plus en détail. Pour notre circuit, nous avons choisi la technique de la puce retournée pour assurer la continuité de la connexion entre les lignes de transmission du circuit imprimé et les pastilles de connexion du circuit intégré.

3.2 La conception du circuit imprimé

Une fois la conception du circuit intégré terminée, nous avons entamé la conception du circuit imprimé, sur lequel la puce a été soudée, grâce au logiciel de CAO *Altium Designer*. La description globale du circuit imprimé est donnée dans la section 1 de l'annexe III. De plus, les différents choix de conception sont décrits en détail dans la section 2 de l'annexe III.

3.2.1 L'adaptation d'impédance pour la partie principale

L'impédance standard des câbles, des circuits et des instruments RF est de $50\ \Omega$. Un circuit avec une impédance différente de cette valeur ne peut pas recevoir ou transmettre toute la puissance à cause de la réflexion des ondes. Or, le circuit intégré ne comporte aucun élément d'adaptation d'impédance. En effet, pour faciliter la conception du circuit intégré, nous avons décidé de n'inclure aucun élément inductif. or, des bobines auraient permis d'annuler la réactance négative des condensateurs et des capacités parasites, pour offrir une impédance réelle. Tout le circuit d'adaptation d'impédance se trouve donc sur le circuit imprimé.

3.2.1.1 Extraction des paramètres S

Le circuit d'adaptation d'impédance, sur le circuit imprimé, a donc été conçu une fois que la conception du circuit intégré fut terminée. D'abord, nous avons extrait les données de la simulation des paramètres S du circuit intégré sous forme de fichier *Touchstone* (.snp) grâce à Cadence Virtuoso. Ensuite, nous avons conçu un circuit d'adaptation d'impédance grâce au logiciel ADS avec les données importées. ADS a été utilisé parce qu'il est beaucoup mieux adapté que Cadence Virtuoso pour simuler des lignes de transmission et pour concevoir des circuits imprimés. La conception dépendait entièrement de la précision des données de simulation de

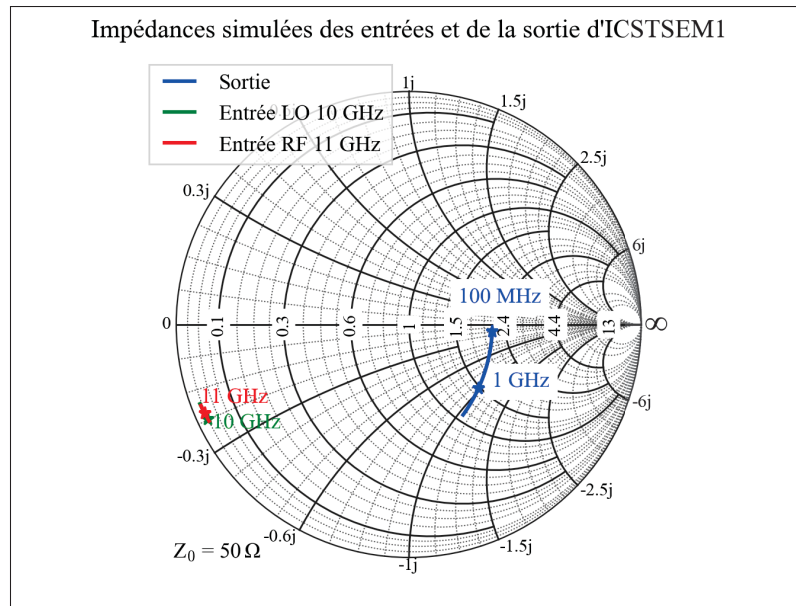


Figure 3.10 Simulation de paramètres S de réflexion des entrées et de la sortie de la puce

Cadence Virtuoso. Or, cette simulation avait été réalisée sans prendre en compte les métaux de remplissage. Une simulation qui prenait en compte les milliers d'éléments carrés métalliques flottants ne pouvait pas être réalisée.

Les coefficients de réflexion du circuit intégré sont représentés sur la figure 3.10. D'après la simulation des paramètres S, les entrées du circuit intégré paraissaient très réfléchissantes. Or, plus une impédance s'éloigne des $50\ \Omega$, plus il est difficile de l'adapter. De plus, selon le critère de Bode-Fanon, la bande passante d'un circuit d'adaptation est d'autant plus limitée que la norme du coefficient de réflexion est grande (p.267 Pozar (2012)). Donc, même avec une adaptation idéale, les entrées du circuit auraient eu une bande passante limitée. Quant à l'impédance du port de sortie, elle était plus proche des $50\ \Omega$ requises. Celle-ci avait une résistance légèrement inférieure à $200\ \Omega$ et une réactance négative.

3.2.1.2 Les lignes de transmission

Les lignes de transmission jouent un rôle clé dans l'adaptation d'impédance grâce à leur impédance caractéristique. Les lignes microruban sont des lignes de transmission simples qui consistent en une trace de cuivre exposée à l'air, d'une couche diélectrique et d'un plan de masse en dessous. L'impédance caractéristique dépend principalement de la largeur de la trace de cuivre, de la distance entre les deux couches de cuivre et de la constante diélectrique. D'autres facteurs, tels que la résistivité de la trace, son épaisseur ou le facteur de dissipation du substrat, affectent les pertes le long de la ligne. Cependant, lors de la conception d'un circuit imprimé, la seule variable que l'on peut véritablement contrôler est la largeur de la trace. Les calculs d'impédance caractéristique des lignes de transmission ont été réalisés grâce au logiciel LineCalc d'ADS.

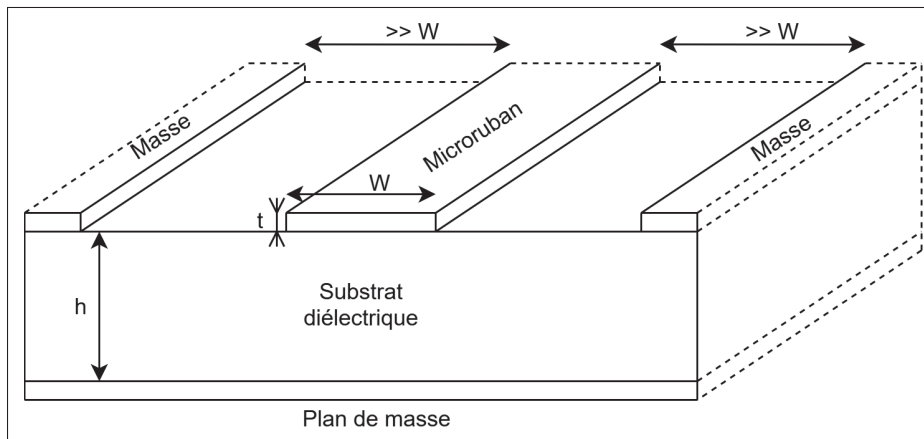


Figure 3.11 Vue en coupe d'une ligne microruban

3.2.1.3 Le circuit d'adaptation pour les entrées

Pour l'adaptation d'impédance des ports d'entrée $\pm$ RF et $\pm$ LO, nous n'avons utilisé que des lignes microrubans. En effet, à de si hautes fréquences, les condensateurs et les bobines discrets entrent en résonance à cause des effets parasites. Nous aurions pu utiliser des bras de réactance

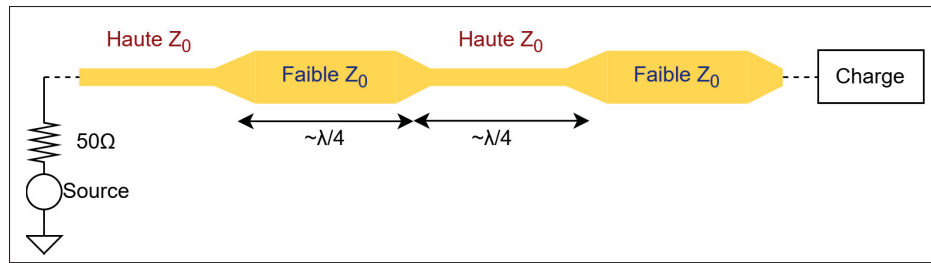


Figure 3.12 Principe de l'adaptation d'impédance avec un ligne microruban à largeurs alternées

(*stub*) qui ajoutent de la susceptance. Cependant, les simulations montraient que la qualité de l'adaptation était très sensible à la moindre variation de la longueur des bras de réactance. Pour ne pas prendre de risque, nous avons donc opté pour des circuits d'adaptation uniquement constitués de lignes en série de différentes largeurs, comme montré dans la figure 3.12. Il y a une alternance entre les sections minces à haute impédance caractéristique ($>50\ \Omega$) et les sections larges de faible impédance ($<50\ \Omega$). Les sections de largeur variable successives ramènent petit à petit l'impédance de la charge vers $50\ \Omega$, comme illustré sur la figure 3.13. Leurs dimensions ont été déterminées grâce à l'outil de CAO.

Cependant, un défaut de conception d'ICSTSEM1 a rendu l'adaptation d'impédance des entrées encore plus compliquée et délicate. On peut remarquer sur la figure 3.1 qu'il n'existe aucune isolation entre les ports $\pm\text{RF}$ et $\pm\text{LO}$. À cause du manque d'isolation, relier quelque chose au port $\pm\text{RF}$ affecte l'adaptation au port $\pm\text{LO}$, et vice versa. Il a donc fallu concevoir des circuits d'adaptation de $\pm\text{RF}$ et de $\pm\text{LO}$ simultanément.

3.2.1.4 Le circuit d'adaptation pour la sortie

Le circuit d'adaptation d'impédance de la sortie d'ICSTSEM1 a été conçu pour une fréquence de 1 GHz, qui était la fréquence IF prévue pour le mélangeur. Bien que la sortie contienne des signaux en bande de base et autour de 1 GHz, il n'était pas possible d'adapter l'impédance à $50\ \Omega$ pour deux fréquences différentes à la fois. Le signal de sortie en bande de base allait inévitablement avoir une adaptation imparfaite.

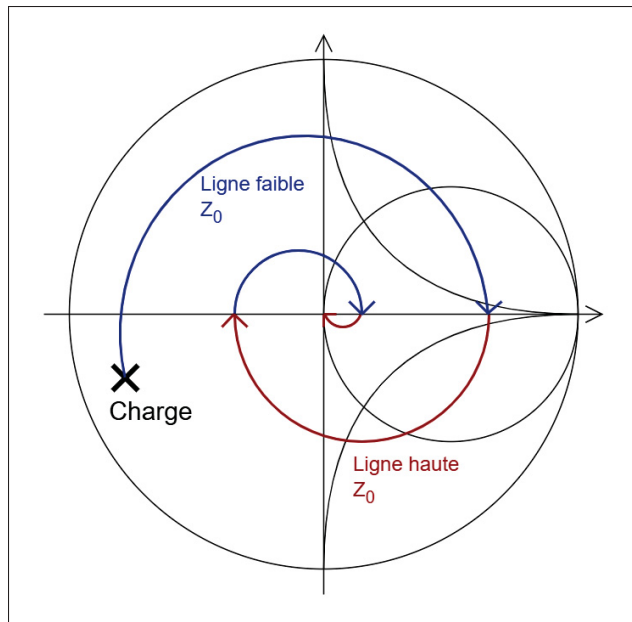


Figure 3.13 Illustration de la technique d'adaptation d'impédance à largeurs de ligne alternant sur une abaque de Smith

L'adaptation de la sortie s'est avérée beaucoup plus simple que celle des entrées, car il n'y avait pas de problématique de couplage. À 1 GHz, nous avons pu nous servir de composants discrets, tels que des condensateurs et des inducteurs. À cette fréquence, la taille des composants discrets est bien inférieure à la longueur d'onde. Donc, ils occupent moins de place qu'un bras de réactance, par exemple. Nous avons placé sur le circuit imprimé une bobine de 7 nH qui suit le condensateur de blocage de courant continu de 10 nF. La bobine est suivie d'une ligne de quart d'onde d'impédance caractéristique légèrement supérieure à 50 Ω .

3.3 Conclusion sur la conception des circuits

Le circuit intégré ICSTSEM1 a été conçu pour fonctionner à la fois en tant que détecteur et en tant que mélangeur. Il est conçu pour recevoir un signal RF à 11 GHz et un signal LO à 10 GHz. À sa sortie, se trouve un signal en bande de base de l'enveloppe de puissance et un signal IF à 1 GHz. Les transistors NMOS, qui ont un rôle central au sein du circuit, ont été conçus pour à la fois maximiser le coefficient de transconductance quadratique g_{m2} et

minimiser celui d'ordre quatre g_{m4} . Les autres éléments du circuit ont été dimensionnés à l'aide du logiciel de CAO pour répondre aux contraintes du gain, de la bande passante, de la réjection des produits d'intermodulation nuisibles et de la simplicité. La puce fait $1\text{ mm} \times 1,25\text{ mm}$. Cependant, la partie active, représentée dans la figure 3.1, occupe une surface de seulement $111,23\text{ }\mu\text{m} \times 83,45\text{ }\mu\text{m}$. Si nous ne comptons pas la surface occupée par les condensateurs, alors le circuit occupe une aire de $24,73\text{ }\mu\text{m} \times 60,695\text{ }\mu\text{m}$.

Le circuit imprimé servait à monter la puce grâce à la technique de la puce retournée. Sur la carte se trouvaient les connecteurs nécessaires pour faire fonctionner le circuit intégré et pour injecter les signaux de deux façons différentes : via des connecteurs SMA ou via des sondes. Puisque le circuit intégré ne comportait aucun élément d'adaptation d'impédance, le circuit imprimé a également servi à adapter les impédances de la puce à $50\text{ }\Omega$.

CHAPITRE 4

SIMULATIONS ET MESURES

Une fois que les circuits intégrés et imprimés furent conçus et assemblés, nous avons effectué les tests et les mesures. Toutes les mesures présentées par la suite ont été effectuées sur la partie principale du circuit imprimé (définie dans la partie 1), sauf les expériences de la section 4.5.1. Quand les mesures n'étaient pas possibles, nous avons utilisé des outils de simulation pour caractériser le circuit ICSTSEM1.

4.1 Les tests en courant continu

Nous avons d'abord effectué les mesures en courant continu, qui sont plus simples à mettre en œuvre et, surtout, qui permettent de vérifier le bon fonctionnement de la puce.

4.1.1 La consommation de courant

Nous avons d'abord mesuré la consommation de courant en fonction de la tension de polarisation au repos (c'est-à-dire sans injection de signal aux entrées). Les résultats de la simulation en courant continu de Cadence Virtuoso et la mesure avec un multimètre sont présentés dans la figure 4.1. D'après la simulation et les mesures, la consommation de courant s'accroît avec la tension de polarisation. Le circuit intégré consomme jusqu'à 6 mA de courant, si $V_p = 1V$. Or, étant donné que le circuit fonctionne le mieux entre 300 mV et 600 mV de polarisation, le circuit consommerait de façon plus réaliste entre environ 430 μA et 3,33 mA de courant.

La mesure du courant avec un multimètre a également servi à vérifier que le circuit intégré est fonctionnel. C'est ainsi que nous avons pu détecter les courts-circuits sous la puce, au niveau des billes de soudure. C'était donc une étape importante du déverminage après l'assemblage d'un circuit intégré sur un circuit imprimé.

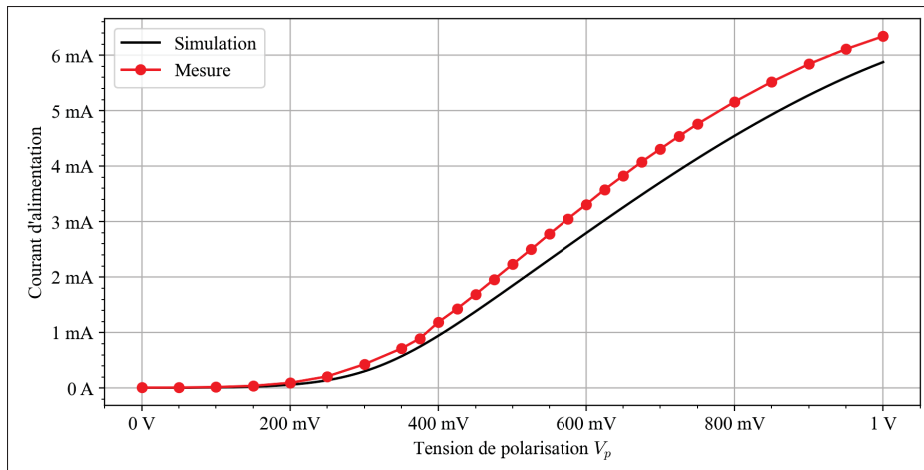


Figure 4.1 Consommation de courant du circuit intégré au repos en fonction de la tension de polarisation

4.1.2 La tension à la sortie au repos

Une autre façon de vérifier le fonctionnement du circuit intégré était la mesure de la tension continue à sa sortie V_r en fonction de la tension de polarisation V_p . Cette mesure a également été réalisée grâce à un multimètre en faisant varier la tension de polarisation sur le générateur de tension continue. De la même façon que la consommation de courant, V_r s'accroît lorsque l'on augmente V_p . Les données de cette mesure, représentées sur la figure 4.2, étaient également utiles plus tard pour évaluer l'augmentation de la tension $|V_s - V_r|$ lorsque l'on stimule le circuit avec un signal à amplitude constante.

4.2 L'adaptation d'impédance de la partie principale

Nous avons ensuite effectué des mesures des paramètres S grâce à un analyseur de réseau vectoriel (ou *vector network analyzer* en anglais) N5222A d'Agilent Technologies à deux ports. Cette mesure nous a notamment permis de vérifier si le circuit d'adaptation d'impédance de la partie principale du circuit imprimé était bien conçu. Trois choses nous intéressaient : l'adaptation des entrées autour de 10 et 11 GHz, l'isolation entre les ports RF et LO et l'adaptation de la sortie entre 0 Hz et 1 GHz.

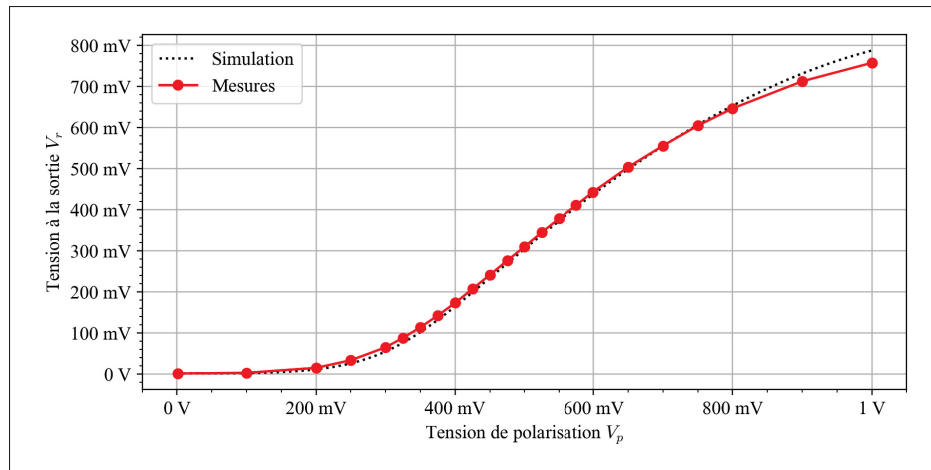


Figure 4.2 Tension à la sortie du circuit intégré au repos en fonction de la tension de polarisation

Les mesures, illustrées dans la figure 4.3 et 4.4, montrent que les taux d'adaptation aux entrées et à la sortie sont différents de ce qui était prévu lors de la conception. L'adaptation d'impédance à la sortie est moins bonne que prévu, mais elle reste centrée autour de 1 GHz. Quant aux ports d'entrée, les différences sont considérables. En effet, le port RF n'est pas adapté à 11 GHz et le port LO n'est pas adapté à 10 GHz. Le port RF semble être le mieux adapté pour des fréquences de 6,7 GHz, 7,72 GHz, 8,94 GHz et 10,5 GHz, tandis que le port LO semble être le mieux adapté pour 6,44 GHz et 9,16 GHz. Par conséquent, pour tester le circuit en tant que détecteur, nous avons utilisé 6,7 GHz ou 7,72 GHz comme fréquences de la porteuse f_c . Pour tester le circuit en mode mélangeur, nous avons dû choisir une fréquence RF de 10,515 GHz f_c et une fréquence LO f_{LO} de 9,163 GHz, ce qui donne une fréquence IF de 1,352 GHz. Ce n'est pas la fréquence IF f_{IF} de 1 GHz initialement prévue, mais cela permettait d'avoir des pertes de réflexion relativement basses sur les différents ports.

Si l'adaptation des impédances ne s'est pas matérialisée comme prévu, c'est probablement parce que les impédances de la puce réelle différaient des données extraites par simulation, que l'impédance du *balun* n'était pas exactement de 50 Ω pour toutes les fréquences, et que le circuit imprimé n'était pas parfaitement conçu. Cependant, l'adaptation d'impédance du circuit imprimé était tout de même en partie efficace pour réduire les réflexions à certaines fréquences.

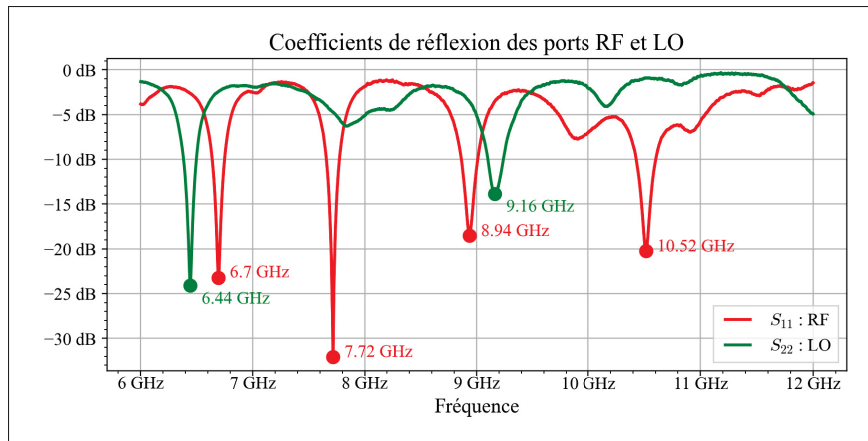


Figure 4.3 Mesure des coefficients de réflexion aux entrées de la partie principale du circuit imprimé

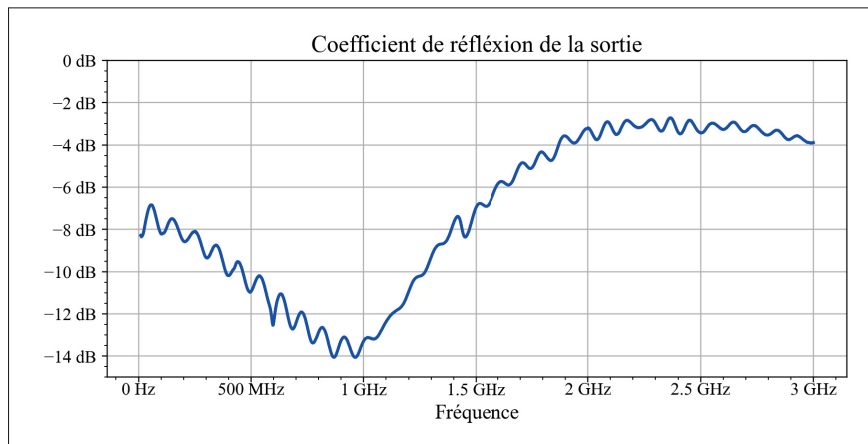


Figure 4.4 Mesure du coefficient de réflexion à la sortie de la partie principale du circuit imprimé

Enfin, nous avons déterminé le taux d'isolation entre les ports RF et LO en mesurant les taux de transfert S_{12} et S_{21} . Les données sont représentées sur la figure 4.5. Selon les mesures, l'isolation semble être d'au moins 20 dB pour des fréquences supérieures à 7 GHz. Si l'isolation semble bonne, malgré l'absence totale de moyens d'isolation au sein du circuit intégré, c'est grâce aux *baluns* reliés aux ports RF et LO.

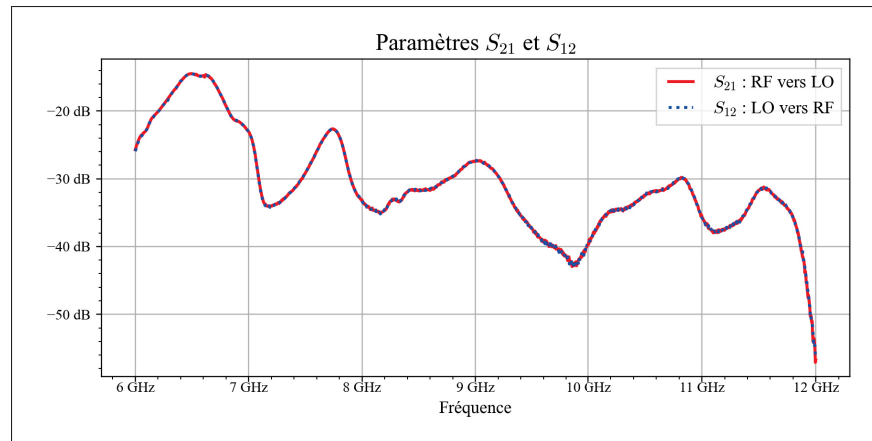


Figure 4.5 Coefficients de transmission entre les ports RF et LO

4.3 Les tests en mode détecteur

Une fois que nous avons appris à quelles fréquences le circuit fonctionne le mieux, nous l'avons d'abord testé en tant que détecteur. Dans ces tests, nous avons seulement injecté un signal RF à l'entrée.

4.3.1 La détection d'un signal à amplitude constante

De la même façon que Li *et al.* (2019), nous avons mesuré la tension continue à la sortie V_s grâce à un voltmètre, tout en faisant varier la puissance du signal RF P_{RF} ¹ à amplitude constante et la tension de polarisation V_p . Le banc de test utilisé est représenté sur la figure 4.6. La fréquence du signal f_c était d'environ 7,7 GHz.

Tout d'abord, nous avons mesuré V_s en fonction de V_p pour une puissance de signal P_{RF} fixe de -5 dBm. Nous avons pu déduire la relation entre la responsivité et la tension de polarisation. La figure 4.7 montre que la responsivité mesurée n'est égale qu'à environ la moitié de celle simulée. Il y a donc environ 3 dB de pertes de signal entre le générateur et la puce. Ces pertes se

¹ Ici et dans toutes les mesures qui suivent, pour les puissances des signaux d'entrée, nous avons enregistré la puissance générée par le générateur de signal. Cette puissance est différente de celle reçue par le circuit intégré, que nous n'avons pas pu mesurer.

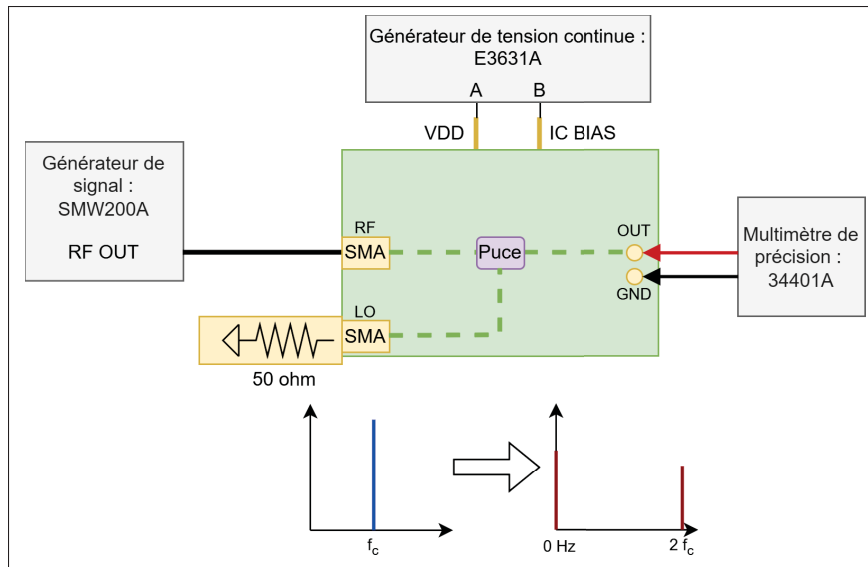


Figure 4.6 Banc d'essai du détecteur avec un signal sinusoïdal

trouvent en partie au niveau du câble reliant le générateur et le circuit imprimé, mais elles sont également dues à la mauvaise adaptation du circuit intégré. Pour une tension de polarisation de 340 mV, nous avons simulé une responsivité de 65,6 V/W et nous avons mesuré une responsivité de 31,6 V/W. Au-delà de cette tension, les transistors commencent à entrer en saturation et la responsivité décroît donc rapidement.

Par la suite, nous avons mesuré la tension continue à la sortie V_s en fonction de la puissance du signal P_{RF} , comme illustré dans la figure 4.8. Or, V_s ne constitue pas en soi une information interprétable ; nous nous intéressons à la différence de tension continue avec et sans signal à l'entrée $V_s - V_r$, telle qu'illustrée dans la figure 4.9. Ces mesures nous ont permis d'obtenir des résultats importants sur la responsivité, le point de compression $IP_{1\text{ dB}}$ et la plage dynamique, tels qu'ils sont enregistrés dans le tableau 4.1.

D'abord, nous remarquons un décalage constant entre les courbes des mesures et celles des simulations causé par les pertes par réflexion qui n'ont pas été simulées. À cause de ces pertes, il a fallu envoyer un signal plus puissant pour obtenir la même réponse du circuit intégré que dans les simulations. Bien que les valeurs mesurées aient été différentes de celles simulées, le

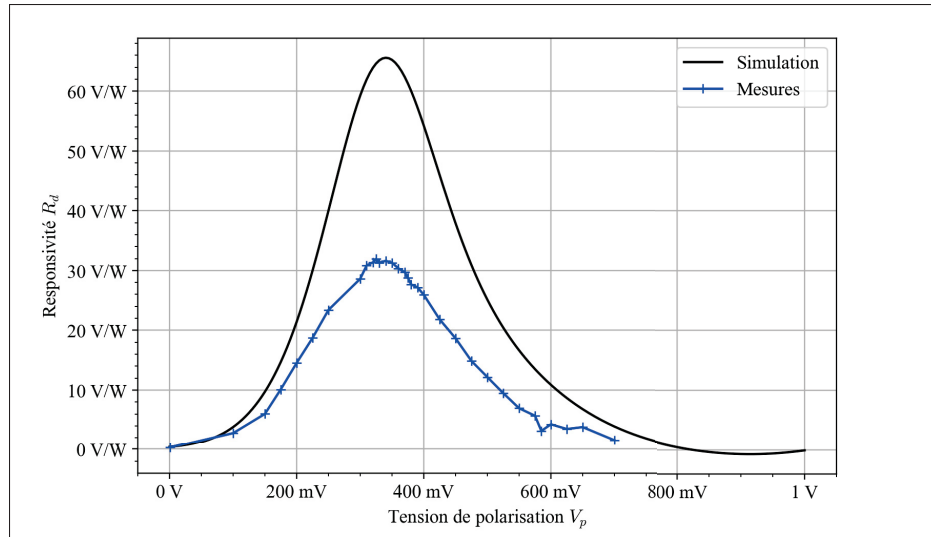


Figure 4.7 Responsivités du détecteur simulée et mesurée

comportement du circuit reste qualitativement le même. Ensuite, nous tenons à préciser que si les points de mesure de la figure 4.9 à basse puissance semblent s'écarter de la pente attendue, c'est parce que le bruit faisait constamment varier la tension lue sur le voltmètre de quelques dizaines de millivolts. Cela affectait la précision de notre estimation de la borne inférieure de la plage dynamique. Nos mesures nous indiquaient tout de même une plage dynamique atteignant ou dépassant les vingt décibels. De plus, elles nous ont permis de comparer la performance de notre circuit à celle des autres travaux. Enfin, nous avons pu observer le phénomène de compression. Le taux d'accroissement de $V_s - V_r$ diminuait à haute puissance P_{RF} . Nous remarquons la tendance selon laquelle, plus la tension de polarisation est élevée, plus le point de compression est repoussé vers de plus hautes puissances. D'après la théorie, le point de compression dépend du rapport entre α_2 et α_4 de la fonction de transfert. Or, nous avons vu dans la partie 3.1.3.1 que, pour V_p allant de 490 mV à 575 mV, le ratio g_{m4}/g_{m2} des NMOS est le plus bas. Dès lors, on peut penser que α_2/α_4 du circuit atteint son maximum sur cette plage de tension. Les résultats paraissent donc cohérents avec la théorie.

Pour conclure, en fonction de la tension de polarisation, le circuit, en tant que détecteur de puissance, peut avoir une responsivité de 31,3 V/W, voire le double d'après les simulations,

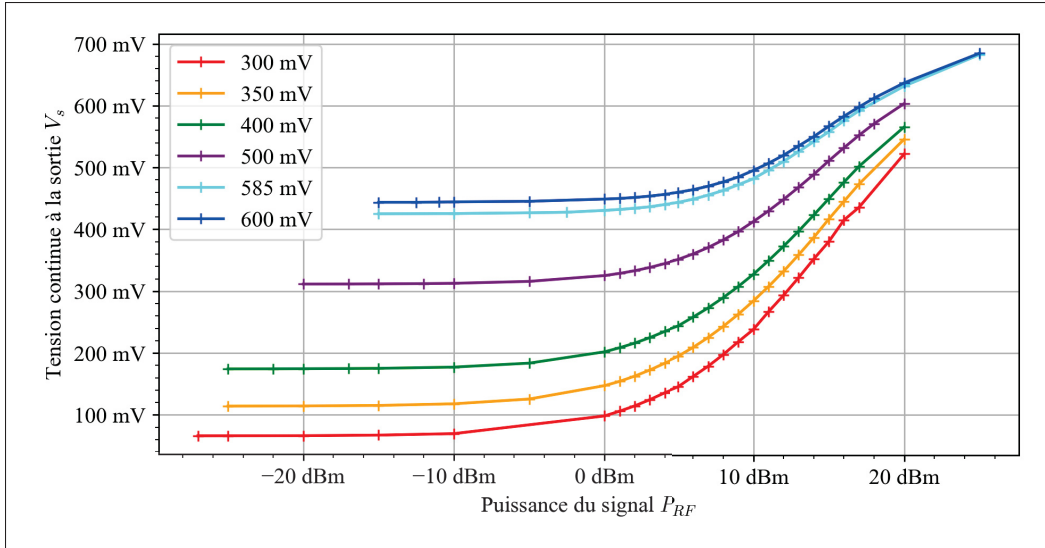


Figure 4.8 V_s mesurée en fonction de la puissance du signal à amplitude constante à l'entrée

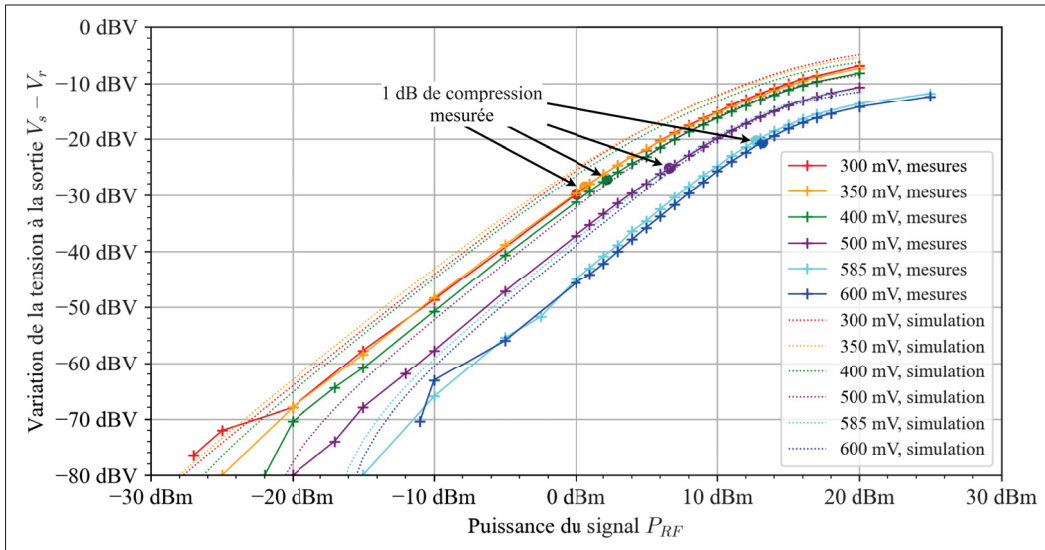


Figure 4.9 $V_s - V_r$ en fonction de la puissance du signal à amplitude constante à l'entrée

un point de compression $IP_{1\text{ dB}}$ allant jusqu'à 13,15 dBm ou une plage dynamique s'étendant jusqu'à 27,72 dB.

Tableau 4.1 Résultats des mesures sur le détecteur avec un signal à amplitude constante à l'entrée

V_p [mV]	R_d [V/W] (mesurée ; simulée)	IP_1 dB [dBm]	PD [dB]
300	(30 ; 59,34)	0,01	[20,01 ; 27,01]
350	(31,3 ; 65,2)	0,67	25,67
400	(25,92 ; 54,28)	2,25	[22,25 ; 24,25]
500	(12,18 ; 25,14)	6,63	26,63
585	(4,28 ; 12,42)	12,72	27,72
600	(3,2 ; 10,91)	13,15	[18,15 ; 24,15]

4.3.2 La détection d'un signal à deux fréquences

Afin de prouver que notre circuit peut fonctionner comme détecteur d'enveloppe de puissance et afin d'observer la présence de produits d'intermodulation, nous avons ensuite injecté un signal à deux fréquences dans le circuit à travers le port RF. Lors des mesures, la fréquence de la porteuse f_c était de 7,717 GHz et f_m était de 50 MHz. Nous nous attendions donc à des produits d'intermodulation à $2f_m = 100$ MHz et à $4f_m = 200$ MHz. Le premier est l'enveloppe de puissance du signal et le second est un produit d'intermodulation d'ordre quatre nuisible. Le banc d'essai utilisé est représenté sur la figure 4.10.

Nous avons étudié la fonction de transfert du circuit en mesurant la puissance sortante à $2f_m$ et à $4f_m$ en fonction de la puissance du signal P_{RF} . Nous avons répété les mesures pour différentes tensions de polarisation. Comme nous pouvons le voir sur la figure 4.11, la puissance sortante à $2f_m$ augmente à un taux de 20 dB par décade : cela confirme bel et bien la nature quadratique de la fonction de transfert. Comme vu précédemment, le gain est meilleur pour une tension de polarisation proche de 340 mV ; au-delà de cette tension, il diminue. Une fois de plus, le changement du point de compression IP_1 dB suit la même tendance que lors des mesures précédentes. En plus de la compression, nous avons pu observer l'apparition du produit d'intermodulation d'ordre quatre à $4f_m$. Ce produit augmente à un taux de 40 dB par décade : il est proportionnel à P_{RF}^4 , comme prévu par la théorie. Pour caractériser le rapport entre le produit d'intermodulation d'ordre deux et celui d'ordre quatre, nous avons mesuré les points

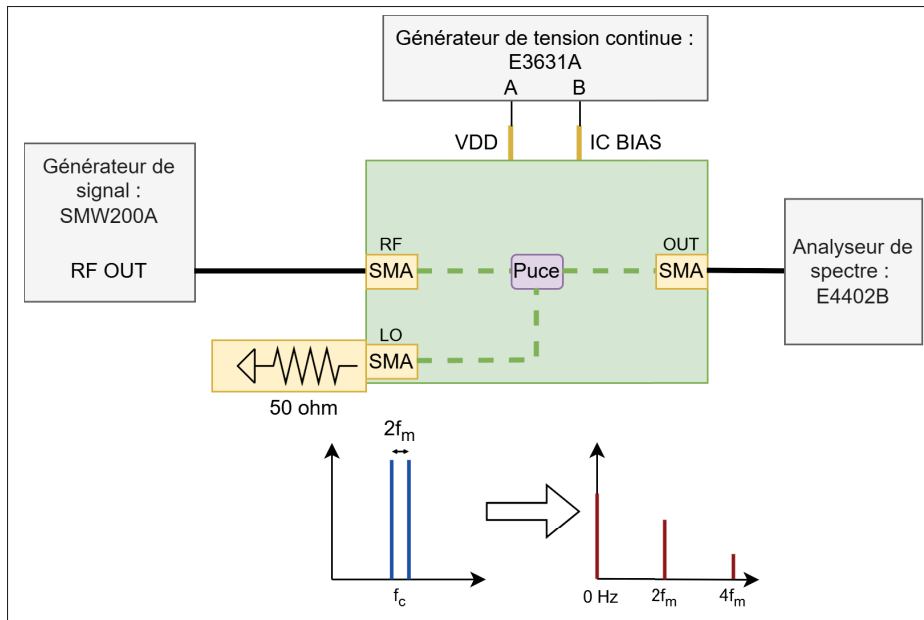


Figure 4.10 Taux d'adaptation d'impédance aux entrées de la partie principale du circuit imprimé

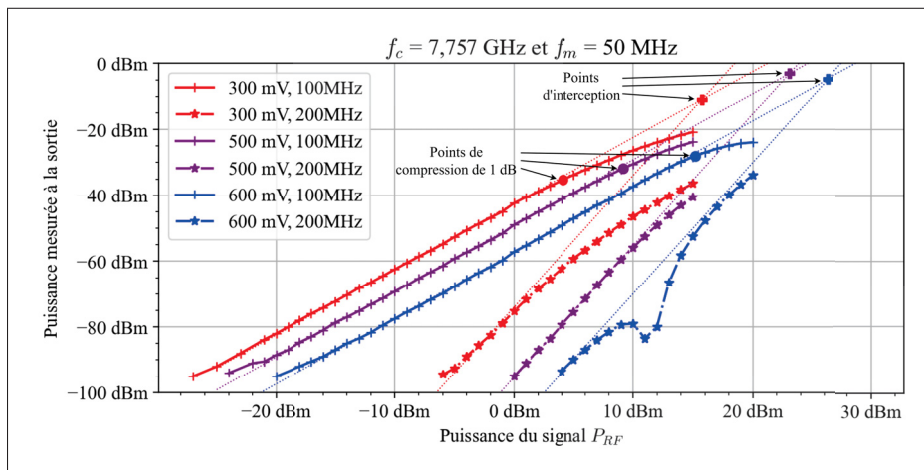


Figure 4.11 Puissances mesurées à $2f_m$ et $4f_m$ à la sortie du détecteur en fonction de P_{RF} et de V_p

d'interception, de la même façon qu'avec un mélangeur de fréquence. Les points d'interception semblaient se trouver à des puissances plus importantes lorsque nous augmentions V_p .

Nous avons recueilli les valeurs des points de compression $IP_{1\text{ dB}}$, des points d'interception IP_{IP} et de leurs ratios $IP_{IP}/IP_{1\text{ dB}}$ pour différentes tensions de polarisation V_p dans le tableau 4.2. En premier lieu, nous remarquons que les valeurs des $IP_{1\text{ dB}}$ dans ce tableau sont différentes de celles du tableau 4.1. Il semble effectivement y avoir un décalage de deux à quatre décibels, alors que, selon nos formules de (A I-19) et (A I-26), ces valeurs sont censées être identiques. Nous pouvons en partie attribuer ces écarts aux erreurs de mesure et aux pertes par réflexion entre la sortie de la puce et l'analyseur de spectre. Ensuite, nous nous apercevons que le ratio $IP_{IP}/IP_{1\text{ dB}}$ reste à peu près constant, à quelques décibels près. Cependant, nous pouvons noter que les ratios mesurés sont inférieurs de quelques décibels au ratio de 15,65 dB estimé par calcul dans la partie 2.6.2.4. Si les valeurs mesurées ne correspondent pas exactement à notre calcul, c'est parce que notre calcul théorique ne tenait pas compte de l'influence des intermodulations des ordres supérieurs à quatre. Néanmoins, malgré les imprécisions, nous pouvons conclure que notre modèle théorique demeure quelque peu cohérent avec la réalité.

Tableau 4.2 Points de compression, points d'interception pour un détecteur et leurs rapport selon V_p

V_p [mV]	$IP_{1\text{ dB}}$ [dBm]	IP_{IP} [dBm]	$IP_{IP}/IP_{1\text{ dB}}$ [dB]
300	4,1	15,75	11,65
350	3,7	16,0	12,3
400	4,3	17,23	12,93
500	9,2	23,07	13,87
600	15,2	26,3	11,1

Nous avons également tracé le rapport de la puissance de sortie à $2f_m$, $P_{s@2f_m}$, sur celle à $4f_m$, $P_{s@4f_m}$. Ce rapport, illustré dans la figure 4.12, indique le taux de réjection du produit d'intermodulation d'ordre quatre. Généralement, plus P_{RF} est importante, plus ce taux de réjection est faible, parce que $P_{s@4f_m}$ s'accroît avec un taux supérieur à celui de $P_{s@2f_m}$. De plus, il apparaît que ce taux s'améliore avec une tension de polarisation s'approchant de 585 mV. Pour $V_p = 585$ mV et $V_p = 600$ mV, ce taux de réjection vaut au moins 40 dB sur une grande partie de la plage dynamique ; naturellement, il se dégrade lorsque le circuit s'approche de son point de compression. Nous observons tout de même un phénomène qui ne peut pas être expliqué par notre modèle théorique : pour les tensions de polarisation de 585 mV et 600 mV, $P_{s@4f_m}$

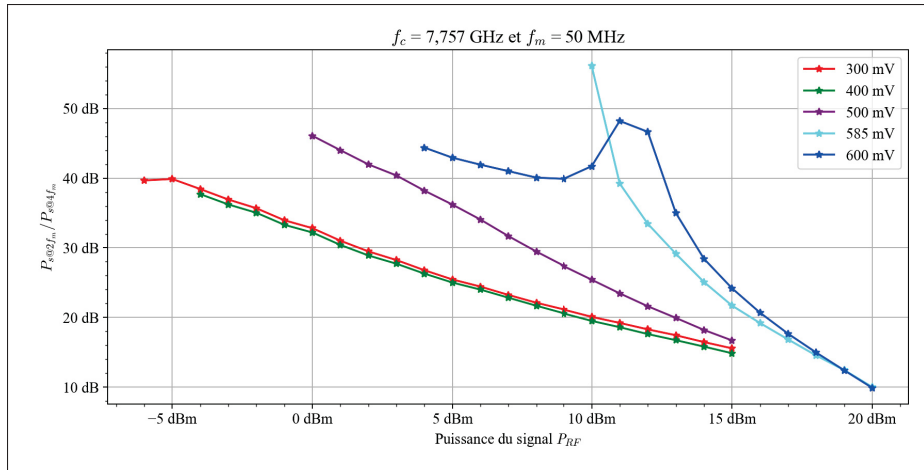


Figure 4.12 Le taux de réjection du produit d'intermodulation d'ordre quatre en fonction de P_{RF}

connaît une chute lorsque $P_{RF} \approx 10$ dBm. Le taux de réjection $P_{s@2f_m}/P_{s@4f_m}$ connaît alors un pic. Ce phénomène a également été observé lors des simulations. Notre observation n'est donc pas une erreur de mesure. Lorsque nous avons mesuré $P_{s@2f_m}$ et $P_{s@4f_m}$ en fonction de V_p pour une puissance P_{RF} fixée à 10 dBm, nous observons de nouveau ce phénomène, comme le montre la figure 4.13. Nous pouvons conjecturer que la distorsion de phase intervient sous ces conditions particulières pour créer une interférence destructive qui réduit l'amplitude du produit d'intermodulation à $4f_m$. La distorsion de phase n'est pas un phénomène pris en compte par notre modèle non-linéaire.

4.3.3 Les simulations transitoires et les mesures à l'oscilloscope

Les simulations transitoires et les mesures à l'oscilloscope nous ont permis non seulement de visualiser le fonctionnement du détecteur, mais également de mesurer la bande passante à la sortie du circuit. Le banc d'essai utilisé est représenté sur la figure 4.14. La sortie de bande de base du générateur, produisant un signal modulant le signal RF, est directement branchée à un canal de l'oscilloscope. La sortie du détecteur est branchée à un autre canal de l'oscilloscope. Nous avons donc pu comparer le signal modulant à l'enveloppe de puissance du détecteur. Aux

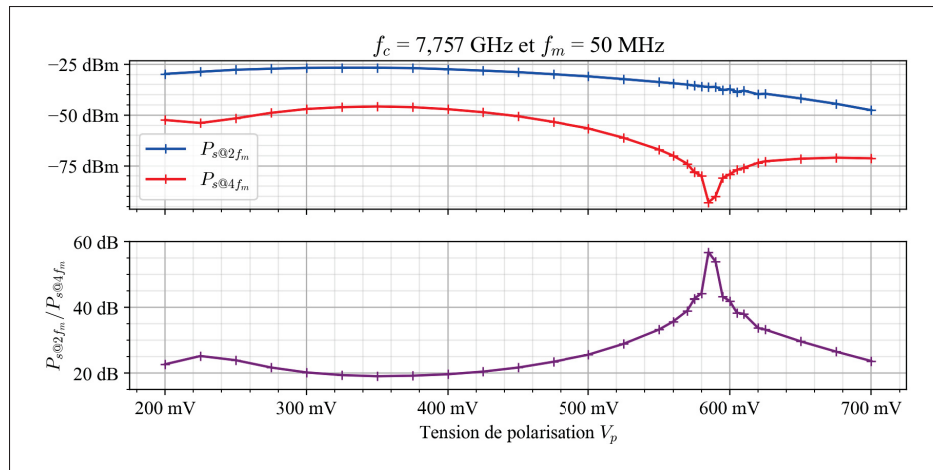


Figure 4.13 $P_{s@2f_m}$ et $P_{s@4f_m}$ en fonction de la tension de polarisation V_p

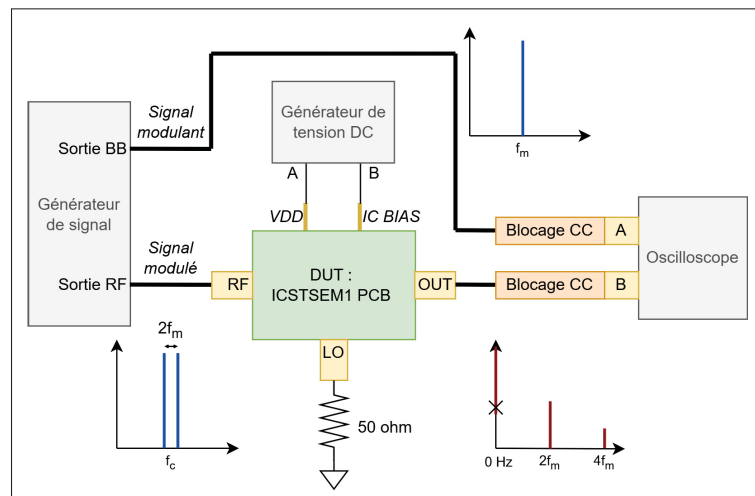


Figure 4.14 Banc d'essai du détecteur branché à l'oscilloscope

entrées de l'oscilloscope, des circuits de blocage de courant continu étaient également branchés. Ils servaient à protéger l'instrument de mesure sensible.

4.3.3.1 Les mesures à l'oscilloscope de l'enveloppe de puissance d'un signal à deux fréquences

Lors d'un premier test avec l'oscilloscope, nous avons de nouveau injecté un signal à deux fréquences dans le détecteur. Le détecteur reçoit donc un signal porté à $f_c = 6,667$ GHz, composé de deux fréquences espacées de $2f_m = 100$ MHz.

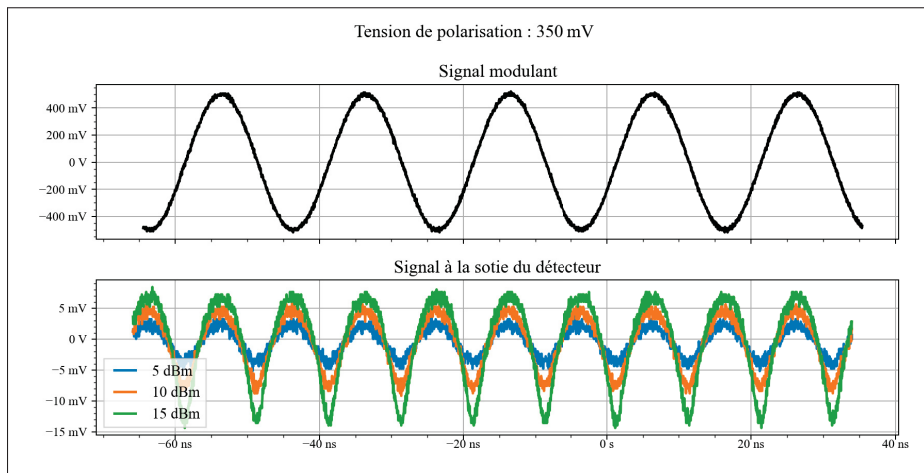


Figure 4.15 Comparaison à l'oscilloscope du signal modulant et de la sortie du détecteur

Sur la figure 4.15, nous pouvons voir le signal sinusoïdal modulant, ainsi que le signal de l'enveloppe de puissance, qui sort du détecteur. À cause du circuit de blocage de courant continu, les signaux sortants du détecteur ont une tension moyenne de 0 V. En effet, il manque la tension continue de la sortie du détecteur V_s qui dépend de la polarisation et de la puissance du signal injecté. Nous pouvons tout de même analyser qualitativement le fonctionnement du circuit. L'enveloppe de puissance du modulant atteint sa tension minimale à l'instant où la tension du signal modulant vaut 0 V, et elle est maximale lorsque le modulant atteint sa crête maximale ou minimale. Par conséquent, la fréquence de l'enveloppe de puissance est le double de celle du signal modulant. Cela est cohérent avec le comportement quadratique du circuit. De même que l'amplitude de l'enveloppe de puissance augmente avec la puissance réglée au générateur, le niveau de distorsion augmente. Nous pouvons constater visuellement l'effet de distorsion sur

l'enveloppe de puissance pour une puissance de 15 dBm. En effet, alors que le signal à la sortie du détecteur est censé être purement sinusoïdal, ce dernier comporte en réalité des harmoniques.

4.3.3.2 La simulation et la mesure de la réactivité du détecteur

Il existe deux façons de tester la bande passante de sortie d'un détecteur. La première méthode consiste à injecter un signal à deux fréquences, puis à écarter les fréquences (*i.e.* augmenter f_m) jusqu'à ce que le signal à $2f_m$ à la sortie atteigne la fréquence de coupure. C'est ce que Bourgault (2023) a simulé pour son détecteur. Mais, il n'a pas pu réaliser ce test en laboratoire. C'est probablement parce que le générateur de signaux ne pouvait pas produire un signal à deux fréquences avec un écartement suffisamment important. La seconde façon de tester la réactivité du détecteur est de procéder comme Serhan *et al.* (2015a). Dans cet ouvrage, les auteurs ont stimulé le détecteur avec un signal RF dont l'amplitude était modulée par un signal carré. Le signal sortant du détecteur, avec sa réactivité limitée, connaît une durée de transition entre les deux niveaux d'amplitude. Comme pour caractériser un filtre RC, on peut mesurer la constante de temps τ , qui est la durée de transition caractéristique. Cette constante est inversement proportionnelle à la bande passante.

Dans un premier temps, nous avons mesuré le temps que le signal à la sortie du détecteur met à réagir à l'apparition soudaine d'un signal RF en simulation transitoire, avant de réaliser l'expérience en laboratoire. Comme illustré dans la figure 4.16, τ correspond au temps que met le signal à atteindre 63,2% de sa valeur finale. Cette valeur varie avec la polarisation. Les valeurs de τ et de la fréquence de coupure en fonction de la polarisation simulée de notre circuit apparaissent sur la figure 4.16. En effet, selon Bourgault (2023), un détecteur dont les transistors sont en saturation est plus réactif grâce à la mobilité supérieure des porteurs de charge dans cette région d'opération. Le temps de réaction diminue en augmentant la tension de polarisation. Entre 200 mV et 600 mV de polarisation, la fréquence de coupure varie de 148 MHz à 1,18 GHz. Étant donné que la responsivité décroît au-delà d'une tension de polarisation de 350 mV, où la bande passante est de 491 MHz, si l'on souhaite un détecteur encore plus réactif, il faut

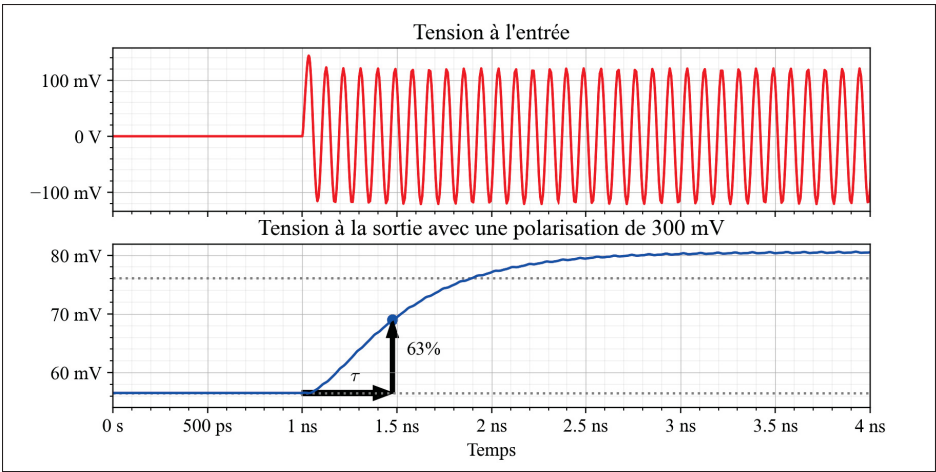


Figure 4.16 Simulation transitoire de la réactivité du détecteur

Tableau 4.3 La bande passante simulée en fonction de la tension de polarisation

V_p [mV]	350	400	450	500	550	600
Bande passante	491 MHz	663 MHz	850 MHz	978 MHz	1,10 GHz	1,18 GHz

sacrifier de la responsivité. Le choix de la tension de polarisation pour le détecteur est donc en partie un compromis entre la bande passante de sortie et la responsivité.

En ce qui concerne le circuit en tant que mélangeur, le choix de la tension de polarisation s’avère plus complexe. En effet, à 350 mV de polarisation, le gain en tension quadratique est maximal, mais la fréquence de coupure est limitée. En augmentant la tension de polarisation, on peut augmenter la fréquence de coupure, mais le gain en tension quadratique diminue. Donc, en fonction de la fréquence IF désirée, une tension de polarisation proche de 340 mV ne garantit pas un gain de conversion maximal. Ajouté à tout cela, la tension de polarisation fait varier le taux de réjection des produits d’intermodulation néfastes. Par conséquent, la tension de polarisation optimale, qui répond aux critères du gain de conversion et de la linéarité, varie avec la fréquence IF désirée.

Nous avons tenté de répliquer expérimentalement cette simulation. Le générateur de signal a généré en interne un signal modulant carré d’un mégahertz qui alterne entre 0 V et 0,5 V : à 0 V,

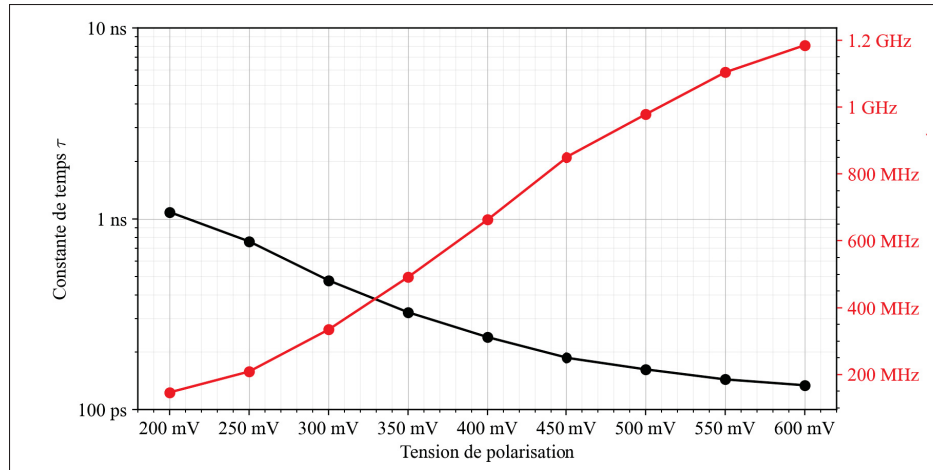


Figure 4.17 Réactivité simulée du détecteur en fonction de la tension de polarisation

le signal modulé a une amplitude nulle et, à 0,5 V, le signal modulé atteint l'amplitude maximale. Cependant, par rapport à la simulation, ces essais furent imparfaits. En effet, contrairement aux simulations, dans la réalité, le signal carré modulant connaît également un temps de transition non nul entre les deux états. De plus, la forme du signal modulant mesurée à l'oscilloscope était probablement affectée par le câble et le circuit de blocage de courant continu. Donc, nous ne pouvions pas être entièrement certain de son temps de transition. Enfin, le circuit imprimé a probablement également influencé la performance du circuit intégré.

Pour mesurer la constante τ grâce aux mesures, nous avons évalué le temps que le signal mettait pour monter de 20% de la tension finale, à l'instant t_a , à 80% de la tension finale, à l'instant t_b . Nous avons choisi ces pourcentages pour éviter que le bruit n'impacte notre mesure de τ . Puisque l'on utilise toujours un modèle de circuit RC, la relation entre τ et les instants de franchissement des seuils, t_a et t_b , est la suivante :

$$\tau = \frac{t_b - t_a}{-\ln(1 - 0.8) + \ln(1 - 0.2)} \quad (4.1)$$

Certains résultats de la mesure sont représentés dans la figure 4.18. Pour des tensions de polarisation de 350 mV et de 600 mV, nous avons trouvé respectivement des valeurs de τ de

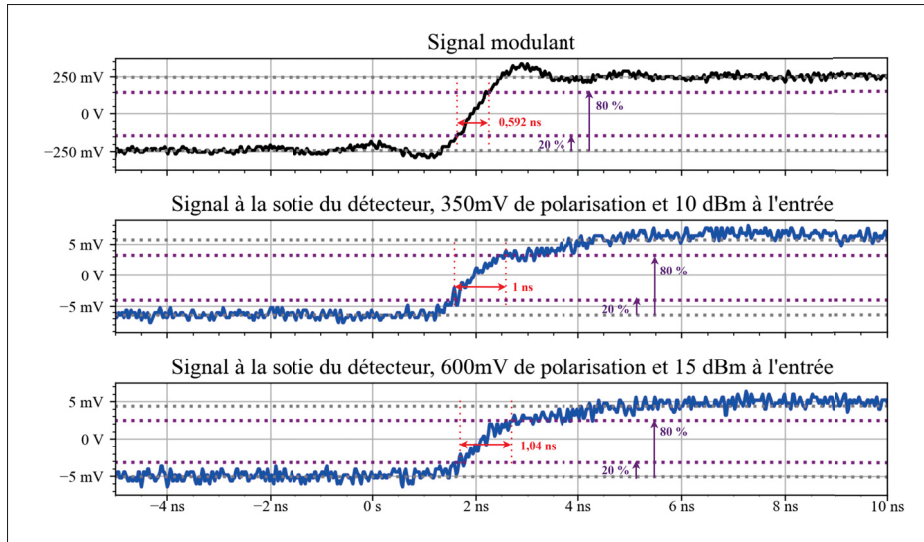


Figure 4.18 Mesure à l'oscilloscope de la réactivité du détecteur

721 ps et de 750 ps, correspondant respectivement aux fréquences de coupure de 221 MHz et de 212 MHz. Même pour des tensions de polarisation très différentes, nous avons obtenu presque le même résultat. Nous pouvons alors déduire qu'ici, contrairement à la simulation, la fréquence de coupure est limitée par un facteur externe à la puce. Les résultats des mesures sont peu concluants. Cependant, ils prouvent que la fréquence de coupure de la puce est **au moins** de 212 MHz.

4.4 Simulations de la réjection de la porteuse

Bourgault (2023) définit le taux de réjection de la porteuse comme le rapport de la puissance du signal RF à l'entrée P_{RF} sur la puissance à cette même fréquence à la sortie $P_{s@f_{RF}}$ ². La topologie du circuit utilisée rejette intrinsèquement les fréquences du signal RF injecté à l'entrée. Par conséquent, idéalement, aucun signal ne doit rester dans la bande de fréquences RF à la sortie du détecteur. Si cette réjection est parfaite en théorie, tel n'est pas le cas en pratique. Dans la section 5 de l'annexe I, nous avons démontré mathématiquement comment des asymétries du signal et celles du circuit causent un rejet imparfait de la porteuse.

² Il l'appelle ainsi, bien que le circuit ne rejette pas seulement la fréquence singulière f_c . Une appellation plus précise serait le «taux de réjection de la partie impaire».

Comme démontré dans l'annexe, un taux de réjection de la porteuse n'est pas nécessaire pour maximiser la responsivité ou le gain de conversion. Ces tensions peuvent ne pas être déphasées du tout et le circuit fonctionnerait quand même comme détecteur et mélangeur. Le circuit est donc résilient vis-à-vis de l'erreur de phase. Néanmoins, le rejet de la porteuse demeure important pour plusieurs raisons. Premièrement, si le détecteur est utilisé dans une boucle de rétroaction, comme dans le projet de Bourgault (2023), alors un signal à haute fréquence rémanent à la sortie du détecteur peut déstabiliser le reste du système. Ensuite, dans notre cas, si M1 et M2 ne rejettent pas la porteuse, alors M3 et M4, qui admettent toujours une part de non-linéarité, peuvent générer à leur tour des produits d'intermodulation néfastes et inattendus. Cela est d'autant moins désirable lorsque notre circuit fonctionne en tant que mélangeur. Enfin, si les signaux RF et LO entrants sont éliminés dans le circuit, alors le taux d'isolation entre les entrées et la sortie du circuit est meilleur.

Grâce aux outils de simulation, nous étudierons le lien entre le taux de réjection de la porteuse et la symétrie du signal, dans un premier temps, et la symétrie du circuit, dans un second temps.

4.4.1 La réjection de la porteuse en fonction du déphasage

La symétrie des signaux aux grilles de M1 et de M2 contribue au rejet du signal RF ; pour cela, les signaux au niveau des grilles doivent être déphasés de 180° avec la même amplitude. Posons ϕ l'erreur de phase, tel que si $\phi = 0$ alors les signaux entrants sont parfaitement déphasés, et si $\phi = 180^\circ$, alors les deux signaux sont en phase. Selon la théorie, $P_{s@f_c}$ est proportionnelle à $\cos(\phi)$. Donc, si $\phi = 0$, la puissance de la porteuse à la sortie est absolument nulle, et si $\phi = 180^\circ$, sa puissance est maximale. La figure 4.19 représente la puissance de la porteuse à la sortie en fonction de ϕ .

Assez logiquement, la puissance de la porteuse à la sortie devient infiniment faible avec ϕ qui tend vers 0° . À 1° d'erreur de phase, la puissance du signal à l'entrée est 92 dB plus puissante que celle à la sortie à la même fréquence. Une partie de ce rejet se fait grâce à la symétrie des signaux aux grilles des NMOS, et une autre partie est due aux pertes de conversion

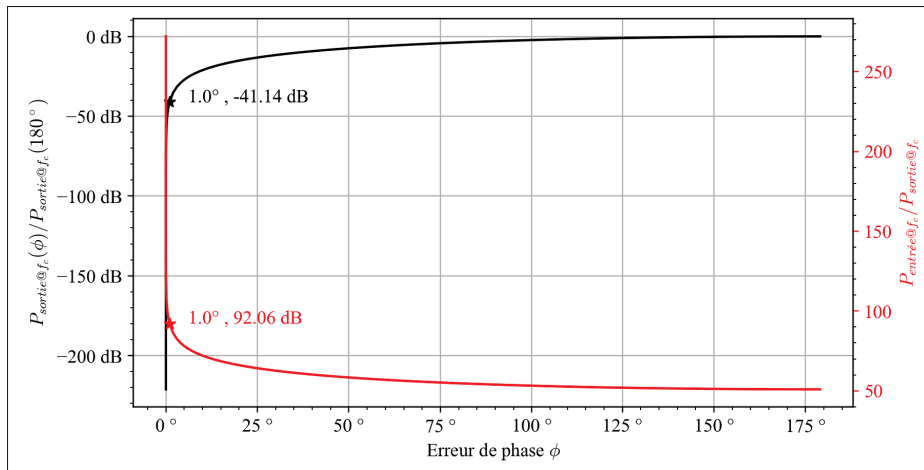


Figure 4.19 Simulation de la réjection de la porteuse en fonction de l'erreur de phase ϕ

inhérentes au circuit. (En effet, le gain en tension linéaire du circuit α_1 est inférieur à 1.) Pour évaluer la contribution au rejet de la porteuse par la symétrie, nous pouvons évaluer le ratio de $P_{\text{sortie}@f_c}(\phi)$ sur $P_{\text{sortie}@f_c}(\phi = 180)$. Pour $\phi = 1$, ce ratio vaut -41,1 dB.

4.4.2 La réjection de la porteuse en fonction de la symétrie du circuit intégré

Tout le circuit ICSTSEM1 a été conçu de façon symétrique, avec M1 et M2 côte à côte au centre de la puce, car, selon la partie 5 de l'annexe I, le rejet de la porteuse dépend également de l'identité entre M1 et M2. Dans la littérature, on appelle les différences involontaires entre des composants jumeaux le *mismatch*. Lors du placement des composants sur la puce, nous avons donc mis toutes les chances de notre côté pour éviter les asymétries dues à des défauts de fabrication. Il aurait également été possible d'optimiser les dimensions de M1 et de M2 afin de minimiser le *mismatch* (Lovett, Welten, Mathewson & Mason (1998)). Cependant, cela aurait inévitablement affecté d'autres caractéristiques plus importantes du circuit.

Grâce à une simulation de Monte Carlo, où les caractéristiques des composants de 1 000 échantillons du circuit ont été variées aléatoirement selon les modèles statistiques du fabricant, nous avons pu déterminer l'influence des défauts de fabrication sur le taux de réjection de la

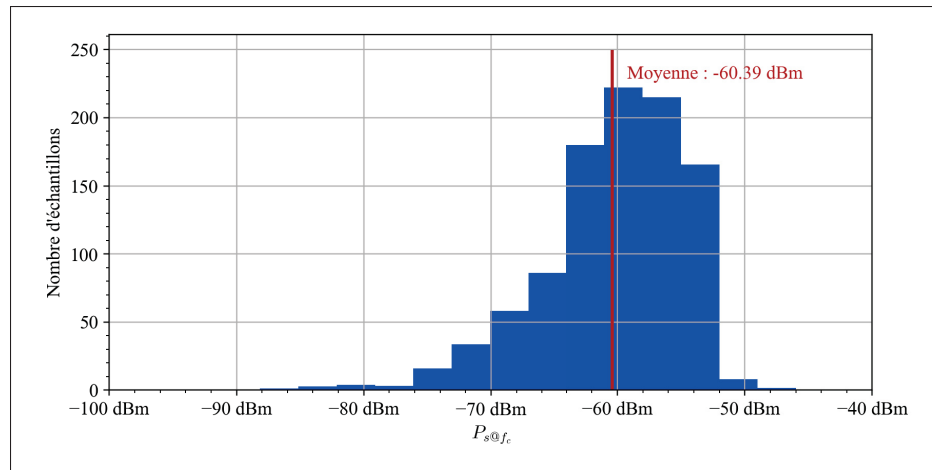


Figure 4.20 Simulation de Monte Carlo de la réjection de la porteuse

porteuse. Dans la simulation, nous avons injecté un signal RF sinusoïdal à $f_c=11$ GHz de 0 dBm parfaitement symétrique aux entrées. Puis, nous avons mesuré la puissance à la sortie à f_c . La figure 4.20 se présente sous la forme d'un histogramme, où nous trouvons une valeur moyenne de $P_{s@f_c}$ de -60,4 dBm et un écart-type de 5,8 dB.

4.5 Les tests en mode mélangeur

Suite aux mesures du circuit dans son rôle de détecteur, nous avons effectué des mesures pour le caractériser en tant que mélangeur de fréquence. Une partie des mesures a été réalisée sur la partie du circuit imprimé dédiée aux mesures sous pointes, tandis que le reste a été effectué sur la partie principale du circuit imprimé, avec ses connecteurs SMA.

4.5.1 Les mesures sous pointes

Au laboratoire de Poly-Grames de Polytechnique Montréal, nous avons pu effectuer des mesures sous pointe, où nous avons testé le circuit en tant que mélangeur de fréquence. Le banc d'essai est illustré dans la figure 4.21 et dans la figure 4.22. Le générateur produisant le signal RF à deux fréquences était branché à la sonde simple. Ce signal était ensuite symétrisé par le *balun*

soudé au circuit imprimé. Deux générateurs produisant des signaux sinusoïdaux étaient branchés à la sonde différentielle. Ces générateurs étaient censés générer des signaux LO de la même amplitude, de la même fréquence et de phases opposées. Un analyseur de spectre était branché à un connecteur SMA relié à la sortie du circuit intégré.

Pour produire un signal LO symétrique, les ports de référence (REF), émettant et recevant un signal de référence, des deux générateurs E8257D et SMR40 étaient reliés. En théorie, les deux générateurs étaient censés produire exactement la même fréquence f_{LO} . Cependant, lors des mesures, malgré la synchronisation des générateurs, il y a eu quelques hertzs de décalage entre les signaux. Par conséquent, un phénomène de battement a eu lieu. Cela se manifestait par une variation des amplitudes mesurées par l'analyseur de spectre au cours du temps. À l'origine, nous souhaitions régler manuellement la phase du signal sortant d'E8257D afin de symétriser les signaux. Cependant, le décalage des fréquences a rendu le réglage de phase impossible. Donc, non seulement le battement rendait difficile les mesures précises, mais il empêchait la symétrisation du signal LO. Le signal RF, quant à lui, était effectivement symétrique aux entrées de la puce grâce au *balun*. Nous avons donc procédé aux mesures sans le générateur SMR40 pour éviter ce phénomène de battement. Le signal LO était donc injecté asymétriquement ; seul un des deux NMOS de la puce recevait un signal LO. Cependant, cela ne nous a pas empêché de montrer le fonctionnement du mélangeur.

Cette partie du circuit imprimé ne comporte aucun élément d'adaptation d'impédance. Par conséquent, les pertes par réflexion ici sont beaucoup plus importantes que dans la partie principale du circuit imprimé. Il a donc fallu que les générateurs de signaux produisent beaucoup plus de puissance rien que pour que le signal IF soit mesurable à la sortie. Il n'a même pas été possible d'observer la compression, car les générateurs atteignaient leurs limites de puissance avant que le circuit intégré n'entrât en compression.

Nous avons tout de même pu observer que le circuit intégré pouvait fonctionner à des fréquences RF, LO et IF variées, tel qu'illustré dans les figures 4.23 et 4.24. Nous pouvons observer sur les spectres que le signal LO rémanent à la sortie n'a pas pu être éliminé par la symétrie. Le signal

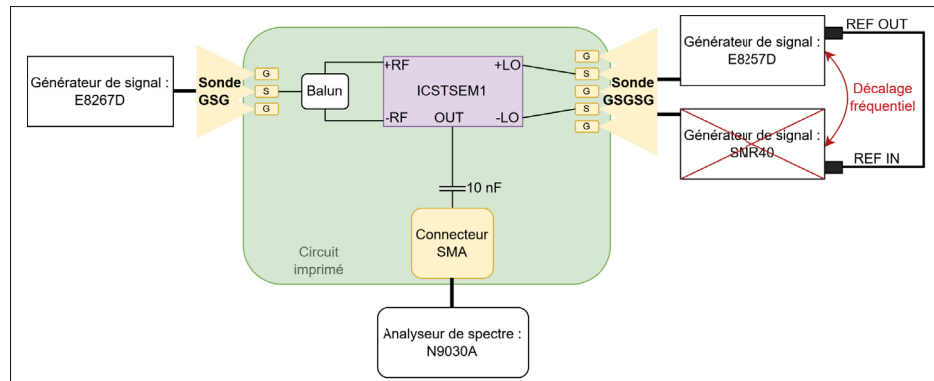


Figure 4.21 Schéma du banc de test pour les mesures sous pointes

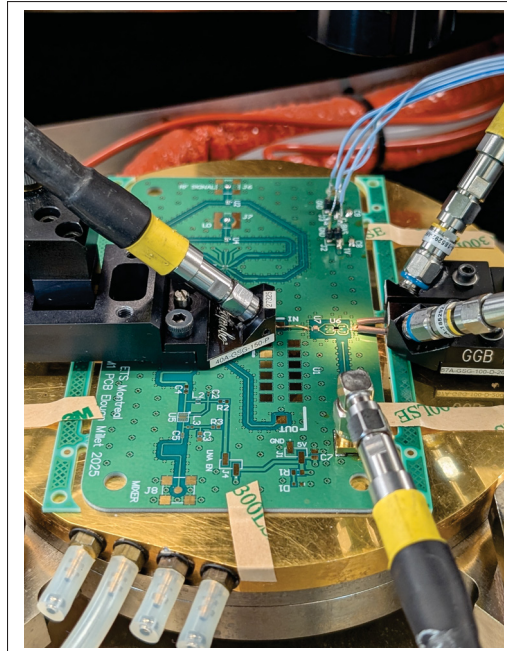


Figure 4.22 Photographie du circuit imprimé sous les sondes de mesure

RF rémanent, quant à lui, est beaucoup plus faible grâce au *balun*. Le circuit a pu quand même fonctionner en tant que mélangeur, avec les signaux RF et LO variant entre 5 GHz et 11 GHz, et le signal IF pouvant aller jusqu'à 1,5 GHz. Cependant, nous ne pouvons pas déduire la bande passante exacte du circuit intégré à cause des défauts du banc d'essai précédemment mentionnés.

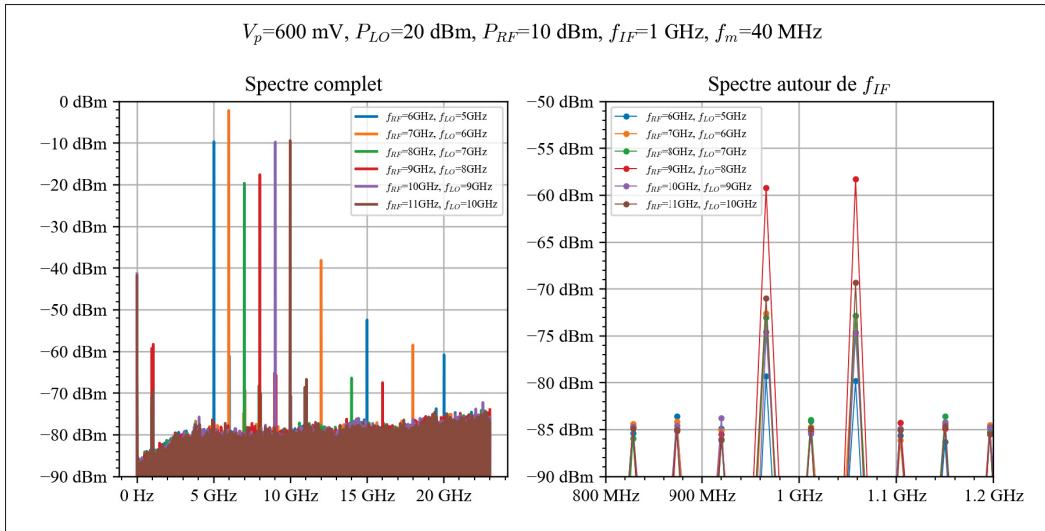


Figure 4.23 Spectres de la sortie du mélangeur avec f_{RF} et f_{LO} qui varient

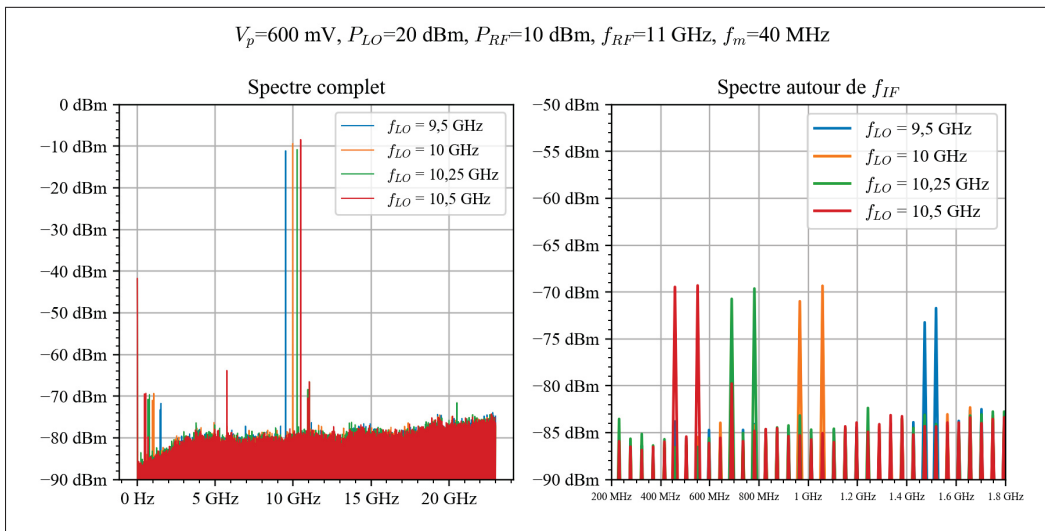


Figure 4.24 Spectres de la sortie du mélangeur avec f_{IF} et f_{LO} qui varient

Finalement, bien que nous n'ayons pas obtenu de valeurs fiables et comparables à d'autres mesures, lors des mesures sous pointe, nous avons démontré la capacité du circuit à mélanger différentes fréquences. De plus, ces mesures prouvent que le circuit peut fonctionner malgré une symétrisation imparfaite des signaux d'entrée.

4.5.2 Les mesures des spectres sur la partie principale

Suite aux mesures sous pointes, nous avons procédé au test du mélangeur sur la partie principale du circuit imprimé afin d'obtenir des données plus fiables. Pour ces essais, nous avons utilisé le banc de test représenté dans la figure 4.25. Le générateur SMW200A générait un signal RF modulé, porté à $f_c = 10,515$ GHz, tandis que l'E8244A générait un signal LO à $f_{LO} = 9,163$ GHz, ce qui donnait une fréquence IF à la sortie de $f_{IF} = 1,352$ GHz. Si ces fréquences ont été choisies, c'était pour minimiser les pertes par réflexion aux entrées du circuit intégré. La sortie du circuit était branchée à un analyseur de spectre.

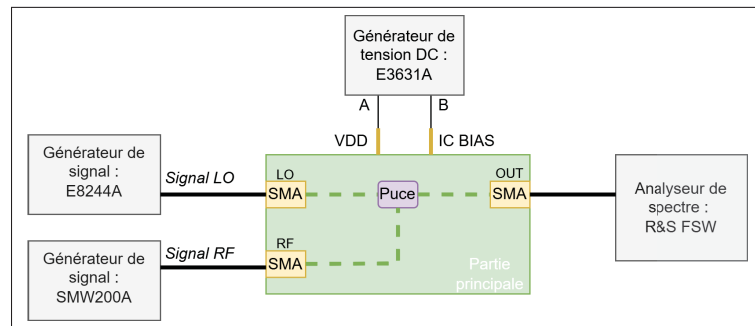


Figure 4.25 Banc d'essai du mélangeur sur la partie principale du circuit imprimé

En injectant un signal à deux fréquences, dont $f_m = 10$ MHz, dans le port RF, nous avons pu mesurer, grâce à l'analyseur de spectre, la puissance du signal IF, ainsi que celles des signaux RF et LO rémanents à la sortie. Sur la figure 4.26, nous pouvons constater que, bien que $P_{LO} = P_{RF}$ aux entrées, les signaux RF et LO rémanents n'ont pas le même niveau à la sortie. Cette différence est probablement due à un écart de performance entre les *baluns* reliés aux entrées RF et à ceux de LO. Le *balun* relié au port RF devait être notamment moins bien soudé.

Ensuite, la figure 4.27 illustre deux spectres de la sortie du mélangeur entre 0 Hz et 1,6 GHz. Un des spectres correspond à une tension de polarisation de 350 mV, où le signal comporte plus de produits d'intermodulation néfastes, et l'autre spectre correspond à une polarisation de 600 mV, où le circuit sacrifie du gain au profit de la linéarité. Sur les deux spectres, nous pouvons apercevoir le signal IF, qui se situe autour de 1,352 GHz, ainsi que le signal du détecteur à la

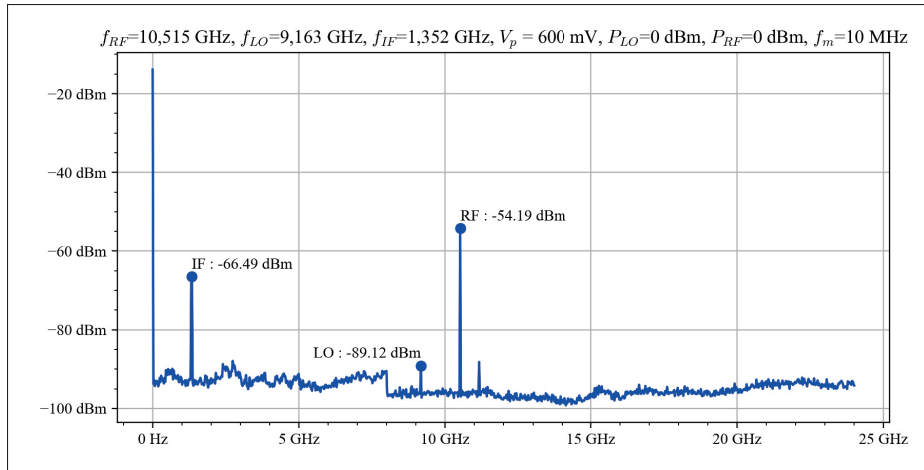


Figure 4.26 Spectre à la sortie du mélangeur avec les signaux RF et LO rémanents

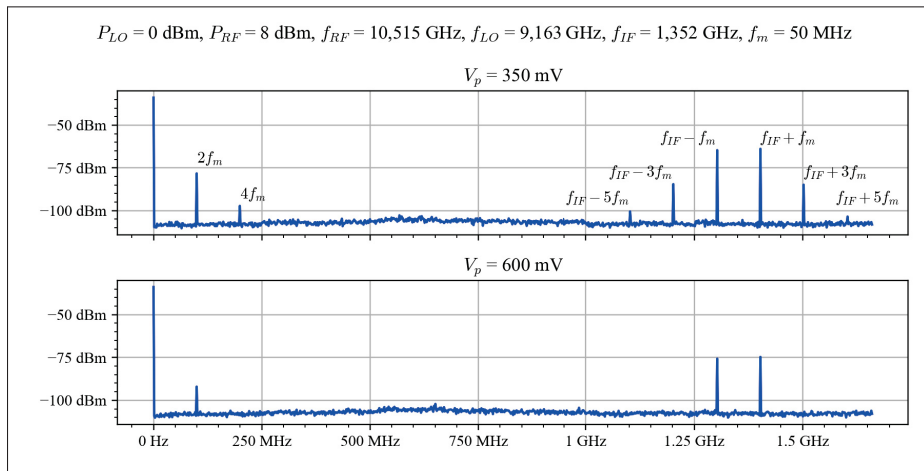


Figure 4.27 Spectres à la sortie du mélangeur illustrant les basses fréquences

bande de base. Cela prouve définitivement que le circuit fonctionne à la fois comme mélangeur et comme détecteur.

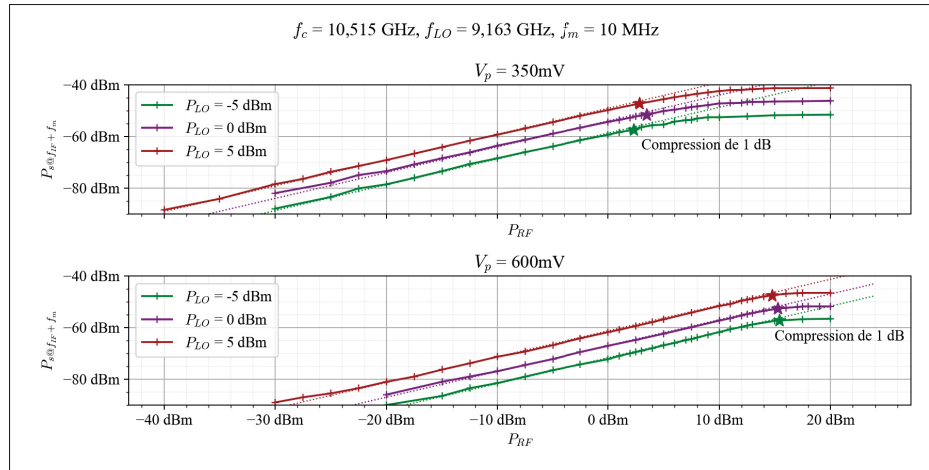


Figure 4.28 Puissance mesurée à la sortie IF en fonction de P_{RF}

4.5.3 La linéarité du mélangeur

L'analyseur de spectre a également servi à évaluer la linéarité du mélangeur. Nous avons mesuré la puissance du signal sortant à $f_{IF} \pm f_m$ et $f_{IF} \pm 3f_m$ pour différents niveaux de puissance et pour différentes tensions de polarisation.

La figure 4.28 montre la relation entre la puissance mesurée à $f_{IF} + f_m$ et P_{RF} pour différentes valeurs de P_{LO} et de V_p . Tout d'abord, nous pouvons observer que la puissance IF à la sortie est proportionnelle à celle des signaux RF et LO, contrairement au signal en bande de base sortant du détecteur qui est quadratique. Cela est cohérent avec le calcul du gain de conversion du mélangeur de la partie 2.4.2. Tout comme pour le détecteur, le gain de conversion à 350 mV de polarisation est meilleur qu'avec 600 mV de polarisation. En effet, le gain en tension quadratique α_2 du circuit reste le même, que le circuit fonctionne comme détecteur ou comme mélangeur. Enfin, nous constatons de nouveau que les points de compression $IP_{1\text{ dB}}$ sont supérieurs avec $V_p = 600\text{ mV}$ qu'avec $V_p = 300\text{ mV}$.

Nous avons pu déduire, grâce aux mesures, les caractéristiques du mélangeur, qui sont résumées dans le tableau 4.4. Il est clair que le mélangeur a un gain de conversion très faible, même lorsque $V_p = 350\text{ mV}$, où α_2 est maximal. Les pertes par réflexion et les pertes résistives

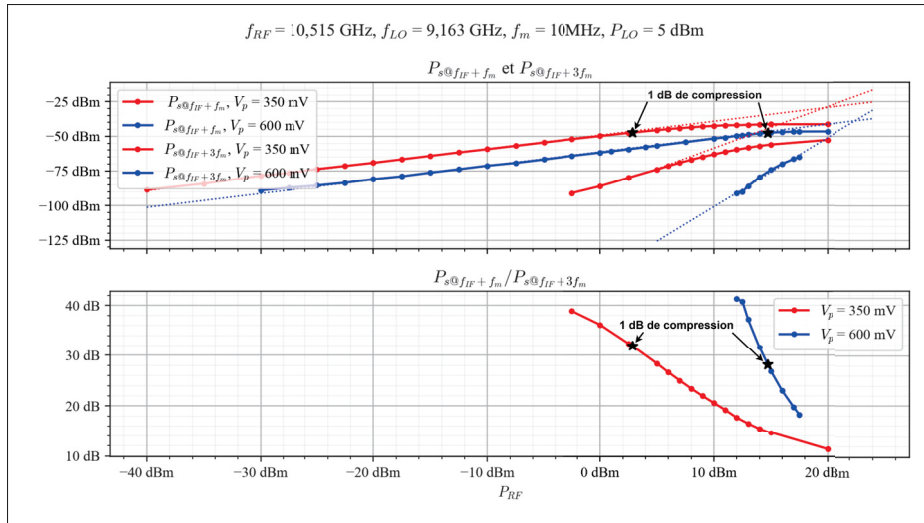


Figure 4.29 Mesures des produits d'intermodulation d'ordre deux et quatre du mélangeur

Tableau 4.4 Performance du mélangeur avec $P_{LO} = 5 \text{ dBm}$

Tension de polarisation V_p	350 mV	600 mV
Gain de conversion G_{cm}	-46,4	-58,2 dB
Point de compression $IP_{1 \text{ dB}}$	2,8 dBm	14,7 dBm
Point d'interception IIP	19,7 dBm	22,5 dBm
Plage dynamique	42,8 dB	44,7 dB

au niveau des câbles et du circuit imprimé ne sont pas les seules causes de ce gain si faible, car même les simulations, sans aucune perte par réflexion, le prévoyaient. Selon elles, avec $V_p = 350 \text{ mV}$, le gain de conversion aurait été de -39 dB et, avec $V_p = 600 \text{ mV}$, il aurait été de -50 dB. Augmenter la puissance du signal LO n'aurait pas permis de compenser ce défaut majeur. Si nous souhaitions 0 dB de gain de conversion mesuré, alors, en théorie, il aurait fallu que $P_{LO} \approx 50 \text{ dBm}$. Or, c'est un niveau de puissance ridiculement élevé qui aurait de toute manière provoqué la compression. Non, le circuit manque intrinsèquement de gain quadratique. Les NMOS M1 et M2 manquent de g_{m2} . Néanmoins, le mélangeur présente une plage dynamique de 42,8 dB pour $V_p = 350 \text{ mV}$ et de 44,7 dB pour $V_p = 600 \text{ mV}$. Comme autre qualité, sur la figure 4.29, nous pouvons apercevoir que la puissance du signal IF à $f_{IF} + f_m$ est au moins de 30 dB plus importante que celle du produit d'intermodulation d'ordre quatre à $f_{IF} + 3f_m$ jusqu'aux

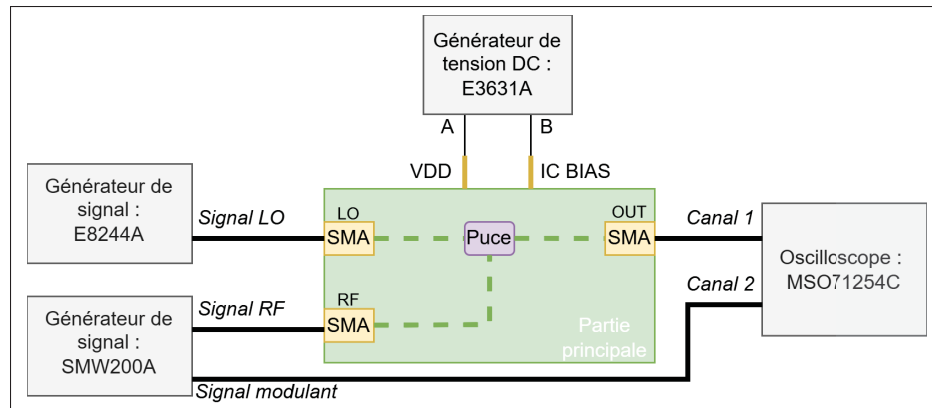


Figure 4.30 Banc de test pour les mesures à l'oscilloscope du mélangeur

points de compression de 1 dB. Le rejet du produit d'intermodulation d'ordre quatre semble meilleur avec $V_p = 600$ mV. En effet, la puissance à $f_{IF} + 3f_m$ s'accroît avec un taux de presque 50 dB par décade, au lieu des 30 dB par décade attendus. Nous pouvons de nouveau inférer que, comme dans la partie 4.3.2, il y a une annulation de l'intermodulation d'ordre quatre causée par la distorsion de phase.

4.5.4 Les mesures à l'oscilloscope sur la partie principale

Nous nous sommes de nouveau servi de l'oscilloscope afin de visuellement vérifier le fonctionnement correct du circuit en tant que mélangeur. Le banc d'essai utilisé est représenté sur la figure 4.30. Nous avons de nouveau mesuré le signal sortant du circuit imprimé, ainsi que celui sortant du générateur de bande de base.

La figure 4.31 illustre la réponse du mélangeur à un signal RF à deux fréquences à son entrée. À la sortie du mélangeur, l'amplitude du signal IF suit une forme sinusoïdale. Ensuite, nous avons répété l'expérience, mais avec un signal WCDMA qui modulait le signal RF. Sur la figure 4.32, nous pouvons voir les signaux I et Q modulateurs, ainsi que le signal IF. De nouveau, l'amplitude de ce dernier suit l'évolution des signaux I et Q.

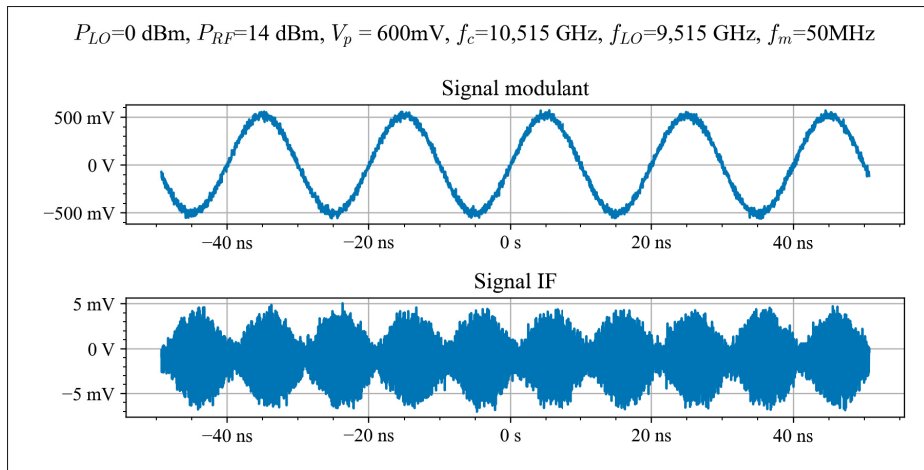


Figure 4.31 Signal IF à la sortie du mélangeur mesuré à l'oscilloscope comparé au signal modulant l'entrée

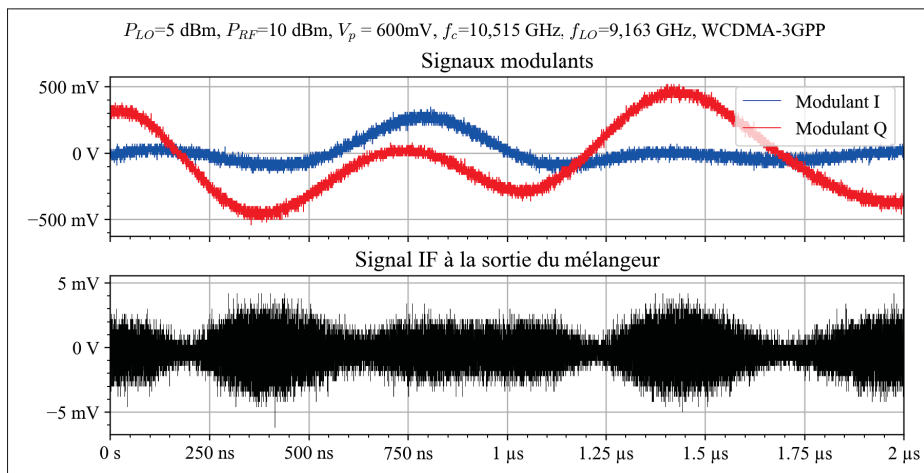


Figure 4.32 Signal IF à la sortie du mélangeur mesuré à l'oscilloscope comparé au signal WCDMA modulant l'entrée

4.6 Conclusion sur la performance du circuit

Nous avons présenté les résultats de nos mesures et de nos simulations du circuit intégré ICSTSEM1. Celui-ci s'avère capable d'opérer à la fois en tant que détecteur d'enveloppe de puissance et en tant que mélangeur de fréquences. Les mesures ont donné des résultats similaires

à ceux des simulations. Cependant, à cause des pertes par réflexions qui n'ont pas pu être prises en compte dans les simulations, le gain mesuré est plus faible que celui simulé.

D'après les mesures, le circuit, en tant que détecteur, a une responsivité maximale de 31,3 V/W. En fonction de la tension de polarisation, son point de compression $IP_{1\text{ dB}}$ peut s'élever jusqu'à 13,15 dBm. La plage dynamique maximale mesurée est de 27,72 dB, s'étendant de -15 dBm à 12,72 dBm, lorsque $V_p = 585$ mV. À cette tension de polarisation particulière, le circuit manifeste un bon taux de réjection des produits d'intermodulation d'ordre quatre. Le tableau 4.5 nous permet de comparer certaines figures de performance de notre circuit et d'autres détecteurs de puissance. Nous constatons que, par rapport à d'autres détecteurs, notre circuit a une plage dynamique qui est respectable. Cependant, la responsivité et le gain de conversion laissent à désirer. De plus, notre circuit consomme plus de puissance par rapport à plusieurs autres détecteurs.

En tant que mélangeur, le circuit s'est montré capable de fonctionner avec des fréquences RF et IF variées lors des mesures sous pointes. Le circuit a présenté un comportement très linéaire sur au moins 42,8 dB de plage dynamique. En fonction de la polarisation, le point de compression $IP_{1\text{ dB}}$ peut varier entre 2,8 et 14,7 dBm. Quant au point d'interception IP_{IP} , il varie entre 19,7 dBm et 22,5 dBm. Cependant, avec $P_{LO} = 5$ dBm, le gain de conversion ne s'élève pas au-delà de -46,4 dB. Le gain de conversion est véritablement le talon d'Achille de ce circuit, qui, autrement, présente un bon taux de réjection des produits d'intermodulation d'ordre quatre. Le tableau 4.6 permet de comparer notre circuit en tant que mélangeur à d'autres circuits de la littérature. La donnée la plus remarquable est le gain de conversion très faible de notre circuit par rapport à celui d'autres mélangeurs. Même les mélangeurs passifs ne subissent pas autant de pertes de conversion que notre circuit. Les points de compression et d'interception ($IP_{1\text{ dB}}$ et IP_{IP}) semblent meilleurs pour notre circuit. Cependant, cela peut tout simplement être une conséquence du très faible gain de conversion qui permet d'opérer à de plus hautes puissances avec peu de distorsion. Enfin, la simplicité de notre circuit permet tout de même de consommer moins de puissance, surtout par rapport aux cellules de Gilbert.

Tableau 4.5 Comparaison des performances de notre circuits et d'autres détecteurs de puissance en technologie CMOS

Ouvrage	Technologie	R_d [V/mW]	PD [dBm]	Puissance consommée [mW]	Fréquence
Qayyum & Negra (2017)	CMOS 130 nm	7*	24	0,16	7 à 70 GHz
Xu <i>et al.</i> (2014)	CMOS 65 nm	0,25*	25 à 30 *	0,015	100 MHz à 44 GHz
Li <i>et al.</i> (2019)	CMOS 65 nm	$\sim 5.10^{-3}$ *	34	0,7	4 à 6 GHz
Serhan <i>et al.</i> (2015a)	BiCMOS 55 nm	0,32*	30	0,09	50 à 66 GHz
Ozeren <i>et al.</i> (2016)	SiGe BiCMOS 250 nm	6,7*	52	7,2	7 à 20 GHz
Yang <i>et al.</i> (2012)	CMOS 130 nm	1,5*	10*	0,116	100 MHz à 40 GHz
Ce travail (2025)	CMOS 65 nm	>0,031	<27,7	[0,45 ; 3,3]	$\leq 10,5$ GHz

* : données extrapolées des figures

Plusieurs caractéristiques n'ont pas pu être mesurées de manière fiable en laboratoire. Premièrement, nous ne savons pas quelle est la fréquence de coupure aux entrées, car les circuits d'adaptation d'impédance aux entrées ont des bandes passantes limitées. Nous n'avons pu tester le circuit qu'avec certaines valeurs précises de f_c et f_{LO} . Ensuite, en ce qui concerne sa réactivité, nous avons simulé une bande passante allant de 491 MHz à 1, 18 GHz selon V_p . Or, nous n'avons pas pu répliquer ces simulations. Enfin, nous ne savons pas quel est le taux de réjection de la porteuse inhérent au circuit intégré. Mesurer cela nécessite des *baluns* parfaits aux entrées et des circuits d'adaptation d'impédance indépendants de la fréquence.

Cependant, ce que les mesures nous révèlent de manière certaine, c'est qu'un détecteur quadratique et un mélangeur quadratique sont essentiellement le même circuit. De plus, ces mesures nous ont permis de confirmer en partie les prédictions du chapitre théorique. Nous avons prouvé que les produits d'intermodulation néfastes en bande de base et en bande IF sont principalement causés par l'intermodulation d'ordre quatre. Enfin, nous avons montré un lien fort entre l'intermodulation d'ordre quatre et la compression du signal désiré.

Tableau 4.6 Comparaison des performances de notre circuits et d'autres mélangeurs en technologie CMOS

Ouvrage	Technologie	G_{cm} [dB]	IP_1 dB [dBm]	IP_{1P} [dBm]	Puissance consommée [mW]	Fréquence RF
Hsiao <i>et al.</i> (2014)	CMOS 180 nm	6	-10	-2	32,8	60 GHz
Bekkaoui (2017)	CMOS 65 nm	10	-2,54	6	2,17	1,9 GHz
Tsai <i>et al.</i> (2007a)	CMOS 90 nm	3	-4,5 à 2,1	?	93	25 à 75 GHz
Michaelsen <i>et al.</i> (2015)	SiGe BiCMOS	4	>-11	13	31,62*	10,5 GHz
Shahriar Rashid & Rashid (2010)	CMOS 90 nm	11,4	-10,8	0	9,25	17,3 à 21,8 GHz
Nouri <i>et al.</i> (2010)	CMOS 130 nm	1,7	?	14,2	6 + 1,4*	60 GHz
Bao <i>et al.</i> (2006)	CMOS 90 nm	-11 à -8	?	3 à 7	9,33* à 18,2*	9 à 31 GHz
Ce travail (2025)	CMOS 65 nm	>-46,4	< 14, 7	< 22, 5	3,3 + 1,5*	≤10,5 GHz

* : puissance du signal LO

CONCLUSION ET RECOMMANDATIONS

Dans un premier temps, nous avons établi le besoin d'antennes peu coûteuses capables de modelage de faisceau. Ces antennes utilisent souvent des déphaseurs à capacité variable. Or, ces dernières ont tendance à causer de la distorsion. Pour corriger cela, nous avons proposé l'usage de détecteurs de puissance et de mélangeurs. Grâce à notre revue de la littérature, nous savons désormais que ces circuits peuvent être réalisés simplement dans un circuit intégré à partir de composants à réponse quadratique. Pour rendre notre solution encore plus simple, nous avons proposé un seul circuit qui peut accomplir les deux tâches.

Tout d'abord, nous avons étudié la théorie des circuits à fonctionnement quadratique. Nous avons notamment expliqué comment l'intermodulation du second ordre permet la détection et le mélange. Puis, nous avons mis en évidence comment les produits d'intermodulation d'ordre quatre étaient nuisibles pour notre circuit. Nous avons posé un cadre théorique généralisé pour les circuits quadratiques, enrichissant ainsi la littérature scientifique. Suite à cette analyse, nous avons entamé la conception du circuit intégré ICSTSEM1 et du circuit imprimé. La puce a été conçue pour minimiser la présence des produits d'intermodulation d'ordre quatre. De plus, nous avons démontré une méthode simple pour effectuer l'analyse d'un circuit non linéaire à l'aide du modèle des petits signaux. Le circuit imprimé, quant à lui, a été conçu pour permettre le test de la puce de deux façons différentes. Lors des essais, nous avons vérifié que le circuit intégré pouvait effectivement fonctionner à la fois comme détecteur d'enveloppe de puissance et comme mélangeur, sans modification du circuit, malgré quelques défauts. Cependant, le manque de gain de conversion est son principal point faible. Une faute de conception du circuit intégré est à l'origine du gain insatisfaisant. De plus, les pertes par réflexion au niveau du circuit imprimé empirent la performance du circuit. Néanmoins, le gain peut être relativement facilement amélioré en redimensionnant les transistors, en modifiant la topologie ou en greffant un étage d'amplification à la sortie.

Nous avons réussi à concevoir un circuit intégré en technologie CMOS de détecteur-mélangeur miniaturisé. La puce occupe $1\text{ mm} \times 1,25\text{ mm}$ de la surface du circuit imprimé. Les composants actifs du circuit intégré n'occupent cependant qu'une aire de $24,73\text{ }\mu\text{m} \times 60,695\text{ }\mu\text{m}$ de la puce. La simplicité lui permet de consommer moins de 3,33 mW de puissance au repos, avec une tension d'alimentation de 1 V. En tant que détecteur, sa responsivité mesurée est de 31,3 V/W maximum. En fonction de la polarisation, sa plage dynamique peut s'étendre jusqu'à 27,7 dB de largeur et son point de compression IP_{1dB} peut s'élever jusqu'à 13,15 dBm. En tant que mélangeur de fréquence, nous avons remarqué un gain de conversion maximal de -46.4 dB lorsque $P_{LO} = 5\text{ dBm}$. Sa plage dynamique peut être 44,7 dB de large, son point de compression IP_{1dB} peut s'élever jusqu'à 14,7 dBm et son point d'interception IP_{IP} peut valoir 22,5 dBm. Dans les deux modes de fonctionnement, la puissance du signal désiré peut être 40 dB plus élevée que celle des produits d'intermodulation nuisibles l'avoisinant sur presque toute la plage dynamique, si le circuit est polarisé correctement. Par ailleurs, nous n'avons pas mesuré toutes les caractéristiques du circuit. Nous n'avons pas pu mesurer les bandes passantes aux entrées et à la sortie, notamment à cause du circuit imprimé. Nous avons seulement prouvé que le circuit pouvait fonctionner avec un signal aux entrées d'au moins 10,5 GHz. Nous avons dû nous limiter à une simulation concernant la bande passante de la sortie, car les mesures nous permettent seulement de savoir qu'elle est d'au moins 212 MHz. Nous avons également dû faire confiance aux simulations quant au taux de réjection de la porteuse.

Les mesures nous ont non seulement permis de caractériser le circuit, mais elles nous ont aussi permis de valider le modèle non linéaire du circuit basé sur une série de Taylor. Un modèle aussi simple nous a permis de comprendre les origines des différents produits d'intermodulation. De plus, il a permis de prédire la relation entre le phénomène de compression et l'intermodulation d'ordre quatre. Enfin, nous sommes désormais certains, grâce au modèle et aux mesures, qu'un circuit avec une réponse quadratique peut tout aussi bien servir de détecteur de puissance que

de mélangeur. Ce travail permet donc de combler le manque d'informations détaillées sur le fonctionnement et la conception des circuits quadratiques.

À notre avis, un circuit intégré ICSTSEM1 révisé et amélioré pourrait véritablement servir à améliorer la polarisation des diodes des déphaseurs des antennes réseau. Puisque la partie active du détecteur-mélangeur est compacte, on peut également envisager d'intégrer plusieurs détecteurs-mélangeurs dans une seule puce. Cela permettrait de réduire le nombre de circuits intégrés sur l'antenne. Ou alors, un circuit de détecteur-mélangeur de petite taille peut être facilement intégré dans un système plus large pour du BIST. La double fonctionnalité du circuit permettrait le test et la caractérisation d'un circuit plus complet. Par exemple, un amplificateur de puissance pourrait être testé lors de sa fabrication grâce à notre circuit. Le détecteur pourrait mesurer son gain, tandis que le mélangeur permettrait de caractériser sa non-linéarité en transposant vers le bas le signal sortant de l'amplificateur.

Recommandations

Nous avons de nombreuses recommandations pour celles et ceux qui souhaiteraient approfondir notre travail sur les circuits quadratiques.

Tout d'abord, en ce qui concerne la théorie, nous recommandons à examiner plus scrupuleusement les effets non linéaires secondaires que nous avons négligés ou simplifiés. Nous avons mis de côté l'effet de modulation de la longueur du canal et l'effet du substrat pour une question de simplicité. Ensuite, nous n'avons pas étudié en profondeur les intermodulations d'ordres supérieurs à quatre. Nous savons que les intermodulations d'ordres six, huit, dix, etc. peuvent avoir des effets similaires à ceux d'ordre quatre. Cependant, les étudier permettrait de savoir s'il existe un moyen d'annuler l'intermodulation d'ordre quatre. Cela donnerait au circuit un comportement plus quadratique. Dans ce but, nous proposons d'étudier des modèles de transistors plus complexes que celui de l'équation (2.1). De plus, développer une formule mathématique généralisée des

produits d'intermodulation faciliterait l'analyse pour les ordres supérieurs, car le développement de la formule d'un signal à deux fréquences, élevé rien qu'à la puissance quatre, est chronophage.

Outre la théorie, il faudrait revoir le circuit. Bien que les entrées symétriques permettent l'annulation de la porteuse et des produits d'intermodulation des ordres impairs, cette topologie ne semble pas offrir un bon gain sans compromettre un critère de performance. Par conséquent, nous proposons de revoir entièrement la conception du circuit intégré en s'aidant de l'analyse en petits-sinaux non linéaire. De plus, il serait critique d'implémenter un circuit d'adaptation d'impédance amélioré pour le prochain circuit. Sans cela, tout effort d'amélioration du gain serait vain.

Pour ceux qui souhaitent continuer le travail sur un circuit de détecteur-mélangeur, nous les encourageons vivement à revoir l'isolation entre les ports RF et LO en s'inspirant de Shahriar Rashid & Rashid (2010). Nous incitons également fortement à développer un circuit avec des sorties séparées pour la bande de base et la bande IF, car ces sorties ont des besoins et des contraintes différents. Nous encourageons également à étudier les méthodes proposées dans la littérature pour minimiser les effets capacitifs parasites des diodes de protection, qui peuvent réduire la bande passante d'un circuit intégré. Enfin, en ce qui concerne la conception de la puce, nous décourageons l'utilisation de la méthode d'assemblage de la puce retournée, sauf si elle s'avère absolument nécessaire, surtout pour les concepteurs de puces débutants. En effet, cette technique impose des contraintes de dimensions strictes pour le circuit intégré et le circuit imprimé, ce qui rend la conception plus exigeante. De plus, avec la puce retournée, il y a plus de chances qu'il y ait un court-circuit sous le circuit intégré qui ne soit pas détecté. Finalement, avec une puce retournée, il n'est pas possible de sonder directement les pastilles de connexion du circuit intégré. Par conséquent, les résultats des mesures dépendent de la qualité du circuit imprimé.

ANNEXE I

CALCULS THÉORIQUES

1. Les produits d'intermodulation

1.1 Les ordres des produits d'intermodulation néfastes

Dans cette partie, nous étudions quels ordres d'intermodulations génèrent des produits dans la bande de base pour le détecteur et dans la bande IF pour le mélangeur.

1.1.1 Détecteur

Soit $(a, b) \in \mathbb{Z}^2$. Si l'on injecte un signal RF à deux fréquences, $f_1 = f_c + f_m$ et $f_2 = f_c - f_m$, les produits d'intermodulation se trouvent aux fréquences suivantes :

$$af_1 + bf_2 = a(f_c + f_m) + b(f_c - f_m) = (a + b)f_c + (a - b)f_m \quad (\text{A I-1})$$

L'ordre d'intermodulation est la valeur de $|a| + |b|$. Les produits d'intermodulation dans la bande de base se trouvent aux fréquences kf_m , où k est un entier non nul.

Alors, deux conditions s'imposent sur le couple (a, b) :

$$\begin{cases} (a + b)f_c = 0 \\ (a - b)f_m = kf_m \end{cases} \quad (\text{A I-2})$$

alors

$$\begin{cases} a = -b \\ -2b = k \end{cases} \quad (\text{A I-3})$$

Par conséquent, $(a, b) = (k/2, -k/2)$. Puis, étant donné que $(a, b) \in \mathbb{Z}^2$, k est nécessairement un nombre pair. Ensuite, si l'on calcule l'ordre d'intermodulation, nous obtenons :

$$|a| + |b| = |k/2| + |-k/2| = |k| \quad (\text{A I-4})$$

Donc tous les produits d'intermodulation en bande de base se trouvent à une harmonique paire de f_m . De plus, l'ordre de l'intermodulation qui génère un produit d'intermodulation à la bande de base est nécessairement pair. Le produit d'intermodulation d'ordre deux à $2f_m$ est désirable pour un détecteur de puissance, tandis que tous les autres ne le sont pas.

Tableau-A I-1 Ordres des intermodulations d'un détecteur d'enveloppe de puissance

k	Fréquence	(a,b)	Ordre
0	0	(0,0)	0
2	$2f_m$	(1,-1)	2
4	$4f_m$	(2,-2)	4
6	$6f_m$	(3,-3)	6
8	$8f_m$	(4,-4)	8

1.1.2 Mélangeur

Soit $(a, b, c) \in \mathbb{Z}^3$. Si l'on injecte un signal RF à deux fréquences, $f_1 = f_c + f_m$ et $f_2 = f_c - f_m$, et un signal LO à $f_{LO} = f_c + f_{IF}$, les produits d'intermodulation se trouvent aux fréquences suivantes :

$$af_1 + bf_2 + cf_{LO} = a(f_c + f_m) + b(f_c - f_m) + c(f_c + f_{IF}) = (a+b+c)f_c + (a-b)f_m + cf_{IF} \quad (\text{A I-5})$$

L'ordre d'intermodulation est la valeur de $|a| + |b| + |c|$. Par exemple, le mélange désirable des fréquences se fait avec l'intermodulation d'ordre 2, où $(a, b, c) = (1, 0, -1)$ ou $(a, b, c) = (0, 1, -1)$, ce qui donne les fréquences $f_{IF} \pm f_m$.

Les produits d'intermodulation néfastes se trouvent aux fréquences $f_{IF} + k f_m$, avec $k \in \mathbb{Z}$. Nous pouvons tenter de déterminer quels ordres d'intermodulation sont les plus nuisibles.

Trois conditions s'imposent sur (a, b, c) :

$$\begin{cases} (a + b + c)f_c = 0 \\ (a - b)f_m = k f_m \\ c f_{IF} = f_{IF} \end{cases} \quad (\text{A I-6})$$

Donc

$$\begin{cases} a + b + 1 = 0 \\ a - b = k \\ c = 1 \end{cases} \quad (\text{A I-7})$$

Puis on obtient que $(a, b, c) = (\frac{k-1}{2}, -\frac{k+1}{2}, 1)$.

Pour que a, b et c soient des entiers, il faut que k soit un entier impair. En effet :

Soit $i \in \mathbb{Z}$,

Si $k = 2i$, alors $(\frac{k-1}{2}, -\frac{k+1}{2}, 1) = (\frac{2i-1}{2}, -\frac{2i+1}{2}, 1) \notin \mathbb{Z}^3$.

Alors que si $k = 2i + 1$, alors $(\frac{k-1}{2}, -\frac{k+1}{2}, 1) = (\frac{2i+1-1}{2}, -\frac{2i+1+1}{2}, 1) = (i, -(i+1), 1) \in \mathbb{Z}^3$.

Nous avons ainsi démontré que $\forall i \in \mathbb{Z}, (a, b, c) = (i, -(i+1), 1)$ engendre les produits d'intermodulation néfastes. Cherchons les ordres de ces intermodulations :

$|a| + |b| + |c| = |i| + |-(i+1)| + 1 = |i| + |i+1| + 1$. Si $i \geq 0$, alors $|a| + |b| + |c| = 2i + 2$. Si $i < 0$, alors $|a| + |b| + |c| = |i| + |i| - 1 + 1 = 2|i|$.

Donc, dans tous les cas, **l'ordre de l'intermodulation est pair**.

Tableau-A I-2 Les ordres des intermodulations des mélangeurs quadratiques

i	Fréquence	(a,b,c)	Ordre
$i < 0$	$f_{IF} + (-2 i + 1)f_m$	$(- i , i - 1, 1)$	$2 i $
-3	$f_{IF} - 5f_m$	$(-3, 2, 1)$	6
-2	$f_{IF} - 3f_m$	$(-2, 1, 1)$	4
-1	$f_{IF} - f_m$	$(-1, 0, 1)$	2
0	$f_{IF} + f_m$	$(0, -1, 1)$	2
1	$f_{IF} + 3f_m$	$(1, -2, 1)$	4
2	$f_{IF} + 5f_m$	$(2, -3, 1)$	6
$i > 0$	$f_{IF} + (2i + 1)f_m$	$(i, -i - 1, 1)$	$2i + 2$

1.2 Développements complets des produits d'intermodulation pour un signal à deux fréquences

Dans cette partie, nous retrouvons les développements quasi-complets des signaux à deux fréquences élevés à des puissances paires. Ces calculs nous permettent de retracer les origines des produits d'intermodulation de différents ordres.

1.2.1 Détecteur

Soient ω_1 et ω_2 deux nombres réels positifs, non nuls et distincts. Soit $V_{RF} \in \mathbb{R}^{+*}$. Posons le signal $v_{RF}(t) = \frac{V_{RF}}{2}(\cos(\omega_1 t) + \cos(\omega_2 t))$ qui est à l'entrée du détecteur. Posons également ω_c et ω_m tels que $\omega_1 = \omega_c + \omega_m$ et $\omega_2 = \omega_c - \omega_m$.

Dans un premier temps, ce signal élevé au carré donne :

$$\begin{aligned}
 v_{RF}^2 &= \left(\frac{V_{RF}}{2} \right)^2 [\cos(\omega_1 t) + \cos(\omega_2 t)]^2 \\
 &= \frac{V_{RF}^2}{4} [\cos^2(\omega_1 t) + 2 \cos(\omega_1 t) \cos(\omega_2 t) + \cos^2(\omega_2 t)] \\
 &= \frac{V_{RF}^2}{4} \left[\frac{1 + \cos(2\omega_1 t)}{2} + 2 \frac{\cos((\omega_1 - \omega_2)t) + \cos((\omega_1 + \omega_2)t)}{2} + \frac{1 + \cos(2\omega_2 t)}{2} \right] \quad (\text{A I-8}) \\
 &= \frac{V_{RF}^2}{4} \left[1 + \underbrace{\cos((\omega_1 - \omega_2)t)}_{\cos(2\omega_m t)} + \underbrace{\cos((\omega_1 + \omega_2)t)}_{\cos(2\omega_c t)} + \frac{\cos(2\omega_1 t) + \cos(2\omega_2 t)}{2} \right]
 \end{aligned}$$

Ce signal élevé à la puissance quatre donne :

$$\begin{aligned}
v_{RF}^4 &= \left(\frac{V_{RF}}{2} \right)^4 [\cos(\omega_1 t) + \cos(\omega_2 t)]^4 \\
&= \left(\frac{V_{RF}^4}{16} \right) [\cos^4(\omega_1 t) + \cos^4(\omega_2 t) + 4 \cos^3(\omega_1 t) \cos(\omega_2 t) + 4 \cos^3(\omega_2 t) \cos(\omega_1 t) \\
&\quad + 6 \cos^2(\omega_1 t) \cos^2(\omega_2 t)] \\
&= \left(\frac{V_{RF}^4}{16} \right) \left[\frac{3 + 4 \cos(2\omega_1 t) + \cos(4\omega_1 t)}{8} + \frac{3 + 4 \cos(2\omega_2 t) + \cos(4\omega_2 t)}{8} \right. \\
&\quad + 4 \frac{3 \cos(\omega_1 t) + \cos(3\omega_1 t)}{4} \cos(\omega_2 t) + 4 \frac{3 \cos(\omega_2 t) + \cos(3\omega_2 t)}{4} \cos(\omega_1 t) \\
&\quad \left. + 6 \frac{1 + \cos(2\omega_1 t)}{2} \frac{1 + \cos(2\omega_2 t)}{2} \right] \\
&= \left(\frac{V_{RF}^4}{16} \right) \left[\frac{3}{4} + \frac{1}{2} (\cos(2\omega_1 t) + \cos(2\omega_2 t)) + \frac{1}{8} (\cos(4\omega_1 t) + \cos(4\omega_2 t)) \right. \\
&\quad + 3 (\cos(\omega_1 t) \cos(\omega_2 t) + \cos(\omega_2 t) \cos(\omega_1 t)) + \cos(3\omega_1 t) \cos(\omega_2 t) \\
&\quad + \cos(3\omega_2 t) \cos(\omega_1 t) + \frac{6}{4} (1 + \cos(2\omega_1 t) + \cos(2\omega_2 t) + \cos(2\omega_1 t) \cos(2\omega_2 t)) \left. \right] \\
&= \left(\frac{V_{RF}^4}{16} \right) \left[\frac{9}{4} + 2 (\cos(2\omega_1 t) + \cos(2\omega_2 t)) + \frac{1}{8} (\cos(4\omega_1 t) + \cos(4\omega_2 t)) \right. \\
&\quad + 6 \underbrace{\cos(\omega_1 t) \cos(\omega_2 t)}_{(\cos(2\omega_m t) + \cos(2\omega_c t))/2} + \cos(3\omega_1 t) \cos(\omega_2 t) + \cos(3\omega_1 t) \cos(\omega_2 t) \\
&\quad + \frac{3}{2} \underbrace{\cos(2\omega_1 t) \cos(2\omega_2 t)}_{(\cos(4\omega_m t) + \cos(4\omega_c t))/2} \left. \right] \\
&= \left(\frac{V_{RF}^4}{16} \right) \left[\frac{9}{4} + 3 \cos(2\omega_m t) + \frac{3}{4} \cos(4\omega_m t) + \dots \right]
\end{aligned} \tag{A I-9}$$

1.2.2 Mélangeur

Soit un mélangeur qui reçoit la somme d'un signal RF et d'un signal LO à son entrée. Le signal RF est le même que celui ci-dessus. Le signal LO est de la forme $v_{LO}(t) = V_{LO} \cos(\omega_{LO} t)$.

Posons ω_{IF} tel que $\omega_c = \omega_{LO} + \omega_{IF}$.

Si nous élevons $v_{RF} + v_{LO}$ au carré, cela donne :

$$\begin{aligned}
 (v_{RF} + v_{LO})^2 &= v_{RF}^2 + v_{LO}^2 + 2v_{RF} v_{LO} \\
 &= \frac{V_{RF}^2}{4} (\cos(\omega_1 t) + \cos(\omega_2 t))^2 + V_{LO}^2 \cos^2(\omega_{LO} t) \\
 &\quad + \frac{2V_{RF}V_{LO}}{2} (\cos(\omega_1 t) + \cos(\omega_2 t)) \cos(\omega_{LO} t) \\
 &= \frac{V_{RF}^2}{4} (\cos(\omega_1 t) + \cos(\omega_2 t))^2 + V_{LO}^2 \cos^2(\omega_{LO} t) \\
 &\quad + V_{RF}V_{LO} (\cos(\omega_1 t) \cos(\omega_{LO} t) + \cos(\omega_2 t) \cos(\omega_{LO} t)) \\
 &= \frac{V_{RF}^2}{4} (\cos(\omega_1 t) + \cos(\omega_2 t))^2 + V_{LO}^2 \cos^2(\omega_{LO} t) \\
 &\quad + V_{RF}V_{LO} \left(\frac{\cos((2\omega_c + \omega_m + \omega_{IF})t) + \cos((\omega_{IF} - \omega_m)t)}{2} \right. \\
 &\quad \left. + \frac{\cos((2\omega_c - \omega_m + \omega_{IF})t) + \cos((\omega_{IF} + \omega_m)t)}{2} \right) \\
 &= \frac{V_{RF}^2}{4} (\cos(\omega_1 t) + \cos(\omega_2 t))^2 + V_{LO}^2 \cos^2(\omega_{LO} t) \\
 &\quad + \frac{V_{RF}V_{LO}}{2} [\cos((\omega_{IF} + \omega_m)t) + \cos((\omega_{IF} - \omega_m)t) + \dots] \\
 &= \frac{V_{RF}^2}{4} \left[1 + \underbrace{\cos((\omega_1 - \omega_2)t)}_{\cos(2\omega_m t)} + \underbrace{\cos((\omega_1 + \omega_2)t)}_{\cos(2\omega_c t)} + \frac{\cos(2\omega_1 t) + \cos(2\omega_2 t)}{2} \right] \\
 &\quad + V_{LO}^2 \frac{1 + \cos(2\omega_{LO} t)}{2} \\
 &\quad + \frac{V_{RF}V_{LO}}{2} [\cos((\omega_{IF} + \omega_m)t) + \cos((\omega_{IF} - \omega_m)t) + \dots]
 \end{aligned}
 \tag{A I-10}$$

Si nous élevons ce signal à la puissance quatre, cela donne :

$$\begin{aligned}
(v_{RF} + v_{LO})^4 &= \dots + 4v_{RF}^3 v_{LO} + 6v_{RF}^2 v_{LO}^2 + 4v_{RF} v_{LO}^3 + \dots \\
&= \dots + 4 \left(\frac{V_{RF}}{2} \right)^3 (\cos(\omega_1 t) + \cos(\omega_2 t))^3 V_{LO} \cos(\omega_{LO} t) + 6 \left(\frac{V_{RF}}{2} \right)^2 (\cos(\omega_1 t) + \cos(\omega_2 t))^2 V_{LO} \cos^2(\omega_{LO} t) \\
&\quad + 4 \left(\frac{V_{RF}}{2} \right) (\cos(\omega_1 t) + \cos(\omega_2 t)) V_{LO} \cos^3(\omega_{LO} t) + \dots \\
&= \dots + \frac{V_{RF}^3 V_{LO}}{2} (\cos^3(\omega_1 t) + 3 \cos^2(\omega_1 t) \cos(\omega_2 t) + 3 \cos^2(\omega_2 t) \cos(\omega_1 t) + \cos^3(\omega_2 t)) \cos(\omega_{LO} t) \\
&\quad + \frac{3V_{RF}^2 V_{LO}^2}{2} \left(\dots + \cos(2\omega_c t) + \frac{\cos(2\omega_1 t) + \cos(2\omega_2 t)}{2} \right) \left(\dots + \cos(2\omega_{LO} t) \right) \\
&\quad + 2V_{RF} V_{LO}^3 (\cos(\omega_1 t) + \cos(\omega_2 t)) \left(\frac{3 \cos(\omega_{LO} t) + \cos(3\omega_{LO} t)}{4} \right) + \dots \\
&= \dots + \frac{V_{RF}^3 V_{LO}}{2} \left(\frac{3 \cos(\omega_1 t) + \dots}{4} + 3 \frac{1 + \cos(2\omega_1 t)}{2} \cos(\omega_2 t) + 3 \frac{1 + \cos(2\omega_2 t)}{2} \cos(\omega_1 t) + \frac{3 \cos(\omega_2 t) + \dots}{4} \right) \cos(\omega_{LO} t) \\
&\quad + \frac{3V_{RF}^2 V_{LO}^2}{2} \left(\dots + \frac{\cos(2\omega_c t) \cos(2\omega_{LO} t)}{2} + \frac{\cos(2\omega_1 t) \cos(2\omega_{LO} t)}{4} + \frac{\cos(2\omega_2 t) \cos(2\omega_{LO} t)}{4} \right) \\
&\quad + 2V_{RF} V_{LO}^3 \left(\frac{3 \cos(\omega_1 t) \cos(\omega_{LO} t) + \dots}{4} + \frac{3 \cos(\omega_2 t) \cos(\omega_{LO} t) + \dots}{4} \right) \\
&= \dots + \frac{V_{RF}^3 V_{LO}}{2} \left(\frac{3 \cos(\omega_1 t)}{4} + \frac{3 \cos(\omega_2 t)}{2} + \frac{3 \cos((\omega_c + 3\omega_m)t) + \dots}{4} + \frac{3 \cos(\omega_1 t)}{2} + \frac{3 \cos((\omega_c - 3\omega_m)t) + \dots}{4} + \frac{3 \cos(\omega_2 t)}{4} + \dots \right) \cos(\omega_{LO} t) \\
&\quad + \frac{3V_{RF}^2 V_{LO}^2}{2} (\dots) \\
&\quad + 2V_{RF} V_{LO}^3 \left(\frac{3 \cos((\omega_1 - \omega_{LO})t) + \dots}{8} + \frac{3 \cos((\omega_2 - \omega_{LO})t) + \dots}{8} \right) + \dots \\
&= \dots + \frac{V_{RF}^3 V_{LO}}{2} \left(\frac{9 \cos(\omega_1 t) \cos(\omega_{LO} t)}{4} + \frac{9 \cos(\omega_2 t) \cos(\omega_{LO} t)}{4} + \frac{3 \cos((\omega_c + 3\omega_m)t) \cos(\omega_{LO} t)}{4} + \frac{3 \cos((\omega_c - 3\omega_m)t) \cos(\omega_{LO} t)}{4} + \dots \right) \\
&\quad + 2V_{RF} V_{LO}^3 \left(\frac{3 \cos((\omega_{IF} - \omega_m)t)}{8} + \frac{3 \cos((\omega_{IF} + \omega_m)t)}{8} + \dots \right) + \dots \\
&= \dots + \frac{V_{RF}^3 V_{LO}}{2} \left(\frac{9 \cos((\omega_{IF} + \omega_m)t) + \dots}{8} + \frac{9 \cos((\omega_{IF} - \omega_m)t) + \dots}{8} + \frac{3 \cos((\omega_{IF} + 3\omega_m)t) + \dots}{8} + \frac{3 \cos((\omega_{IF} - 3\omega_m)t) + \dots}{8} + \dots \right) \\
&\quad + 2V_{RF} V_{LO}^3 \left(\frac{3 \cos((\omega_{IF} - \omega_m)t)}{8} + \frac{3 \cos((\omega_{IF} + \omega_m)t)}{8} + \dots \right) + \dots \\
&= \frac{V_{RF}^3 V_{LO}}{16} (9 \cos((\omega_{IF} \pm \omega_m)t) + 3 \cos((\omega_{IF} \pm 3\omega_m)t)) + \frac{3V_{RF} V_{LO}^3}{4} (\cos((\omega_{IF} \pm \omega_m)t)) + \dots
\end{aligned}$$

(A I-11)

Par ailleurs, les équations (A I-8), (A I-9), (A I-10) et (A I-11) de cette partie ont été vérifiées numériquement via un outil de transformation de Fourier de la librairie NumPy sur Python.

2. La compression

Puisque un circuit de détecteur de puissance n'est pas purement quadratique en réalité, à un certain niveau de puissance d'entrée, il atteint un point de compression. Nous pouvons calculer analytiquement ces points de compression en nous inspirant de la méthode utilisée par Pozar (2012) (p.512) pour étudier la non linéarité d'un amplificateur.

2.1 La compression du détecteur avec un signal à amplitude constante

Pour un détecteur de puissance, nous pouvons mesurer ce point en évaluant la tension continue à sa sortie.

Soit un signal $v_{RF}(t) = V_{RF} \cos(\omega t)$ à l'entrée d'un détecteur. Alors, la tension à la sortie est composée de V_s qui est constante et v_s qui dépend du temps.

$$V_s + v_s(t) = V_r(V_p) + \sum_{i=1}^{\infty} \alpha_i v_{RF}^i(t) \quad (\text{A I-12})$$

Les termes $\alpha_i v_{RF}^i$ d'ordre impair de la série ne génèrent pas de composant à 0 Hz. Seuls les termes d'ordre pair influencent la tension continue. Donc, en ce qui concerne la tension de sortie continue, nous pouvons ignorer tous les termes d'ordre impair. Puis, nous pouvons supposer que les termes d'ordre 2 et 4 sont prépondérants.

$$V_s \approx V_r(V_p) + (\alpha_2 v_{RF}^2 + \alpha_4 v_{RF}^4(t))|_{f=0 \text{ Hz}} \quad (\text{A I-13})$$

$$\approx V_r(V_p) + \frac{\alpha_2 V_{RF}^2}{2} + \frac{3\alpha_4 V_{RF}^4}{8} \quad (\text{A I-14})$$

Le phénomène de compression a lieu à cause de l'existence de α_4 . Si le gain diminue, c'est parce que α_2 et α_4 sont de signes opposés. Dans le cas d'ICSTSEM1, α_2 est positif, donc α_4 est négatif. Posons $V_{1 \text{ dB,d}}$ comme l'amplitude V_{RF} qui engendre 1 dB de compression de $|V_s - V_r|$. C'est-à-dire :

$$\frac{|V_s - V_r|_{\alpha_4=0}}{|V_s - V_r|_{\alpha_4 \neq 0}} = 10^{1/20} \quad (\text{A I-15})$$

$$10^{-1/20} \alpha_2 \frac{V_{1 \text{ dB,d}}^2}{2} = \alpha_2 \frac{V_{1 \text{ dB,d}}^2}{2} + \alpha_4 \frac{3V_{1 \text{ dB,d}}^4}{8} \quad (\text{A I-16})$$

$$-\alpha_4 \frac{3V_{1 \text{ dB,d}}^2}{8} = \frac{\alpha_2}{2} (1 - 10^{-1/20}) \quad (\text{A I-17})$$

$$V_{1 \text{ dB,d}}^2 = \frac{8\alpha_2}{6|\alpha_4|} (1 - 10^{-1/20}) \quad (\text{A I-18})$$

$$V_{1 \text{ dB,d}} = 2\sqrt{\frac{\alpha_2}{3|\alpha_4|} (1 - 10^{-1/20})} \quad (\text{A I-19})$$

(Cette formule a été numériquement vérifiée sur Python.) D'après ce résultat, plus le ratio α_2/α_4 est important, plus $V_{1 \text{ dB}}$ est grande. Donc, pour maximiser la plage dynamique utilisable du circuit, il faut donc maximiser α_2/α_4 .

2.2 La compression du détecteur avec un signal à deux fréquences

Nous pouvons également tenter de prédire le point de compression en évaluant la réponse d'un détecteur de puissance à un signal à deux fréquences. Soit un signal de la forme $v_{RF} = \frac{V_{RF}}{2} (\cos((\omega_c + \omega_m)t) + \cos((\omega_c - \omega_m)t))$. L'enveloppe de la puissance de sortie se situe à la fréquence $2f_m$. Idéalement, la puissance à la sortie du détecteur à $2f_m$ est proportionnelle au carré de la puissance du signal injecté. Cependant, nous avons vu dans la partie 1.2.1 que v_{RF}^4 génère également un terme à $2f_m$. Ceci est la cause de la compression de notre enveloppe de puissance du signal. Lorsque $V_{RF} = V_{1 \text{ dB,d}}$, la puissance à la sortie à $2f_m$, qui dépend de α_2 et de α_4 , est 1 dB plus faible que si elle ne dépendait que de α_2 .

Selon la partie 1.2.1, la tension à la sortie à $2f_m$ est :

$$v_s(t)|_{f=2f_m} \approx \left[\alpha_2 \frac{V_{RF}^2}{4} + \alpha_4 \frac{3V_{RF}^4}{16} \right] \cos(2\omega_m t) \quad (\text{A I-20})$$

Par conséquent, au point de compression, on a :

$$10^{-1/10} \frac{1}{2R} \left(\alpha_2 \frac{V_{1\text{ dB,d}}^2}{4} \right)^2 = \frac{1}{2R} \left(\alpha_2 \frac{V_{1\text{ dB,d}}^2}{4} + \frac{3}{16} \alpha_4 V_{1\text{ dB,d}}^4 \right)^2 \quad (\text{A I-21})$$

$$10^{-1/20} \left| \alpha_2 \frac{V_{1\text{ dB,d}}^2}{4} \right| = \left| \alpha_2 \frac{V_{1\text{ dB,d}}^2}{4} + \frac{3}{16} \alpha_4 V_{1\text{ dB,d}}^4 \right| \quad (\text{A I-22})$$

Nous supposons toujours que $\alpha_2 > 0 > \alpha_4$. Pour $V_{RF} = V_{1\text{ dB,d}}$, l'effet quadratique reste prépondérant devant celui de la puissance quatre. Donc, on peut affirmer que $\alpha_2 V_{1\text{ dB,d}}^2/4 > |\frac{3}{16} \alpha_4 V_{1\text{ dB,d}}^4| > 0$. Donc, $\alpha_2 \frac{V_{1\text{ dB,d}}^2}{4} + \frac{3}{16} \alpha_4 V_{1\text{ dB,d}}^4 > 0$.

$$10^{-1/20} \alpha_2 \frac{V_{1\text{ dB,d}}^2}{4} = \alpha_2 \frac{V_{1\text{ dB,d}}^2}{4} - \frac{3}{16} |\alpha_4| V_{1\text{ dB,d}}^4 \quad (\text{A I-23})$$

$$\alpha_2 \frac{V_{1\text{ dB,d}}^2}{4} (1 - 10^{-1/20}) = \frac{3}{16} |\alpha_4| V_{1\text{ dB,d}}^4 \quad (\text{A I-24})$$

$$\frac{16}{12} \frac{\alpha_2}{|\alpha_4|} (1 - 10^{-1/20}) = V_{1\text{ dB,d}}^2 \quad (\text{A I-25})$$

$$V_{1\text{ dB,d}} = 2 \sqrt{\frac{\alpha_2}{3|\alpha_4|} (1 - 10^{-\frac{1}{20}})} \quad (\text{A I-26})$$

(Ce résultat a été vérifié numériquement via Python.)

Le résultat ci-dessus est identique à celui que nous avons vu précédemment.

La mesure du point de compression du détecteur peut donc se faire de deux manières. Une première façon serait d'injecter un signal à amplitude constante et de mesurer la croissance de la tension continue à la sortie en fonction de l'amplitude du signal. Une seconde façon serait d'injecter un signal à deux fréquences, puis de mesurer la croissance de la tension à $2f_m$ en fonction de l'amplitude du signal.

2.3 La compression du mélangeur avec un signal à deux fréquences

Le phénomène de compression a également lieu avec les mélangeurs quadratiques. Le mécanisme est le même que celui que nous avons vu précédemment. Dans la partie 1.2.2, nous avons vu

qu'élever le signal à la puissance quatre générerait des termes à $f_{IF} \pm f_m$. Ces termes sont à l'origine de la compression du signal IF à la sortie du mélangeur. En effet, selon la partie 1.2.2, la tension à la sortie à $f_{IF} + f_m$ (ou $f_{IF} - f_m$) est :

$$v_s(t)|_{f=f_{IF}+f_m} \approx \left[\alpha_2 \frac{V_{RF} V_{LO}}{2} + \alpha_4 \left(\frac{9V_{RF}^3 V_{LO}}{16} + \frac{3V_{RF} V_{LO}^3}{4} \right) \right] \cos((\omega_{IF} + \omega_m)t) \quad (\text{A I-27})$$

Au point de compression à 1 dB, la puissance à la sortie est 1 dB inférieure à ce qu'elle serait si le système était idéal (i.e. pour tout $i > 2$, $\alpha_i = 0$). (On suppose toujours que $\alpha_2 > |\alpha_4| > 0$ et que $\alpha_4 < 0$).

$$10^{\frac{-1}{20}} \frac{1}{2R} \left(\frac{\alpha_2 V_{1 \text{ dB,m}} V_{LO}}{2} \right)^2 = \frac{1}{2R} \left(\frac{\alpha_2 V_{1 \text{ dB,m}} V_{LO}}{2} + \frac{3\alpha_4 V_{1 \text{ dB,m}}^3 V_{LO}}{8} + \frac{3\alpha_4 V_{1 \text{ dB,m}} V_{LO}^3}{4} \right)^2 \quad (\text{A I-28})$$

$$10^{\frac{-1}{20}} \frac{\alpha_2 V_{1 \text{ dB,m}} V_{LO}}{2} = \left| \frac{\alpha_2 V_{1 \text{ dB,m}} V_{LO}}{2} + \frac{3\alpha_4 V_{1 \text{ dB,m}}^3 V_{LO}}{8} + \frac{3\alpha_4 V_{1 \text{ dB,m}} V_{LO}^3}{4} \right| \quad (\text{A I-29})$$

Supposons que $\frac{\alpha_2 V_{1 \text{ dB,m}} V_{LO}}{2} > \left| \frac{9\alpha_4 V_{1 \text{ dB,m}}^3 V_{LO}}{19} \right| + \left| \frac{3\alpha_4 V_{1 \text{ dB,m}} V_{LO}^3}{4} \right|$.

$$10^{\frac{-1}{20}} \frac{\alpha_2 V_{1 \text{ dB,m}} V_{LO}}{2} = \frac{\alpha_2 V_{1 \text{ dB,m}} V_{LO}}{2} - \frac{9|\alpha_4| V_{1 \text{ dB,m}}^3 V_{LO}}{16} - \frac{3|\alpha_4| V_{1 \text{ dB,m}} V_{LO}^3}{4} \quad (\text{A I-30})$$

$$\left(10^{\frac{-1}{20}} - 1 \right) \frac{\alpha_2}{2} = -\frac{9|\alpha_4| V_{1 \text{ dB,m}}^2}{16} - \frac{3|\alpha_4| V_{LO}^2}{4} \quad (\text{A I-31})$$

$$\left(1 - 10^{\frac{-1}{20}} \right) \frac{\alpha_2}{2|\alpha_4|} - \frac{3V_{LO}^2}{4} = \frac{9V_{1 \text{ dB,m}}^2}{16} \quad (\text{A I-32})$$

$$V_{1 \text{ dB,m}}^2 = \frac{16}{9} \left[\left(1 - 10^{\frac{-1}{20}} \right) \frac{\alpha_2}{2|\alpha_4|} - \frac{3V_{LO}^2}{4} \right] \quad (\text{A I-33})$$

$$V_{1 \text{ dB,m}} = \frac{4}{3} \sqrt{\left(1 - 10^{\frac{-1}{20}} \right) \frac{\alpha_2}{2|\alpha_4|} - \frac{3V_{LO}^2}{4}} \quad (\text{A I-34})$$

(Ce résultat a été numériquement vérifié sur Python.)

Cette fois-ci, le point de compression dépend également de V_{LO} . Donc, si l'amplitude du signal LO est excessive, alors cela peut causer une compression prématurée.

3. Le point d'interception

Les produits d'intermodulation d'ordre quatre, qui se retrouvent à des fréquences adjacentes à celles des produits d'intermodulation d'ordre deux, ont un taux d'accroissement supérieur à celui de ces derniers. Par conséquent, il existe un point théorique où les puissances de ces deux types de produits d'intermodulation sont égales. Procédons avec la même méthode que Pozar (2012) (p.515) pour examiner ces points d'interception.

3.1 Pour le détecteur

3.1.1 Calcul de $V_{IP,d}$

Soit un détecteur qui reçoit un signal à deux fréquences. Posons la puissance à la sortie du circuit de détecteur à la fréquence $2f_m$ supposant qu'il n'y ait que l'intermodulation d'ordre 2 :

$$P_{s@2f_m | i>2, \alpha_i=0} = \frac{1}{2R} \left(\alpha_2 V_{RF}^2 / 4 \right)^2 \quad (\text{A I-35})$$

Puis posons la puissance à la sortie du détecteur à la fréquence $4f_m$ supposant qu'il n'ait que l'intermodulation d'ordre 4 :

$$P_{s@4f_m | i>4, \alpha_i=0} = \frac{1}{2R} \left(\alpha_4 \frac{3}{64} V_{RF}^4 \right)^2 \quad (\text{A I-36})$$

En théorie, il existe une amplitude de tension $V_{RF} = V_{IP,d}$ telle que les deux puissances soient égales :

$$\frac{1}{2R}(\alpha_2 V_{IP,d}^2/4)^2 = \frac{1}{2R}(\alpha_4 \frac{3}{64} V_{IP,d}^4)^2 \quad (\text{A I-37})$$

$$\alpha_2 V_{IP,d}^2/4 = |\alpha_4| \frac{3}{64} V_{IP,d}^4 \quad (\text{A I-38})$$

$$V_{IP,d} = 4 \sqrt{\frac{\alpha_2}{3|\alpha_4|}} \quad (\text{A I-39})$$

(Ce résultat a été vérifié numériquement via Python.)

Pour minimiser la distorsion de l'enveloppe de puissance, il faut maximiser le ratio α_2/α_4 . Si nous injectons la formule de

3.1.2 La relation entre $V_{1\text{ dB,d}}$ et $V_{IP,d}$

Nous pouvons également trouver la relation entre $V_{1\text{ dB}}$ et V_{IP} pour le détecteur :

$$\begin{cases} V_{1\text{ dB,d}} &= 2 \sqrt{\frac{\alpha_2}{3|\alpha_4|} (1 - 10^{-1/20})} \\ V_{IP,d} &= 4 \sqrt{\frac{\alpha_2}{3|\alpha_4|}} \end{cases} \quad (\text{A I-40})$$

$$\begin{cases} \frac{\alpha_2}{|\alpha_4|} &= \frac{3V_{1\text{ dB,d}}^2}{4(1-10^{-1/20})} \\ V_{IP,d} &= 4 \sqrt{\frac{3V_{1\text{ dB,d}}^2}{3 \times 4(1-10^{-1/20})}} \end{cases} \quad (\text{A I-41})$$

Donc

$$\frac{V_{IP,d}}{V_{1\text{ dB,d}}} = \frac{2}{\sqrt{1 - 10^{-1/20}}} \approx 6,065 = 15,66 \text{ dB} \quad (\text{A I-42})$$

(Ce résultat a été numériquement vérifié sur Python.)

3.2 Pour le mélangeur

3.2.1 Calcul de $V_{IP,m}$

Au point d'interception, les produits d'intermodulation d'ordre quatre néfastes à $f_{IF} \pm 3f_m$ ont la même amplitude que le signal transposé à $f_{IF} \pm f_m$. Il existe une amplitude du signal RF $V_{IP,m}$ telle que $P_{f_{IF} \pm f_m | i>2, \alpha_i=0}(V_{RF} = V_{IP,m}) = P_{f_{IF} \pm 3f_m | i>4, \alpha_i=0}(V_{RF} = V_{IP,m})$.

$$\frac{1}{2R} \left(\frac{\alpha_2 V_{IP,m} V_{LO}}{2} \right)^2 = \frac{1}{2R} \left(\frac{3\alpha_4 V_{IP,m}^3 V_{LO}}{16} \right)^2 \quad (\text{A I-43})$$

$$\frac{\alpha_2 V_{IP,m} V_{LO}}{2} = \frac{3|\alpha_4| V_{IP,m}^3 V_{LO}}{16} \quad (\text{A I-44})$$

$$V_{IP,m} = 4 \sqrt{\frac{\alpha_2}{6|\alpha_4|}} \quad (\text{A I-45})$$

(Ce résultat a été numériquement vérifié sur Python.)

3.2.2 La relation entre $V_{1 \text{ dB,m}}$ et $V_{IP,m}$

Ensuite, la relation entre $V_{IP,m}$ et V_{1dB} pour le mélangeur peut être calculé. D'abord, on peut exprimer $\alpha_2/|\alpha_4|$ en fonction de $V_{IP,m}$:

$$\frac{V_{IP,m}^2}{16} = \frac{\alpha_2}{6|\alpha_4|} \quad (\text{A I-46})$$

$$\frac{6V_{IP,m}^2}{16} = \frac{\alpha_2}{|\alpha_4|} \quad (\text{A I-47})$$

Ensuite, nous pouvons injecter ceci dans (A I-39) :

$$V_{1 \text{ dB,m}} = \frac{4}{3} \sqrt{\left(1 - 10^{\frac{-1}{20}}\right) \frac{6V_{IP,m}^2}{2 \times 16} - \frac{3}{2} V_{LO}^2} \quad (\text{A I-48})$$

$$\frac{9}{16} V_{1 \text{ dB,m}}^2 = \left(1 - 10^{\frac{-1}{20}}\right) \frac{6V_{IP,m}^2}{32} - \frac{3}{2} V_{LO}^2 \quad (\text{A I-49})$$

$$\frac{32}{6} \left[\frac{9}{16} V_{1 \text{ dB,m}}^2 + \frac{3}{2} V_{LO}^2 \right] = V_{IP,m}^2 (1 - 10^{\frac{-1}{20}}) \quad (\text{A I-50})$$

$$V_{IP,m} = \sqrt{\frac{3V_{1 \text{ dB,m}}^2 + 8V_{LO}^2}{1 - 10^{\frac{-1}{20}}}} \quad (\text{A I-51})$$

(Ce résultat a été numériquement vérifié sur Python.)

4. L'analyse du circuit ICSTSEM1

4.1 La tension continue à la sortie du circuit

Exprimons la tension V_r à la sortie du circuit lorsqu'il n'y a aucune stimulation par un signal à l'entrée :

$$V_r = R_1 I_B \quad (\text{A I-52})$$

avec

$$I_B = \frac{1}{2} C_{ox} \mu_p \frac{W_4}{L_4} (V_{gs,4} - V_{th})^2 \quad (\text{A I-53})$$

or $V_{gs,4} = V_{gs,3}$. Rappelons-nous que

$$I_A = \frac{1}{2} C_{ox} \mu_p \frac{W_3}{L_3} (V_{gs,3} - V_{th})^2 \quad (\text{A I-54})$$

donc

$$V_{gs,3} = \sqrt{\frac{2I_A}{C_{ox} \mu_p W_3 / L_3}} + V_{th} \quad (\text{A I-55})$$

Or, $I_A = I_{ds,1} + I_{ds,2} = 2I_{ds,1}$, donc :

$$V_{gs,3} = \sqrt{\frac{2C_{ox} \mu_n W_1 / L_1 (V_{gs,1} - V_{th})^2}{C_{ox} \mu_p W_3 / L_3}} + V_{th} \quad (\text{A I-56})$$

Finalement,

$$V_r = R_1 \frac{1}{2} C_{ox} \mu_p \frac{W_4}{L_4} \left(\sqrt{\frac{2\mu_n W_1 / L_1 (V_{gs,1} - V_{th})^2}{\mu_p W_3 / L_3}} + V_{th} - V_{th} \right)^2 \quad (\text{A I-57})$$

$$= R_1 C_{ox} \underbrace{\frac{W_4 / L_4}{W_3 / L_3}}_K \frac{W_1}{L_1} \underbrace{(V_{gs,1} - V_{th})^2}_{\text{Polarisation}} \quad (\text{A I-58})$$

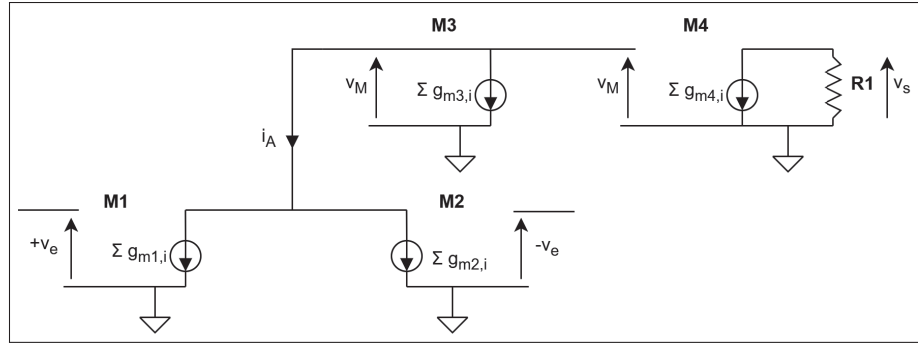


Figure-A I-1 Schéma du circuit intégré en modèle de petits signaux sans modulation de la longueur du canal

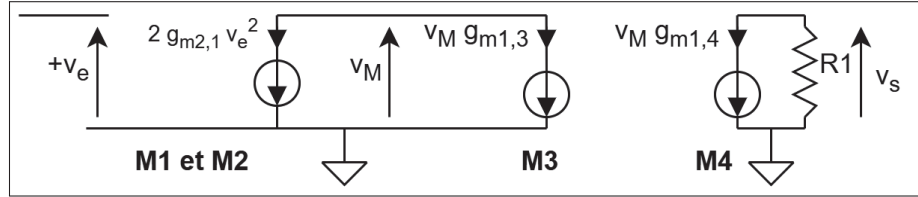


Figure-A I-2 Schéma simplifié du circuit intégré en modèle de petits signaux sans modulation de la longueur du canal

4.2 Le gain en tension quadratique idéal

Dans un premier temps, nous négligerons l'effet de modulation de canal. Donc, nous supposons que $r_o \rightarrow \infty$. Sur la figure I-1, nous avons l'équivalent en modèle de petits signaux du schéma de la figure 3.2. Nous pouvons encore simplifier le schéma grâce à la symétrie de l'étage d'entrée et aux hypothèses, que nous avons posées dans la partie 3.1.2.2. Cette simplification est représentée dans la figure I-2. D'abord, nous pouvons calculer le courant i_A tiré par M1 et M2 :

$$i_A = i_{gm,1} + i_{gm,2} = \sum_{i=1}^{\infty} g_{mi,1} v_e^i + g_{mi,2} (-v_e)^i \quad (\text{A I-59})$$

Grâce à la symétrie de M1 et de M2, nous pouvons annuler les termes d'ordres impairs :

$$i_A = 2 \sum_{i=1}^{\infty} g_{m(2i),1} v_e^{2i} \quad (\text{A I-60})$$

Ensuite, nous pouvons supposer que le terme d'ordre deux est prépondérant devant ceux d'ordres supérieurs :

$$i_A \approx 2g_{m2,1} v_e^2 \quad (\text{A I-61})$$

Ensuite, nous pouvons affirmer, grâce au schéma simplifié ci-dessus, que

$$\begin{cases} v_s &= -R_1 g_{m1,4} v_M \\ g_{m1,3} v_M &= -2g_{m2,1} v_e^2 \end{cases} \quad (\text{A I-62})$$

Par conséquent,

$$v_s = +2R_1 g_{m1,4} \frac{g_{m2,1} v_e^2}{g_{m1,3}} \quad (\text{A I-63})$$

$$\alpha_2 = \frac{v_s}{v_e^2} = 2R_1 K g_{m2,1} \quad (\text{A I-64})$$

Où $K = g_{m1,4}/g_{m1,3}$. Nous avons ainsi obtenu une expression simplifiée du gain en tension quadratique.

4.3 Le gain en tension quadratique avec la modulation de la longueur du canal

Désormais, supposons que r_o est une grandeur finie pour chaque transistor. Alors, nous pouvons dessiner le circuit équivalent en petits signaux, tel qu'il est illustré dans la figure I-3. Ce schéma peut être encore simplifié : cela donne la figure I-4.

Grâce au circuit équivalent petit-signal simplifié, nous pouvons exprimer le gain du circuit.

Tout d'abord,

$$v_s = -(r_{o4} || R_1) g_{m1,4} v_M \quad (\text{A I-65})$$

Avec

$$v_M = -\left(\frac{r_{o1}}{2} || r_{o3}\right) (2 g_{m2,1} v_e^2 + g_{m1,3} v_M) \quad (\text{A I-66})$$

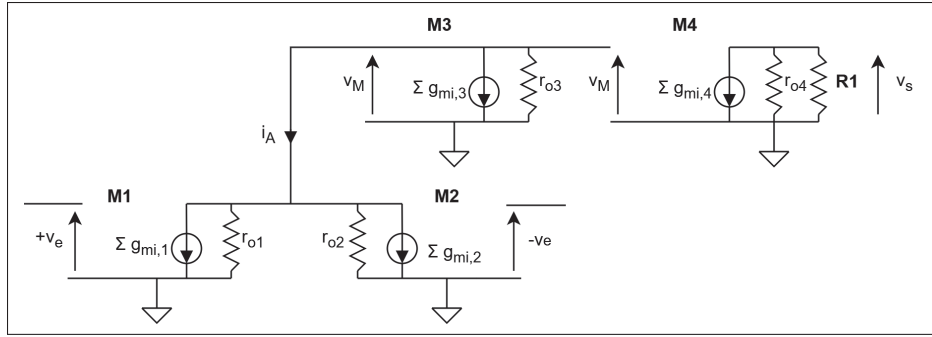


Figure-A I-3 Circuits équivalents du circuits ICSTSEM1 en petit-signal

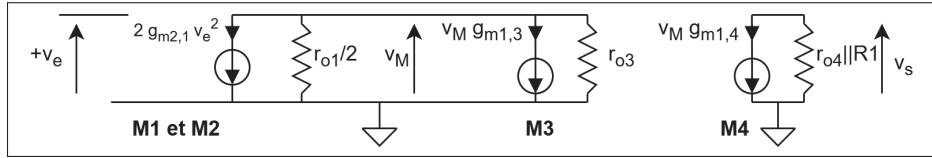


Figure-A I-4 Circuit équivalent petit-signal d'ICSTSEM1 simplifié

Cela donne :

$$v_M = \frac{-2\left(\frac{r_{o1}}{2} \parallel r_{o3}\right) g_{m2,1} v_e^2}{1 + \left(\frac{r_{o1}}{2} \parallel r_{o3}\right) g_{m1,3}} \quad (\text{A I-67})$$

En injectant (A I-67) dans (A I-65), nous obtenons :

$$v_s = \frac{2 (r_{o4} \parallel R_1) \left(\frac{r_{o1}}{2} \parallel r_{o3}\right) g_{m1,4} g_{m2,1}}{1 + \left(\frac{r_{o1}}{2} \parallel r_{o3}\right) g_{m1,3}} v_e^2 \quad (\text{A I-68})$$

Le gain en tension quadratique du circuit vaut donc :

$$\alpha_2 = \frac{v_s}{v_e^2} = \frac{2 (r_{o4} \parallel R_1) \left(\frac{r_{o1}}{2} \parallel r_{o3}\right) g_{m1,4} g_{m2,1}}{1 + \left(\frac{r_{o1}}{2} \parallel r_{o3}\right) g_{m1,3}} \quad (\text{A I-69})$$

L'équation ci-dessus est cohérente avec (A I-64), car si les r_o tendent vers l'infini, alors nous obtenons la même formule que (A I-64). De plus, nous avons comparé cette formule à la simulation dans la partie 3.1.4.

Tentons d'exprimer (A I-69) en fonction des dimensions du circuit et des paramètres de polarisation :

$$\alpha_2 = \frac{2 (r_{o4} || R_1) \left(\frac{\frac{1}{2\lambda_1 I_A/2} \frac{1}{\lambda_3 I_A}}{\frac{1}{2\lambda_1 I_A/2} + \frac{1}{\lambda_3 I_A}} \right) \frac{2I_B}{V_{ov,4}} g_{m2,1}}{1 + \left(\frac{\frac{1}{2\lambda_1 I_A/2} \frac{1}{\lambda_3 I_A}}{\frac{1}{2\lambda_1 I_A/2} + \frac{1}{\lambda_3 I_A}} \right) \frac{2I_A}{V_{ov,3}}} \quad (\text{A I-70})$$

$$= \frac{2 (r_{o4} || R_1) \frac{\frac{1}{\lambda_1 \lambda_3 I_A^2}}{\frac{\lambda_1 I_A + \lambda_3 I_A}{\lambda_1 \lambda_3 I_A^2}} \frac{2I_B}{V_{ov,4}} g_{m2,1}}{1 + \frac{\frac{1}{\lambda_1 \lambda_3 I_A^2}}{\frac{\lambda_1 I_A + \lambda_3 I_A}{\lambda_1 \lambda_3 I_A^2}} \frac{2I_A}{V_{ov,3}}} \quad (\text{A I-71})$$

$$= \frac{2 (r_{o4} || R_1) \frac{1}{I_A (\lambda_1 + \lambda_3)} \frac{2I_B}{V_{ov,4}} g_{m2,1}}{1 + \frac{1}{I_A (\lambda_1 + \lambda_3)} \frac{2I_A}{V_{ov,3}}} \quad (\text{A I-72})$$

$$= \frac{4 (r_{o4} || R_1) \frac{K g_{m2,1}}{V_{ov,4} (\lambda_1 + \lambda_3)}}{1 + \frac{2}{V_{ov,3} (\lambda_1 + \lambda_3)}} \quad (\text{A I-73})$$

Ensuite, la polarisation de M1 et M2 détermine le courant I_A qui est tiré. Ce courant dicte la tension des grilles de M3 et de M4. En effet :

$$I_A/2 = \frac{1}{2} \mu_n C_{ox} \frac{W_1}{L_1} V_{ov,1}^2 = \frac{1}{2} \mu_n C_{ox} \frac{W_2}{L_2} V_{ov,2}^2 \quad (\text{A I-74})$$

et

$$I_A = \frac{1}{2} \mu_p C_{ox} \frac{W_3}{L_3} V_{ov,3}^2 \quad (\text{A I-75})$$

Puisque M3 et M4 partagent la même tension de grille :

$$V_{ov,4} = V_{ov,3} = \sqrt{\frac{2I_A}{\mu_p C_{ox} \frac{W_3}{L_3}}} \quad (\text{A I-76})$$

$$= \sqrt{\frac{2\mu_n C_{ox} \frac{W_1}{L_1} V_{ov,1}^2}{\mu_p C_{ox} \frac{W_3}{L_3}}} \quad (\text{A I-77})$$

$$= \sqrt{\frac{2\mu_n \frac{W_1}{L_1}}{\mu_p \frac{W_3}{L_3}}} V_{ov,1} \quad (\text{A I-78})$$

Nous pouvons injecter cette dernière expression dans l'expression du gain :

$$\alpha_2 = \frac{4 (r_{o4} || R_1) \frac{K_{gm2,1}}{(\lambda_1 + \lambda_3) V_{ov,1}} \sqrt{\frac{\mu_p \frac{W_3}{L_3}}{2\mu_n \frac{W_1}{L_1}}}}{1 + \frac{2}{(\lambda_1 + \lambda_3) V_{ov,1}} \sqrt{\frac{\mu_p \frac{W_3}{L_3}}{2\mu_n \frac{W_1}{L_1}}}} \quad (\text{A I-79})$$

Enfin exprimons r_{o4} en fonction des paramètres des composants, sachant que :

$$r_{o4} = \frac{1}{\lambda_4 I_B} = \frac{1}{\lambda_4 K I_A} \quad (\text{A I-80})$$

$$= \frac{1}{\lambda_4 K \mu_n C_{ox} \frac{W_1}{L_1} V_{ov,1}^2} \quad (\text{A I-81})$$

Par substitution, cela donne :

$$\alpha_2 = \frac{4 \left(\frac{1}{\lambda_4 K \mu_n C_{ox} \frac{W_1}{L_1} V_{ov,1}^2} || R_1 \right) \frac{K g_{m2,1}}{(\lambda_1 + \lambda_3) V_{ov,1}} \sqrt{\frac{\mu_p \frac{W_3}{L_3}}{2 \mu_n \frac{W_1}{L_1}}}}{1 + \frac{2}{(\lambda_1 + \lambda_3) V_{ov,1}} \sqrt{\frac{\mu_p \frac{W_3}{L_3}}{2 \mu_n \frac{W_1}{L_1}}}} \quad (\text{A I-82})$$

$$= \frac{4 \left(\frac{R_1}{1 + R_1 \lambda_4 K \mu_n C_{ox} \frac{W_1}{L_1} V_{ov,1}^2} \right) \frac{K g_{m2,1}}{(\lambda_1 + \lambda_3) V_{ov,1}} \sqrt{\frac{\mu_p \frac{W_3}{L_3}}{2 \mu_n \frac{W_1}{L_1}}}}{1 + \frac{2}{(\lambda_1 + \lambda_3) V_{ov,1}} \sqrt{\frac{\mu_p \frac{W_3}{L_3}}{2 \mu_n \frac{W_1}{L_1}}}} \quad (\text{A I-83})$$

Pour vérifier l'équation ci-dessus, nous pouvons calculer la limite lorsque les λ tendent vers zéro, c'est-à-dire lorsque la modulation de la longueur du canal devient négligeable.

$$\lim_{\lambda \rightarrow 0} \frac{4 \left(\frac{R_1}{1 + R_1 \lambda_4 K \mu_n C_{ox} \frac{W_1}{L_1} V_{ov,1}^2} \right) \frac{K g_{m2,1}}{(\lambda_1 + \lambda_3) V_{ov,1}} \sqrt{\frac{\mu_p \frac{W_3}{L_3}}{2 \mu_n \frac{W_1}{L_1}}}}{1 + \frac{2}{(\lambda_1 + \lambda_3) V_{ov,1}} \sqrt{\frac{\mu_p \frac{W_3}{L_3}}{2 \mu_n \frac{W_1}{L_1}}}} = 2 R_1 K g_{m2,1} \quad (\text{A I-84})$$

Nous retrouvons la même formule que (A I-64), où nous avons négligé l'effet de modulation de la longueur du canal. Nous avons ainsi vérifié la validité de (A I-83).

5. Le taux de réjection de la porteuse

Étudions ce qui se passe dans le circuit lorsque les signaux injectés ne sont pas parfaitement symétriques et lorsqu'il y a une asymétrie dans le circuit.

Soient $v_{e,1}(t)$ et $v_{e,2}(t)$ les tensions des signaux respectivement injectés aux grilles de M1 et de M2. Pour simplifier l'analyse, supposons que les signaux soient sinusoïdaux sans modulation d'amplitude. Supposons qu'ils ne soient pas exactement déphasés à 180° , qu'ils n'aient pas

exactement la même amplitude et que les transistors NMOS soient légèrement asymétriques. Posons d'abord l'erreur de phase ϕ en radians entre les signaux et l'erreur relative d'amplitude entre les signaux α .

$$v_{e,1}(t) = V_{RF} \cos(\omega_c t) \quad (\text{A I-85})$$

$$v_{e,2}(t) = -V_{RF}(1 + \alpha) \cos(\omega_c t + \phi) \quad (\text{A I-86})$$

Si le *balun* avant les entrées du circuit est parfait, alors $\phi = 0$ et $\alpha = 0$. Mais en réalité, ϕ et α sont non nuls.

Puis M1 et M2 connaissent de légères différences de fabrication, donc $\forall i \in \mathbb{N}^*$, $g_{mi,2} = (1 + \gamma_i)g_{mi,1}$. Voyons quels courants tirent M1 et M2. Cependant, pour des raisons de simplicité, négligeons les termes de transconductance supérieurs à l'ordre 2.

$$i_{gm,1} = \sum_{i=1}^{\infty} g_{mi,1} v_{in,1}^i \quad (\text{A I-87})$$

$$\approx g_{m1,1} v_{in,1} + g_{m2,1} v_{in,1}^2 \quad (\text{A I-88})$$

$$\approx g_{m1,1} V_{RF} \cos(\omega_c t) + g_{m2,1} [V_{RF} \cos(\omega_c t)]^2 \quad (\text{A I-89})$$

$$\approx g_{m1,1} V_{RF} \cos(\omega_c t) + g_{m2,1} V_{RF}^2 \left[\frac{1 + \cos(2\omega_c t)}{2} \right] \quad (\text{A I-90})$$

et

$$i_{gm,2} = \sum_{i=1}^{\infty} g_{mi,2} v_{in,2}^i \quad (\text{A I-91})$$

$$\approx g_{m1,2} v_{in,2}^+ + g_{m2,2} v_{in,2}^2 \quad (\text{A I-92})$$

$$\approx -g_{m1,2}(1 + \alpha)V_{RF} \cos(\omega_c t + \phi) + g_{m2,2} [-(1 + \alpha)V_{RF} \cos(\omega_c t + \phi)]^2 \quad (\text{A I-93})$$

$$\begin{aligned} &\approx -g_{m1,1}(1 + \gamma_1)(1 + \alpha)V_{RF} \cos(\omega_c t + \phi) \\ &+ (1 + \gamma_2)g_{m2,1}(1 + \alpha)^2 V_{RF}^2 \left[\frac{1 + \cos(2\omega_c t + 2\phi)}{2} \right] \end{aligned} \quad (\text{A I-94})$$

Puis, les drains de M1 et de M2 étant reliés, la somme des courants donne i_A .

$$i_A = i_{gm,1} + i_{gm,2} \quad (\text{A I-95})$$

$$\begin{aligned} &= g_{m1,1} V_{RF} [\cos(\omega_c t) - (1 + \gamma_1)(1 + \alpha) \cos(\omega_c t + \phi)] \\ &+ g_{m2,1} V_{RF}^2 \left[\frac{1 + \cos(2\omega_c t)}{2} + (1 + \gamma_2)(1 + \alpha)^2 \frac{1 + \cos(2\omega_c t + 2\phi)}{2} \right] \end{aligned} \quad (\text{A I-96})$$

Posons $u = (1 + \gamma_1)(1 + \alpha)$ et $v = (1 + \gamma_2)(1 + \alpha)^2$.

$$\begin{aligned} i_A &= g_{m1,1} V_{RF} [\cos(\omega_c t) - u \cos(\omega_c t + \phi)] \\ &+ g_{m2,1} V_{RF}^2 \left[\frac{1 + \cos(2\omega_c t)}{2} + v \frac{1 + \cos(2\omega_c t + 2\phi)}{2} \right] \end{aligned} \quad (\text{A I-97})$$

$$\begin{aligned} &= g_{m1,1} V_{RF} [\cos(\omega_c t) - u (\cos(\omega_c t) \cos(\phi) - \sin(\omega_c t) \sin(\phi))] \\ &+ g_{m2,1} V_{RF}^2 \left[\frac{1}{2}(1 + v) + \frac{\cos(2\omega_c t) + v \cos(2\omega_c t + 2\phi)}{2} \right] \end{aligned} \quad (\text{A I-98})$$

$$= g_{m1,1} V_{RF} [\cos(\omega_c t)(1 - u \cos(\phi)) - u \sin(\omega_c t) \sin(\phi)] + g_{m2,1} V_{RF}^2 \left[\frac{1}{2}(1 + v) + \dots \right] \quad (\text{A I-99})$$

Posons R et ψ tel que

$$i_A = g_{m1,1} V_{RF} R \cos(\omega_c t - \psi) + g_{m2,1} V_{RF}^2 \left[\frac{1}{2}(1 + v) + \dots \right] \quad (\text{A I-100})$$

avec

$$\begin{cases} R &= \sqrt{(1 - u \cos(\phi))^2 + (-u \sin(\phi))^2} = \sqrt{1 - 2u \cos(\phi) + u^2} \\ \tan(\psi) &= \frac{-u \sin(\phi)}{1 - u \cos(\phi)} \end{cases} \quad (\text{A I-101})$$

(Les équations ci-dessus ont également été vérifiées numériquement sur Python)

La tension à la sortie du signal est proportionnelle à i_A donc il n'y a pas besoin de développer davantage nos équations. Nous constatons qu'il reste un terme à la fréquence f_c dont l'amplitude est proportionnelle à V_{RF} . C'est le signal injecté rémanent à cause de la symétrie imparfaite des

signaux et des transistors. L'amplitude de ce terme est d'autant plus importante que ϕ , l'erreur de phase, α , l'erreur d'amplitude et γ_1 , l'erreur de fabrication, sont importantes. Cependant, la bonne nouvelle est que l'amplitude de l'enveloppe de puissance, à 0 Hz, n'est pas affectée par ϕ l'erreur de phase. En effet, le circuit continue de fonctionner comme un détecteur, peu importe l'erreur de phase. En théorie, ce circuit fonctionnerait même si les tensions aux entrées n'avaient pas de déphasage du tout, et le gain de conversion ou la responsivité ne seraient quand même pas affectés.

LE CIRCUIT INTÉGRÉ

1. Un aperçu du circuit intégré au complet

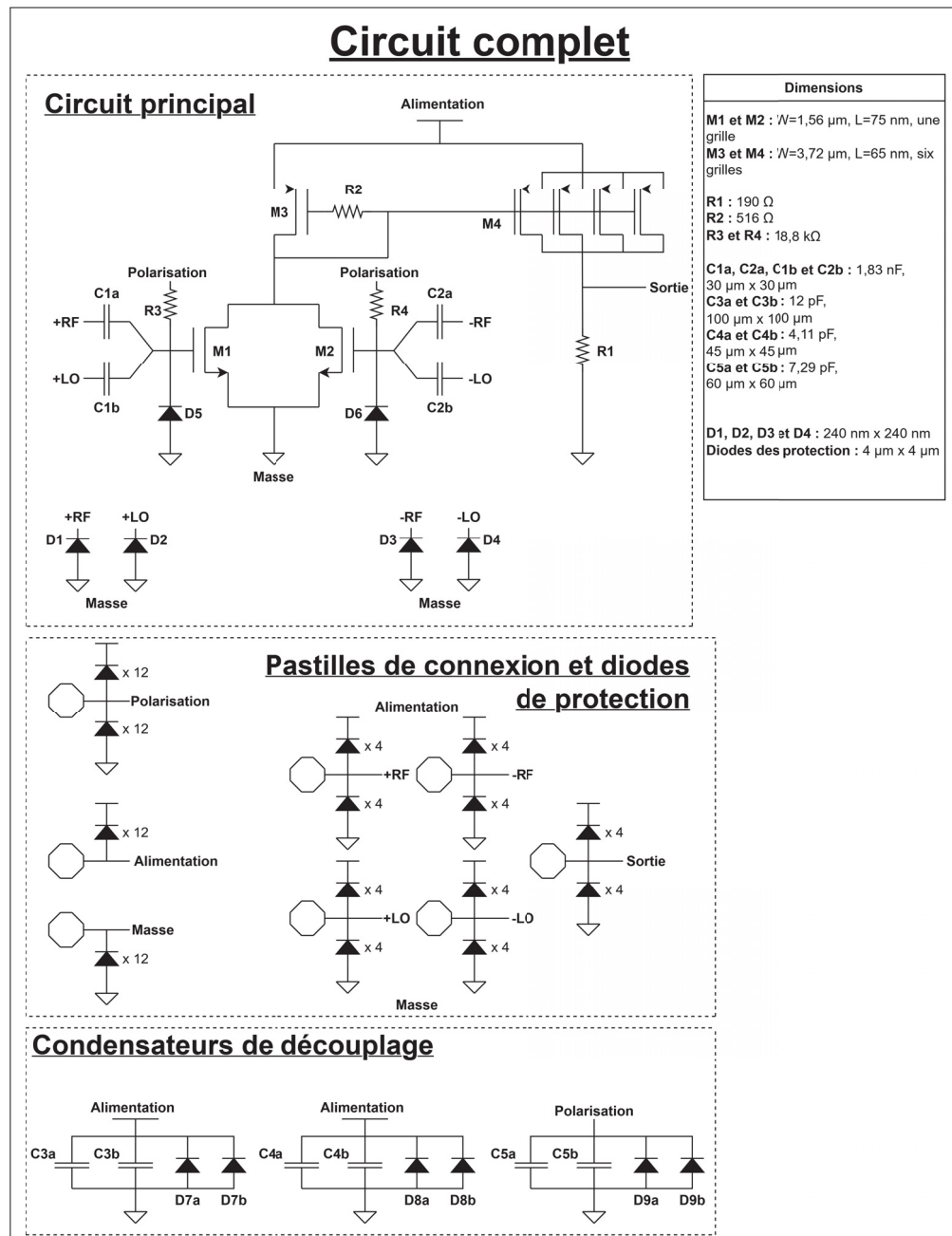


Figure-A II-1 Schéma du circuit intégré ICSTSEM1 au complet

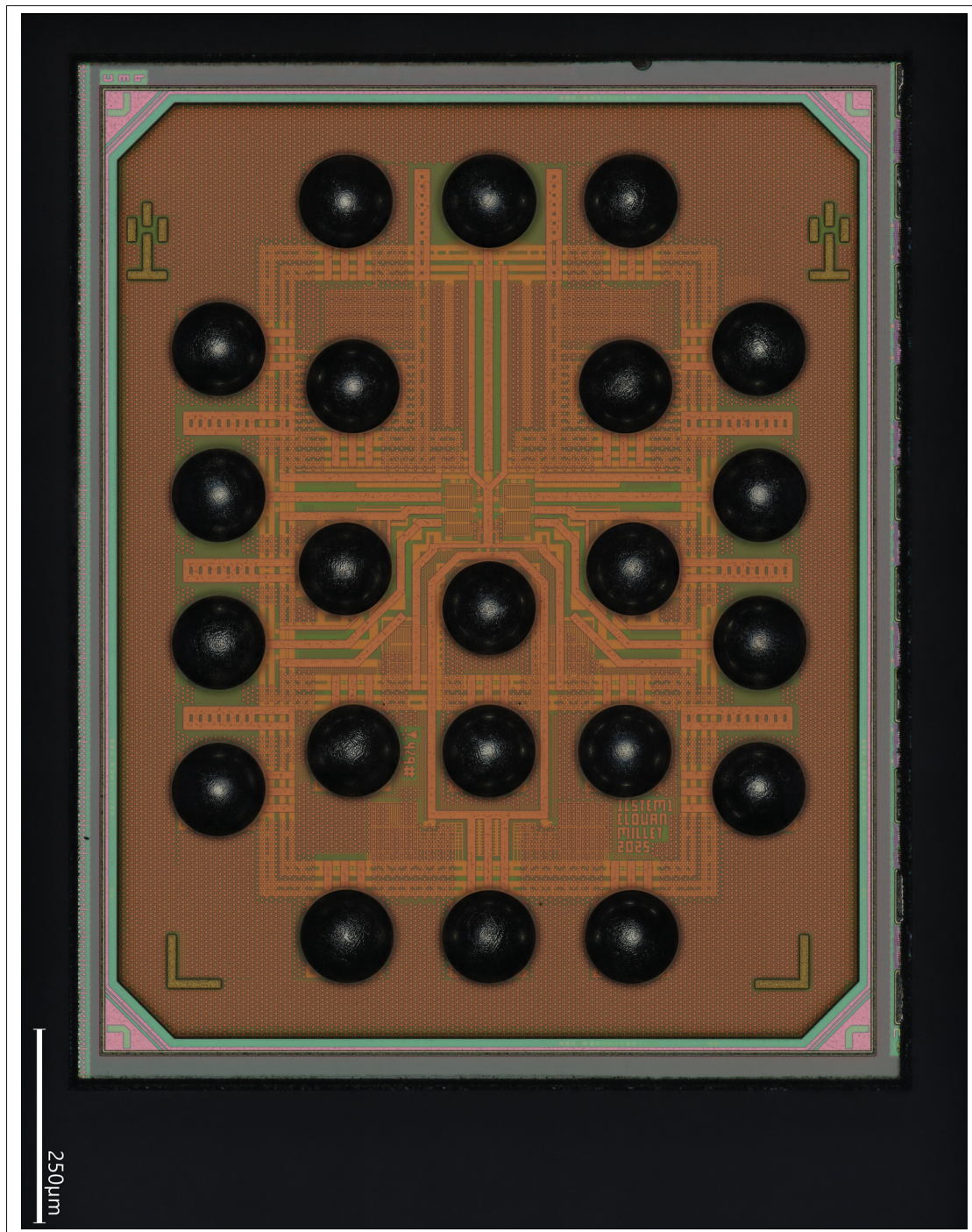


Figure-A II-2 Photographie au microscope du circuit intégré ICSTSEM1

2. La structure physique du circuit intégré

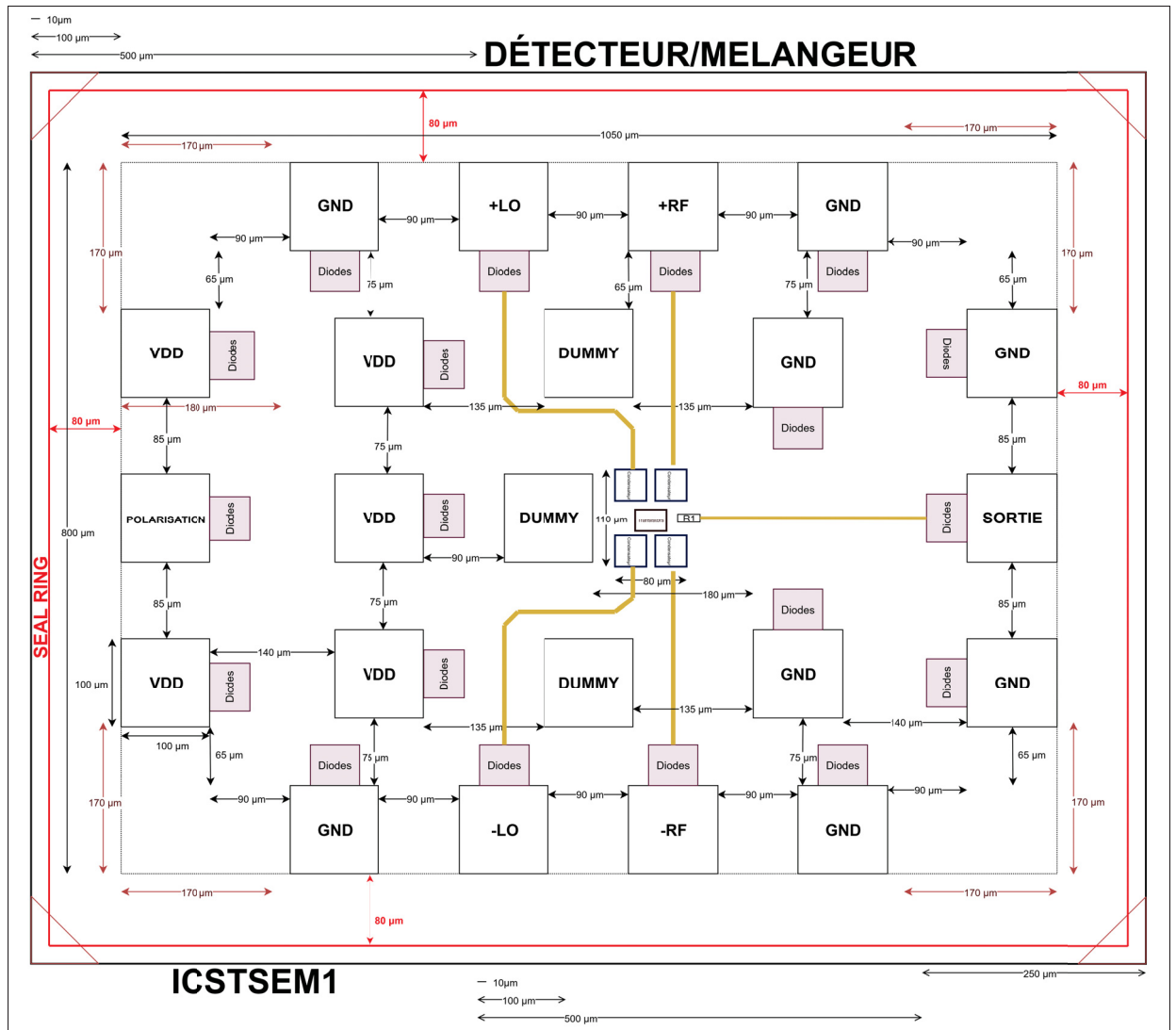


Figure-A II-3 Le plan du circuit intégré ICSTSEM1

2.1 Les couches métalliques

Le procédé utilisé est celui du 65 nm GP (*general purpose*) de TSMC, avec une couche de silicium polycristallin (*polysilicium* en anglais) et neuf couches de métal (1P9M), comme représentées sur la figure II-4. La neuvième couche de métal, celle la plus haute, est dite ultra épaisse (*ultra thick metal*). Elle admet moins de résistivité et permet la création de spirales inductives.

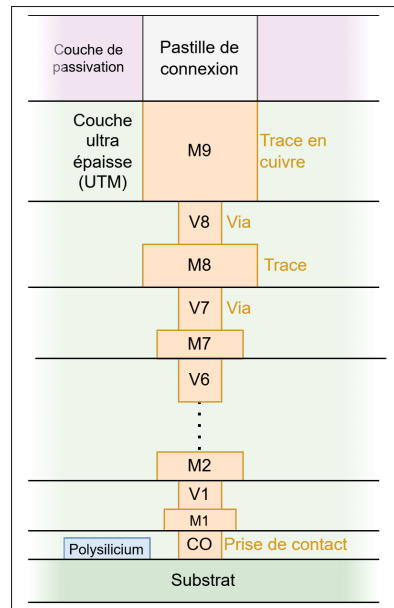


Figure-A II-4 Les différentes couches présentes dans la technologie de fabrication de 65 nm de TSMC utilisée

2.2 La structure des transistors

Comme illustré sur la figure II-5, les transistors NMOS et PMOS sont formés de zones de dopage positive et négative, formant le drain et la source, séparés par un canal contrôlé par la grille. Ces zones de dopage se trouvent dans un puits dont le type de dopage est complémentaire.

Les *guard rings* (ou les anneaux protecteurs) sont des structures ayant pour but de protéger les transistors du bruit ou de l'effet *latchup*. Ce dernier est un phénomène potentiellement catastrophique pour les transistors, car il peut entièrement détruire le circuit. Il a tendance à avoir lieu au sein des transistors directement reliés à l'extérieur de la puce. Les *guard rings* permettent de l'éviter en offrant une connexion à basse résistance entre les puits, l'alimentation et la masse. La structure est constituée d'un puits avec des prises de contact encerclant le circuit à protéger. Le puits en question doit être de dopage opposé à celui du substrat, et son potentiel doit être tel que le puits et le substrat forment une diode bloquante. Le procédé de fabrication utilisé

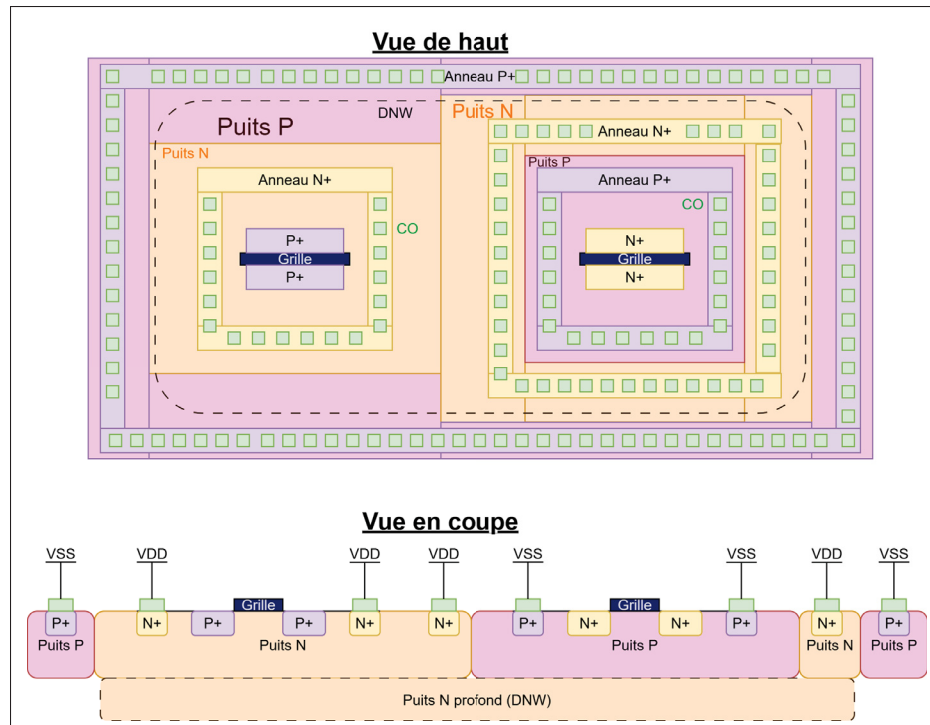


Figure-A II-5 La structure des transistors protégés par des puits

permet l'emploi de *Deep N-Well* (DNW) : des puits de fort dopage négatif sous-jacents. Le DNW permet de protéger les transistors du bruit venant d'en dessous (Weste & Harris (2011)).

2.3 La structure des diodes de protection

Les diodes, comme les transistors, sont formées de zones de dopage positif et négatif. Leur structure est décrite par la figure II-6. Deux types de diodes peuvent être formés, selon le dopage du puits entourant la diode. Elles sont appelées dans la librairie de composants de TSMC « PDIO » et « NDIO ». Le concepteur du circuit doit choisir entre ces deux types de diode, selon la tension de polarisation prévue. Enfin, les nombreuses surfaces de contact entre ces zones dopées et les couches métalliques offrent au courant un chemin à faible résistance lorsque le circuit subit un choc électrostatique.

2.4 Les carrés de remplissage

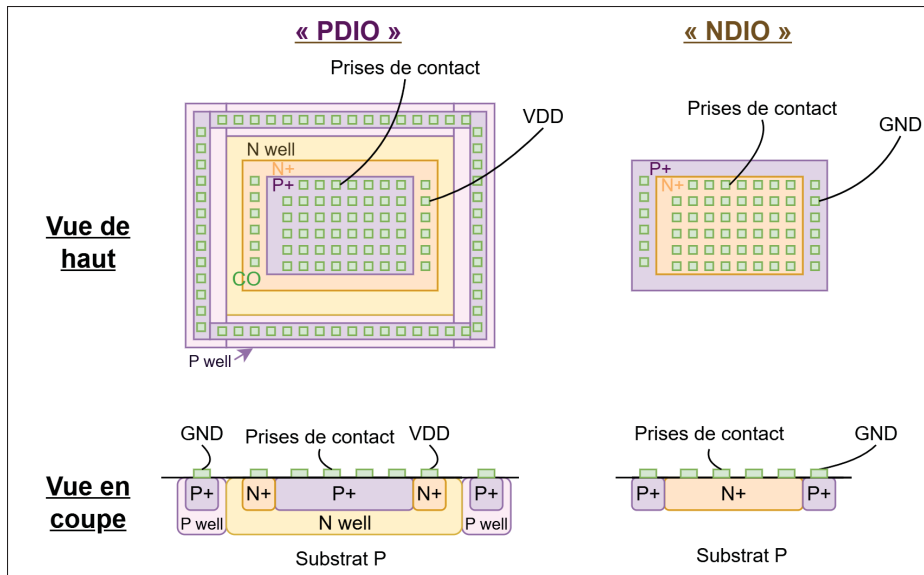


Figure-A II-6 Structure physique des diodes de protection

Pour assurer un bon rendement de fabrication, TSMC exige que chaque couche de silicium polycristallin et de métal atteigne un certain seuil de densité : un certain pourcentage de la surface totale de la puce doit être recouvert sur chacune des couches. Cependant, il n'est pas possible de simplement recouvrir chaque couche par un plan solide et continu de silicium polycristallin ou de métal, comme on le ferait sur un circuit imprimé. Par conséquent, tout l'espace qui n'est pas utilisé par des composants ou des traces est occupé par un motif régulier de carrés d'à peu près $1\ \mu\text{m} \times 1\ \mu\text{m}$. Ces remplissages en métal s'appellent en anglais des *dummies*, car ils ne servent à rien électriquement parlant. Les traces transportant des signaux de haute fréquence sont tout de même vulnérables à un effet de couplage avec ces carrés. Comme illustré sur la figure II-8, les traces conduisant des signaux de haute fréquence ont donc été séparées des carrés par des traces reliées à la masse, bien que cela engendre une capacité parasite.

D'ailleurs, Xu *et al.* (2014) note qu'ils ont dimensionné leurs pastilles de connexion grâce à des simulations électromagnétiques, parce que la présence de ces carrés de remplissage autour des pastilles dégrade la performance de leur circuit.

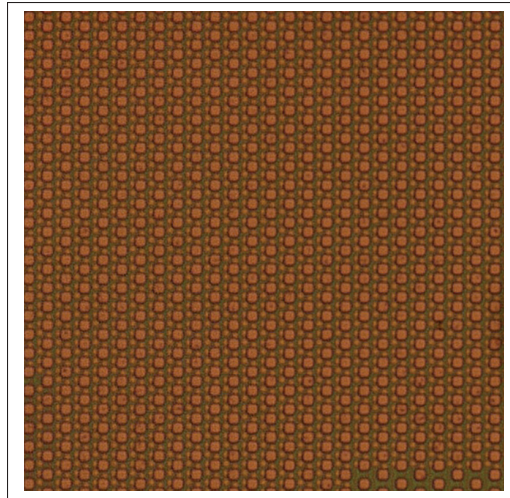


Figure-A II-7 Photographie au microscope des carrés de remplissage

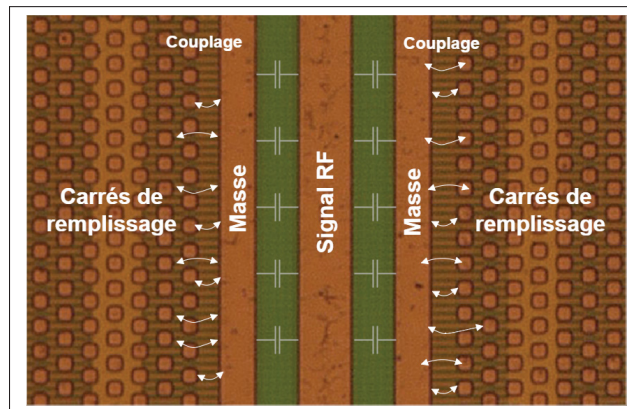


Figure-A II-8 Photographie au microscope d'une trace sensible protégée par la masse

3. Les techniques de montage d'une puce sur un circuit imprimé

Le câblage par fil, représenté dans la figure II-9, consiste à relier les pastilles du circuit intégré à celles du circuit imprimé grâce à des fils en or très fins. Ces fils sont soudés grâce à des ultrasons ou par chauffage. Cette technique nécessite que les pastilles de connexion de la puce se trouvent sur ses bords. Ces fils ajoutent cependant une certaine inductance parasite de l'ordre de quelques

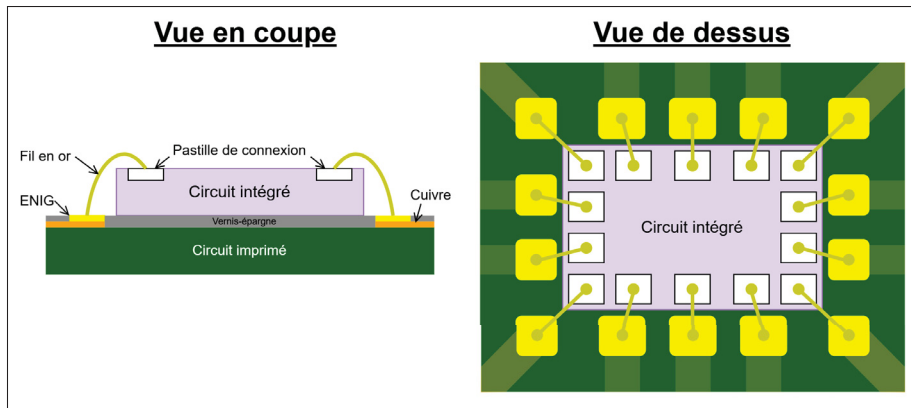


Figure-A II-9 Technique du câblage par fils

dixièmes de nanohenrys, qui devient non négligeable à des fréquences micro-ondes. De plus, l'effet de peau rend ces fils très fins résistifs. (Cui, Yuan, Liu, Cui & Liu (2022))

La technique de la puce retournée consiste à placer des billes de soudure sur les pastilles de connexion qui couvrent toute la surface de la puce. Ces billes de soudure sont notamment visibles sur la photo de la figure II-2. Puis, on aligne précisément ces billes avec les pastilles du circuit imprimé, qui sont illustrées dans les figures III-2 et III-3. Une fois le circuit intégré placé sur le circuit imprimé, pour souder les circuits ensemble, on chauffe le circuit imprimé par dessous, grâce à une plaque chauffante, et on souffle également de l'air chaud sur le dessus de la puce. Cette étape est illustrée dans la figure II-10. Cette technique permet (et oblige) de placer les pastilles de connexion sur toute la surface du circuit intégré. Elle permet donc à la puce d'avoir un très grand nombre de ports d'entrée et de sortie. De plus, cette technique permet d'assurer une continuité entre les lignes de transmission du circuit imprimé et les traces du circuit intégré, car à aucun moment un signal ne traverse un fil entièrement exposé. Il est donc très intéressant pour l'intégrité du signal à haute fréquence. Au moment de la conception du circuit intégré, l'intégrité du signal fut la raison principale pour laquelle nous avons choisi la technique de la puce retournée.

Pour la technique de la puce retournée, le fabricant de circuits TSMC exigeait que toutes les pastilles de connexion soient en forme octogonale. De plus, il imposait des règles strictes

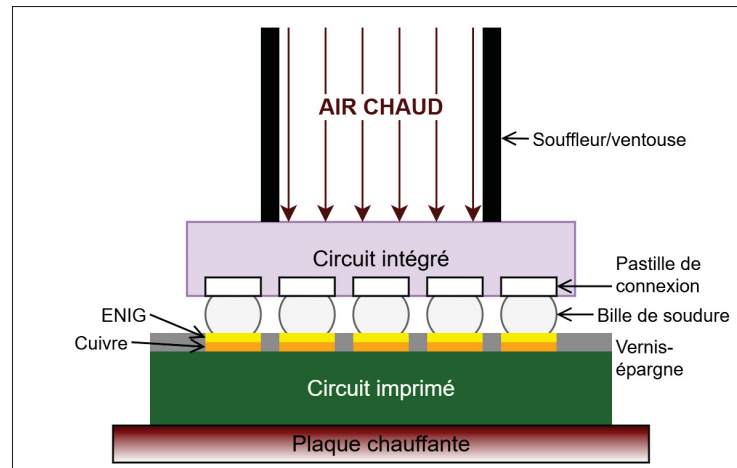


Figure-A II-10 Technique de la puce retournée

concernant l'espace entre les pastilles et leur taille. Étant donné que nous ne pouvions pas placer trop peu de pastilles, nous avons placé des pastilles redondantes. La redondance des pastilles et leur placement à des intervalles réguliers servent à garantir un bon rendement de fabrication. La figure II-3 illustre le placement de toutes les pastilles sur notre puce.

Après la conception du circuit intégré, il a fallu concevoir le circuit imprimé avec des pastilles de connexion de dimensions identiques à celles de la puce. Or, de nombreux fabricants de circuits imprimés n'offraient tout simplement pas la précision nécessaire pour réaliser ce dont nous avions besoin. Ceci est un aspect du *flip-chip* que nous n'avons pas pris en compte lors du choix de la technique d'assemblage. En plus de cela, la très petite taille des pastilles et leur placement serré rendaient l'assemblage une affaire non triviale. En effet, lors de la soudure de notre puce sur le circuit imprimé, nous avons souvent accidentellement court-circuité les billes de connexion en exerçant trop de pression sur la puce. Puisque les connexions se situaient en dessous de la puce, il était impossible de vérifier visuellement la soudure, tandis que le câblage par fils nous aurait permis de vérifier la connexion au microscope. Avec la puce retournée, seule une vérification électrique à l'aide d'un multimètre nous permettait de savoir rapidement si la connexion était bonne. Enfin, l'espace serré entre les pastilles et la largeur minimale des traces réalisable par le fabricant de circuits imprimés nous empêchaient d'insérer une trace de circuit imprimé entre deux pastilles de connexion. De plus, nous ne pouvions pas placer de

vias en-dessous du circuit intégré, car la surface sous la puce devait être parfaitement plate. Placer des vias aurait ajouté du relief sous le circuit intégré. Par conséquent, nous avons placé les pastilles pour le signal sur les bords, où la connexion est simple, et les pastilles redondantes au milieu de la puce.

Pour résumer, outre les considérations d'intégrité du signal, le choix de la technique de montage de la puce affecte non seulement la conception du circuit intégré, mais également celle du circuit imprimé et les étapes de vérification du circuit.

ANNEXE III

LE CIRCUIT IMPRIMÉ

1. Les différentes parties du circuit imprimé

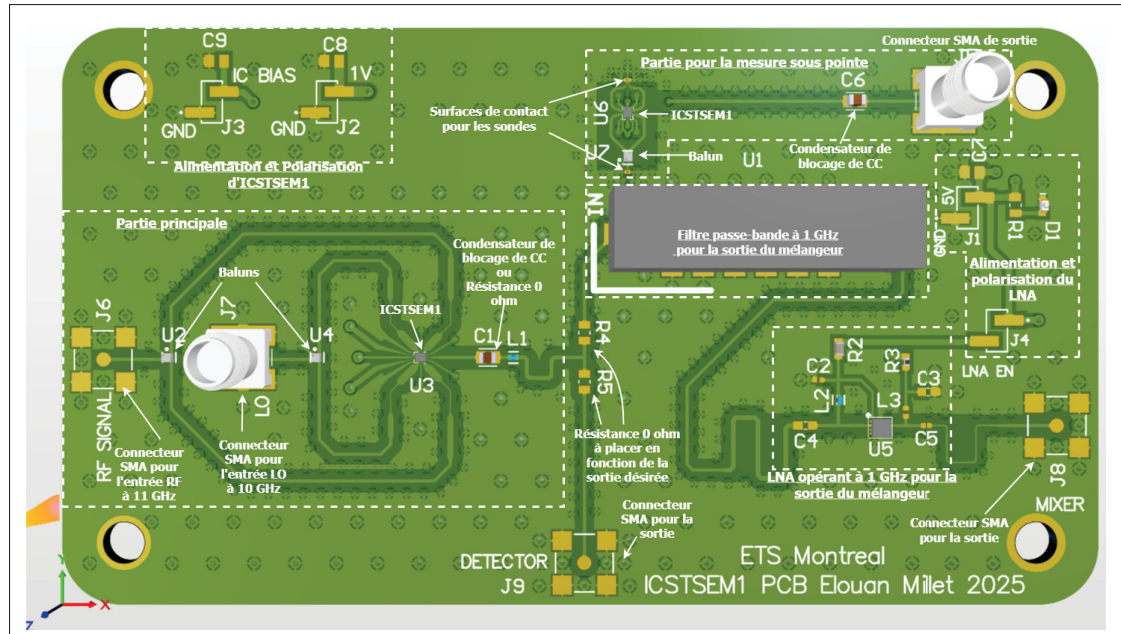


Figure-A III-1 Un aperçu du circuit imprimé pour l'ICSTSEM1

1.1 La partie principale

Le circuit imprimé a une partie principale qui a été utilisée pour la plupart des mesures. Sur cette partie, l'emplacement U3, dédié à la puce ICSTSEM1, est relié à des connecteurs de type SMA pour les entrées et les sorties. Entre les entrées de la puce et les connecteurs SMA qui reçoivent les signaux RF et LO, se trouvent des *baluns* BD60120N50100AHF aux emplacements U2 et U4. Entre la sortie de la puce et le connecteur SMA à l'emplacement J9, se trouvent un condensateur de 10 nF à C1 et une bobine de 7 nH à L1. Ces deux composants servent à adapter l'impédance de la sortie de la puce à 50 Ω . L'alimentation et la polarisation sont assurées par les connecteurs J2 et J3. C8 et C9 sont des condensateurs de découplage pour l'alimentation et la polarisation.

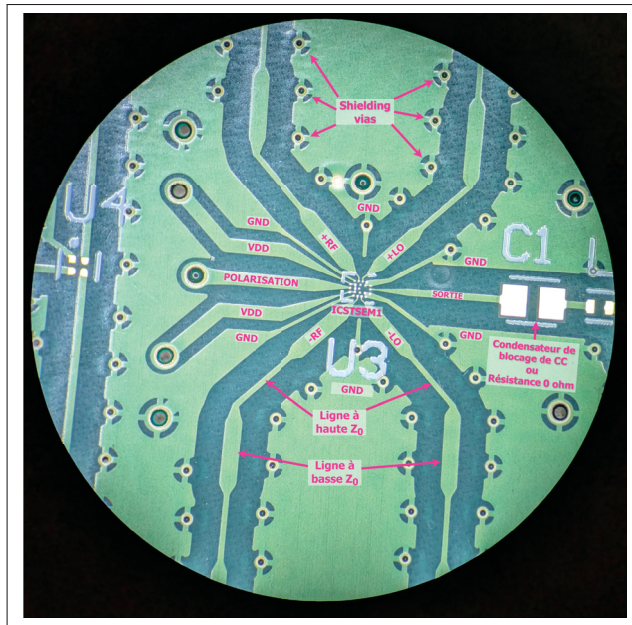


Figure-A III-2 Photographie au microscope du circuit d'adaptation d'impédance de la partie principale du circuit imprimé

1.2 La partie pour les mesures sous pointe

En plus de cela, nous pouvons trouver une partie pour effectuer des mesures sous pointe, représentée dans la figure III-3. À l'origine, l'idée derrière cette partie était de réaliser des mesures indépendantes de la performance des *baluns* et du circuit d'adaptation d'impédance. En effet, nous souhaitions utiliser la station de mesure sous pointe et les sondes appartenant à Polytechnique Montréal (parce qu'il en manquait au LaCIME). Une paire d'entrées de la puce est directement reliée à une pastille de connexion dédiée à la sonde différentielle. L'autre paire d'entrées de la puce est reliée à un *balun* monté sur le circuit imprimé. Ce *balun* est lui-même relié à une pastille dédiée à une sonde simple. S'il se trouve encore un *balun*, et si nous n'avons pas placé deux pastilles adaptées aux sondes différentielles, c'est parce que deux sondes différentielles auraient requis quatre générateurs. Or, nous n'avons que trois générateurs de signaux à notre disposition.

1.3 La partie pour amplifier le signal IF

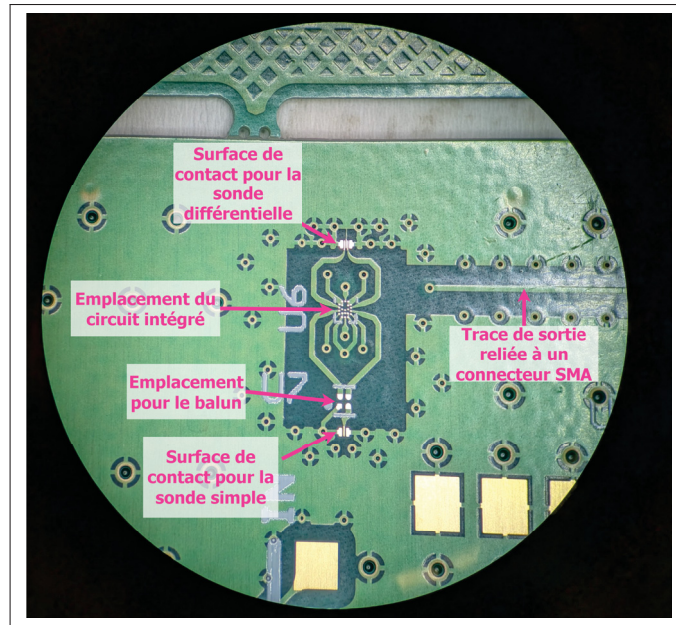


Figure-A III-3 Photographie au microscope de la partie du circuit imprimé dédiée aux mesures sous pointe

Enfin, en anticipant un faible gain de conversion du circuit intégré, nous avons dédié une partie du circuit à un chemin alternatif pour le signal de sortie. En effet, la sortie de la puce à U3 peut être reliée à deux chemins différents selon le placement d'une résistance de $0\ \Omega$. Si on la soude à R5, alors le signal sort directement par un connecteur SMA. Si on la soude R4, alors le signal peut être acheminé vers un filtre passe-bande *BPF-A1140+*, à U1, puis vers un LNA *BLB01*, à U5. Le filtre et l'amplificateur sont censés opérer autour de 1 GHz pour amplifier le signal IF. Ils ne laissent pas passer de courant continu. Donc, cette partie n'aurait pas pu amplifier l'enveloppe de puissance sortant du détecteur. Au final, cette partie n'a pas été utilisée, car nous avons tout de même pu effectuer des mesures du mélangeur.

2. L'anatomie du circuit imprimé

2.1 Le choix du matériau du circuit imprimé

Nous avons opté pour un matériau FR-4 en raison de son prix bas. Plus précisément, notre carte est faite de FR408HR de l'entreprise Isola. Un matériau spécialement conçu pour les hautes fréquences, tel que le Rogers 4003, aurait beaucoup augmenté le prix du circuit imprimé. Si les matériaux FR-4 sont en général moins adaptés aux hautes fréquences que ceux de Rogers, c'est parce que leur constante diélectrique (notée D_k ou ϵ_r) varie davantage avec la fréquence, et parce que leur facteur de dissipation (noté D_f ou $\tan(\delta)$) est relativement important. La constante diélectrique détermine la longueur d'onde ainsi que les dimensions des traces pour une impédance caractéristique souhaitée. Le facteur de dissipation quantifie la perte de puissance due au substrat diélectrique. Plus le facteur de dissipation est faible, plus la perte d'insertion d'une ligne de transmission est faible. Parmi les différents matériaux FR-4, nous avons choisi le FR408HR en particulier, car, selon Isola Group (2021), sa constante diélectrique varie peu entre une fréquence de 1 GHz et 10 GHz. Cela a simplifié la conception des circuits d'adaptation d'impédance des entrées et de la sortie.

2.2 La surface du circuit imprimé

Les pastilles de connexion du circuit imprimé ont une finition en *electroless nickel immersion gold* (ENIG). Cela consiste en du nickel recouvert d'une fine couche d'or. Une surface en or était nécessaire afin de souder le circuit intégré grâce à la technique de la puce retournée.

2.3 Les couches

La carte comporte quatre couches de cuivre, comme dans la figure III-4. Entre chaque couche de cuivre se trouve une couche de matériau diélectrique. Une couche de vernis-épargne recouvre la surface de la carte sur chaque côté. À la couche de surface supérieure se trouvent les lignes microruban pour les signaux de haute fréquence. Juste en dessous, enfouie dans la carte, se trouve une couche dédiée à la masse. Cette couche ne comporte aucune trace pour que les lignes microruban aient un plan de masse continu. La continuité du plan de masse minimise l'impédance du chemin de retour du courant. Encore en dessous, la seconde couche de cuivre

intérieure sert à l'alimentation et à la polarisation du circuit intégré. Enfin, la couche de surface inférieure est encore une couche de masse.

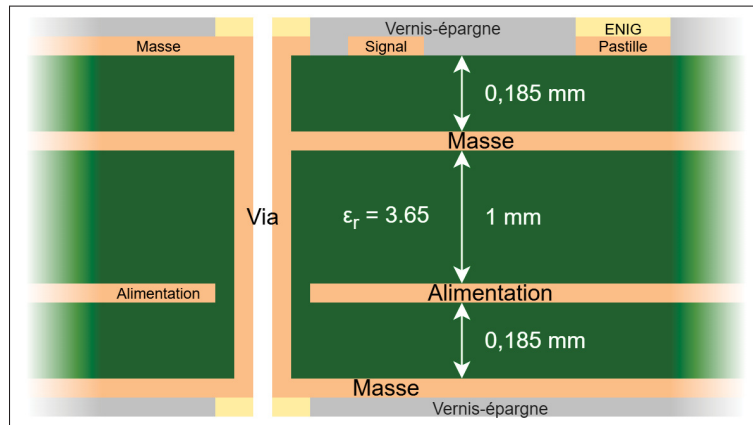


Figure-A III-4 Illustration des couches du circuit imprimé

2.4 Le placement des vias

Le *via-stitching* (ou littéralement la « couture de vias ») est une pratique qui consiste à placer de nombreux vias de façon régulière sur toute la surface du circuit imprimé. Ces vias permettent de relier les plans de masse qui se trouvent sur différentes couches. Cela donne un chemin à basse impédance pour le courant de retour. Cela assure ainsi un potentiel égal sur tous les plans de masse (Altium (2025)). Sur la carte, nous avons placé des vias de 0,7 mm de diamètre, espacés de 5 mm entre eux.

En plus de cela, afin de protéger les signaux de haute fréquence des interférences, il a fallu utiliser des vias de protection (*shielding vias* en anglais). Ces vias, qui relient deux plans de masse, sont placés à intervalles réguliers autour des traces sensibles, assurant l'isolation des flancs des traces. Selon Altium (2025), il est conseillé de placer ces vias avec un espacement d'un dixième, voire d'un vingtième, de la longueur d'onde. Sur la carte, nous n'avons pas pu atteindre ce niveau de densité pour des raisons de fabrication. Des vias de protection, de 0,3 mm de diamètre, ont été placés à un intervalle de 2,1 mm, à 1,5 mm des traces sensibles. Sachant

que la fréquence maximale est de 11 GHz et que la constante diélectrique du substrat est de 3,65, l'espacement des vias vaut presque un septième de la longueur d'onde minimale.

BIBLIOGRAPHIE

- Altium. (2025). Adding Via Stitching & Via Shielding. Repéré le 2025-09-17 à <https://www.altium.com/documentation/altium-designer/pcb/via-stitching-via-shielding>.
- Ansys. (2025). What is a Phased Array Antenna ? Repéré à <https://www.ansys.com/simulation-topics/what-is-a-phased-array-antenna>.
- Aparici, J. (1988). A wide dynamic range square-law diode detector (for radioastronomy). *IEEE Transactions on Instrumentation and Measurement*, 37(3), 429-433. doi : 10.1109/19.7469.
- Bao, M., Jacobsson, H., Aspemyr, L., Carchon, G. & Sun, X. (2006). A 9–31-GHz Subharmonic Passive Mixer in 90-nm CMOS Technology. *IEEE Journal of Solid-State Circuits*, 41(10), 2257-2264. doi : 10.1109/JSSC.2006.881551.
- Bekkaoui, M. O. (2017). Gilbert cell Mixer design in 65nm CMOS technology. *2017 4th International Conference on Electrical and Electronic Engineering (ICEEE)*, pp. 67-72. doi : 10.1109/ICEEE2.2017.7935794.
- Berthiaume, D. (2022). *Highly linear RF amplifying transmitarray unit-cell with varactor-based reconfigurable phase and amplitude*. (phd, École de technologie supérieure). Repéré à <https://espace.etsmtl.ca/id/eprint/2978/>.
- Berthiaume, D., Laurin, J.-J. & Constantin, N. G. (2021). Anti-Series Varactor Network With Improved Linearity Performances in the Presence of Inductive and Capacitive Parasitics. *IEEE Access*, 9, 48325-48340. doi : 10.1109/ACCESS.2021.3069090.
- Bigdeli, Y., Burasa, P. & Wu, K. (2025). Extending the Dynamic Range of Square-Law Power Detectors for Large-Scale Receiver Arrays. *IEEE Microwave and Wireless Technology Letters*, 1-4. doi : 10.1109/LMWT.2025.3566330.
- Bourgault, P. (2023). *Amplificateur et détecteur d'enveloppe large-bande pour la polarisation dynamique des amplificateurs RF CMOS en rétroaction positive d'enveloppe*. (Mémoire de maîtrise, École de technologie supérieure).
- Brunet, J., Ayling, A. & Hajimiri, A. (2024). Transmitarrays for Wireless Power Transfer on Earth and in Space. *IEEE Journal of Microwaves*, 4(4), 934-945. doi : 10.1109/JMW.2024.3459859.
- BSIM Group. (2025). BSIM4 Model. Repéré à <https://www.bsim.berkeley.edu/models/bsim4/>.

- Buisman, K., de Vreede, L. C. N., Larson, L., Spirito, M., Akhnoukh, A., Scholtes, T. & Nanver, L. K. (2005). “Distortion-Free” Varactor Diode Topologies for RF Adaptivity. *IEEE MTT-S International Microwave Symposium Digest, 2005.*, pp. 157-160. doi : 10.1109/MWSYM.2005.1516547.
- Carvalho, N. & Pedro, J. C. (1999). Compact Formulas to Relate ACPR and NPR to Two-Tone IMR and IP3. *Microwave Journal*, 42(12), 0-0.
- Cha, J., Woo, W., Cho, C., Park, Y., Lee, C.-H., Kim, H. & Laskar, J. (2009). A highly-linear radio-frequency envelope detector for multi-standard operation. *2009 IEEE Radio Frequency Integrated Circuits Symposium*, pp. 149-152. doi : 10.1109/RFIC.2009.5135510.
- Chen, W., Lin, X., Lee, J., Toskala, A., Sun, S., Chiasserini, C. F. & Liu, L. (2023). 5G-Advanced Toward 6G : Past, Present, and Future. *IEEE Journal on Selected Areas in Communications*, 41(6), 1592-1619. doi : 10.1109/JSAC.2023.3274037.
- Cui, G., Yuan, W., Liu, J., Cui, X. & Liu, X. (2022). Extraction and Compensation of Bonding Wires Parasitic Parameter. *2022 Cross Strait Radio Science & Wireless Technology Conference (CSRSWTC)*, pp. 1-3. doi : 10.1109/CSRSWTC56224.2022.10098402.
- de Dieuleveult, F. & Romain, O. (2008). *Électronique appliquée aux hautes fréquences - 2e éd. : Principes et applications*. Dunod. Repéré à <https://books.google.ca/books?id=I0I6dWQk5UsC>.
- Di Palma, L., Clemente, A., Dussopt, L., Sauleau, R., Potier, P. & Pouliguen, P. (2016). Radiation Pattern Synthesis for Monopulse Radar Applications With a Reconfigurable Transmitarray Antenna. *IEEE Transactions on Antennas and Propagation*, 64(9), 4148-4154. doi : 10.1109/TAP.2016.2586491.
- Doczkat, M., Etemad, K. & Gosain, A. (2025). *FCC Technology Advisory Council 6G Working Group Report August 5, 2025*. FCC. Repéré à <https://www.fcc.gov/sites/default/files/FCC-TAC-6G-Working-Group-Report-2025-Final.pdf>.
- Díaz Canales, F. & Abbasi, M. (2013). A 75–90 GHz high linearity MMIC power amplifier with integrated output power detector. *2013 IEEE MTT-S International Microwave Symposium Digest (MTT)*, pp. 1-4. doi : 10.1109/MWSYM.2013.6697600.
- European 5G Observatory. (2025). *5G Observatory report 2025*. Repéré le 2025-10-20 à <https://digital-strategy.ec.europa.eu/en/policies/5g-observatory>.
- European Commission. (2025). *6G networks in Europe*. Repéré à <https://digital-strategy.ec.europa.eu/en/policies/6g>.

- Fischer, A., Steinmetz, C. & Adel, H. (2025). Varactor Based Reconfigurable Reflectarray Antenna at Ka-Band. *2025 19th European Conference on Antennas and Propagation (EuCAP)*, pp. 1-5. doi : 10.23919/EuCAP63536.2025.10999706.
- Fong, K. L. & Meyer, R. (1999). Monolithic RF active mixer design. *IEEE Transactions on Circuits and Systems II : Analog and Digital Signal Processing*, 46(3), 231-239. doi : 10.1109/82.754857.
- Frater, R. H. (1964). Accurate Wideband Multiplier—Square-Law Detector. *Review of Scientific Instruments*, 35(7), 810-813. doi : 10.1063/1.1746797.
- Fu, Z., Zou, X., Liao, Y., Lai, G., Li, Y. & Chung, K. L. (2022). A Brief Review and Comparison Between Transmitarray Antennas, Reflectarray Antennas and Reconfigurable Intelligent Surfaces. *2022 IEEE Conference on Telecommunications, Optics and Computer Science (TOCS)*, pp. 1192-1196. doi : 10.1109/TOCS56154.2022.10016145.
- Gibiino, G. P. (2016). Nonlinear Characterization and Modeling of Radio-Frequency Devices and Power Amplifiers with Memory Effects. Repéré à <https://api.semanticscholar.org/CorpusID:115087858>.
- Government of Canada. (2022). Canadian Table of Frequency Allocations (2022). Repéré à <https://ised-isde.canada.ca/site/spectrum-management-telecommunications/en/learn-more/key-documents/consultations/canadian-table-frequency-allocations-sf10759#s2.3>.
- Gray, P., Hurst, P., Lewis, S. & Meyer, R. (2009). *Analysis and Design of Analog Integrated Circuits*. Wiley. Repéré à <https://books.google.ca/books?id=VBB4QAAACAAJ>.
- Guo, Y. J. & Ziolkowski, R. W. (2022). A Perspective of Antennas for 5G and 6G. Dans *Advanced Antenna Array Engineering for 6G and Beyond Wireless Communications* (pp. 1-22). doi : 10.1002/9781119712947.ch1.
- Hsiao, Y.-C., Meng, C., Wang, T.-W. & Huang, G.-W. (2014). 60-GHz 0.18- μm CMOS Schottky-diode ring-mixer down-converter. *2014 Asia-Pacific Microwave Conference*, pp. 1202-1204.
- Huang, J., Zhang, R., Xie, R., Shi, S., Ding, J. & Zhai, G. (2019). Beam-Steerable Reflectarray Antenna for C-Band Radar. *2019 IEEE International Conference on Computational Electromagnetics (ICCEM)*, pp. 1-3. doi : 10.1109/COMPEM.2019.8778906.
- Huang, J. & Encinar, J. A. (2008). Introduction to Reflectarray Antenna. Dans *Reflectarray Antennas* (pp. 1-7). doi : 10.1002/9780470178775.ch1.

- Hum, S. V. & Perruisseau-Carrier, J. (2014). Reconfigurable Reflectarrays and Array Lenses for Dynamic Antenna Beam Control : A Review. *IEEE Transactions on Antennas and Propagation*, 62(1), 183-198. doi : 10.1109/TAP.2013.2287296.
- Isola Group. (2021). FR408HR Lead Free, Mid Loss Laminate and Prepreg. Repéré le 2025-10-22 à <https://www.isola-group.com/pcb-laminates-prepreg/fr408hr-laminate-and-prepreg/>.
- Ivanov, S. I., Lavrov, A. P. & Matveev, Y. A. (2015). A square-law microwave transistor detector for wideband radiometers. *2015 International Siberian Conference on Control and Communications (SIBCON)*, pp. 1-4. doi : 10.1109/SIBCON.2015.7147076.
- Ker, M.-D., Chen, T.-Y. & Chang, C.-Y. (2001). ESD protection design for CMOS RF integrated circuits. *2001 Electrical Overstress/Electrostatic Discharge Symposium*, pp. 344-352.
- Ku, H. & Kenney, J. (2001). Estimation of error vector magnitude using two-tone intermodulation distortion measurements [power amplifier]. *2001 IEEE MTT-S International Microwave Symposium Digest (Cat. No.01CH37157)*, 1, 17-20 vol.1. doi : 10.1109/MWSYM.2001.966829.
- Lee, T. (2003). *The Design of CMOS Radio-Frequency Integrated Circuits*. Cambridge University Press. Repéré à <https://books.google.ca/books?id=mlYhAwAAQBAJ>.
- Li, C., Yi, X., Boon, C. C. & Yang, K. (2019). A 34-dB Dynamic Range 0.7-mW Compact Switched-Capacitor Power Detector in 65-nm CMOS. *IEEE Transactions on Power Electronics*, 34(10), 9365-9368. doi : 10.1109/TPEL.2019.2908283.
- Liu, S., Niu, B., Wang, B. & Yu, W. (2024). A 340-GHz Upconversion Mixer MMIC Using an Antiparallel Series Schottky Diode Pair With High Output Power. *IEEE Microwave and Wireless Technology Letters*, 34(7), 927-930. doi : 10.1109/LMWT.2024.3401579.
- Liu, Y., Niu, Z., Zhang, B., Hu, Y., Dai, B., Qiao, C., Feng, Y., Fan, Y. & Chen, X. (2022). A High-Performance 330-GHz Subharmonic Mixer Using Schottky Diodes. *IEEE Microwave and Wireless Components Letters*, 32(6), 571-574. doi : 10.1109/LMWC.2022.3149378.
- Lovett, S., Welten, M., Mathewson, A. & Mason, B. (1998). Optimizing MOS transistor mismatch. *Solid-State Circuits, IEEE Journal of*, 33, 147 - 150. doi : 10.1109/4.654947.
- Ludwig, R. & Bretchko, P. (2000). *RF Circuit Design : Theory and Applications*. Pearson Education. Repéré à <https://books.google.ca/books?id=4D1iPgAACAAJ>.
- Meyer, R. (1995). Low-power monolithic RF peak detector analysis. *IEEE Journal of Solid-State Circuits*, 30(1), 65-67. doi : 10.1109/4.350192.

- Meyer, R. & Stephens, M. (1975). Distortion in variable-capacitance diodes. *IEEE Journal of Solid-State Circuits*, 10(1), 47-54. doi : 10.1109/JSSC.1975.1050553.
- Michaelsen, R. S., Johansen, T. K., Tamborg, K. M. & Zhurbenko, V. (2015). A SiGe BiCMOS double-balanced mixer with active balun for X-band Doppler radar. *2015 SBMO/IEEE MTT-S International Microwave and Optoelectronics Conference (IMOC)*, pp. 1-5. doi : 10.1109/IMOC.2015.7369090.
- Narayanan, S. (1967). Transistor distortion analysis using volterra series representation. *The Bell System Technical Journal*, 46(5), 991-1024. doi : 10.1002/j.1538-7305.1967.tb01723.x.
- Nayeri, P., Yang, F. & Elsherbeni, A. Z. (2018). Introduction to Reflectarray Antennas. Dans *Reflectarray Antennas : Theory, Designs, and Applications* (pp. 1-8). doi : 10.1002/9781118846728.ch1.
- Nguyen, V. (2025). Global 5G Adoption Skyrockets to 2.25 Billion, Four Times Faster Than 4G. Repéré à <https://www.5gamerica.org/global-5g-adoption-skyrockets-to-2-25-billion-four-times-faster-than-4g/>.
- Norwood, M. & Shatz, E. (1968). Voltage variable capacitor tuning : A review. *Proceedings of the IEEE*, 56(5), 788-798. doi : 10.1109/PROC.1968.6408.
- Nouri, N., Nezhad-Ahamdi, M.-R., Mirabbasi, S. & Safavi-Naeini, S. (2010). A double-balanced CMOS mixer with on-chip balun for 60-GHz receivers. *Proceedings of the 8th IEEE International NEWCAS Conference 2010*, pp. 321-324. doi : 10.1109/NEWCAS.2010.5604029.
- Office québécois de la langue française. (2025). responsivité. Repéré le 2025-12-19 à <https://vitrlinguistique.oqlf.gouv.qc.ca/fiche-gdt/fiche/8985830/responsivite>.
- Ou, J. & Farahmand, F. (2012). Transconductance/drain current based distortion analysis for analog CMOS integrated circuits. *10th IEEE International NEWCAS Conference*, pp. 61-64. doi : 10.1109/NEWCAS.2012.6328956.
- Ozeren, E., Kalyoncu, I., Ustundag, B., Cetindogan, B., Kayahan, H., Kaynak, M. & Gurbuz, Y. (2016). A High Dynamic Range Power Detector at X-Band. *IEEE Microwave and Wireless Components Letters*, 26(9), 708-710. doi : 10.1109/LMWC.2016.2597266.
- Packard, H. (1981). Square Law and Linear Detection. Hewlett Packard.
- Pan, W., Huang, C., Ma, X. & Luo, X. (2014). An Amplifying Tunable Transmitarray Element. *IEEE Antennas and Wireless Propagation Letters*, 13, 702-705. doi : 10.1109/LAWP.2014.2313596.

- Peng, G., Wang, S., Li, G., Huang, T., Huang, Y. & Liu, Y. (2025). Space-Terrestrial Integrated Time-Sensitive Networks : Architecture, Challenges, and Open Issues. *IEEE Network*, 39(3), 140-147. doi : 10.1109/MNET.2025.3551322.
- Peterson, Z. (2023). RF Microstrips and Ground Plane Clearance : How Close is Too Close ? Repéré le 2025-09-17 à <https://resources.altium.com/p/microstrip-ground-clearance-how-close-too-close>.
- Pozar, D. M. (2012). *Microwave engineering* (éd. Fourth edition). Hoboken, NJ : John Wiley & Sons, Inc.
- Qayyum, S. & Negra, R. (2017). 0.16 mW, 7–70 GHz distributed power detector with 75 dB voltage sensitivity in 130 nm standard CMOS technology. *2017 12th European Microwave Integrated Circuits Conference (EuMIC)*, pp. 13-16. doi : 10.23919/EuMIC.2017.8230648.
- Ravanne, J. G., Then, Y. L., Su, H. T. & Hijazin, I. (2021). Design and Analysis of a Practical RMS Power Detector Using Power MOSFET. *2021 IEEE 12th Annual Information Technology, Electronics and Mobile Communication Conference (IEMCON)*, pp. 0744-0748. doi : 10.1109/IEMCON53756.2021.9623258.
- Razavi, B. (1997). Design considerations for direct-conversion receivers. *IEEE Transactions on Circuits and Systems II : Analog and Digital Signal Processing*, 44(6), 428-435. doi : 10.1109/82.592569.
- Razavi, B. (2013). *Fundamentals of Microelectronics, 2nd Edition*. Wiley. Repéré à <https://books.google.ca/books?id=Lvr8ngEACAAJ>.
- Razavi, B. (2012). *RF microelectronics* (éd. 2. edition). Upper Saddle River, NJ Munich : Prentice Hall.
- Reid, M. S., Gardner, R. A. & Stelzried, C. T. (1975). *A new broadband square law detector* (Rapport n°NASA-CR-143459). Repéré le 2025-07-08 à <https://ntrs.nasa.gov/citations/19750023271>.
- Richier, C., Salome, P., Mabboux, G., Zaza, I., Juge, A. & Mortini, P. (2000). Investigation on different ESD protection strategies devoted to 3.3 V RF applications (2 GHz) in a 0.18 μm CMOS process. *Electrical Overstress/Electrostatic Discharge Symposium Proceedings 2000 (IEEE Cat. No.00TH8476)*, pp. 251-259. doi : 10.1109/EOESD.2000.890084.

- Romanofsky, R., Mueller, C. & Chandrasekar, C. V. (2009). Concept for a low cost, high efficiency precipitation radar system based on ferroelectric reflectarray antenna. *2009 IEEE Radar Conference*, pp. 1-6. doi : 10.1109/RADAR.2009.4977135.
- Saeed, M., Hamed, A., Uzlu, B., Baskent, E., Otto, M., Wang, Z. & Negra, R. (2021). Compact V-Band MMIC Square-law Power Detector with 70 dB Dynamic Range exploiting State-of-the-art Graphene diodes. *2021 IEEE MTT-S International Microwave Symposium (IMS)*, pp. 888-891. doi : 10.1109/IMS19712.2021.9574843.
- Sarkas, I., Mavridis, D. & Papadopoulos, G. (2006). Large and Small Signal Distortion Analysis Using Modified Volterra Series. *2006 NORCHIP*, pp. 63-66. doi : 10.1109/NORCHIP.2006.329245.
- Sarkas, I., Mavridis, D., Papamichail, M. & Papadopoulos, G. (2007). Volterra Analysis Using Chebyshev Series. *2007 IEEE International Symposium on Circuits and Systems (ISCAS)*, pp. 1931-1934. doi : 10.1109/ISCAS.2007.378353.
- Serhan, A., Lauga-Larroze, E. & Fournier, J.-M. (2015a). A 700MHz output bandwidth, 30dB dynamic range, common-base mm-wave power detector. *2015 IEEE MTT-S International Microwave Symposium*, pp. 1-3. doi : 10.1109/MWSYM.2015.7166764.
- Serhan, A., Lauga-Larroze, E. & Fournier, J.-M. (2015b). Common-Base/Common-Gate Millimeter-Wave Power Detectors. *IEEE Transactions on Microwave Theory and Techniques*, 63(12), 4483-4491. doi : 10.1109/TMTT.2015.2496220.
- Shahriar Rashid, S. M. & Rashid, A. H. (2010). Design of a double balanced square law CMOS up-conversion mixer with improved input isolation technique for high frequency applications. *2010 17th IEEE International Conference on Electronics, Circuits and Systems*, pp. 1088-1091. doi : 10.1109/ICECS.2010.5724705.
- Sharma, S. & Constantin, N. G. (2019). Nonlinear Three-Port Representation of PAs for Embedded Self-Calibration of Envelope-Dependent Dynamic Biasing Implementations. *IEEE Access*, 7, 172796-172815. doi : 10.1109/ACCESS.2019.2956721.
- Su, Y.-b., Xia, Q., Geng, L. & Liu, X.-K. (2019). A Highly Linear Low Power Envelope Detector. *2019 IEEE International Conference on Electron Devices and Solid-State Circuits (EDSSC)*, pp. 1-2. doi : 10.1109/EDSSC.2019.8754090.
- Tang, J., Xu, S., Yang, F. & Li, M. (2021). Design and Measurement of a Reconfigurable Transmitarray Antenna With Compact Varactor-Based Phase Shifters. *IEEE Antennas and Wireless Propagation Letters*, 20(10), 1998-2002. doi : 10.1109/LAWP.2021.3101891.

- Terrovitis, M. & Meyer, R. (2000). Intermodulation distortion in current-commutating CMOS mixers. *IEEE Journal of Solid-State Circuits*, 35(10), 1461-1473. doi : 10.1109/4.871323.
- Toka, M., Lee, B., Seong, J., Kaushik, A., Lee, J., Lee, J., Lee, N., Shin, W. & Poor, H. V. (2024). RIS-Empowered LEO Satellite Networks for 6G : Promising Usage Scenarios and Future Directions. *IEEE Communications Magazine*, 62(11), 128-135. doi : 10.1109/MCOM.002.2300554.
- Tsai, J.-H., Wu, P.-S., Lin, C.-S., Huang, T.-W., Chern, J. G. J. & Huang, W.-C. (2007a). A 25–75 GHz Broadband Gilbert-Cell Mixer Using 90-nm CMOS Technology. *IEEE Microwave and Wireless Components Letters*, 17(4), 247-249. doi : 10.1109/LMWC.2007.892934.
- Tsai, J.-H., Wu, P.-S., Lin, C.-S., Huang, T.-W., Chern, J. G. J. & Huang, W.-C. (2007b). A 25–75 GHz Broadband Gilbert-Cell Mixer Using 90-nm CMOS Technology. *IEEE Microwave and Wireless Components Letters*, 17(4), 247-249. doi : 10.1109/LMWC.2007.892934.
- Tsividis, Y., Suyama, K. & Vavelidis, K. (1995). Simple ‘reconciliation’ MOSFET model valid in all regions. *Electronics Letters*, 31, 506-508. doi : 10.1049/el:19950256.
- Valdes-Garcia, A., Venkatasubramanian, R., Silva-Martinez, J. & Sanchez-Sinencio, E. (2008). A Broadband CMOS Amplitude Detector for On-Chip RF Measurements. *IEEE Transactions on Instrumentation and Measurement*, 57(7), 1470-1477. doi : 10.1109/TIM.2008.917196.
- Voo, T. & Toumazou, C. (1995). A novel high speed current mirror compensation technique and application. *1995 IEEE International Symposium on Circuits and Systems (ISCAS)*, 3, 2108-2111 vol.3. doi : 10.1109/ISCAS.1995.523841.
- Weste, N. & Harris, D. (2011). *CMOS VLSI Design : A Circuits and Systems Perspective*. Addison Wesley. Repéré à <https://books.google.ca/books?id=sv8OQgAACAAJ>.
- Xia, J. & Boumaiza, S. (2015). A Novel Broadband Linear-in-Magnitude RF Envelope Detector With Enhanced Detection Speed and Accuracy. *IEEE Microwave and Wireless Components Letters*, 25(5), 325-327. doi : 10.1109/LMWC.2015.2409796.
- Xu, K., Yang, X., Wang, W. & Yoshimasu, T. (2014). 44-GHz, 0.5-V compact power detector IC in 65-nm CMOS. *2014 IEEE International Wireless Symposium (IWS 2014)*, pp. 1-3. doi : 10.1109/IEEE-IWS.2014.6864216.
- Yang, X., Uchida, Y., Liu, Q. & Yoshimasu, T. (2012). Low-power ultra-wideband power detector IC in 130 nm CMOS technology. *2012 IEEE MTT-S International Microwave Workshop Series on Millimeter Wave Wireless Technology and Applications*, pp. 1-4. doi : 10.1109/IMWS2.2012.6338206.

- Yao, H., Wang, L., Wang, X., Lu, Z. & Liu, Y. (2018). The Space-Terrestrial Integrated Network : An Overview. *IEEE Communications Magazine*, 56(9), 178-185. doi : 10.1109/MCOM.2018.1700038.
- Zhang, C., Gharpurey, R. & Abraham, J. A. (2007). Built-In Test of RF Mixers Using RF Amplitude Detectors. *8th International Symposium on Quality Electronic Design (ISQED'07)*, pp. 404-409. doi : 10.1109/ISQED.2007.44.
- Zhang, T., Eisenstadt, W., Fox, R. & Yin, Q. (2006). Bipolar Microwave RMS Power Detectors. *IEEE Journal of Solid-State Circuits*, 41(9), 2188-2192. doi : 10.1109/JSSC.2006.880592.
- Zhou, Y. & Chia, M. Y.-W. (2008). A Low-Power Ultra-Wideband CMOS True RMS Power Detector. *IEEE Transactions on Microwave Theory and Techniques*, 56(5), 1052-1058. doi : 10.1109/TMTT.2008.921299.
- Zhu, Y., Vilenskiy, A. R., Iupikov, O. A., Krasov, P. S., Emanuelsson, T., Lasser, G. & Ivashina, M. V. (2024). A Varactor-Based Reconfigurable Intelligent Surface Concept for 5G/6G mm-Wave Applications. *2024 18th European Conference on Antennas and Propagation (EuCAP)*, pp. 1-5. doi : 10.23919/EuCAP60739.2024.10501072.
- Zirath, H. & He, Z. (2010). Power detectors and envelope detectors in mHEMT MMIC-technology for millimeterwave applications. *The 5th European Microwave Integrated Circuits Conference*, pp. 353-356.