

Titre: Étude et réalisation de composantes et algorithme pour les
Title: récepteurs directs en ondes millimétriques

Auteur: Tiberiu Marian Visan
Author:

Date: 2002

Type: Mémoire ou thèse / Dissertation or Thesis

Référence: Visan, T. M. (2002). Étude et réalisation de composantes et algorithme pour les
Citation: récepteurs directs en ondes millimétriques [Ph.D. thesis, École Polytechnique de
Montréal]. PolyPublie. <https://publications.polymtl.ca/7081/>

 **Document en libre accès dans PolyPublie**
Open Access document in PolyPublie

URL de PolyPublie: <https://publications.polymtl.ca/7081/>
PolyPublie URL:

**Directeurs de
recherche:** Rénato Bosisio, & Jacques Beauvais
Advisors:

Programme: Unspecified
Program:

INFORMATION TO USERS

This manuscript has been reproduced from the microfilm master. UMI films the text directly from the original or copy submitted. Thus, some thesis and dissertation copies are in typewriter face, while others may be from any type of computer printer.

The quality of this reproduction is dependent upon the quality of the copy submitted. Broken or indistinct print, colored or poor quality illustrations and photographs, print bleedthrough, substandard margins, and improper alignment can adversely affect reproduction.

In the unlikely event that the author did not send UMI a complete manuscript and there are missing pages, these will be noted. Also, if unauthorized copyright material had to be removed, a note will indicate the deletion.

Oversize materials (e.g., maps, drawings, charts) are reproduced by sectioning the original, beginning at the upper left-hand corner and continuing from left to right in equal sections with small overlaps.

**ProQuest Information and Learning
300 North Zeeb Road, Ann Arbor, MI 48106-1346 USA
800-521-0600**

UMI[®]

UNIVERSITÉ DE MONTREAL

**ÉTUDE ET RÉALISATION DE COMPOSANTES ET ALGORITHME POUR LES
RÉCEPTEURS DIRECTS EN ONDES MILLIMÉTRIQUES**

**TIBERIU MARIAN VISAN
DÉPARTEMENT DE GÉNIE ÉLECTRIQUE
ÉCOLE POLYTECHNIQUE DE MONTRÉAL**

**THÈSE PRÉSENTÉE EN VUE DE L'OBTENTION
DU DIPLÔME DE PHILOSOPHIAE DOCTOR (Ph.D.)
(GÉNIE ÉLECTRIQUE)
AVRIL 2002**

@Tiberiu Marian Visan, 2002



**National Library
of Canada**

**Acquisitions and
Bibliographic Services**

**385 Wellington Street
Ottawa ON K1A 0N4
Canada**

**Bibliothèque nationale
du Canada**

**Acquisitions et
services bibliographiques**

**385, rue Wellington
Ottawa ON K1A 0N4
Canada**

Your file Votre référence

Our file Notre référence

The author has granted a non-exclusive licence allowing the National Library of Canada to reproduce, loan, distribute or sell copies of this thesis in microform, paper or electronic formats.

The author retains ownership of the copyright in this thesis. Neither the thesis nor substantial extracts from it may be printed or otherwise reproduced without the author's permission.

L'auteur a accordé une licence non exclusive permettant à la Bibliothèque nationale du Canada de reproduire, prêter, distribuer ou vendre des copies de cette thèse sous la forme de microfiche/film, de reproduction sur papier ou sur format électronique.

L'auteur conserve la propriété du droit d'auteur qui protège cette thèse. Ni la thèse ni des extraits substantiels de celle-ci ne doivent être imprimés ou autrement reproduits sans son autorisation.

0-612-75940-7

Canada

UNIVERSITÉ DE MONTREAL

ÉCOLE POLYTECHNIQUE DE MONTRÉAL

Cette thèse intitulée :

**ÉTUDE ET RÉALISATION DE COMPOSANTES ET ALGORITHME POUR LES
RÉCEPTEURS DIRECTS EN ONDES MILLIMÉTRIQUES**

présentée par : **VISAN Tiberiu Marian**

en vue de l'obtention du diplôme de : **Philosophiae Doctor**

a été dûment acceptée par le jury d'examen constitué de :

M. **WU Ke**, Ph.D., président

M. **BOSISIO G. Renato**, M. Sc., membre et directeur de recherche

M. **BEAUVAIS Jacques**, Ph.D., membre et codirecteur de recherche

M. **GAGNON, François**, Ph.D., membre

M. **DEMERS, Yves**, Ph.D., membre

DÉDICACE

À mon frère, Traian qui a été toujours là pour moi.

REMERCIEMENTS

Je voudrais exprimer mes chaleureux sentiments pour mon directeur de thèse, prof. Renato G. Bosisio, et le remercier pour sa confiance et son soutien inconditionnel pendant toute la période passée au PolyGrames. Il n'y a pas de qualificatif suffisant pour mesurer la qualité des relations humaines dont il a fait preuve à l'égard de ses étudiants et ses collaborateurs.

Ce travail n'aurait pas été possible sans la collaboration de mon codirecteur de l'Université de Sherbrooke, prof. Jacques Beauvais, et son groupe du laboratoire GRM. Je les remercie pour leurs efforts et leur patience. La contribution importante de Bruce Faure et Pierre Langlois dans le travail du développement des procédés MMIC est largement appréciée.

Bien sûr, je ne peux pas oublier l'âme de PolyGrames, M. Jules Gauthier qui, par son effort et sa bonne humeur a largement facilité mon travail.

Je remercie également les membres du jury pour leur participation et pour le temps dédié à la révision de mon manuscrit.

Enfin, un gros merci à mes collègues de l'École Polytechnique avec lesquels j'ai eu de nombreuses discussions intéressantes.

RÉSUMÉ

Cette thèse présente une approche innovatrice visant l'amélioration de l'architecture de récepteurs utilisés dans les systèmes de télécommunications. Les aspects relatifs à la topologie sont considérés tant au niveau du traitement analytique de l'information qu'au niveau de la caractérisation des paramètres technologiques. Un des nos objectifs a été de montrer la possibilité de construire un récepteur direct à faible coût et de concevoir de nouvelles techniques pour contourner les désavantages de cette architecture par rapport à celle des récepteurs super-hétérodyne. Ce projet fait suite aux travaux de recherche précédents, dans le domaine des récepteurs directs basés sur la jonction six-ports et on a analysé cette solution en soulignant les différences entre l'application initiale comme réflectomètre et la spécificité d'un récepteur.

Après une analyse critique des techniques proposées antérieurement pour la démodulation dans les récepteurs directs, on a introduit une nouvelle méthode de compensation des erreurs physiques. Le nouvel algorithme, utilisant la théorie des réseaux de neurones, modifie de façon adaptative la réponse du récepteur et corrige les erreurs linéaires de phase et d'amplitude. La méthode est indépendante du type de modulation et peut être appliquée à n'importe quel schéma de modulation, usuellement utilisée dans les systèmes de télécommunications actuels (n-QAM, n-PSK etc.). Dans ce cas, l'algorithme est complètement portable sur les processeurs de signaux numériques commerciaux (DSP) et il peut s'adapter au signal modulé d'entrée, pouvant être classifié ainsi comme une variante de "soft-radio".

Le deuxième volet de cette thèse présente le développement d'une infrastructure au sein du laboratoire permettant la fabrication, la caractérisation et la conception de composantes MMIC. Ce travail a été réalisé en étroite collaboration avec le groupe de recherche GMS de l'université de Sherbrooke. Tous les aspects de la technologie MMIC développée sont commentés et nous présentons les résultats qui, à notre connaissance,

sont les premiers essais universitaires au Canada, visant le développement d'une telle infrastructure dans le domaine des micro-ondes.

Le Centre de Recherche PolyGrames ne possédait aucun moyen de fabrication et caractérisation MMIC au commencement de ce projet. Un effort spécial a été dédié au développement d'une infrastructure permettant l'extraction automatique des paramètres de dispositifs à l'aide de mesures sous pointe. Des techniques différentes de calibration et des modèles alternatifs sont présentés, offrant ainsi une vue d'ensemble sur les moyens de caractérisation et de mesures élaborés dans le cadre de ce travail. Les principales limitations relatives aux procédés de fabrication (attribuables à un budget limité et une équipe peu nombreuse) sont analysées et certaines solutions sont proposées afin d'améliorer les résultats.

Nous présentons un mélangeur sous-harmonique qui offre une meilleure isolation entre le port d'oscillateur local et le port RF par rapport aux précédentes architectures de récepteurs directs, basées sur la jonction six-ports, fonctionnant dans la bande des ondes millimétriques et qui utilisent des détecteurs de puissance comme élément convertisseur de fréquence. Cette amélioration est très importante pour les récepteurs directs. En effet, l'atténuation du signal de fuite de la tête réceptrice est toujours très faible dans le cas d'un système homodyne, entraînant ainsi l'émission d'un signal parasite par l'antenne dans la bande utile.

ABSTRACT

This thesis presents a new approach towards an improved architecture of the telecommunications receivers. The topological issues are covered both in terms of analytical treatment of information and characterization of technological parameters. One of our research goals was to show the possibility of obtaining a cost-effective receiver, in a direct conversion variant, and designing special techniques to alleviate some disadvantages of this architecture compared to the classical super-heterodyne receiver.

From a chronological point of view, our project followed some works on direct receivers based on six-ports junction and we analyzed this approach respective to an evident difference between the initial application as reflectometer and the specification of a receiver.

After making a critical review of previous published techniques for demodulation in direct receivers, we introduced a new method to compensate the hardware imperfections. The novel algorithm based on neural networks theory, adaptively changes the actual response of the receiver and corrects the linear errors of gain & phase imbalance. The method is not dependent on modulation type and can be applied to any orthogonal modulation scheme, usually used in today communication systems (n-QAM, n-PSK etc.). In this case the software is completely portable on commercial available signal processors (DSP) and can actually adapt itself to the incoming signal, proving itself as a member of the soft-radios class.

A second area of development was to build an in-house facility of fabrication, characterization and design for MMIC components in collaboration with University of Sherbrooke. All three aspects of the MMIC technology are discussed and we show our

results, which in our knowledge are the first tries at academic levels in Canada to obtain such facility at microwave range.

At the beginning of this project, no MMIC fabrication and characterization means were available at PolyGrames Research Center. A special effort was dedicated to build the infrastructure for automatic parameters extraction for on-wafer measurement. Different calibration techniques and models alternatives are presented to give a complete picture of the measurement capabilities developed during this work.

Most of the limitations met during the fabrication process, due mostly of man-power and budget constraints are discussed in the presentation of the transistors characteristics and some possible solution are proposed at manufacturing level.

In contrast with the previous architecture of direct receivers based on six-port junction and operating in millimetric-wave range, who employed power detectors as down-conversion devices, we present a sub-harmonic mixer design who improves the overall isolation between the local oscillator and the RF ports. This condition is mandatory for practical applications of direct receivers as the RF front-end provides minimum attenuation of the leakage signal. Consequently, the antenna can transmit the in-band spurious signal, becoming a noise source for the wireless system.

TABLE DES MATIÈRES

Dédicace	IV
Remerciements	V
Résumé	VI
Abstract	VIII
Table des Matières.....	X
Liste des Tableaux.....	XII
Liste des Figures.....	XIII
Liste des Symboles et des Abréviations	XVII
Chapitre 1 Introduction	1
Chapitre 2 Récepteur Direct (Homodyne). Analyse Théorique	11
2.1 Récepteur Direct (Homodyne) versus Récepteur Super-Hétérodyne.....	11
2.2 Architecture I&Q.....	16
2.3 Analyse des Techniques de Correction d'Erreurs.....	22
2.3.1 Jonction Six-Ports.	22
2.3.2 Récepteur Direct.....	29
2.3.3 La Calibration de Type "Dual-Tone".	31
2.3.4 La Calibration à l'Aide d'un Réseau des Neurones.	32
2.3.5 L'Autocalibration.....	35
Chapitre 3 POLYCOM. Nouvelle Architecture de Récepteur Direct.....	37
3.1 Polycom.....	37
3.2 La Technique de Résonance Adaptative.	41
3.3 Nouvelle Technique de Correction d'Erreurs.	52
3.4 Résultats des Tests pour la Nouvelle Technique de Correction.....	60
Chapitre 4 Mesures et Fabrication de Circuits MMIC	74
4.1 Mesures des Dispositifs en Hautes Fréquences.....	76
4.1.1 Méthode Multi-lignes.....	79
4.1.2 Calibration en Domaine Temporel.....	84

4.1.3 Changement de Plan de Mesure pour les Structures Non-uniformes.....	91
4.2 Fabrication des Circuits MMIC.....	96
Chapitre 5 Modélisation des Composantes MMIC	105
5.1 Capacités Métal-Isolant-Métal (MIM).	107
5.2 Les Inductances Spirales.	111
5.3 Résistances Déposées	116
5.4 Le Transistor PHEMT	119
5.4.1 Modèle Petit Signal.	123
5.4.2 Modèle Grand Signal.....	134
Chapitre 6 Résultats Finaux et Conclusions.....	146
6.1 Caractéristiques du Transistor PHEMT.	147
6.2 Mélangeur à Transistor PHEMT.	156
6.3 Transition CPWFG à CPS.....	165
6.4 Technologie Homodyne. Évaluation et Synthèse.....	170
6.5 Conclusions et Travaux Futurs.....	172
Bibliographie	177

LISTE DES TABLEAUX

Tableau 1.1 Réseau Sans Fils à Services Intégrés.....	3
Tableau 1.2 Systèmes de Communications par Satellite.....	5
Tableau 3.1 Analyse de Convergence.	62
Tableau 3.2 Capacité de correction.	64
Tableau 5.1 Masque de capacités mesurées.	109
Tableau 5.2 Paramètres électriques du modèle de l'inductance	115
Tableau 5.3 Valeurs des composantes électriques parasites.	131

LISTE DES FIGURES

Figure 2.1 Conversion de fréquence.....	12
Figure 2.2 Mélangeur à rejection du canal image.	15
Figure 2.3 Démodulateur I&Q.	17
Figure 2.4 Suppression du signal d'intercorrélation vs l'erreur d'amplitude	19
Figure 2.5 Suppression du signal d'intercorrélation vs. l'erreur de phase	20
Figure 2.6 Le vecteur d'erreur cumulatif (phase, amplitude et écart DC).....	21
Figure 2.7 Réflectomètre à base de circuit six-ports.	23
Figure 2.8 Le Points Q_i de la jonction six-ports.....	25
Figure 2.9 Démodulateur direct basé sur le principe du six-ports.	29
Figure 2.10 Le principe "Dual-Tone".....	31
Figure 2.11 Répartition des états simulés.....	34
Figure 3.1 Architecture POLYCOM.....	38
Figure 3.2 Réseau de neurones.....	43
Figure 3.3 Architecture ART ("Adaptive Resonance Technique").....	44
Figure 3.4 Architecture anneau du circuit six-ports.....	52
Figure 3.5 Constellation QPSK distorsionnée.....	53
Figure 3.6 Modèle SPW pour les erreurs de phase, gain et écarts DC.....	54
Figure 3.7 Convergence de l'algorithme.	63
Figure 3.8 Constellation I&Q de sortie avec le bruit gaussien.....	65
Figure 3.9 La caractéristique de probabilité d'erreur.....	67
Figure 3.10 Convergence de la constellation 8-PSK.....	68
Figure 3.11 Constellation de sortie 8-PSK avec le bruit gaussien	69
Figure 3.12 Énergie de la constellation 16-QAM	70
Figure 3.13 Convergence de l'algorithme pour une constellation 16-QAM.....	72
Figure 4.1 Sonde de mesure de type GSG.....	77
Figure 4.2 Schéma simplifié de mesure.	78
Figure 4.3 Ligne de transmission en réflexion après la calibration multi-lignes	83

Figure 4.4 Standards de calibration.....	85
Figure 4.5 Le principe de fenêtrage.....	86
Figure 4.6 Cheminement du signal pour la connexion "THRU".	87
Figure 4.7 Cheminement du signal pour la connexion "OPEN".	88
Figure 4.8 Dessin des structures de calibration.....	92
Figure 4.9 Modèle électrique de la monture.....	93
Figure 4.10 Schéma électrique des connexions.	94
Figure 4.11 Vérification de la correction.	95
Figure 4.12 Structure PHEMT sur GaAs, épitaxie par jet moléculaire.....	97
Figure 4.13 Gravure MESA (Définition Zones Actives)	98
Figure 4.14 Contacts Ohmiques	98
Figure 4.15 Définition de la grille par lithographie à faisceau d'électrons	98
Figure 4.16 Première couche de métal d'interconnexions	99
Figure 4.17 Résistance NiCr déposée par évaporation.....	99
Figure 4.18 Passivation Si_3N_4	100
Figure 4.19 Métal 2 interconnexions type pont-à-air	100
Figure 4.20 Système de fabrication de masques à balayage électro-optique	102
Figure 4.21 Détail du porte-échantillon.	102
Figure 4.22 Masque pour lithographie	103
Figure 5.1 Système de mesure sous-pointes.....	106
Figure 5.2 Capacité MIM en vue de haut et vue latérale.	108
Figure 5.3 Modèle de la capacité.....	108
Figure 5.4 Comparaison entre les valeurs mesurées et modélisées.....	110
Figure 5.5 Inductance surélevée.....	111
Figure 5.6 Modèle électrique de l'inductance.....	113
Figure 5.7 Masque pour une inductance spirale.....	114
Figure 5.8 Coefficient de réflexion mesuré et modélisé de l'inductance série.	115
Figure 5.9 Structure test Kelvin pour résistance NiCr.	117
Figure 5.10 L'extraction de la résistance spécifique ($\Omega/\text{carré}$).....	118

Figure 5.11 Structure transversale du transistor PHEMT	120
Figure 5.12 L'architecture en "T" (gauche) et "multi-doigts" (droite).	122
Figure 5.13 Modèle électrique du transistor en mode "petit signal".	124
Figure 5.14 Extraction de capacités intrinsèques C_{gs} et C_{ds}	128
Figure 5.15 Transconductance et le temps de transit.	128
Figure 5.16 Coefficients de réflexion S_{11} et S_{22} (mesurés et simulés).	129
Figure 5.17 Paramètres de transfert S_{21} (mesurés et simulés).....	130
Figure 5.18 Paramètres de transfert S_{12} (mesurés et simulés).....	130
Figure 5.19 Modèle du PHEMT en mode grand signal.	136
Figure 5.20 Coefficients de réflexion S_{11} et S_{22} (mesurés et modèle Root).	139
Figure 5.21 Paramètres de transfert S_{21} (mesurés et modèle ROOT).	140
Figure 5.22 Paramètres de transfert S_{12} (mesurés et modèle ROOT).	140
Figure 5.23 Modèle EEHEMT1.	141
Figure 5.24 Transconductance en mesure DC.....	143
Figure 5.25 S_{21} modèle EEHEMT1	144
Figure 5.26 S_{12} modèle EEHEMT1.	144
Figure 6.1 Grille du transistor PHEMT (~200 nm).....	147
Figure 6.2 Transistor PHEMT (Vue du haut).....	148
Figure 6.3 Détail de la grille.....	149
Figure 6.4 Courant de la grille en connexion diode.	150
Figure 6.5 Courant de grille en représentation logarithmique.....	150
Figure 6.6 Caractéristique de sortie I_D - V_D pour deux tensions de grille.....	151
Figure 6.7 Courant de fuite.....	152
Figure 6.8 Transistor PHEMT en réflexion.	153
Figure 6.9 Gain Maximal pour le transistor	154
Figure 6.10 Profil de la grille champignon.....	155
Figure 6.11 Mélangeur à transistor PHEMT.....	157
Figure 6.12 Réseaux biais pour la grille (a) et source (b).	158
Figure 6.13 Impédance de la sonde.....	160

Figure 6.14 Plage de variation de la tension de grille	161
Figure 6.15 Pertes d'insertion	162
Figure 6.16 Pertes par réflexion du mélangeur.	163
Figure 6.17 L'isolation LO-RF.	164
Figure 6.18 Mélangeur différentiel.	165
Figure 6.19 Transition CPWFG vers CPS – dessin (a) et modèle électrique (b).....	166
Figure 6.20 Pertes par réflexion de la transition CPWFG-CPS.	168
Figure 6.21 Pertes d'insertion pour la transition CPWFG-CPS-CPWFG	169

LISTE DES SYMBOLES ET DES ABRÉVIATIONS

ADC	Analog to Digital Converter
ADS	Advanced Design System
AGC	Automatic Gain Control
AM	Amplitude Modulation
ART	Adaptive Resonance Technique
ATM	Asynchronous Transfer Mode
BER	Bit Error Rate (Taux d'erreurs Binaires)
BJT	Bipolar Junction Transistor
CAD	Computer Aided Design
Cdg	Capacité Drain-Grille dans un FET
CDMA	Code Division Multiple Acces
Cds	Capacité Drain-Source dans un FET
Cgs	Capacité Grille-Source dans un FET
CPW	Coplanar Waveguide (Ligne Coplanaire)
CPWFG	Coplanar Waveguide with Finite Ground
CS	Coplanar Strips
DSP	Digital Signal Processor
DUT	Device Under Test
FET	Field Effect Transistor
$f_{\max}$	Fréquence Maximale d'Oscillation
f_t	Fréquence de Coupure pour un FET
GaAs	Gallium Arsenide (Arséniure de Gallium)
gm	Transconductance du Modèle Petit Signal dans un FET
$G_{\max}$	Maximum Stable Gain
GRAMES	Groupe de Recherche Avancé en Microondes et Électronique Spatiale
GMS	Groupe de Recherche en Microélectronique de Sherbrooke
G_T	Transducer Gain

HBT	Heterojunction Bipolar Transistor
HP8510C	Analyseur de Réseau Vectoriel de la Série 8510
HPIB	Hewlett Packard Interface Bus (IEEE 488)
ICCAP	Integrated Circuit
I_{ds}	Courant Drain-Source dans un FET
IF	Intermediate Frequency
I_g	Courant de Grille dans un FET
InP	Indium Phosphate (Phosphate d'Indium)
LASER	Light Amplification Stimulated by Emission Radiation
L_d	Inductance de Drain pour un FET
L_g	Inductance de Grille pour un FET
LMDS	Local to Multi-Points Distribution System
LNA	Low Noise Amplifier
LO	Local Oscillator
LPF	Low-Pass Filter
LRM	Line-Reflect-Match
L_s	Inductance de Source pour un FET
MDS	Microwave Design System
MESFET	Metal-Semiconductor FET
MHMIC	Miniaturized Hybrid Microwave Integrated Circuit
MIM	Metal-Insulated-Metal
MMIC	Microwave Monolithic Integrated Circuit
NNMS	Nonlinear Network Measurement System
PHEMT	Pseudomorphic High Electron Mobility Transistor
PM	Phase Modulation
RF	Radio Frequency
SNR	Signal to Noise Ratio
TDMA	Time Division Multiple Acces
TDR	Time Domain Reflectometry

TRL	Thru-Line-Reflect
UDS	Université de Sherbrooke
VCO	Voltage Controlled Oscillator
Vds	Tension Drain-Source dans un FET
Vgs	Tension Grille-Source dans un FET

CHAPITRE 1 INTRODUCTION

La demande accrue pour des taux de transmission toujours plus élevés a profondément modifié les contraintes de conception des circuits radios traditionnels. *N'importe où, n'importe quand et par n'importe quel moyen* sont les mots d'ordre pour les ingénieurs qui conçoivent les générations futures des systèmes de télécommunications à hauts débits. Chacun d'entre eux, est cependant conscient d'une contradiction: une bande passante extrêmement large nécessite l'utilisation de fréquences de plus en plus élevées, ce qui a des répercussions évidentes sur le choix de technologies, sur les coûts d'implémentation et sur la fiabilité des composantes. Ce paradoxe de la réalité quotidienne d'un ingénieur actif dans le domaine des micro-ondes, représente un défi pour les chercheurs qui proposent des nouvelles solutions pour les systèmes de télécommunications. Le travail présenté dans ce mémoire porte sur l'étude d'un récepteur radio pour les communications sans fils numériques. Comme on le montrera dans ce mémoire, ce récepteur pourrait être intégré dans des systèmes à hauts débits, et son architecture a été conçue dans l'optique de réduire les coûts de fabrication.

Dans ce chapitre d'introduction, nous proposons de faire tout d'abord une courte revue de l'architecture de certains systèmes de télécommunications en service, ou qui seront déployés prochainement. Nous nous limitons ici à des systèmes qui permettent d'illustrer l'intérêt de l'architecture de récepteur étudiée lors de ce doctorat.

Commençons par le système LMDS (Local-to-Multi-Point Distribution System), connu aussi, au Canada, sur le nom de LMCS (Local-to-Multi-Point Communication System). Ce système, dont le déploiement a déjà commencé dans plusieurs endroits au Canada, aux États-Unis et en Europe, pourrait s'avérer être un compétiteur redoutable des systèmes de distribution par fibre optique. Fondé sur une architecture de type cellulaire, il peut offrir un service de communication sans fils sur une bande passante extrêmement

large (1.3 GHz) pour un coût associé de 500-800\$ par unité (l'unité de l'abonné comprenant seulement le récepteur et transmetteur).

Les principales caractéristiques techniques du système nous montrent les spécifications pour des sous-composantes qui pourraient satisfaire les contraintes d'implémentation. On donne comme référence les valeurs standardisées par l'organisme gouvernemental américain dans le domaine, FCC (Federal Commision for Communications). L'assignation de la bande utile, couvre deux licences (pour de raisons des lois antitrust) :

- Licence A (1150 MHz) – Primaire (27,5-28,35 et 31,075-31,225 GHz HS,SH,HH)
Co-Primaire (29,1-29,25 GHz HS,HH).
- Licence B (150 MHz) Co-Primaire (31,0-31,075 et 31,225-31,3 GHz).

HS = station de base vers abonné, HH = station de base vers station de base,
SH = abonné vers station de base.

L'architecture du système a des caractéristiques très intéressantes :

- Réseau ATM (Mode de Transfert Asynchrone) pour haut débit.
- TDMA/FDMA schémas d'accès multiple.
- Commutation dynamique entre les types de modulation. Le mode d'accès FDMA (Accès Multiple à Répartition en Fréquence) est utilisé pour le services aux abonnés rapprochés alors que le mode d'accès TDMA (Accès Multiple à Répartition en Temps) est utilisé pour les usagers éloignés. Des réseaux en anneaux sont créés dans les cas de disparité entre les services requis.
- QPSK (Modulation de Phase en Quadrature) pour les débits variables (max. de 10 Mbs) et QAM (Modulation d'Amplitude en Quadrature) pour hauts taux de transmission à débit fixe (plus que 10 Mbs).

Une telle architecture réseau permet l'obtention d'une flexibilité accrue pour les divers types des services offerts. Dans ce contexte on peut pratiquement parler d'une intégration des services pour un réseau sans fils à large bande. Le tableau suivant décrit une partie des applications (services) qui peuvent être embarquer sur un système LMDS.

Multimédia résidentielle	Accès Internet	Tv/Vidéo Distribution
Travail à distance	Bureau Personnel	Connections LAN
Télé-Éducation	Surveillance Vidéo	Réseau Métropolitaine
Réseau privé virtuel	Réseau public des données	Passerelle PCS
Câble/Fibre interconnexions	Télémédecine	Réseaux Temporaires
Vidéo	Téléphonie	

Tableau 1.1 Réseau Sans Fils à Services Intégrés.

Soulignons ici une des caractéristiques principales de ce système, qui justifie que nous l'ayons pris comme exemple de référence dans le développement de notre architecture générique. Les modulations numériques utilisées sont de type QPSK et QAM. Pratiquement, plus que 90% des systèmes de radio numérique actuels (ou que l'on prévoit de déployer) fonctionnent avec ces types de modulation. Dans ce cas, les tests développés et l'analyse théorique du premier chapitre vont utiliser un type générique de modulation orthogonale (soit N-PSK ou N-QAM). À l'exception de quelques rares schémas de modulations du type spectre étalé ("spread spectrum"), notamment la modulation à saut de fréquence ("frequency hopping"), les modulations de fréquence peuvent aussi être traitées analytiquement comme une modulation du même type que celle analysée dans nos travaux (ex. GMSK peut être générée et démodulée par une architecture de type I&Q).

Une deuxième caractéristique importante porte sur l'aspect dynamique du traitement de l'information. Dans la majorité de cas, il y a une grande disparité entre les canaux d'une

liaison duplex. Habituellement, le débit vers l'abonné est de quelques ordres des grandeurs plus important que celui de l'abonné vers la station de base (par exemple, la télévision payante, accès Internet etc.). Dans ce cas, on doit utiliser des transmetteurs intelligents qui peuvent modifier de façon dynamique leur taux de transmission. Ainsi, pour le système LMDS, on doit être capable de changer le type de modulation en temps réel. Le type de modulation est ajusté suivant la valeur effective du débit, et ce pour optimiser les paramètres de la communication (taux d'erreurs, efficacité spectrale, coût des circuits). Nous verrons dans le chapitre 3, comment notre algorithme original nous permet de changer le schéma de modulation sans répercussions sur la complexité de la tête réceptrice ("front-end"). Dans ce cas, nous parlons d'une radio intelligente ("soft-radio"), dont la composante logiciel est très importante. Une grande partie du développement du système consistera alors en du développement de logiciel et en son intégration avec les composantes matérielles (hardware).

Le deuxième système de communication pris comme référence ici représente actuellement un effort énorme en termes de coût d'implémentation et de développement de solutions techniques. Il s'agit de la grande famille des systèmes de communications mobiles par satellites, capable de supporter les services de liaison mobile de 3^{ème} génération. Contrairement au système LMDS, qui apporte seulement une liaison fixe, ces systèmes apportent l'avantage de la mobilité aux usagers. Alors, ils peuvent aussi répondre à la contrainte "n'importe où". Malgré les bandes de fréquences différentes par rapport au système LMDS, les deux systèmes partagent les mêmes fonctionnalités : En effet, les deux utilisent des communications du type 'en ligne de vue' ("line of sight") et, ils offrent le même type de services (les critères du marché imposent les fonctionnalités des systèmes de communications); Ainsi les services accessibles, présentés dans le tableau 1, reste un dénominateur commun.

Même si le premier système de communications par satellite mis en orbite (IRIDIUM) a connu un échec financier, les caractéristiques techniques de ces systèmes sont une bonne

illustration d'applications possibles pour ce travail de doctorat. Il faut dire que chaque usager est pourvu d'un système récepteur-transmetteur ("transceiver"). Ces systèmes s'adressent donc à un marché de masse. On comprend alors l'intérêt pour un progrès au niveau du coût et/ou des performances en termes de taux d'erreurs par bit (BER).

		GEO Satellite Géostationnaires	LEO Satellite à Orbite Basse
Système	Fréquence Couverture Angle d'élévation Délai Réserve Commutation Entre satellite	Bande Ka Ondes millimétriques Globale sauf les pôles Plus que 5 degrés 270 msec. Réserve en orbite Non	800 MHz, Bande L Bande S Globale 10-20 degrés 10-30 msec. Mutuelle Possible
Satellite	Gabarit Traitement de l'info Antennes EIRP Orbite	Très Large (2 tonnes) Nécessaire 10-30 m. Élevé Géostationnaire (36000 Km)	Petits. Sophistique Max. 5 m. Faible Variable (Iridium, 750 Km)
Terminal Usager	Volume Antennes Protection contre la déviation Doppler	Moyen, Directionnelles Verrouillage directionnel Nécessaire (fréquence élevée utilisée)	Portatif Omnidirectionnelles Sans verrouillage Obligatoire (satellites mobiles)
Notes	Liaison Couverture Faisceau Exemple	Deux Satellites Fixe ACTS (ÉU) COMETS (JP)	À besoin Mobile Iridium, ~5 Km/sec. Iridium, Globalstar, Odyssey

Tableau 1.2 Systèmes de Communications par Satellite.

Pour chaque famille des satellites (géostationnaire, à orbite basse ou moyenne) il y a plusieurs solutions envisageables au niveau de l'architecture du système. Indépendamment de la solution choisie, on peut observer (voir les caractéristiques du

systèmes IRIDIUM et GLOBALSTAR) qu'il existe des caractéristiques communes aux systèmes. Par exemple, en ce qui a trait à notre sujet de recherche, la tête réceptrice pour les récepteurs des terminaux usagers, les deux systèmes doivent répondre à des contraintes similaires : une excellente figure de bruit, une bande passante suffisamment large, un fonctionnement en hautes fréquences (même en ondes millimétriques) et un faible coût en raison du nombre important des clients potentiels.

- **IRRIDIUM**

- Système à basse altitude (785 Km), utilisant 66 satellites actifs plus 6 en réserve, déployés en 6 orbites polaires.
- 3168 Cellules avec 2150 nécessaires pour la couverture de l'entière surface terrestre. Chaque satellite possède 3 antennes en bande L.
- 50 Kbps TDMA pour un duplexage temporel pour les deux liaisons (ascendante et descendante). Liaisons inter-satellite en bande Ka, avec 4 liaisons croisées (avant, arrière, et 2 pour les orbites adjacentes).
- La bande passante est divisée en 12 sous-bandes, réutilisées 4 fois pour chaque satellite. Chaque faisceau a 80 canaux.
- La capacité totale est $2150 \times 80 = 172\ 000$ canaux.

- **GLOBALSTAR**

- Système à moyenne altitude (1414 Km), utilisant 48 satellites actifs plus 6 en réserve, déployés en 8 orbites polaires, couvrant les régions comprises entre 70 degrés sud et 70 degrés Nord.
- Liaison ascendante en bande L (1610-1625,5 MHz) et en bande S pour la liaison descendante (2483,5-2500 MHz). Les stations de base terrestres ont les connexions en bande C (5 GHz pour la section ascendante et 7 GHz pour a version descendante).
- Accès multiple à répartition en code ("CDMA") avec contrôle de la puissance.

- Matrices des antennes actives à balayage de phase pour le contrôle du faisceau et pour une meilleure figure de bruit.
- Diversité spatiale pour une fiabilité accrue (récepteur de type RAKE).
- Commutation progressive ("soft handoff") entre les faisceaux des satellites.

Les deux systèmes sont à la fine pointe de la technologie en ce qui concerne les techniques de routage pour les réseaux sans fils et l'architecture de la section RF. Cependant, il est clair que la qualité et le prix du service désiré, sont directement dépendants de facteurs technologiques, et ceci est vrai pour n'importe quel système considéré. Il ne suffit pas d'avoir une architecture optimale si les composantes n'arrivent pas à fournir les performances demandées. Malheureusement, il va toujours exister une contradiction entre le besoin d'augmenter la qualité des composantes (par exemple la plage de fonctionnement de transistors, leur gain et leur figure de bruit) et le prix de la technologie utilisée.

Un des buts de notre recherche est de montrer comment, tout en utilisant une technologie donnée, on peut par des améliorations au niveau de l'architecture interne de composantes, améliorer les performances et diminuer le coût d'un récepteur. Cependant, l'aspect technologique forme aussi, une partie importante de ce mémoire. L'élément commun de n'importe quel récepteur performant reste la vitesse de traitement de l'information. Dans la plage des micro-ondes et ondes millimétriques il y a des contraintes physiques (la mobilité de porteurs dans les composantes actives) qui limitent fortement l'utilisation des composantes à base de silicium. Une grande partie de notre travail va porter sur la mise au point de moyens de fabrication de composantes intégrées monolithiques pour les micro-ondes et les ondes millimétriques (MMIC). Un aspect phare de la recherche actuelle dans le domaine des récepteurs pour les hautes fréquences, est la tendance à la miniaturisation des circuits. Cela a pour avantage d'augmenter la fiabilité des circuits et de permettre d'intégration de plusieurs fonctions sur la même puce. On parle maintenant de "tête réceptrice", "d'ensemble transmetteur–

récepteur”, et de plus en plus, on trouve sur le marché, des circuits qui intègrent des fonctions habituellement utilisées dans le traitement analogique du signal (amplificateur, mélangeur, oscillateur, multiplicateur de fréquence, commutateur, détecteur de puissance, etc.).

Alors, le projet présenté ici a évolué suivant deux axes. Un premier volet du projet porte sur l'analyse théorique d'un nouveau récepteur de type homodyne (direct) et sera couvert en détails dans la première partie de ce mémoire. L'idée du récepteur direct représente un sujet de recherche important pour le groupe du professeur Renato Bosisio au sein de laboratoire PolyGrames de l'École Polytechnique de Montréal (EPM) et elle est née suite au développement des circuits six-ports, modifiés pour les besoins spécifiques d'une architecture de récepteur. Une nouvelle approche de récepteur direct, en particulier avec une stratégie de traitement des signaux novatrice, va être présentée minutieusement dans le chapitre 3. De plus, nous allons comparer, grâce à des simulations de fonctionnement, les performances de la nouvelle version du récepteur homodyne, avec celles du récepteur homodyne basée sur l'architecture six-ports. Pour référence, nous comparerons aussi les performances de notre architecture à celles obtenues avec un récepteur I&Q classique (hétérodyne).

Dans un deuxième temps, nous nous sommes efforcés de concevoir et fabriquer des sous systèmes (amplificateur, mélangeur, etc.) de notre récepteur original. La partie "fabrication" a été développée en étroite collaboration avec le groupe du professeur Jacques Beauvais de l'Université de Sherbrooke (UDS). Les installations du GMS (Groupe de Recherche en Micro-électronique) à UDS permettent la fabrication de circuits MMIC's. Au laboratoire Polygrames de l'EPM, nous avons utilisé les outils de conception et de modélisation par ordinateur, les outils de mesures micro-ondes, et les appareils pour la fabrication de masques. Tout ceci a permis la mise au point des techniques de conception, de mesure, et de caractérisation pour des composantes monolithiques à hautes fréquences (fréquences millimétriques). À notre connaissance,

c'est la première fois que l'on a pu ainsi développer des circuits MMIC fabriqués à 100% dans le milieu universitaire canadien. L'expérience acquise tout au long de ces trois années de collaboration a apporté deux acquis majeurs. Au premier rang, évidemment, on a eu l'occasion de mettre en place un ensemble unitaire qui permet l'intégration des 3 étapes cruciales dans le développement des circuits : la modélisation la fabrication et la caractérisation. Une telle plate-forme unitaire a permis dans la cadre de ce projet de quantifier l'effort (au niveau de prix et performances) d'implémentation d'une architecture MMIC pour un récepteur homodyne. En deuxième lieu, la nouvelle réalité créée par l'existence d'une unité de fabrications de circuits MMIC dans nos laboratoires, constitue une base de départ pour les activités de recherche sur ces circuits.

Finalement, comme un guide dédié au lecteur, nous présentons la structure de ce mémoire. Le chapitre deux fait une revue de différentes techniques de correction d'erreurs utilisées pour l'amélioration du comportement d'un récepteur directe qui utilise la jonction six-ports. De plus, quelques caractéristiques importantes qui différencient un récepteur homodyne par rapport à l'architecture super-hétérodyne sont misent en relief.

Le troisième chapitre fait une analyse comparative entre l'architecture classique du récepteur super-hétérodyne avec démodulateur de type I&Q et une nouvelle architecture de récepteur direct. De plus, les avantages de la nouvelle approche sont analysés de façon qualitative et il sont comparés aux résultats obtenus avec les 'anciens' récepteurs utilisant la jonction six-ports. Les simulations de l'algorithme proposé (en considérant certaines erreurs par rapport aux conditions de réception optimales, ce qui correspond aux conditions réelles de fonctionnement) ont permis la validation quantitative du fonctionnement du nouveau récepteur qui utilise une nouvelle technique de correction.

Le quatrième chapitre présente les techniques de mesure à hautes fréquences et les systèmes de mesures mis en place à PolyGrames pour caractériser des composantes

MMIC. On décrit aussi brièvement les installations pour la fabrication des circuits MMIC et la technologie choisie.

La cinquième partie du mémoire approfondie les détails de la modélisation de composantes MMIC développées au long de notre travail. Le lecteur aura l'occasion de voir les différents compromis faits entre diverses solutions, les modèles obtenus et les limites de nos moyens de fabrication en termes de possibilités et aussi de qualité.

Le dernier chapitre décrit les résultats finaux sur la conception d'un mélangeur fonctionnant dans la plage de fréquence ciblée. Le mélangeur est la composante la plus importante de la tête réceptrice. Nous décrivons aussi, une transition de type "BALUN" qui est utilisée dans l'architecture différentielle de notre récepteur, et qui est potentiellement supérieure au niveau des écarts DC à l'architecture 'non-balancée'. En conclusion, le sommaire de l'ensemble de résultats obtenus est présenté, et nous faisons quelques recommandations pour les futurs travaux de caractérisation de composantes des circuits intégrés, opérants dans la bande millimétrique.

CHAPITRE 2 RÉCEPTEUR DIRECT (HOMODYNE). ANALYSE THÉORIQUE

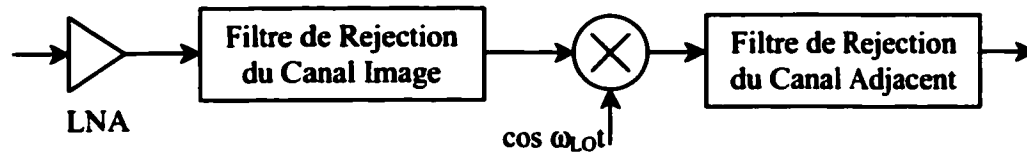
Dans la plupart des cas, l'étude d'un système déjà existant, sous une ou plusieurs versions (en l'occurrence ici un récepteur), doit répondre à une question relativement simple : Comment peut-on diminuer le coût du produit, tout en gardant les paramètres de fonctionnement à l'intérieur de la plage requise ? Au cours de ce chapitre, nous allons essayer de présenter de façon progressive, les principales caractéristiques d'une nouvelle architecture, capable de répondre à cette contrainte.

2.1 Récepteur Direct (Homodyne) versus Récepteur Super-Hétérodyne.

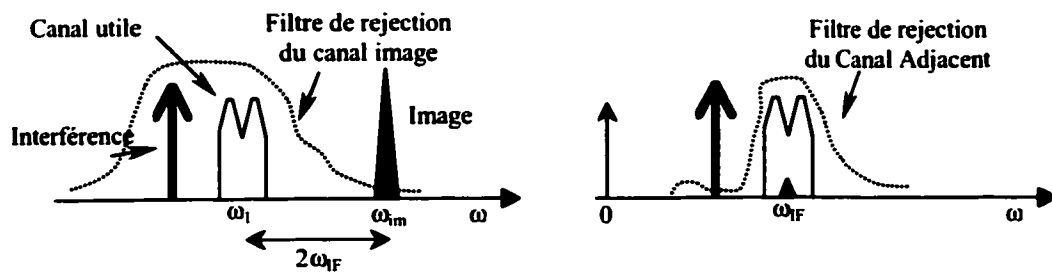
Indépendamment de topologie choisie et la technologie utilisée pour la réalisation physique d'une tête réceptrice ("front-end"), on peut distinguer deux classes des récepteurs : la version super-hétérodyne et la version directe ou homodyne. Il faut dire qu'il existe un déséquilibre entre le nombre d'applications qui utilisent l'une ou l'autre architecture. Actuellement, les versions super-hétérodyne dominent le marché, étant présentées dans plus de 90% de cas. La figure 2.1 (a et b) présente le principe de fonctionnement de chaque classe. Nous allons maintenant expliquer comment le choix de l'une ou l'autre architecture peut influencer les performances obtenues et les coûts de fabrication.

L'architecture super-hétérodyne a été introduite dans les années 20, par l'américain Armstrong, avec le but précis de palier certains défauts des récepteurs existants à l'époque. Le principal défaut était la résistance aux interférences, principalement le "brouillage" introduit par le canal adjacent. La référence était le canal de radio de type AM (modulation en amplitude avec un facteur de modulation de 30%) avec une bande passante de 9 kHz. Ce type de modulation reste en place même aujourd'hui pour les postes de radio émettant dans le domaine de fréquences moyennes. Les possibilités

limitées de l'époque en ce qui concerne la technologie de filtrage, ne permettaient pas la construction d'un filtre passe-bas avec un facteur de forme suffisant, pour bien



2.1a



2.1b

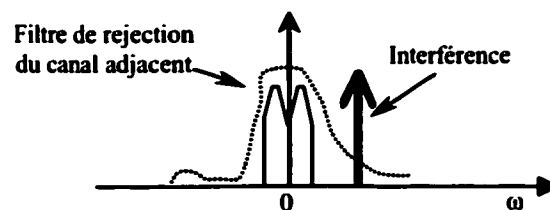


Figure 2.1 Conversion de fréquence
dans les architectures super-hétérodyne (2.1a) et directe (2.1b)

sélectionner le canal utile. En conséquence, le processus de démodulation est réalisé à une fréquence intermédiaire, dont la valeur permet plus facilement la construction d'un filtre passe-bande. Le filtre va présenter un ordre largement inférieur par rapport au filtre passe-bas.

Dans les deux cas, on a besoin d'un amplificateur à faible bruit (LNA), qui est la composante la plus importante pour la sensibilité du circuit. La différence majeure entre les deux architectures, d'où le nom de chaque catégorie, est la fréquence de l'oscillateur local. Pour la variante super-hétérodyne, elle a une valeur supérieure à la valeur de la porteuse ($\omega_{LO} = \omega_{RF} + \omega_{IF}$). Dans la version directe (homodyne) les deux signaux (de l'oscillateur local et la porteuse du signal RF) ont la même fréquence. Ces multiples conversions de fréquences ont comme conséquence l'augmentation de la complexité du circuit pour la variante super-hétérodyne. Le mélangeur qui réalise la conversion vers la fréquence intermédiaire produit un spectre complexe à sa sortie. On y retrouve toutes les combinaisons de type $\omega_{sortie} = m \cdot \omega_{LO} + n \cdot \omega_{RF}$ où les m, n sont des nombres entiers. La présence dans le signal d'entrée, d'une composante de fréquence $\omega_{image} = \omega_{LO} + \omega_{IF}$ va générer un signal de sortie dans la bande utile. Cette composante spectrale, nommée la fréquence image doit être absolument éliminée avant la conversion. Dans ce cas, on doit utiliser un filtre passe-bande, dans le domaine RF, pour obtenir une atténuation suffisante sur le canal image. Au contraire, l'architecture de conversion directe, utilisant une fréquence de l'oscillateur local identique à la porteuse, ne présente pas un canal image. Dans ce cas on parle d'une fréquence intermédiaire nulle.

Ces effets sur les fonctionnalités de chaque circuit à utiliser, montrent deux avantages de l'architecture directe :

- *La complexité.* L'inexistence de la fréquence intermédiaire a comme conséquence immédiate l'absence de l'électronique sur ce niveau. Un récepteur à conversion directe va présenter moins de circuits que son analogue super-hétérodyne. Cette caractéristique représente une façon directe de réduire le coût du circuit.
- *Spécifications élargies.* Le filtre de sélection du canal : Dans le cas du récepteur direct peut être un simple filtre passe-bande, sans un facteur de qualité très élevé, ayant la seule fonctionnalité de sélectionner la bande RF du système. Dans le cas d'un système à bande modérée, les caractéristiques résonantes d'une antenne de

réception ou le filtrage amené par l'amplificateur faible bruit (étant toujours conçu avec une caractéristique de type passe-bande) peuvent suffire.

Une autre différence réside dans le fait que pour un système de communication à plusieurs canaux, l'architecture super-hétérodyne doit faire un alignement entre la fréquence centrale du filtre de rejection sur le canal image et la synthèse de la fréquence de l'oscillateur local. Dans les récepteurs modernes, le signal de l'oscillateur local est habituellement généré par un synthétiseur de fréquence, donc l'alignement doit être réalisé par un filtre variable au niveau du signal RF. Ce facteur augmente le coût de fabrication.

Pour les récepteurs en ondes millimétriques, une différence, plus délicate à mettre en évidence, présente des répercussions au niveau de la conception et du coût global du circuit. La tendance actuelle est à l'intégration des composantes, dans le but de diminuer le coût de fabrication et surtout augmenter la fiabilité des circuits. Le domaine de micro-ondes ou des ondes millimétriques, nécessite, dans la plupart des cas, l'utilisation d'une technologie de haute qualité en termes de figure de bruit et de gain associé. Dans ce cas les circuits utilisant des semi-conducteurs de type III-V, comme le GaAs ou InP sont principalement utilisés. Dans le cas d'utilisation de fréquences intermédiaires relativement basses (ex. 1200, 900 et 70 MHz pour les systèmes DBS), l'optimisation de performances montre un avantage pour les technologies à base de silicium. L'intégration de la partie RF et de la partie IF du récepteur ne va alors pas de soit alors qu'une architecture directe utilisera la même technologie pour l'intégralité de la tête réceptrice.

L'architecture super-hétérodyne peut utiliser des composantes spéciales pour éviter le désavantage induit par la présence du canal image. Il s'agit d'un type spécial de mélangeur à rejection d'image (figure 2.2), qui élimine l'influence d'un potentiel signal d'interférence sur la fréquence image. Malgré l'amélioration apportée par cette solution, la conception d'un tel système augmente la complexité du circuit et consécutivement le

coût. De plus, en hautes fréquences, ce circuit présente des erreurs de phase et d'amplitude, causant une diminution des performances. Comme ces erreurs sont similaires avec celles produites par le récepteur de type I&Q, on va détailler leurs effets d'une manière quantitative dans la section dédiée (2.2).

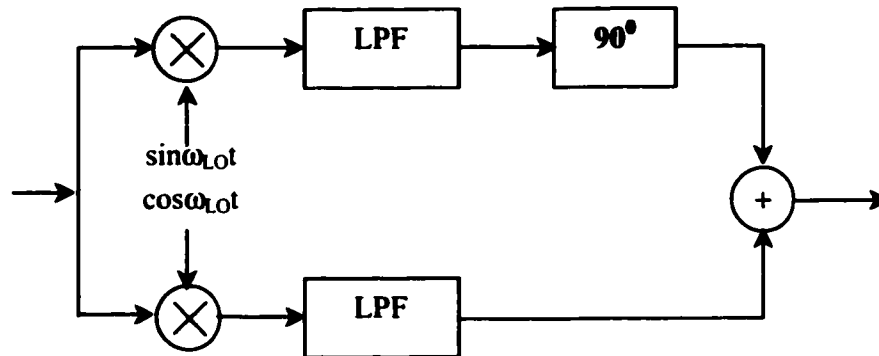


Figure 2.2 Mélangeur à rejection du canal image.

Il n'en reste pas moins que le récepteur super-hétérodyne présente certains avantages sur sa contrepartie homodyne. On ne doit pas "démoniser" cette architecture qui est finalement le standard pour les récepteurs actuels. Il y a deux facteurs importants qui jouent en faveur de la version super-hétérodyne :

- **Sélectivité.** Son principal avantage demeure de la raison pour laquelle on l'a choisi à l'aube de l'ère électronique. Il est plus facile à sélectionner le canal utile à une fréquence intermédiaire. La protection contre les interférences dans les canaux adjacents va être toujours meilleure que la variante directe. Dans l'analyse de notre architecture proposée, nous allons parler en détail comment on peut contrôler quand même les effets de ces interférences. Les deux niveaux de filtrage et l'amplification sur des fréquences différentes (RF et IF) permettent un fonctionnement supérieur dans un milieu hostile (brouillage, signaux faible de type radar, etc.).

- **Erreurs DC.** Un des principaux problèmes associés avec la conversion directe reste la présence d'un écart DC ("offset"). Pour une architecture super-hétérodyne le traitement analogique ne se fait pas au niveau de la composante continue, donc il n'y a pas un besoin de considérer ce paramètre dans la conception du système.

L'influence des erreurs DC sur la qualité du signal dans le cas de la conversion directe sera analysée quantitativement dans la section suivante. L'ensemble des caractéristiques présentés dans cette section forme les paramètres d'analyse pour l'évaluation du système que nous nous proposons d'étudier. Les méthodes développées pour contourner les désavantages de la variante homodyne seront décrites en détail, ceci dans le but d'apporter des arguments solides à l'effet que l'architecture directe peut présenter un avantage important en ce qui concerne le coût de développement et la flexibilité d'implémentation.

2.2 Architecture I&Q.

Que l'architecture choisie soit directe ou non, la structure du démodulateur influence de façon très importante les performances du récepteur au niveau de la qualité du signal. La majorité des démodulateurs actuellement utilisés sont de type I&Q (figure 2.3) en raison de la simplicité d'implémentation et la flexibilité de la topologie pour accommoder plusieurs schémas de modulation.

Nous allons commencer notre analyse par la présentation du principe de démodulation I&Q, et expliquer comment les imperfections dans réalisation physique du circuit peuvent introduire des erreurs de démodulation.

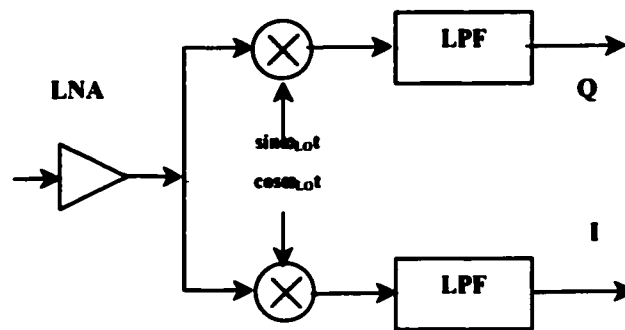


Figure 2.3 Démodulateur I&Q.

Soit le signal modulé en représentation complexe:

$$Z_m = \Re\{X(t)\exp(j\omega t)\} \quad 2.2.1$$

où $X(t)$ = le signal modulé dans la bande de base et ω la pulsation de la porteuse. Le signal Z_m est ici sous la forme la plus générale, on peut avoir une modulation de phase ou une modulation d'amplitude. Si on écrit :

$$X(t) = I(t) + jQ(t) \quad 2.2.2$$

l'équation 2.2.1 devient:

$$Z_m = \Re\{(I(t) + jQ(t))(\cos(\omega t) + j\sin(\omega t))\} = I(t)\cos(\omega t) - Q(t)\sin(\omega t) \quad 2.2.3$$

Les signaux I et Q sont les signaux dits en phase et en quadrature. Si les signaux I et Q sont générés par des états discrets (dont leur totalité forme un alphabet) on parle d'une modulation numérique. Sinon on a une modulation analogique. Il faut préciser qu'une source générale de confusion reste cette différence entre les deux types de modulations (analogique et numérique). Indépendant du type utilisé, le signal de sortie d'un modulateur est un signal continu (donc analogique). L'analogie du terme numérique provient du fait que dans le cas d'un modulateur à états discrets, dénommés symboles, on peut exécuter des opérations numériques sur les signaux en bande de base. Dans ce cas

on peut appliquer toutes les techniques de filtrage numérique, du codage et décodage, dans le but d'augmenter la robustesse du signal face au bruit.

Au niveau du démodulateur, on fait l'opération inverse du processus de modulation dans le but de retrouver les signaux I et Q porteurs d'information :

$$\begin{aligned} O_1 &= Z_m \cos(\omega t) = I(t) \cos^2(\omega t) - Q(t) \sin(\omega t) \cos(\omega t), \\ O_2 &= Z_m \sin(\omega t) = I(t) \sin(\omega t) \cos(\omega t) - Q(t) \sin^2(\omega t), \end{aligned} \quad 2.2.4$$

$$\begin{aligned} O_1 &= \frac{1}{2} I(t) + \frac{1}{2} I(t) \cos(2\omega t) - \frac{1}{2} Q(t) \sin(2\omega t), \\ O_2 &= \frac{1}{2} I(t) \sin(2\omega t) - \frac{1}{2} Q(t) + \frac{1}{2} Q(t) \cos(2\omega t), \end{aligned} \quad 2.2.5$$

Les composantes de haute fréquence (2ω) seront éliminées par les filtres passe-bas, et les signaux dans la bande de base (filtrés en sortie) seront proportionnels avec les signaux I, respectivement Q.

Il faut dire que le principe de démodulation suppose qu'au niveau du récepteur, on connaît avec exactitude la phase et la fréquence de la porteuse. Ce principe de démodulation cohérente ne tient pas compte des erreurs physiques introduites par les circuits non-idéaux. En réalité, il y a une erreur de phase, entre les deux chemins de traitement du signal et aussi une erreur d'amplitude, due au gain différent de l'ensemble diviseur de puissance, mélangeur, filtre passe-bas et déphaseur.

Reprenons l'équation 2.2.4, mais cette fois-ci, on considère une erreur de phase α et une différence de gain entre les deux mélangeurs, correspondant aux deux chemins O_1 et O_2 :

$$\begin{aligned} O_1 &= Z_m \mu \cos(\omega t + \alpha) = I(t) \mu \cos(\omega t) \cos(\omega t + \alpha) - Q(t) \mu \sin(\omega t) \cos(\omega t + \alpha), \\ O_2 &= Z_m \sin(\omega t) = I(t) \sin(\omega t) \cos(\omega t) - Q(t) \sin^2(\omega t), \end{aligned} \quad 2.2.6$$

Le développement de l'équation 2.2.6 suit facilement si on garde seulement les termes dans la bande de base :

$$\begin{aligned}\tilde{O}_1 &= \frac{1}{2}I(t)\mu \cos(\alpha) + \frac{1}{2}\mu Q(t)\sin(\alpha), \\ \tilde{O}_2 &= \frac{1}{2}Q(t),\end{aligned}\tag{2.2.7}$$

Dans l'expression 2.2.6 on a défini les erreurs en prenant le canal 'hors de phase' ' $\tilde{O}_2$ ' comme référence. On peut observer ainsi le terme d'intercorrélation sur le canal I, donc les données transmises sur le canal en quadrature (Q) vont modifier le canal en phase. Pour analyser quantitative l'effet négatif sur le processus de démodulation, on peut considérer comme paramètre l'expression de la suppression de la bande latérale non-désirée [98].

$$S[\text{dBc}] = 10 \log \left[\frac{\mu^2 + 2\mu \cos(\alpha) + 1}{\mu^2 - 2\mu \cos(\alpha) + 1} \right],\tag{2.2.8}$$

La figure 2.4 montre la valeur de la suppression du signal d'intercorrélation par rapport au niveau utile en fonction des erreurs d'amplitude (rapport de gain).

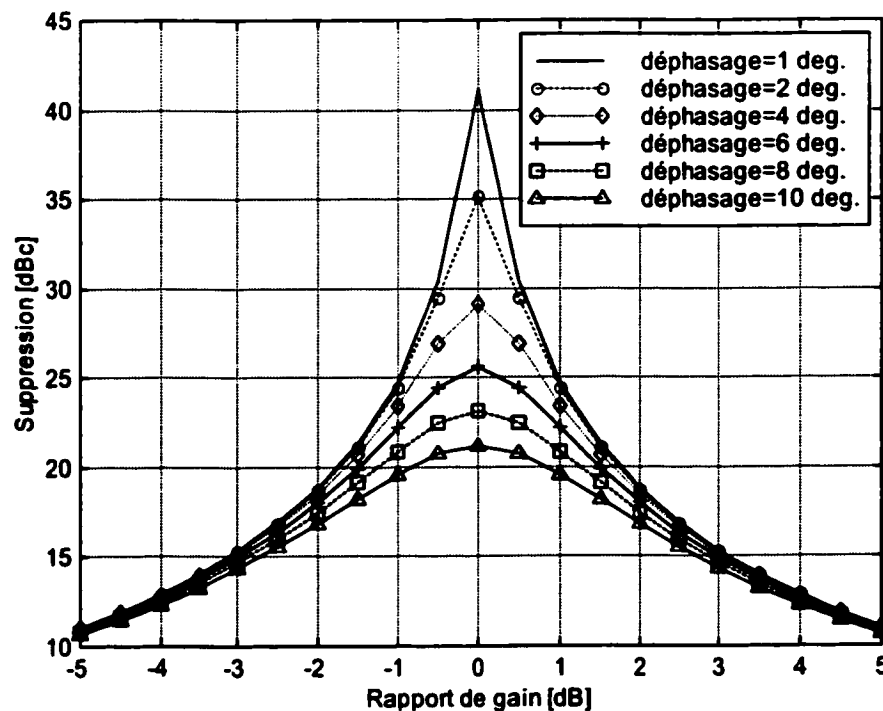


Figure 2.4 Suppression du signal d'intercorrélation vs l'erreur d'amplitude

L'analyse du graphique ci-dessus donne une valeur quantitative sur les effets induits par les erreurs d'amplitude, pour différentes valeurs du déphasage. Pour des erreurs de phase supérieure à 3 degrés et un rapport de gain entre les deux mélangeurs plus grand que 0.5 dB, l'atténuation du signal d'intercorrélation descend à un niveau relativement bas. La même expression peut être utilisée avec le rapport de gain comme paramètre. La famille de courbes résultante est présentée dans la figure 2.5.

Pour des circuits à très hautes fréquences, les contraintes deviennent très serrées, et la conception d'un circuit qui respecte ces performances est difficile et coûteuse en terme de temps de réalisation et le coût.

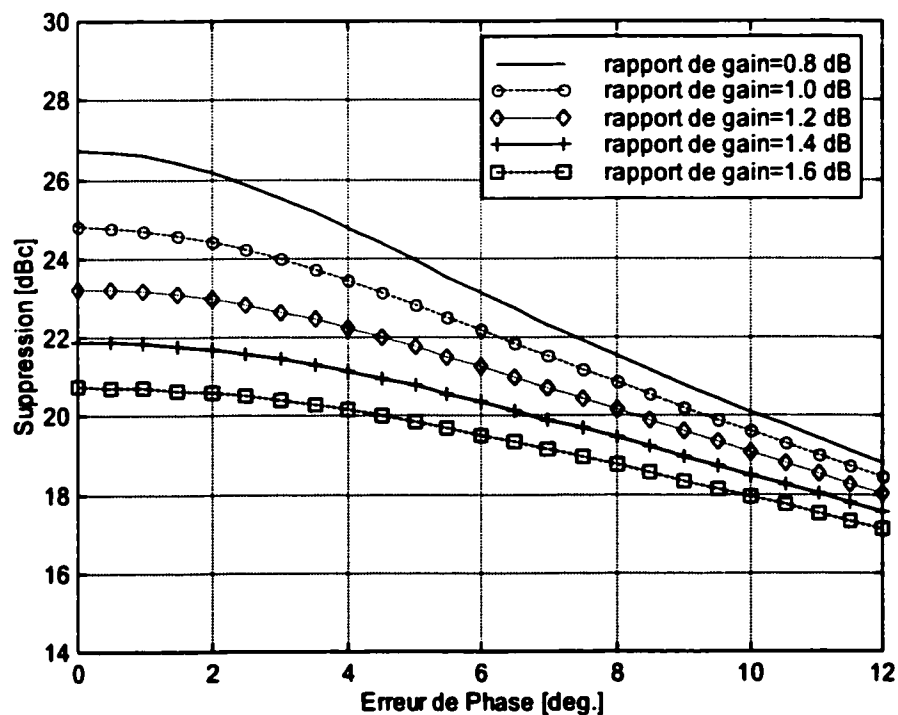


Figure 2.5 Suppression du signal d'intercorrélation vs. l'erreur de phase

Jusqu'au présent, on a analysé l'influence de ce type particulier d'erreurs. Les résultats nous indiquent que l'écart entre un circuit idéal (dont les erreurs de phase et amplitude entre les deux canaux sont nulles) et une réalisation réelle, peut influencer la qualité de

la réception. Cependant, il faut dire qu'il est difficile de quantifier avec précision, l'effet de ces erreurs sur le taux d'erreurs de transmission (BER). En effet, ce type d'erreurs affecte non seulement la prise de décision par rapport à la détermination des états I et Q, mais aussi le comportement du circuit de récupération de la porteuse. Ainsi, une erreur de phase sur le canal de référence (dans notre cas sur le canal O_2) modifie aussi l'écart entre la phase de la porteuse et la phase du signal de l'oscillateur local. Donc le phénomène est amplifié dans une certaine mesure. Pour les récepteurs directs, il faut aussi considérer un autre type d'erreur qui est dû à la présence d'une composante continue dans les signaux de bande de base qui permettent de déterminer les états I et Q. Cette composante fait en sorte que la position de l'origine de la constellation vectorielle des signaux est déplacée. Notons aussi que l'erreur de gain qui peut être produite par le diviseur de puissance doit être considérée en plus de l'erreur provenant des mélangeurs.

La figure 2.6 présente de façon qualitative les erreurs dont on a parlé.

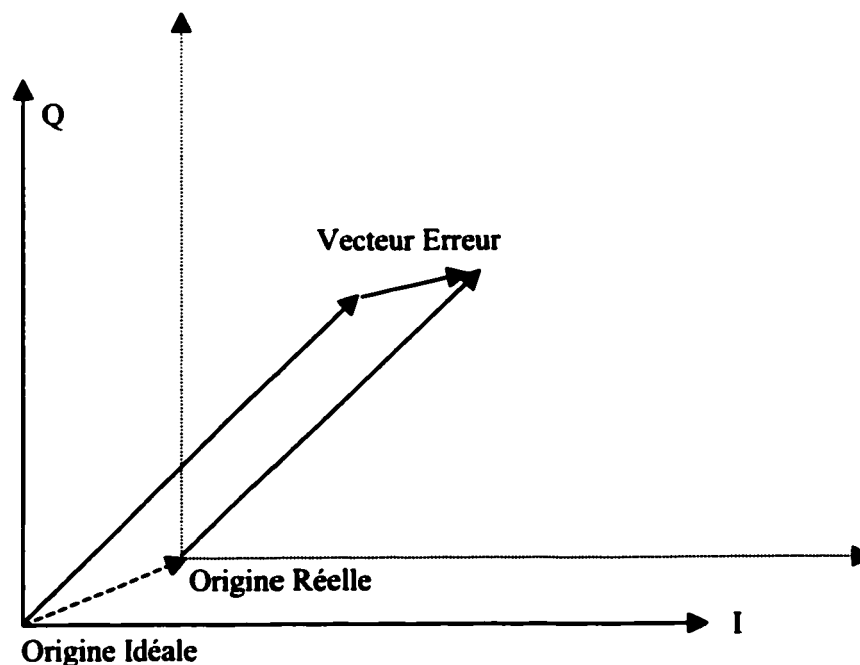


Figure 2.6 Le vecteur d'erreur cumulatif (phase, amplitude et écart DC).

De façon pratique ces erreurs traduisent la constellation des états modulés, et la correction d'erreur consiste à trouver 'l'origine réelle' de cette constellation. Ce processus de correction est le principal objet des améliorations que nous allons proposer dans le chapitre 3. Nous proposons maintenant de faire une revue des principales méthodes de correction décrites dans la littérature [2,95,98]. Ces techniques utilisent différentes méthodes de correction pour ce type d'erreurs linéaires. En particulier, nous allons expliquer quels paramètres ont été améliorés, et nous allons comparer les performances de ces méthodes.

2.3 Analyse des Techniques de Correction d'Erreurs.

Tout au long de la présente section, nous allons adopter une approche essentiellement chronologique, ce qui permettra de voir comment notre projet a évolué par rapport aux résultats publiés antérieurement dans la littérature. Parallèlement, anticipant ainsi la présentation de notre solution, nous allons commenter exclusivement les techniques dédiées à une classe particulière de récepteurs directs, celle qui a fait l'objet des plusieurs études au sein du groupe de recherche du professeur Bosisio : les récepteurs six-ports.

2.3.1 Jonction Six-Ports.

Les systèmes à base de la jonction six-ports ont fait leur apparition au début des années 70 [41-46,60]. La première application ciblée par ce type de circuit était la réflectométrie. Les idées qui ont poussé vers une telle structure sont tout à fait communes avec notre objectif : éliminer les erreurs de phase et d'amplitude dans la mesure d'un signal. Dans le cas du dispositif six-ports, le signal analysé était l'onde réfléchie par une charge, dont on voulait mesurer le coefficient de réflexion. Indépendamment du circuit utilisé, les erreurs physiques du circuit influençaient les résultats de mesure. La nouvelle technique 'six-port' permettait de corriger ces erreurs en utilisant un traitement mathématique similaire à celui d'un analyseur de réseau. L'avantage

principal de cette approche réside dans le faible coût du circuit six-ports et dans la possibilité de faire des mesures multi-ports et multi-harmoniques sans une augmentation considérable de la complexité des appareils de mesure. Une excellente source d'information pour les mesures de type multi-ports et multi-harmoniques est la référence [94]. Le lecteur peut y trouver une analyse comparative des avantages et des désavantages des réflectomètres de type six-ports.

La figure 2.7 présente le schéma typique d'un circuit six-ports dans sa configuration de base (réflectomètre):

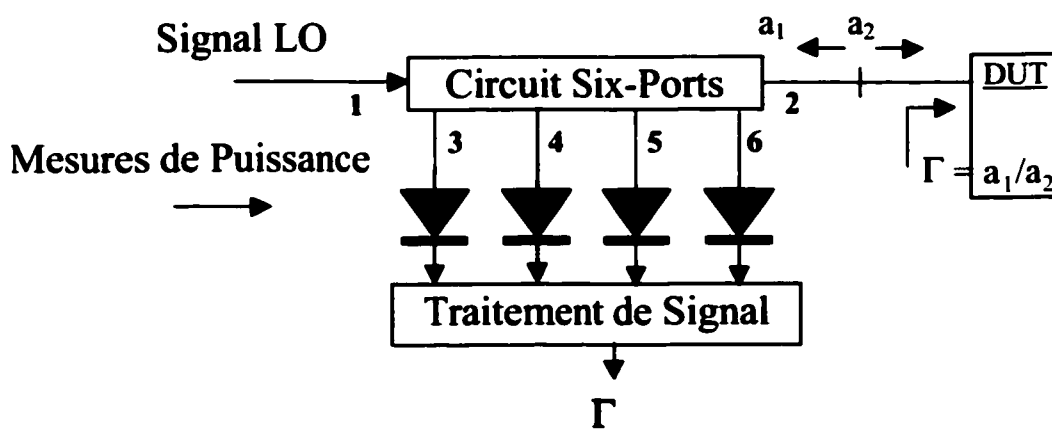


Figure 2.7 Réflectomètre à base de circuit six-ports.

Comme le nom de la jonction le suggère, le circuit possède six-ports, dont un port d'entrée (le signal de l'oscillateur local), un port bi-directionnel, relié au dispositif sous test et quatre ports de sortie dédiés aux mesures de puissance (scalaires). Le principe de fonctionnement est relativement simple. Le signal généré localement est injecté à la fois dans le dispositif sous test (en occurrence une charge inconnue) et dans la jonction. Sur chaque canal de sortie on retrouve une combinaison linéaire entre le signal réfléchi par la charge et le signal de l'oscillateur local. On comprend ici, d'un point de vue général,

que la mesure des quatres puissances sur les ports de sortie permettra de résoudre un système linéaire et calculer l'amplitude et la phase des vecteurs a_1 et a_2 . De façon usuelle, l'un des ports de mesure est pris comme port de référence. Une des propriétés essentielles de la jonction six-ports est que le signal de sortie du canal de référence doit être proportionnel seulement au signal de l'oscillateur local. Autrement dit, le port de référence est isolé par rapport au port de mesure. La valeur d'isolation constitue un critère de performance du circuit. L'unité de calcul du système 'six-port' va effectuer les opérations mathématiques suivantes : La puissance mesurée sur les ports de sorties est divisée par la puissance sur le canal de référence. Les valeurs résultant de cette opération sont utilisées dans le traitement mathématique pour calculer le coefficient de réflexion 'à mesurer'.

Les équations mathématiques qui gèrent le fonctionnement de la jonction peuvent être trouvées dans les premiers articles dédiés à ce principe de mesure [59-61]. Pour la simplicité de la présentation, nous ne présentons pas ici les outils mathématiques associés au 'six-ports'. Nous allons plutôt expliquer quels sont les principaux paramètres de fonctionnement du réflectomètre. Ceci permettra de mieux comprendre les améliorations amenées par notre architecture de récepteur et par notre algorithme de traitement des signaux.

Comme la figure 2.8 le suggère, le coefficient de réflexion de la charge peut être représenté comme un vecteur dans l'espace cartésien, caractérisé par son amplitude et sa phase. Pour une famille des vecteurs d'amplitude constante, le lieu géométrique décrit par leur vertex est un cercle. Chaque port de sortie en faisant une combinaison linéaire entre le vecteur inconnu et le vecteur fixe, représenté par le signal de l'oscillateur local, va transformer ce cercle en un autre, suite au principe des transformées conformes.

Pour chaque canal de sortie, les paramètres de la transformée conforme associée sont différents, donc la position des cercles et leur rayon, seront différents. La figure 2.8

montre dans le même plan de référence les cercles de sortie. Comme chaque cercle provient de la même famille de vecteurs, leur point d'intersection dans le plan cartésien doit forcément caractériser un vecteur unique! Ce point représente le vertex du vecteur inconnu.

Cette approche schématique a permis d'expliquer comment une jonction six-ports peut mesurer le coefficient de réflexion. Cette analyse qualitative nous indique aussi quelles précautions on doit prendre pour avoir un système de mesure efficace et performant.

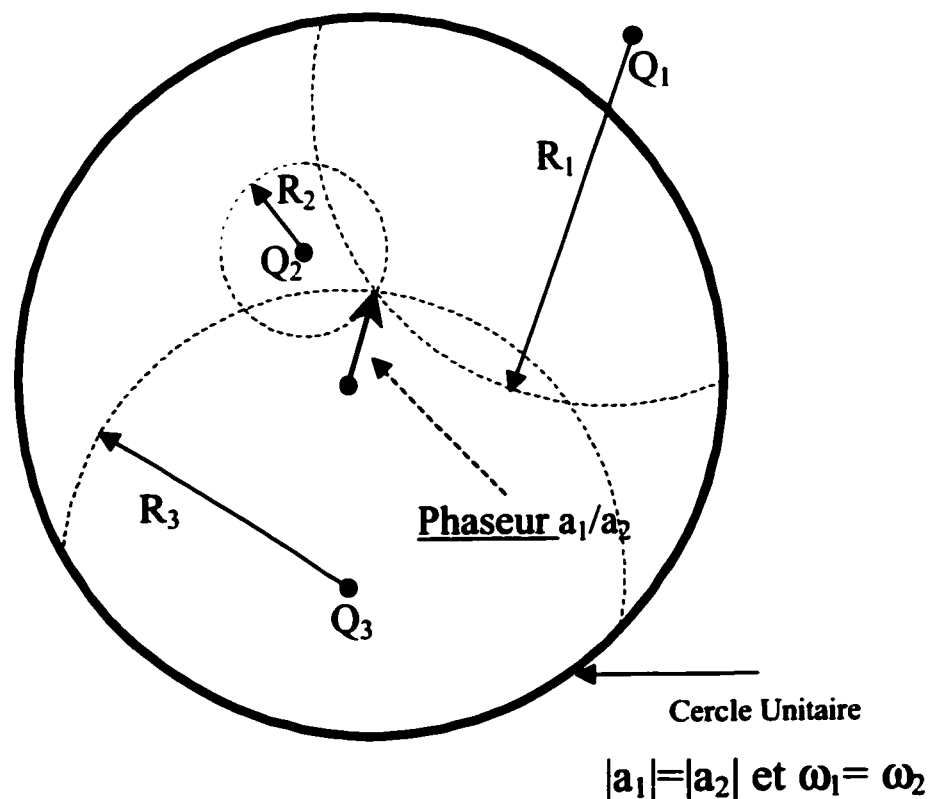


Figure 2.8 Les Points Q_i de la jonction six-ports.

En premier lieu, on doit s'assurer que le point d'intersection est unique. Ceci veut dire que les centres de cercles (les points Q_i) devraient être non-colinéaires. En termes pratiques, les déphasages subits par le signal réfléchi et le signal de référence ne doivent

pas être proportionnels entre les différents canaux (doivent former un système d'équations déterminé).

En deuxième lieu, il reste bien sûr la nécessité d'avoir un point d'intersection ! À ceci s'ajoute la nécessité d'avoir une précision constante, indépendante de la position du vecteur inconnu (sa phase). Ces conditions ont conduit à des critères de conception très précis, pour la jonction six-ports :

- La position des points Q_i par rapport à l'origine du cercle unitaire doit avoir un rayon normalisé de 1,5.
- Les trois points Q_i doivent avoir un déphasage réciproque de 120 degrés.
- L'isolation entre le port de mesure et le port de référence doit être très importante (~ 40 dB).

Aussi dans la figure 2.8 on observe que la position du vecteur inconnu est déterminée comme une position relative par rapport aux points Q_i . Pour un système de mesure capable de donner la valeur absolue du coefficient de réflexion, on doit déterminer la position absolue de points Q_i par rapport au cercle unitaire absolue. Cette procédure porte le nom de processus de calibration et représente le point fondamental de la théorie de mesure de la jonction six-ports. La procédure de calibration n'est pas seulement nécessaire pour la valeur absolue mais aussi pour la correction des erreurs de fabrication par rapport à la position idéale des points Q_i décrite ci-dessus. En d'autres termes, la procédure de calibration permet de référencer les positions exactes du centre des cercles, Q_1 , Q_2 , et Q_3 . Elle permet aussi de déterminer comment le rayon de ces cercles change et fonction du coefficient de réflexion à mesurer.

La procédure de calibration est similaire aux techniques de calibration d'un analyseur de réseau vectoriel. On mesure des dispositifs connus et on extrait les paramètres caractéristiques de la jonction six-ports. La procédure est extrêmement performante,

étant équivalente au niveau conceptuel avec la procédure de calibration de type TLR ("Thru-Line-Reflect"). Les deux techniques ont été développées par les mêmes auteurs [44,46] dans le but précis d'obtenir une technique de mesure de grande précision du coefficient de réflexion.

En 1983, Hodgetts et Griffin [59] ont présenté de façon unifiée la théorie de calibration sous formes des équations paramétriques :

$$\begin{aligned} & pQ_1^2 + qA^2Q_2^2 + rB^4Q_3^2 + (r-p-q)A^2Q_1Q_2 + (q-p-r)B^2Q_1Q_3 + \\ & + (p-q-r)A^2B^2Q_2Q_3 + p(p-q-r)Q_1 + q(q-p-r)A^2Q_2 + \\ & + r(r-p-q)B^2Q_3 + pqr = 0 \end{aligned} \quad 2.3.1$$

avec un changement de variable on peut écrire:

$$\begin{aligned} & X_1Q_1^2 + X_2Q_2^2 + X_3Q_3^2 + X_4Q_1Q_2 + X_5Q_1Q_3 + X_6Q_2Q_3 + \\ & + X_7Q_1 + X_8Q_2 + X_9Q_3 = -1 \end{aligned} \quad 2.3.2$$

Dans les équations 2.3.1 et 2.3.2, Q_i représente le rapport de puissance entre le port i et le port de référence, p, q, r, A, B sont les paramètres de la jonction (dépendent des paramètres S de la jonction et de la réponse des détecteurs de puissance). Les auteurs ont prouvé que, en mesurant 13 standards différents, on peut résoudre un système sur-déterminé par la méthode de moindres carrées.

En principe, la procédure de calibration est relativement simple pour l'utilisateur. L'utilisateur mesure les rapports de puissance de sortie en connectant consécutivement 13 charges standard, au port de mesure. Normalement, les charges sont choisies pour couvrir uniformément l'abaque de Smith. Une solution pratique peut être la suivante :

1 charge adaptée, 1 circuit ouvert sans atténuation et avec atténuateurs de 3 et 6 dB, 1 court-circuit sans atténuation et avec atténuateurs de 3 et 6 dB, 1 court circuit coulissant

avec atténuateur de 3 dB en mesurant en 6 positions (45, 90, 135, 225, 270 et 315 degrés).

À la fin de cette section, dédiée aux circuits de type six-ports en tant que réflectomètre, on doit souligner certains aspects liés aux lectures de puissance. Dans leurs premières versions, les appareils utilisaient des puissance-mètres externes. Suite aux progrès de la miniaturisation, des diodes ont été intégrées aux six-ports pour effectuer les mesures de puissance. Cette technique, oblige l'utilisateur à faire un traitement supplémentaire des signaux de sortie. En effet, la tension de sortie des diodes n'est proportionnelle à la puissance d'entrée que dans une plage limitée. Donc un processus de linéarisation est nécessaire avant l'étape de calibration décrite ci-dessus. L'avantage de la solution avec diodes intégrées est qu'elle est peu coûteuse, utilisée avec un processeur de signal (DSP), elle sera la solution de choix pour un système portable, peu volumineux et flexible.

Par rapport à un analyseur vectoriel standard, le circuit six-ports présente certains avantages :

- Possibilité de mesure multi-harmoniques.
- Mesures de charges actives ($|\Gamma| > 1$).
- Possibilité de mesure multi-ports (mélangeurs , diviseurs de puissance etc.).
- Très faible coût.

Bien sûr, le gain n'est pas total, il existe aussi certaines limites associées aux six-ports:

- Bande de fréquences moins large.
- Bande dynamique limitée (due à la linéarisation).
- La nécessité d'un utilisateur avisée.

2.3.2 Récepteur Direct.

Au-delà de l'application en réflectométrie, on note que d'un point de vue général le principe de fonctionnement de la jonction six-ports est le suivant : Le système 'six-ports' permet de mesurer des rapports de puissance entre deux vecteurs (dans le plan complexe). Si on considère un des vecteurs connus (le signal de l'oscillateur local), le système mesure l'amplitude et la phase de l'autre. Finalement, on peut reconnaître la similitude avec le processus de démodulation qui demande au récepteur de calculer la phase ou (et) l'amplitude du signal d'arrivé, porteur d'information. La figure 2.9 illustre ce principe de démodulation :

Dans la figure 2.9 et l'équation 2.3.3 a_1 représente un signal proportionnel au signal d'oscillateur local (moins un déphasage fixe) et a_2 le signal RF modulé en phase ou amplitude.

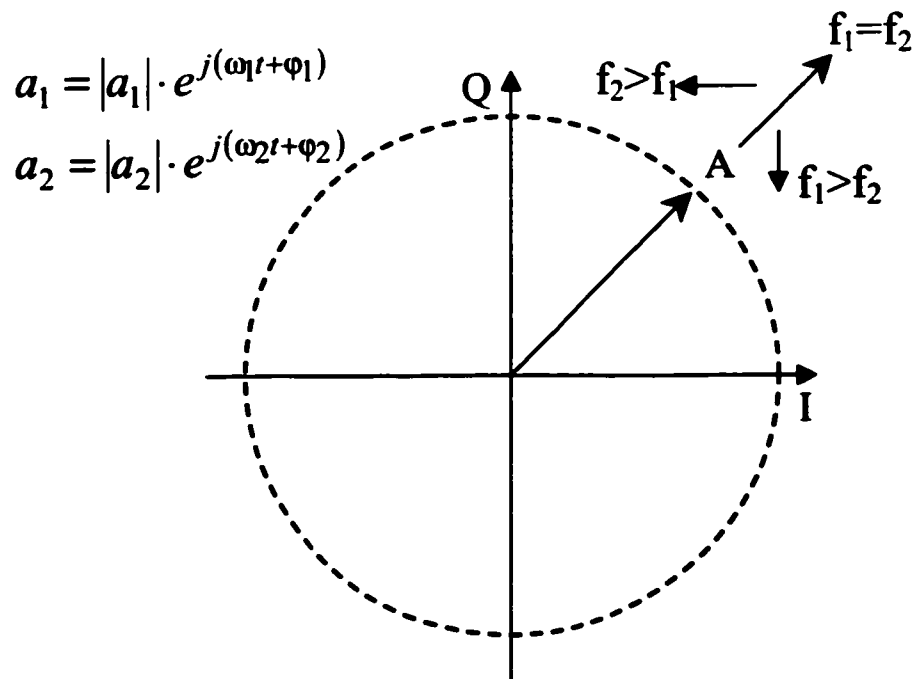


Figure 2.9 Démodulateur direct basé sur le principe du six-ports.

$$\frac{a_1}{a_2} = A = |A| \exp(j\varphi) = \frac{\sum_{i=1}^4 (R_i + jI_i) P_i}{\sum_{i=1}^4 D_i P_i} \quad 2.3.3$$

Supposons que le vecteur d'arrivé (signal à démoduler) est modulé en fréquence. Alors le vertex du vecteur 'A' décrit un cercle, en tournant dans le sens horaire ou anti-horaire, suivant la différence entre les fréquences instantanées de deux signaux. Si notre jonction est capable de mesurer la phase et l'amplitude d'un tel vecteur, évidemment nous obtiendrons la démodulation du signal porteur d'information. Ce principe, exposé pour la première fois par Li et Bosisio [75] a montré la possibilité d'utiliser la jonction six-ports dans une application de récepteur direct. La raison d'un comportement de type homodyne est très simple. Pour la mesure du coefficient de réflexion, la fréquence de signal réfléchi est identique avec celle de l'oscillateur local. Par l'analogie, si on désire démoduler un signal RF complexe, la fréquence de la porteuse doit être strictement identique avec celle de l'oscillateur local, donc il s'agit d'une architecture homodyne.

Malgré la simplicité du principe, la mise en pratique d'un récepteur direct basé sur la jonction six-ports restait un défi. Pour implémenter un système de réception de bonne qualité, certaines contraintes doivent être rencontrées. La plus importante de ces contraintes est la procédure de calibration. Comme on l'a vu, celle ci permet de corriger les erreurs physiques du circuit rf ce qui est un avantage par rapport à une architecture I&Q. Malheureusement la technique classique 'à 13 charges' ne peut pas être utilisée de façon viable pour un récepteur. En effet, on ne peut pas demander à l'utilisateur d'un téléphone cellulaire de faire l'opération de calibration avant d'avoir accès au réseau. De plus la calibration doit être faite dans un temps très court, négligeable par rapport au temps de communications. Ces nouvelles contraintes ont poussé les chercheurs à trouver de nouvelles techniques de calibration, dites 'en temps réel', qui feront l'objet de notre étude critique dans les prochaines sections.

2.3.3 La Calibration de Type "Dual-Tone".

Quelle que soit la méthode de calibration, une contrainte devra toujours être rencontrée. : Nous avons besoin de standards connus, pour résoudre le système d'équations de la jonction. Pour satisfaire la condition 'temps réel' une solution est de synthétiser ces standards de façon 'électronique'. Dans cette optique, Li et Bosisio [77] ont développé une méthode robuste. Leur méthode trouve ses fondements dans la méthode classique 'à 13 charges' : Le vecteur incident (le signal RF monochromatique) a une fréquence proche de celle de l'oscillateur local. Alors les signaux de sortie au niveau de la jonction sont des signaux purement sinusoïdaux. En regardant la figure 2.10, représentant 'A' le signal de sortie, on peut associer certains niveaux de signal au signal que produirait une charge fictive située sur le port où arrive le signal RF. Cette charge fictive varie avec le temps ou de façon équivalente avec le déphasage relatif entre le signal RF et le signal LO. Ainsi par exemple, 'A' avec une amplitude maximale correspond à un circuit ouvert, alors que 'A' avec une amplitude nulle correspond à une charge adaptée.

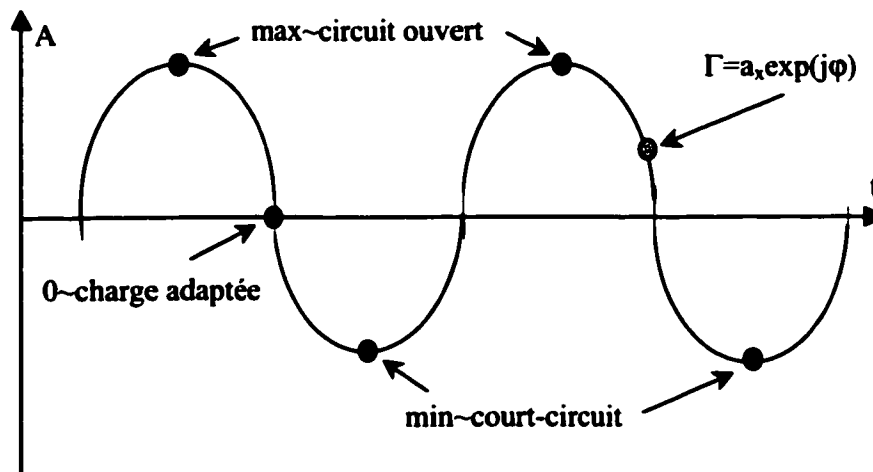


Figure 2.10 Le principe "Dual-Tone".

Un récepteur utilisant cette technique de calibration, devrait en principe fournir un signal monochromatique (simple sinusoïde), qui a une fréquence différente par rapport à la fréquence de l'oscillateur local. Le produit de battement des signaux sera mesuré et un nombre suffisant des valeurs sera échantillonné et gardé en mémoire. En temps réel, on analysera les signaux de sortie de la jonction, et on choisira au moins 13 positions différentes sur une période. Alors on va résoudre le système d'équations 2.3.2 pour déterminer l'origine et l'amplitude des vecteurs Q_i .

Dans les résultats de la référence [77], Li et Bosisio ont utilisé les mêmes positions correspondants aux critères optimaux, énoncées dans la présentation de la méthode à 13 charges.

Une explication simple du principe de cette technique montre clairement qu'elle répond aux contraintes énoncées pour le fonctionnement d'un récepteur direct: C'est une méthode automatisée, une méthode exacte (elle a les mêmes performances que la procédure standard) et elle est relativement rapide (la vitesse dépend de la façon dont l'algorithme de calcul est implanté).

Cependant, elle présente un désavantage : on doit fournir un signal de haute fréquence supplémentaire par rapport au signal de l'oscillateur local. En plus la différence de fréquence doit être relativement faible (~ 50 -200 Hz), alors la précision de synthèse de fréquence doit avoir une résolution comparable avec ces valeurs. Ces deux facteurs appliqués à un système opérant dans le domaine des micro-ondes ou des ondes millimétriques, peuvent augmenter de façon considérable le coût du circuit.

2.3.4 La Calibration à l'Aide d'un Réseau des Neurones.

Une autre méthode de calibration, qui peut être classifiée dans la catégorie des méthodes 'temps réel', a été introduite par Liu [79]. Même si l'auteur ne l'a pas utilisée dans le contexte d'un récepteur direct, mais plutôt dans une variante du réflectomètre, cette technique peut être facilement utilisée dans un système de transmission/réception. Liu a

proposé d'utiliser un réseau de neurones qui est spécialement entraîné dans le but de mesurer la valeur complexe du coefficient de réflexion. Un générateur vectoriel génère un signal dont chaque état correspond à un coefficient de réflexion précis. L'idée réside dans le fait que, au lieu de résoudre un système d'équations, comme dans la méthode standard, on va générer un ensemble important d'états. Cet ensemble permet au réseau de neurones de reconnaître une valeur de signal inconnue.

La figure 2.11 montre le principe de cette technique. Les triangles noirs représentent les états générés par la source vectorielle de test qui seront utilisés pour l'entraînement du réseau. Après cette étape, un signal correspondant à un coefficient de réflexion (représenté en gris) sera "démodulé" par un processus d'interpolation fait par le réseau de neurones. L'avantage majeur de cette technique est le fait que l'on n'a pas besoin d'une technique de linéarisation des détecteurs de puissance. Le réseau de neurones intègre l'aspect non-linéaire de la réponse des diodes. Bien sûr, quelqu'un peut argumenter que les circuits additionnels nécessaires, en l'occurrence le générateur vectoriel de test est largement plus coûteux qu'un simple générateur d'onde sinusoïdale. Cependant lorsqu'on transpose cette méthode au récepteur direct, les vecteurs d'entraînement peuvent être vus comme une séquence d'entraînement dans le protocole de transmission. Chaque trame d'information va avoir un en-tête composé par la séquence désirée.

Une analyse plus rigoureuse montre les limitations d'un tel système : La précision de mesure est directement proportionnelle avec le nombre de valeurs d'entraînement. Plus la séquence d'entraînement est grande, plus la résolution de mesure est grande. Malheureusement, on ne peut pas augmenter trop la longueur de l'en-tête sinon, l'efficacité de transmission restera faible.

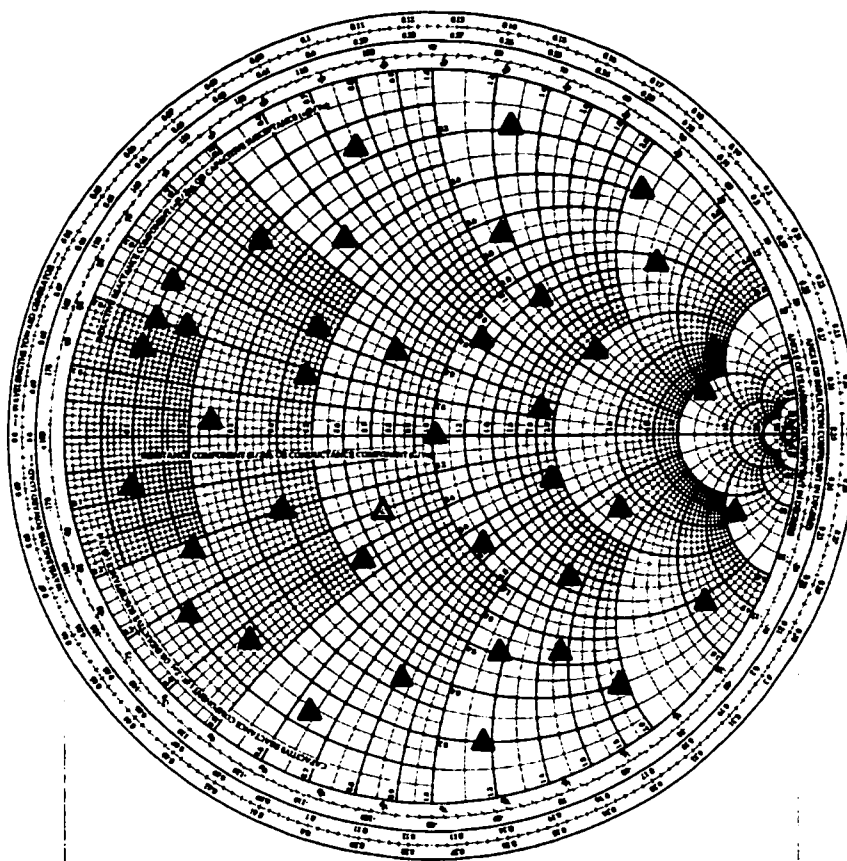


Figure 2.11 Répartition des états simulés.

Même si on pouvait augmenter la taille de l'en-tête, il reste qu'avec cette méthode, il est difficile de dissocier les erreurs linéaires de la jonction par rapport aux erreurs complexes produites par les mécanismes d'évanouissement liés aux multi-trajets. Dans ce cas, la méthode proposée dans cette section n'est pas utilisable. Il reste alors la solution d'intégrer le générateur vectoriel de signal dans le récepteur avec l'augmentation de complexité correspondante.

2.3.5 L'Autocalibration.

Une technique plus performante de calibration a été récemment introduite, spécialement dédié aux applications de récepteur direct [58,63,64]. Les chercheurs du Centre de Recherche en Communication de Canada (Ottawa) ont utilisé une procédure originale à base d'analyse de signaux d'arrivée I et Q. En considérant que la corrélation entre les signaux I et Q à la réception doit être nulle et que les éventuelles erreurs vont modifier cette propriété, on peut développer une procédure exacte de calcul de coefficients de calibration.

Sans insister sur l'aspect mathématique de la procédure de calibration, résumons ici les avantages de cette technique :

- Une technique temps réel.
- On utilise le même équipement électronique que le récepteur, sans ajout des éléments supplémentaires.
- Indépendante du type de modulation si on respecte la condition de corrélation nulle entre les canaux I et Q.

Il y en a aussi de limites pour cette technique, dont les plus importantes sont :

- La plage dynamique limitée due au fait qu'on suppose la réponse des diodes linéaire.
- La précision de la calibration dépend de la taille de la fenêtre d'échantillonnage. La condition d'inter-corrélation nulle demande un ensemble statistiquement grand.

Pour conclure ce chapitre sur l'analyse des méthodes de calibration existantes, on rappelle les caractéristiques qu'une architecture de récepteur direct doit présenter pour être un compétiteur viable comparé aux récepteurs super-heterodynes :

- **Il doit exister d'une méthode de calibration pour la correction des erreurs de phase et amplitude.**
- **La calibration doit corriger les erreurs générées par la présence des écarts DC, ou bien (et) l'architecture retenue doit minimiser ces erreurs.**
- **La méthode de calibration doit être embarquée sur le même matériel que celui utilisé pour la réception de données.**
- **La méthode de calibration doit être transparente pour l'utilisateur.**
- **La technique de correction d'erreurs doit être indépendante du type de modulation.**

CHAPITRE 3 POLYCOM. NOUVELLE ARCHITECTURE DE RÉCEPTEUR DIRECT

Dans le chapitre précédent, nous avons défini les critères associés à une nouvelle architecture, pour que celle-ci puisse présenter des avantages intéressants par rapport à l'architecture classique I&Q super-hétérodyne. Le principal critère s'est avéré l'existence d'une méthode de calibration, similaire ou non avec les techniques de la technologie six-ports. Le traitement mathématique doit rester à un niveau simple pour que le coût d'implémentation soit faible, sans conséquence majeure sur l'architecture du circuit. Dans le présent chapitre, nous allons présenter une nouvelle topologie de circuit qui rencontre ces critères. Les avantages et inconvénients de cette nouvelle architecture sont discutés, et le fonctionnement d'une technique de correction originale est décrit à l'aide de résultats de simulations numériques.

3.1 Polycom.

La figure 3.1 présente un diagramme block de notre système de réception complet, qui intègre aussi la tête réceptrice. On peut observer les caractéristiques principales de notre architecture : Notre architecture utilise des mélangeurs; comme on l'expliquera plus loin dans ce chapitre, ceci permet d'obtenir une bonne isolation entre les ports LO et RF. De plus, le récepteur a une structure à trois canaux, ce qui est différent par rapport au système I&Q (à deux canaux de sortie) et par rapport à la jonction six-ports (à quatre canaux de sortie). La motivation de ce choix est simple. Nous désirons avoir une structure qui maintient les avantages de la structure six-ports si possible, tout en diminuant le nombre de dispositifs nécessaires.

Considérons un signal modulé avec une modulation de type QPSK. L'information est transmise sur les canaux I et Q. Un récepteur I et Q a deux canaux de sortie donc on pourrait construire un système de deux équations et deux inconnues pour retrouver l'état I,Q à partir des mesures de voltage sur les deux canaux. Évidemment, si en réalité les

paramètres des circuits sont différents de ceux pris en compte dans le calcul (erreurs de phase et amplitude), on ne peut pas extraire l'information exacte; celle-ci sera modifiée par les erreurs du circuit physique.

Au contraire, la jonction six-ports présente 4 canaux de sortie, donc à-priori, il y a une redondance de l'information, suffisante pour qu'on puisse extraire les coefficients de calibration (la correction des erreurs de mesure) et aussi pour qu'on puisse mesurer de façon précise la valeur des données reçues. Donc les canaux supplémentaires apportent des degrés de liberté supplémentaires au système, qui seront utilisés pour la correction des erreurs.

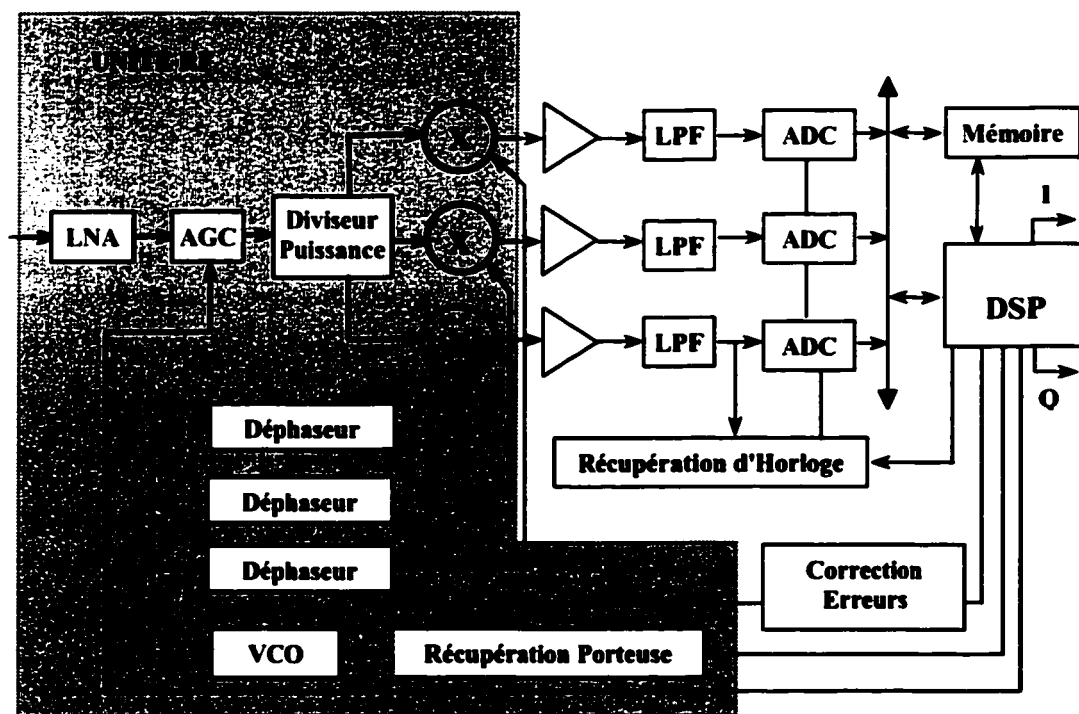


Figure 3.1 Architecture POLYCOM.

On a vu dans le chapitre 2, que le port de référence du six-ports est généralement bien isolé du porte de mesure. Dans ce cas le signal du port de référence ne contient aucune information sur le signal à mesurer, il est seulement dépendant du signal généré

localement. En conséquence, dans le cas du fonctionnement d'un démodulateur, on peut éliminer un canal sans perdre l'information nécessaire au processus de démodulation, ceci explique pourquoi notre architecture compte 3 canaux de sortie.

L'élimination du canal de référence est aussi faite dans le but de simplifier le traitement numérique de l'information. Comme on a déjà vu dans le chapitre précédent, une jonction six-ports extrait l'information à partir des rapports de puissance. L'opération est faite au niveau numérique par le processeur de signal, après l'échantillonnage. La vitesse de calcul du processeur doit être comparable au taux de transmission, ce qui demande d'augmenter la vitesse de traitement des données si la bande de base du signal est augmentée. Habituellement, l'opération de division est plus "gourmande" au niveau temps d'exécution par rapport aux opérations de base d'un processeur de signal (l'addition et multiplication). Réduire la complexité de l'algorithme de calcul et gagner ainsi en terme de vitesse de traitement a été un des facteurs considéré pour notre méthode de correction d'erreurs.

Évidemment, on peut se poser la question suivante : si le canal de référence est éliminé, comment défini-t-on les paramètres de calibration. La réponse est simple. Le canal de référence est nécessaire seulement dans un contexte de type réflectomètre. Dans ce cas, la puissance de référence (une mesure de la puissance de l'oscillateur local) permet de normaliser le coefficient de réflexion. Une image de ce processus est que l'on établit ainsi le rayon unitaire sur l'abaque de Smith. Au contraire, dans notre application démodulation, connaître la valeur absolue des vecteurs mesurés n'est pas nécessaire. Il est suffisant de connaître les valeurs relatives entre les différents états de modulation. Autrement dit le facteur d'échelle dans l'espace cartésien I&Q n'est pas important tant qu'il reste constant (tant que la puissance de l'oscillateur local ne varie pas). Cela amène une contrainte sur la forme de la séquence d'apprentissage : elle doit inclure tous les états de modulation dans un ensemble statistiquement représentatif. Dans notre cas, nous avons utilisé une séquence pseudo-aléatoire.

Dans les grandes lignes, le traitement du signal dans notre architecture (figure 3.1) est divisé en deux parties. Dans l'unité RF, le signal d'entrée est amplifié par un amplificateur faible bruit (LNA) et par un amplificateur à gain variable (AGC). Ensuite il est divisé sur les trois canaux de la jonction et mélangé avec trois versions déphasées du signal produit par l'oscillateur local. Les trois déphaseurs doivent assurer un déphasage relatif de 120 degrés. Pour un système multi-canaux, l'oscillateur local est remplacé par un oscillateur contrôlé en tension (VCO) relié au système de synchronisation de la porteuse. Une analyse qualitative de ce système met en évidence le fait qu'on peut intégrer sur le même circuit intégré toutes ces fonctions. Dans ce cas on pourrait parler d'une tête réceptrice monobloc. En principe tous les sous-systèmes de l'unité RF sont aussi présents dans une architecture classique I&Q directe ou super-hétérodyne. Cependant certaines spécifications sont différentes et certaines contraintes sont modifiées. Par exemple, dans notre architecture, les mélangeurs dans la jonction transposent le spectre directement dans la bande de base. On verra dans le dernier chapitre les implications de ceci au niveau du circuit. De plus, le déphaseur dans notre circuit peut être plus complexe, le déphasage pouvant être ajusté (contrôlé par le DSP) pour compenser les erreurs de phase.

En ce qui concerne le contrôle de la bande dynamique, la fonction de l'AGC peut être divisée en deux niveaux. Un niveau plus grossier dans le domaine RF et un niveau plus fin et performant dans la bande de base. Le premier niveau d'amplification sert seulement à éviter la saturation du mélangeur. L'ensemble des circuits AGC sert à augmenter de façon significative la bande dynamique du signal.

Finalement la deuxième partie du traitement est faite dans son intégralité au niveau de la bande de base du signal modulé. Le filtrage est similaire, en terme de performances, au filtrage sur la fréquence intermédiaire d'un récepteur super-hétérodyne. Cependant, dans notre cas on parle d'un filtre passe-bas et non d'un filtre passe-bande. Le concepteur peut

choisir de faire le filtrage au niveau analogique ou numérique. En principe il doit exister un filtrage minimal sur le signal analogique pour diminuer le bruit d'échantillonnage au niveau du convertisseur analogique-numérique (ADC).

Après l'échantillonnage, l'algorithme de correction entre en jeu. Celui-ci va déterminer de façon exacte les valeurs des données I et Q. Ce puissant algorithme adaptatif va analyser les signaux échantillonnés sur les trois canaux, et va extraire les termes d'erreurs. Il va aussi modifier sa forme en fonction du type de modulation. Cet algorithme, totalement numérique, utilise seulement des opérations primitives (additions et multiplications) et il peut être embarqué sur un processeur de signal numérique standard. Le reste de ce chapitre y sera dédié. Nous allons maintenant présenter ses fondements mathématiques et des résultats de simulations avec des signaux modulés.

3.2 La Technique de Résonance Adaptative.

Les critères énoncés dans le chapitre 2 nous ont conduit à rechercher une technique de calibration mieux adaptée à l'application du récepteur direct. L'utilisation des "standards" de calibration étant a priori écartée, nous avons pensé que la source d'information supplémentaire doit être le signal inconnu lui-même. Le signal inconnu respecte en effet une condition très puissante: il est limité à un ensemble d'états finis (la constellation). Sans connaître la valeur de l'information, on sait que le signal modulé doit appartenir à un ensemble prédéfini par le type de modulation. Dans ce cas, en analysant consécutivement le signal RF, le processus de démodulation peut être vu comme un problème classique de reconnaissance de forme ("pattern recognition"). Dans un sens plus large, la démodulation peut être définie comme une méthode par laquelle on choisit une partition parmi un ensemble d'éléments. Il est très facile, dans ce cas d'associer ce processus avec une technique classique de partition ("clustering"). Pour les signaux modulés numériquement, on va associer une partition à un des états vectoriels de la constellation.

Les algorithmes à base de réseaux de neurones se sont avérés être une solution robuste pour résoudre les problèmes de reconnaissance de forme. Leur simplicité d'implémentation, soit sur des composantes de type ASIC soit sur des architectures génériques (comme les processeurs de signaux) représente leur principal avantage. Dans notre cas, on devait choisir un algorithme qui possède une particularité remarquable : Il ne doit pas y avoir de différence entre la séquence d'entraînement et la séquence de test. Pratiquement les symboles inconnus forment à la fois l'échantillon pour l'apprentissage et la trame de test (dans ce cas la réponse du réseau forme la séquence démodulée).

La technique de résonance adaptative, introduite par Carpenter et Grossberg [20,21,24] à la fin des années '80, constitue un candidat prometteur pour notre algorithme de correction d'erreurs. Nous allons maintenant présenter les fondements mathématiques de cette technique tout en relevant les modifications que nous y avons apporté dans le cadre de notre application spécifique. Les modifications proviennent en général du fait que la technique de Carpenter, a été conçue et optimisée pour la reconnaissance des caractères. Certaines différences de fond, dans le traitement mathématique, doivent être apportées pour qu'on puisse l'utiliser avec succès dans un contexte de démodulation de signaux.

La figure 3.2 présente un réseau de neurones très simple. En principe, un neurone peut être défini comme un nœud du réseau qui effectue une opération primitive sur l'ensemble des éléments d'un vecteur d'entrée, nommée excitation. L'opération primitive dans la majorité de cas est extrêmement simple, par exemple une multiplication par une constante (poids) ou une addition. Dans d'autres cas cette opération primitive peut être une opération non-linéaire comme une comparaison, le calcul d'une fonction non-linéaire (ex. $\ln(x)$) ou un écrêtage. L'idée du réseau de neurones est de remplacer les opérations mathématiques lourdes (en temps de calcul) par une quantité importante de centres de traitement numérique qui peuvent s'adapter au type de signal transmis. Ce calcul distribué est largement plus simple à embarquer sur une architecture numérique et de plus, il peut présenter une courbe d'apprentissage pour des systèmes à paramètres

variables. Dans notre cas, les erreurs physiques du circuit forment l'ensemble des paramètres variables à étudier.

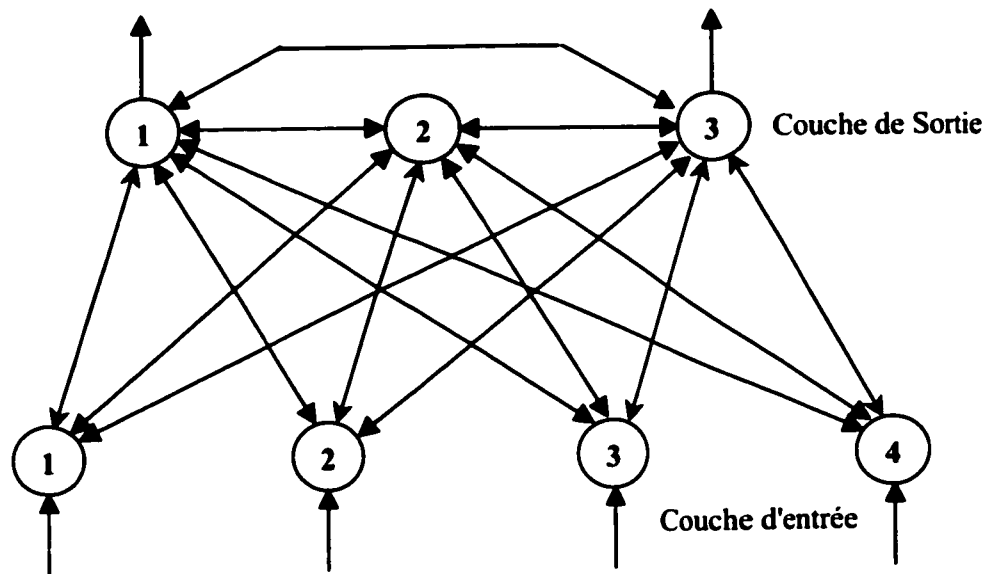


Figure 3.2 Réseau de neurones

La figure 3.2 a présenté un réseau à deux couches avec contre-réaction, dont tous les nœuds sont inter-connectés. Notons que ce n'est pas la seule façon de construire un réseau de neurones. Le nombre de couches, et le nombre de nœuds restent des paramètres de choix pour le concepteur. En principe, seulement le nombre d'entrées et de sorties est déterminé par le système à analyser. En particulier, pour notre problème nous considérerons un réseau à 3 entrées analogiques (mais sous représentation numérique) et à deux sorties (I&Q) représentant le signal démodulé. Pour éviter toute confusion précisons que l'algorithme fonctionne sur une plate-forme numérique (les signaux sont représentés sous forme binaire), et que les signaux d'entrée proviennent de l'échantillonnage des trois signaux de sorties de la jonction (voir figure 3.1). Pour

conserver une bonne précision nous utilisons une représentation numérique en virgule flottante.

Pour mieux comprendre notre technique de calibration originale, nous allons commencer par présenter l'algorithme de la technique "résonance adaptative" (ART). La figure 3.3 présente la structure d'un réseau ART dans sa forme généralisée introduite par Carpenter et Grossberg [21,24]. Comme on peut l'observer, on peut décomposer la structure en 4 sous-systèmes avec des rôles précis : le niveau de pré-traitement (F_0), la première couche, ou champ d'attention (F_1), la couche de décision (ou champ F_2) et le sous-système d'orientation.

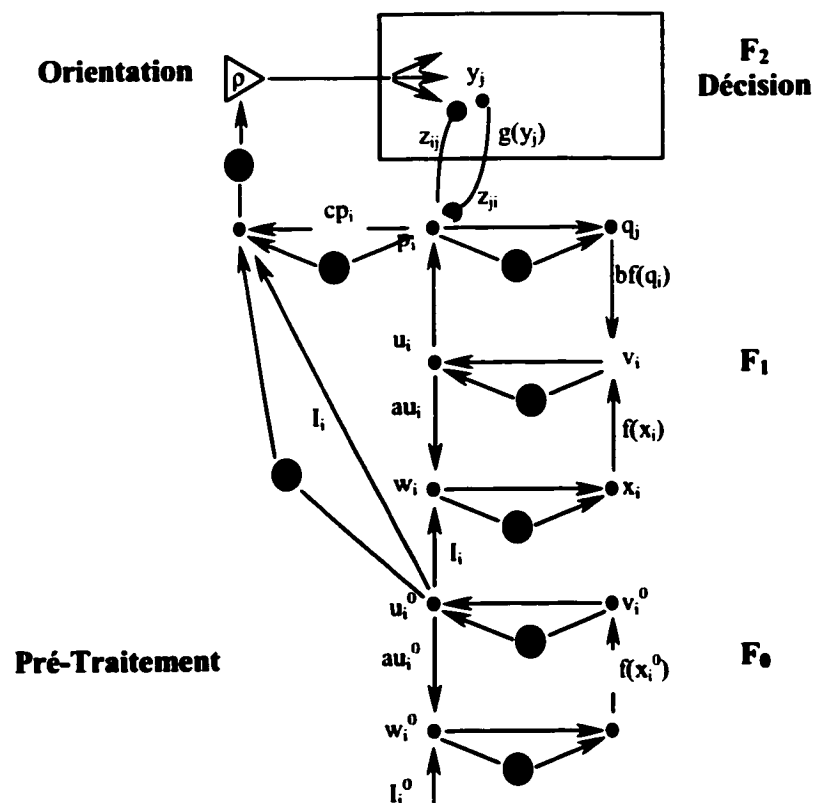


Figure 3.3 Architecture ART ("Adaptive Resonance Technique").

Dans la figure 3.3 on a désigné par les cercles pleins les opérations de normalisation. Par défaut, chaque nœud dans le réseau (neurone), effectue une opération primitive (l'addition de chaque composante du vecteur présenté à son entrée). Même si la figure 3.3 utilise un format scalaire pour la simplicité du graphisme, on doit considérer que sur chaque niveau on parle d'un traitement vectoriel. Ainsi, chaque nœud représente une composante du vecteur et l'opération de normalisation agit sur l'ensemble des composantes de même niveau.

Dans le développement mathématique suivant, les entrées du système numérique (I) sont les valeurs à la sortie des convertisseurs ADC (voir figure 3.1).

Les équations qui décrivent la réponse du réseau sont:

$$w^0 = I^0 + au^0, \quad 3.1$$

la normalisation de type euclidien nous donne:

$$x^0 = \mathfrak{I} \{w^0\}, \quad \text{où } \mathfrak{I} \{x\} = x / \|x\|, \quad 3.2$$

ensuite le vecteur x^0 est transformé via une fonction non-linéaire:

$$v^0 = f(x^0), \quad 3.3$$

$$f(x_i) = \begin{cases} x_i^0 & \text{si } x_i^0 > \theta \\ 0 & \end{cases}, \quad 3.4$$

Dans ce cas θ , est un paramètre choisi sous la contrainte:

$$0 < \theta < \frac{1}{\sqrt{M}}, \quad \text{où } M \text{ est le nombre de dimensions du vecteur d'entrée} \quad 3.5$$

Les opérations au niveau de chaque nœud sont similaires à celles présentées dans les équations 3.1-3.5, avec des paramètres propres à chaque niveau.

Parfois, on utilise un système ayant, comme vitesse d'échantillonnage, un multiple entier du taux de transmission pour avoir plusieurs échantillons par symbole. Dans ce cas, le

champ de pré-traitement présente une caractéristique très importante pour l'architecture : Il permet au vecteur d'entrée d'être transmis de façon stable au niveau F_1 . Nous expliquons maintenant pourquoi :

Initialement le vecteur u^0 est nul donc :

$$v_i^0 = \begin{cases} \frac{I_i^0}{\|I^0\|} & \text{si } I_i^0 > \theta \|I^0\| \\ 0 & \text{sinon} \end{cases}, \quad 3.6$$

Si on prend Ω comme l'ensemble des indices de chaque élément du vecteur I^0 qui est supérieur au seuil θ :

$$\Omega = \{i: I_i^0 > \theta \|I^0\|\}, \quad 3.7$$

on peut déduire qu'il existe une constante unique $K > 1 / \|I^0\|$ dont:

$$u_i^0 = \begin{cases} KI_i^0 & \text{si } i \in \Omega \\ 0 & \text{si } i \notin \Omega \end{cases}, \quad 3.8$$

Après la première itération on trouve que:

$$w_i^0 = \begin{cases} I_i^0(1 + aK) & \text{si } i \in \Omega \\ I_i^0 & \text{si } i \notin \Omega \end{cases} \quad 3.9$$

En analysant l'équation 3.9, on voit que, suite aux premières transformations, les éléments du vecteur I^0 qui sont plus grands que le seuil θ seront amplifiés alors que les éléments inférieurs à θ seront atténués (a est une constante positive). Dans ce cas, l'ensemble Ω reste constant d'itérations en itération! Cette caractéristique est utilisée dans l'algorithme initial de ART pour éliminer le bruit associé au système de lecture (les senseurs). Par la suite, nous verrons que pour satisfaire les contraintes de calibration et aussi préserver cette condition de stabilité on va modifier le traitement numérique du premier niveau dans notre algorithme.

De même que le champ de pré-traitement, le niveau F_1 additionne et normalise le vecteur résultant. Le signal de sortie de cette couche est le vecteur "p" qui est le résultat de l'addition des signaux internes du niveau F_1 et des signaux filtrés par la fonction non-linéaire " $g(x)$ ", appartenant à la couche F_2 de décision :

$$p_i = u_i + \sum_j g(y_j) z_{ji} , \quad 3.10$$

dans cette expression, z_{ji} représente le coefficient de transfert (ou poids) du nœud j du niveau F_2 vers nœud i du niveau F_1 . L'indice j correspond aux différents choix possibles.

Le champ F_2 est l'endroit où l'on prend la décision. Sous une forme générale, on peut considérer que le vecteur d'excitation correspond à une certaine partition. Dans la forme originale, utilisée pour la reconnaissance de forme, le choix du niveau F_2 correspondait à la reconnaissance d'une image (une lettre de l'alphabet). Pour nous, le choix est la détermination de l'état de modulation. Pour chaque nœud du niveau F_2 on peut écrire :

$$y_j = \sum_i p_i z_{ij} , \quad 3.11$$

La décision du F_2 est représentée mathématiquement par:

$$y_k = \max_j \left\{ \sum_i p_i z_{ij} \right\} , \quad 3.12$$

où y_k est l'état déclaré gagnant (la représentation de choix)

En conséquence on peut calculer la contre-réaction du niveau F_2 vers F_1 par l'intermédiaire de la fonction vectorielle $g(x)$:

$$g_j(x) = \begin{cases} d & \text{si } j = k \\ 0 & \text{si } j \neq k \end{cases} , \text{ où } k \text{ est l'index de l'état gagnant.} \quad 3.13$$

Si le champ de décision est activé (un état a été déclaré gagnant) le système va modifier les coefficients de transferts entre les deux niveaux pour que l'état gagnant corresponde au vecteur d'entrée, la source d'excitation. Autrement dit, de façon adaptative, le système va assigner un vecteur semblable au vecteur d'entrée pour les éléments des coefficients de transfert. On reconnaît ici la propriété d'apprentissage d'un réseau de neurones. La catégorie choisie correspond, à un certain degré, au vecteur inconnu présenté à l'entrée du système. Les équations du système d'apprentissage deviennent :

$$\frac{dz_{ji}}{dt} = p_i - z_{ji} \quad , \quad 3.14$$

en considérant 3.10 et 3.13, on obtient :

$$\frac{dz_{ji}}{dt} = (1-d) \left[\frac{u_i}{1-d} - z_{ji} \right] \dots, \quad 3.15$$

Dans les équations 3.14 et 3.15, d/dt représente la dérivée par rapport au temps. Dans notre cas l'unité du temps est l'itération, donc l'opération de dérivation est effectuée numériquement par la modification des éléments entre deux symboles consécutifs.

L'équation 3.15 représente la propriété fondamentale de la technique de résonance adaptative. Elle a comme conséquence immédiate la proportionnalité entre le vecteur d'entrée dans le champ de courte mémoire $F_1(I)$ et le vecteur des coefficients de transfert z_{ij} , dans le cas d'un seuil θ nul pour le niveau F_1 .

Pour prouver ce puissant théorème, reprenons les équations suggérées par la figure 3.3 pour le vecteur u :

$$u = \mathfrak{I}[\mathfrak{I}(I + au) + b\mathfrak{I}(u + d/(1-d)z_j^*)] \quad , \quad 3.16$$

Si le champ F_2 prend une décision, par l'équation 3.15:

$$u = z_j^* = (1-d)z_j, \quad 3.17$$

Notons $z_j^* = z$ pour unifier la notation, 3.16 devient:

$$z = \mathfrak{I}[\mathfrak{I}(I + az) + b\mathfrak{I}(z + d/(1-d)z)] = \mathfrak{I}[\mathfrak{I}(I + az) + bz], \quad 3.18$$

$$z = \left(\frac{I + az}{\|I + az\|} + bz \right) \left\| \frac{I + az}{\|I + az\|} + bz \right\|^{-1}, \quad 3.19$$

donc,

$$(I - A(Ba + b))z = ABI, \quad 3.20$$

$$\text{où } A = \left\| \frac{I + az}{\|I + az\|} + bz \right\|^{-1} \text{ et } B = \|I + az\|^{-1}$$

comme $A \neq 0$ et $B \neq 0$ on obtient:

$$I = \frac{(I - A(Ba + b))}{AB} z, \quad 3.21$$

Mais $\|I\| = \|z\| = 1$ en tant que vecteurs normalisés donc $I = z$.

Ceci complète le développement mathématique du réseau de type ART. Jusqu'à maintenant on n'a rien dit sur le rôle du niveau d'orientation. Son action est importante dans l'application originale de reconnaissance de forme. À priori, on ne connaît pas le nombre de catégories, alors que se passe-t-il quand un vecteur correspondant à une nouvelle catégorie se présente? Une façon de faire serait la suivante : chaque fois qu'un

nouveau vecteur se présente, le système essaie de modifier les coefficients de transfert associés à chaque état déjà connus pour expliquer la nouvelle mesure, ce qui donnera un résultat faux si cette mesure correspond à une nouvelle catégorie. Le sous-système d'orientation peut inhiber le processus de modification des coefficients de transfert si une grande différence entre le vecteur d'entrée et le vecteur caractéristique de l'état associé est détectée. Ainsi, une nouvelle catégorie associée au nouveau vecteur inconnu est générée. Dans un système à états finis ce processus sera actionné seulement pendant le cycle d'apprentissage.

Pour donner au lecteur une image qualitative de ce processus, supposons qu'on utilise un réseau ART pour la reconnaissance de caractères. Supposons aussi que jusqu'à présent on ait appris toutes les lettres de l'alphabet sauf la lettre F. Si à l'entrée, on excite le système avec le vecteur correspondant à la lettre F, la tendance du système de décision est d'associer la lettre soit à la catégorie de la lettre E soit à la catégorie correspondant à P. Évidemment, si on n'empêche pas une telle partition on va commettre une erreur. Dans ce cas là, le sous-système d'orientation va détecter les différences (la proportionnalité demandée par l'équation 3.21 n'est plus valide) entre la décision et le vecteur d'entrée. En conséquence il va annuler les modifications et va demander au système de décision F_2 de générer une nouvelle catégorie similaire à la lettre F. On va revenir sur cet aspect, au cours de la section suivante, dans laquelle les modifications que nous avons apportées pour améliorer le comportement de cette classe d'algorithmes sont expliquées.

Pour résumer cette courte présentation sur la méthode ART, rappelons la séquence d'actions effectuées au cours de cet algorithme:

- Le vecteur d'entrée est soumis à un processus de normalisation. Le vecteur peut être constitué par les échantillons d'une image numérisée, par les composantes

spectrales d'un enregistrement vocal, ou comme dans notre cas, par les niveaux de sortie de la jonction.

- Le vecteur normalisé dans le champ de pré-traitement, devient l'excitation pour le champ d'attention. Le champ d'attention calcule un vecteur résultant entre l'entrée et le vecteur de contre-réaction. Son fonctionnement est similaire à un système à contre-réaction positive. Il va amplifier la somme entre son vecteur d'entrée et son vecteur de sortie qui, en condition d'équilibre, est proportionnel au vecteur d'excitation. Ce comportement a donné le nom à ce type de réseau, parce que, dans ce cas, on dit qu'on a trouvé la résonance du système.
- Le champ de décision choisit parmi les catégories mémorisées celle qui correspond au vecteur d'excitation en résonance. Suite à une décision, il va modifier les vecteurs caractéristiques de cette catégorie pour obtenir une meilleure représentation du vecteur d'entrée. En termes mathématiques, l'algorithme associe le vecteur caractéristique avec le centre de gravité de la partition correspondant à l'ensemble des vecteurs d'entrées analysées qui ont appartenus à cette catégorie. Ce comportement est parfois décrit dans la littérature comme la mémoire à long terme.
- Si aucune catégorie ne représente de manière satisfaisante le vecteur d'entrée, le sous-système d'orientation va inhiber l'activité au niveau de la mémoire à long terme et va imposer soit la création d'une nouvelle catégorie, soit l'annonce de la présence d'un état inconnu (noyée dans le bruit). Au début du fonctionnement du système, l'activité du champ d'orientation sert à l'apprentissage des états. Dans ce cas on peut dire que la technique de la résonance est totalement adaptative, elle peut être vue comme un système expert, capable à se reconfigurer en fonction de la forme de signal reçu.

3.3 Nouvelle Technique de Correction d'Erreurs.

On a vu dans la précédente section, comment on peut construire un algorithme qui peut résoudre un problème typique de partition. Comme la démodulation du signal RF fait partie de la même catégorie de problèmes, la technique ART est un candidat robuste pour être utilisé dans un démodulateur avec la structure présentée dans la figure 3.1. La nouvelle technique proposée, a été conçue avec le but précis d'améliorer les caractéristiques des précédentes techniques de calibration publiées dans la littérature [42,63,77] et elle est adaptée au contexte spécial du récepteur homodyne.

Nous avons analysé dans le chapitre 2, l'analogie entre un récepteur direct et la jonction six-ports. Une jonction six-ports existante a été utilisée comme modèle générique pour notre système. Il s'agit d'une jonction en anneau à 5 ports à laquelle on a ajouté un coupleur distribué pour le canal de référence. Le circuit de test, réalisé en technologie de circuits hybrides MHMIC est représenté dans la figure 3.4.

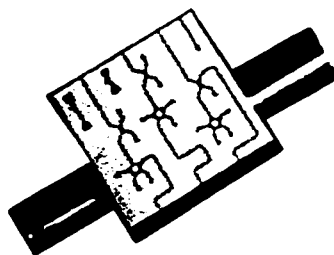


Figure 3.4 Architecture anneau du circuit six-ports

Les erreurs de phase et de gain ont été extraites, en mesurant les paramètres S, pour différentes fréquences dans la bande 24-31 GHz (plus large que les limites considérées dans la conception du circuit). Ces mesures ont été utilisées pour construire une base de données pour notre modèle de récepteur qui a été implémenté sous le logiciel SPW™ (Signal Processing Workstation) de la compagnie Cadence. Il faut mentionner que, dans

ce cas, le domaine de variance des erreurs simulées est plus large que les valeurs moyennes mesurées. Alors, cette hypothèse nous a permis de valider le comportement de l'algorithme dans les cas extrêmes.

Comme nous l'avons expliqué au chapitre 2, la présence des erreurs de phase et de gain dans le signal modulé QPSK va générer des distorsions. Pour une meilleure compréhension du phénomène, la figure 3.5 donne une image qualitative sur les effets des ces distorsions sur la constellation vectorielle de signal.

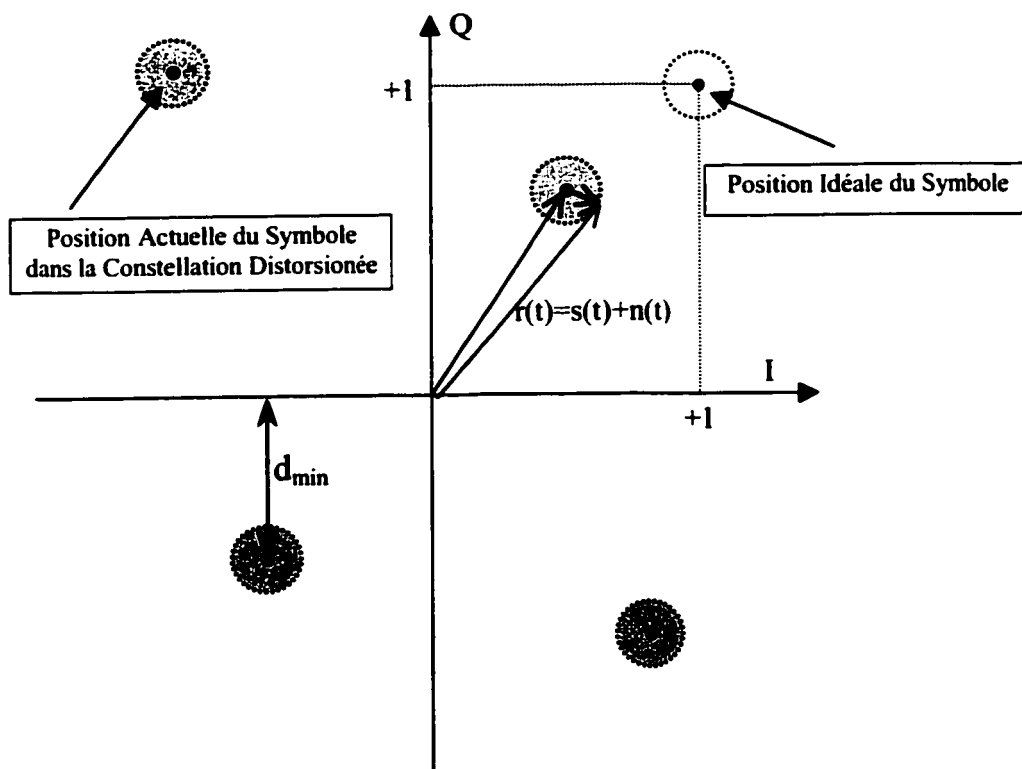


Figure 3.5 Constellation QPSK distorsionnée.

Le centre de chaque cercle représente l'état de modulation. Les changements de phase et de gain sur les chemins différents, dus aux erreurs vont modifier la distance minimale entre le centre d'un tel cercle et la région voisine de décision. Cette diminution de la

distance minimale correspond à une croissance du taux d'erreurs binaires (BER) du système de communication. Ces résultats sont similaires aussi pour les modulations de type M-PSK comme pour celles de type M-QAM. Comme nous l'avons déjà dit, le processus de démodulation est un problème de partition, où l'on doit obtenir le centre de gravité de chaque état de modulation, dans le plan bi-dimensionnel I-Q.

En l'absence d'effets non-linéaires, les erreurs de phase et de gain peuvent être vues comme des transformées linéaires (rotation, changement d'échelle, etc.). Si on connaît la position réelle du vecteur de sortie dans la constellation de signal, on peut trouver la transformation inverse qui va corriger les erreurs, prévenant ainsi la dégradation du BER.

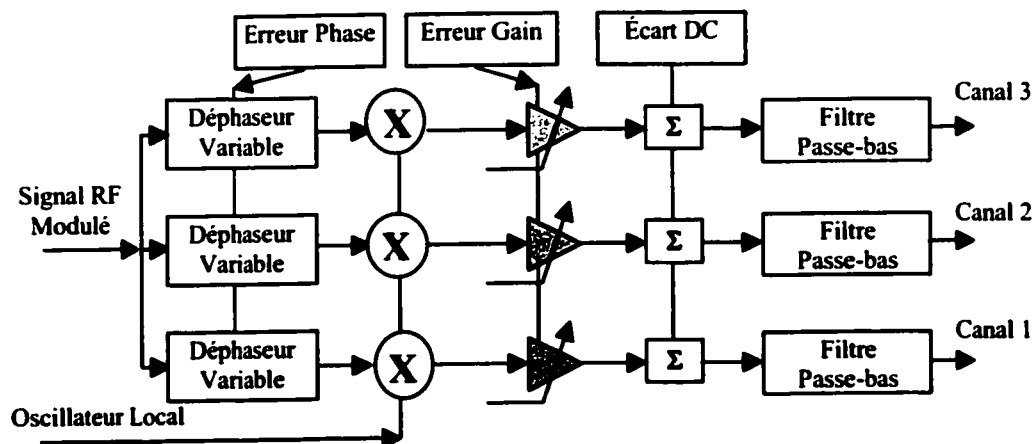


Figure 3.6 Modèle SPW pour les erreurs de phase, gain et écarts DC

Le modèle utilisé, donne la possibilité aux utilisateurs de modifier les erreurs de phase, gain et écarts DC indépendamment sur chaque canal. Dans notre modèle SPW, les erreurs sont assimilées aux paramètres de simulation. En imposant différentes valeurs pour ces paramètres "d'entrée", l'utilisateur peut observer le comportement de notre algorithme pour une plage très large de variation. Éventuellement, il peut aussi extraire des résultats

statistiques, si les caractéristiques de probabilités pour les composantes de la tête réceptrice sont connus (valeur moyenne, écart type, etc.).

Notre procédure de correction respecte en grandes lignes les principes de la méthode ART, mais on a modifié certains fonctions par rapport à la solution typique de ART2 [24]. La première couche du réseau reste un niveau de pré-traitement, mais on n'utilise pas la normalisation vectorielle sur les composantes. Soit la représentation vectorielle de nos sorties échantillonnées :

$$y_o = y_1 \hat{u}_1 + y_2 \hat{u}_2 + y_3 \hat{u}_3 , \quad 3.22$$

où $\hat{u}_i$ sont les vecteur unitaire de a base vectorielle

la normalisation du vecteur rend le résultat :

$$\|y_o\| = \sqrt{y_1^2 + y_2^2 + y_3^2} , \quad 3.23$$

Si on considère la composante $y_i = y_{i0} + \varepsilon_i$, $i=1,3$ et ε_i le terme d'erreur, on y retrouve:

$$\|\tilde{y}_o\| = \sqrt{(y_1 + \varepsilon_1)^2 + (y_2 + \varepsilon_2)^2 + (y_3 + \varepsilon_3)^2} , \quad 3.24$$

Dans ce cas on observe que le rapport entre le vecteur y_o et $\|y_o\|$ ne correspond pas à une opération linéaire. La présence d'erreurs entraîne plus qu'une simple modification du facteur d'échelle, elle change l'énergie du vecteur de sortie de façon non-linéaire. Le phénomène est très connu et dans le langage usuel on dit que le centre de gravité de la constellation a une énergie non-nulle.

Pour éviter ce phénomène et préserver l'aspect linéaire du pré-traitement on a modifié l'opération associée à cette couche. En fait on applique un vrai facteur d'échelle:

$$\tilde{y}^k = \frac{y_o^k}{\max_j \{y_i^j\}} , j = \overline{0, K} , i = \overline{1, 3} , \quad 3.25$$

Dans l'équation 3.25, K représente le nombre d'états reçus durant la fenêtre de temps 'utile' au traitement.

La transformée proposée par 3.25 modifie le traitement numérique par un simple ajustement d'échelle vers une plage prédéfinie (ex. $-1,1$). L'opération est similaire à celle produite par un amplificateur à gain variable dans une boucle de contrôle automatique du gain (AGC). Lors de l'implémentation réelle de notre récepteur, la normalisation pourra se faire au niveau de l'AGC, pour préserver la vitesse de calcul. L'idée fondamentale est de préserver à tout prix, la gamme dynamique du récepteur.

Le concept provient de quelques essais faits dans le but de calibrer une jonction six-ports avec des diodes comme détecteur de puissance. Pour avoir une meilleure sensibilité les diodes ont été utilisées dans une configuration avec polarisation. Comme le récepteur direct doit être couplé à la structure d'échantillonnage en connexion DC, la sortie présentait un écart DC identique au point de polarisation de la diode. Cet écart était même plus grand que la variation de signal AC (0.63V contre ± 0.2 V). L'algorithme à 13 standards ou celui de type dual-tone corrigeait par calcul les écarts DC, cependant la plage dynamique du récepteur était réduite à 0.4 V seulement (0.43-0.83V). En comparaison, la plage totale des convertisseurs analogiques-numériques (ADC) était de -1 à $+1$ V. En comparant au cas d'un point de polarisation nul, ceci équivaut à une diminution d'un facteur 5 fois de la plage dynamique, ou en d'autres termes à 3 bits de résolutions. Comme l'ADC avait 8 bits de résolution, nous avons dans les faits 5 bits de résolution réelle, ce qui a généré une impossibilité d'atteindre la précision nécessaire au processus de calibration.

Avec la modification apportée par l'équation 3.25 et l'architecture à mélangeur direct sans polarisation ce phénomène adverse est contourné. Il faut voir aussi que la valeur du maximum dans l'équation 3.25 peut être utilisée dans le système pour contrôler un amplificateur AGC réel.

La deuxième modification de l'algorithme original est la procédure mathématique de traitement du sous-système de décision. Contrairement à l'application initiale de ART2 dans la reconnaissance des caractères, on connaît le nombre de catégories (quatre pour QPSK ou ses variantes, seize pour 16-QAM etc.). De plus on peut considérer que la réalisation physique du circuit n'est pas complètement hors spécifications (pour un circuit idéal le déphasage nominal entre chaque canal est de 120 degrés, voir le chapitre 2). Dans ce cas, on peut initialiser les coefficients de transfert de chaque nœud interne du champ F_1 (la mémoire de courte durée) vers le champ de décision (mémoire de longue durée) avec les valeurs d'un récepteur idéal (sans erreurs de phase et de gain). Cette modification apportée à l'architecture originale permet au récepteur de fonctionner comme un démodulateur I&Q non-compensé, dans le cas d'un fonctionnement défectueux du logiciel. Cette initialisation est équivalent à l'utilisation d'un biais dans le niveau décisionnel. De plus, cette hypothèse augmente de façon significative la vitesse de convergence de notre algorithme dans le cas d'une jonction correctement conçue, et ce sans conséquence sur la capacité de correction ; nous pouvons toujours corriger des erreurs de phase de l'ordre de 25 degrés et des variations de gain de 40%.

Une autre modification a été apportée au mécanisme du système d'orientation. De la même façon que ART2, notre algorithme vérifie la vraisemblance entre le vecteur d'entrée et les coefficients de transfert associés au nœud gagnant dans la couche de décision. Normalement, une éventuelle asymétrie détectée par le sous-système d'orientation va désactiver le nœud gagnant et le système va continuer la recherche d'une autre solution parmi les catégories encore valides ou va générer un nouvel état. Dans notre cas, on procède seulement à l'inhibition du nœud actif (on ne modifie pas les coefficients de transfert) ou on ré-initialise la valeur des coefficients de transfert à leur valeur 'par-défaut'. Le signal d'orientation sera utilisé au niveau système, pour signaler une erreur, due probablement à un niveau élevé de bruit ou une désynchronisation des symboles et/ou de la porteuse. La synchronisation n'est pas contrôlée par cet algorithme, mais la rotation de la constellation, conséquence d'un manque de synchronisme au

niveau de symbole et/ou porteuse peut être détecté par notre méthode de correction comme un chevauchement de zones de décision, quand deux états de modulation ne sont plus différentiables.

Finalement nous allons maintenant décrire le formalisme utilisé par notre nouvelle méthode de correction pour démoduler les signaux. En particulier, le calcul des états de modulation à partir des mesures à la sortie de la jonction est expliqué.

Considérons la matrice composée par les valeurs numériques, échantillonnées à la sortie de la jonction (unité RF dans la figure 3.1) :

$$\begin{bmatrix} U_1 \\ U_2 \\ U_3 \end{bmatrix} = \begin{bmatrix} K_1 \bullet \alpha_1 & K_1 \bullet \beta_1 \\ K_2 \bullet \alpha_2 & K_2 \bullet \beta_2 \\ K_3 \bullet \alpha_3 & K_3 \bullet \beta_3 \end{bmatrix} \begin{bmatrix} x_n \\ y_n \end{bmatrix} + \begin{bmatrix} O_1 \\ O_2 \\ O_3 \end{bmatrix}, \quad 3.26$$

où K_i sont les valeurs de gain relatif, α_i et β_i sont les coordonnées du déphasage équivalent pour le canal $i = 1 \dots 3$. $\langle x_n, y_n \rangle$ est le signal reçu correspondant à l'espace bi-dimensionnel I&Q (le signal contient du bruit). Les termes O_i sont les niveaux d'écart DC pour chaque canal respectivement. Suite au théorème 3.21, les coefficients de transfert vers la couche de décision sont proportionnelles avec l'excitation d'entrée:

$$\begin{bmatrix} U_1 \\ U_2 \\ U_3 \end{bmatrix} = a \bullet \begin{bmatrix} W_{x1} \\ W_{x2} \\ W_{x3} \end{bmatrix}, \quad 3.27$$

ici "a" est une constante et W_{xi} sont les coefficients de transfert entre la i -ème coordonnée du vecteur d'entrée et le nœud gagnant x . Dans le cas de la modulation QPSK, chaque paire $\langle x, y \rangle$ correspond à un état de modulation et elle va générer une matrice W_x différente dans l'équation 3.27.

Les écarts DC ne sont pas corrélés avec les données et de plus, ils ont une valeur constante, au moins pour un temps court par rapport à la vitesse de transmission. En estimant la moyenne en temps (symbole après symbole), on peut soustraire les valeurs des écarts et on obtient:

$$\begin{bmatrix} \tilde{U}_1 \\ \tilde{U}_2 \\ \tilde{U}_3 \end{bmatrix} = \begin{bmatrix} T_{11} & T_{12} \\ T_{21} & T_{22} \\ T_{31} & T_{32} \end{bmatrix} \begin{bmatrix} x_n \\ y_n \end{bmatrix} \quad 3.28$$

On note que le système décrit par (3.28), est sur-dimensionné, on dispose de trois mesures pour déterminer (x,y). Pour résoudre le système, il suffit en principe de choisir deux équations.

Nous définissons $\|C\|$, la matrice obtenue par la troncation de la matrice $\|T\|$;

$$\|C\| = \begin{bmatrix} C_{q1} & C_{q2} \\ C_{p1} & C_{p2} \end{bmatrix} \quad 3.29$$

avec $p,q=1\dots 3$ et $p \neq q$

Alors, la paire démodulée devient :

$$\begin{bmatrix} \hat{x} \\ \hat{y} \end{bmatrix} = \|C\|^{-1} \bullet \begin{bmatrix} \tilde{U}_q \\ \tilde{U}_p \end{bmatrix} \quad 3.30$$

En remplaçant dans l'équation 3.28, le vecteur d'entrée par la matrice des coefficients de transfert W_x pour chaque nœud gagnant, on va extraire les coefficients de la matrice $\|T\|$. On comprend alors comment notre algorithme permet de démoduler le signal (équation 3.30) tout en ajustant de façon adaptative les paramètres de $\|T\|$. Dans la version de l'algorithme développée sous l'environnement SPW, pour la modulation QPSK, nous

avons utilisé les 3 combinaisons possibles par la partition de la matrice $\|T\|$ et nous avons construit un système surdéterminé. Cette approche augmente la précision numérique de la solution. Pour d'autres types de modulation comme 8-PSK et 16-QAM, nous pouvons utiliser autant d'équations que désiré pour l'extraction de la solution du système. Le nombre d'équation découle d'un compromis entre la vitesse d'exécution, la taille de la mémoire embarquée et la précision voulue. Cependant, il faut mentionner que dans un système réel, la précision numérique sera aussi limitée par la précision des convertisseurs analogique-numérique et par le mode de représentation numérique dans le processeur (virgule flottante ou non).

3.4 Résultats des Tests pour la Nouvelle Technique de Correction.

Maintenant que les détails mathématiques de l'algorithme ont été décrits, nous proposons de présenter les résultats de simulations numériques du fonctionnement d'un récepteur basé sur notre nouvelle architecture. Le système complexe de réseau de neurones a été simulé dans l'environnement SPW. Ce logiciel permet de simuler, point par point (temporellement), le récepteur au niveau du traitement du signal. Pour ces tests un système complet transmetteur-récepteur a été défini pour simuler des différentes performances de la nouvelle méthode de correction d'erreurs (voir figure 3.1 et 3.6). Bien qu'aucune limitation n'ait été imposée pour la vitesse de transmission, les tests ont été effectués avec un taux de transmission de 1Mbaud. Des résultats similaires peuvent être extraits pour différentes vitesses de symbole. La limitation de vitesse de transmission, sera imposée seulement par l'architecture physique du récepteur nécessaire pour l'exécution de l'algorithme. Comme on l'a précisé, dans notre présentation sur la simplicité des opérations primitives à l'intérieur de l'architecture ART (section 3.2), il n'est pas difficile d'atteindre des vitesses de traitement du signal extrêmement grandes (dans l'ordre de 100 Mbaud/s.) avec de processeurs commerciaux avec représentation numérique en virgule flottante.

Le spectre du canal de base a été défini par l'introduction d'un filtre passe-bas de type "cosinus surélevé" avec une fréquence de coupure de 0,8 MHz (voir figure 3.6). Ce filtre respecte la deuxième condition de Nyquist pour éviter l'interférence inter-symbole. Cette classe de filtre est presque un standard pour les transmissions numériques, et une variante, avec sa réponse temporelle, est définie dans le logiciel SPW.

Le tableau 3.1 montre la convergence de l'algorithme par rapport à un paramètre du réseau de neurones, usuellement nommé taux d'apprentissage. Ce paramètre contrôle la vitesse de variation des coefficients de transfert à l'intérieur de chaque neurone du réseau. Les erreurs de phase et de gain sont imposées par l'utilisateur. Le test de convergence mesure le nombre de symboles requis pour que la distance minimale ($d_{\min}$) associée à la constellation de signaux (voir figure 3.5) dépasse 98% de la valeur d'un système sans erreurs. Les termes d'erreurs sont des valeurs relatives par rapport au canal de référence utilisé pour la synchronisation de la porteuse et la synchronisation de symbole (voir figure 3.6).

Comme on peut s'y attendre, les résultats dépendent de la valeur initiale des erreurs et ils montrent une convergence plus rapide si le taux d'apprentissage est augmenté. En pratique, on préfère des valeurs faibles pour le taux d'apprentissage, ce afin d'assurer une meilleure protection au niveau de bruit. Les résultats du tableau 3.1 montrent que l'on peut faire un tel choix; en effet, une différence de 100 symboles représente seulement un retard de 0.1 ms dans le domaine temporel, valeur largement négligeable en pratique. On doit savoir que l'algorithme a besoin de modifier les termes d'erreurs seulement au début de la communication ou après une désynchronisation (changement de la porteuse ou perte de synchronisme). De plus, une valeur basse pour le taux d'apprentissage augmente la plasticité du système face à un signal bruité, ce qui signifie que les coefficients de transfert seront moins modifiés par le bruit. Tout comme les résultats de Carpenter et Grossberg [21,24], la vitesse de convergence dépend faiblement

de l'ordre d'arrivée des symboles. Des séquences différentes ont des vitesses d'apprentissage différentes, mais la différence est très faible (~15 symboles au maximum, pour une modulation à haut degré de modulation 64 QAM).

Taux d'apprentissage		d=0.05	d=0.075	d=0.1	d=0.15	d=0.2	d=0.5
Erreur Phase 1	+5°						
Erreur Phase 2	-5°						
Erreur Gain 1	-5%	17	17	17	17	17	17
Erreur Gain 2	+5%						
Erreur Phase 1	+20°						
Erreur Phase 2	-20°						
Erreur Gain 1	+25%	29	27	26	26	26	26
Erreur Gain 2	-25%						
Erreur Phase 1	+12°						
Erreur Phase 2	-10°						
Erreur Gain 1	-20%	58	55	42	40	39	39
Erreur Gain 2	+20%						
Erreur Phase 1	+15°						
Erreur Phase 2	-15°						
Erreur Gain 1	-25%	102	97	96	95	93	93
Erreur Gain 2	+25%						

Tableau 3.1 Analyse de Convergence.

Nous avons aussi testé notre architecture en la comparant avec un récepteur classique de type I&Q. Au fur et à mesure que les symboles subissent le traitement numérique de l'algorithme à base de réseau de neurones, la constellation distorsionnée par les erreurs physiques du circuit est changée vers la position corrigée. Ce comportement dynamique indique la puissance de la méthode et il est représenté sur la figure 3.7

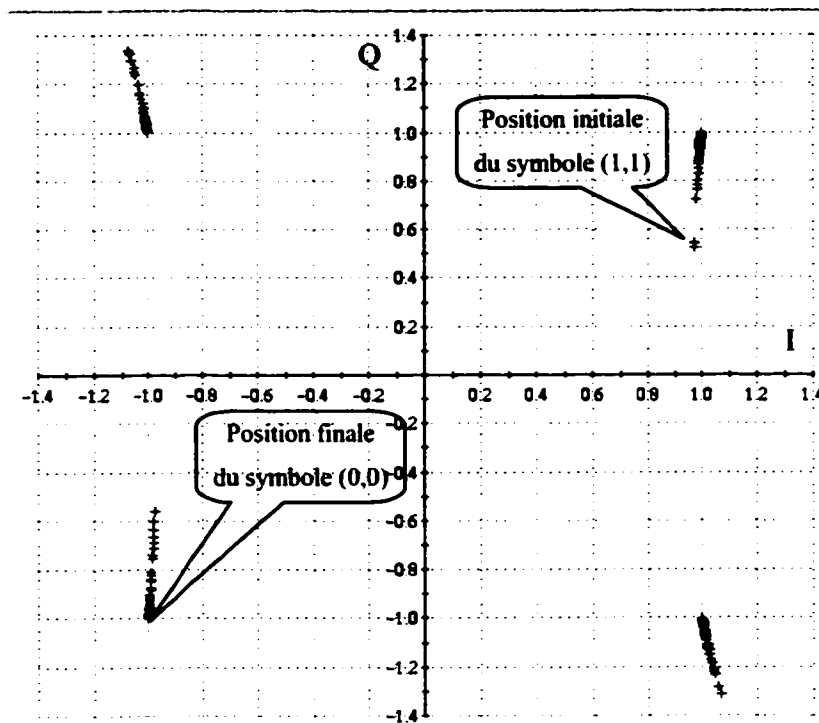


Figure 3.7 Convergence de l'algorithme.

Comme on peut l'observer dans le graphique ci-dessus, après la période de convergence, la position idéale est obtenue si on considère l'absence de bruit. Les vecteurs corrigés par les valeurs des coefficients de transfert assurent la proportionnalité entre le signal d'entrée et la position dans la constellation. Autrement dit, notre algorithme qui minimise l'erreur quadratique entre les vecteurs mesurés et les vecteurs prédéfinis, est équivalent à un récepteur optimal. En conséquence, un paramètre qui quantifie la qualité de la réception est la distance minimale entre deux zones de décision voisines (le plus proche quadrant dans le cas de la modulation QPSK). Les résultats sont présentés dans le tableau 3.2.

Erreur Phase Canal 1 (deg.)	Erreur Phase Canal 2 (deg.)	Erreur Gain Canal 1 (%)	Erreur Gain Canal 2 (%)	d_{min} , Architecture I&Q (%)	d_{min} , Nouvel Algorithme après 100 symboles	d_{min} , Nouvel Algorithme après 500 symboles
10	8	0	0	86.1	98.8	100.0
10	-8	0	0	91.0	99.2	99.99
5	-5	10	-10	92.0	98.9	100.0
5	5	10	-10	97.0	99.2	99.99
15	15	-40	-40	51.0	98.6	99.99
25	-30	-30	+20	39.7	98.7	100

Tableau 3.2 Capacité de correction.

Les valeurs trouvées dans l'analyse de la capacité de correction, nous montrent que l'algorithme possède les performances désirées. Toutefois, on doit préciser qu'en pratique, il y a toujours un bruit de quantification, induit par la résolution finie de la représentation numérique sur le processeur de signal. Cet effet est ici ignoré, nous nous sommes concentrés sur l'étude des caractéristiques intrinsèques de notre méthode en utilisant une représentation numérique en virgule flottante. Ce choix, explique la valeur exacte de 100% de la distance obtenue après la convergence.

L'analyse quantitative montre clairement la supériorité de l'architecture "multiports - méthode adaptative de correction" par rapport à l'architecture classique I&Q. On peut corriger des valeurs extrêmes d'erreur de phase, qui correspondent à une baisse de presque 8 dB dans le rapport signal sur bruit pour sa contrepartie I&Q. Tout en conservant le comportement linéaire du récepteur, notre algorithme amène graduellement, de façon adaptative, la constellation vers la position normale.

Maintenant que notre approche théorique est validée par ces tests, nous étudions le comportement du récepteur adaptatif dans un environnement bruité. Nous avons

considéré un bruit blanc additif de type Gaussien (AWGN) pour nos simulations. Le graphique ci-dessous, présente les résultats pour une constellation QPSK dans les deux cas (version I&Q et nouvelle méthode de correction) :

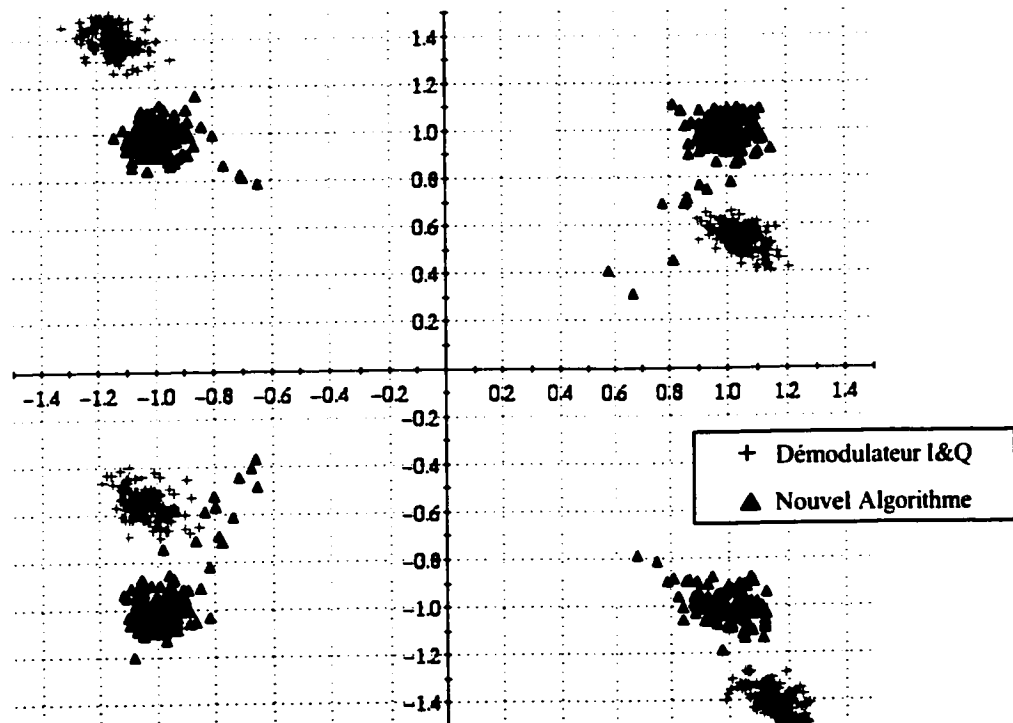


Figure 3.8 Constellation I&Q de sortie avec le bruit gaussien.

Comme il est suggéré de façon qualitative dans la figure 3.8, l'algorithme va déplacer graduellement le centre de gravité de la constellation vers la position unitaire dans le plan cartésien I-Q.

Notons que le processus de calcul des coefficients d'erreurs et d'ajustement des coefficients de transfert ne doit pas nécessairement être fait sur tous les symboles reçus. Ainsi, on peut exciter le réseau de neurones avec une version décimée de la séquence de symbole d'arrivée; Par exemple ce sera le cas si on veut corriger des processus qui

engendrent des erreurs variant lentement. Cet avantage nous permet de faire un compromis entre la vitesse de convergence de l'algorithme et la résolution : Ceci peut permettre d'utiliser des processeurs de signal à haute résolution numérique, plus lents.

D'autre part, si on considère le procédé de récupération de la porteuse, une différence entre la phase instantanée du signal de l'oscillateur locale et la porteuse RF sera transposée dans une rotation par rapport à la constellation actuelle. Un traitement temporel sur ce type de rotation peut conduire à la conception d'un système de synchronisation assisté par le réseau de neurones. Cette caractéristique peut s'avérer une amélioration supplémentaire au fonctionnement des récepteurs homodyne. Le problème est de considérer un espace à trois dimensions, la différence de fréquence instantanée entre la porteuse et le signal de l'oscillateur local étant l'inconnue supplémentaire.

Le test final pour la validation de l'algorithme est le calcul du taux d'erreurs binaires (BER). La figure 3.9 présente les courbes de BER en fonction du rapport signal bruit pour un récepteur utilisant notre algorithme de correction et pour un démodulateur I&Q ayant les mêmes erreurs de phase et de gain. Comme référence de calcul, le graphique incorpore aussi la courbe du récepteur idéal I&Q dans le cas d'absences d'erreurs physiques. L'algorithme proposé dans ce chapitre améliore de manière significative les performances du récepteur. La différence par rapport à la courbe théorique peut être expliquée par le fait que, dans le calcul de BER, tous les symboles sont considérés. Dans, la phase initiale, avant que notre algorithme ait convergé, le taux d'erreur sera relativement élevé, du même ordre que celui obtenu avec le récepteur I&Q avec des erreurs. Maintenant, si on considère seulement la séquence des symboles après convergence de notre algorithme, les performances de notre récepteur approchent celles du récepteur I&Q sans erreurs de phase et de gain. La différence résiduelle provient seulement de la représentation numérique finie utilisée par le logiciel et du seuil de bruit imposé par le taux d'apprentissage.

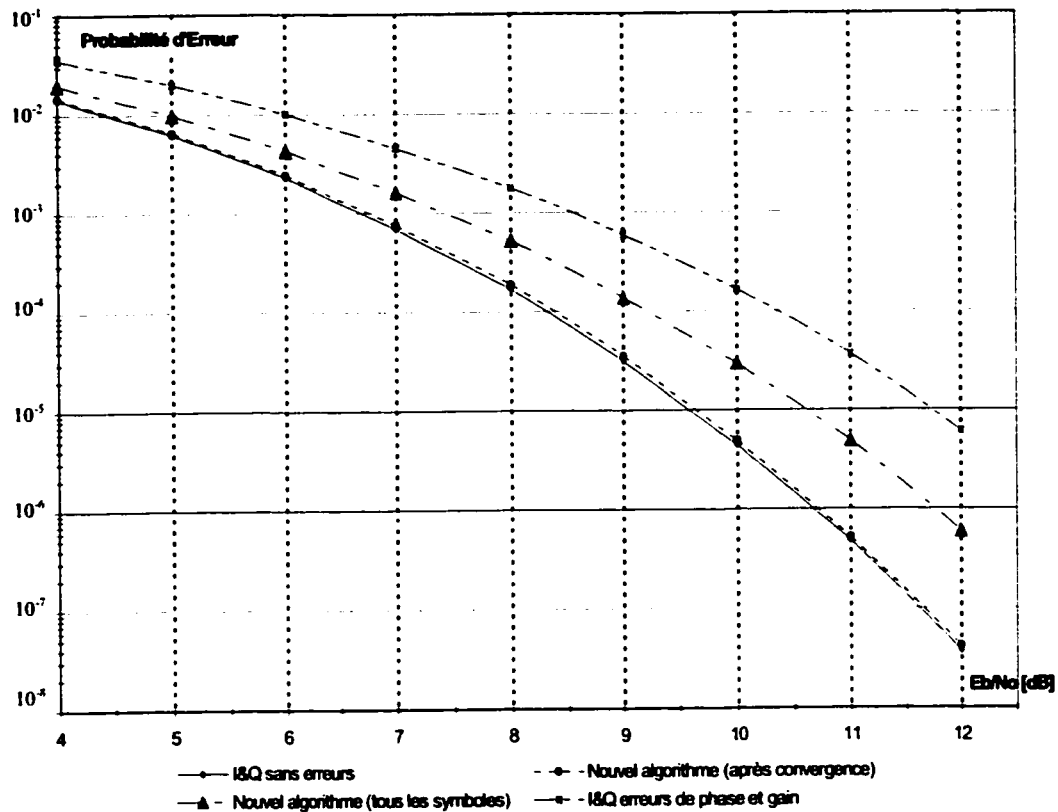


Figure 3.9 La caractéristique de probabilité d'erreur

Finalement, le graphique 3.9 peut être interprété de la façon suivante. Les résultats obtenus avec notre algorithme évoluent, d'une réponse de type "I&Q avec erreurs de phase et de gain" vers une réponse de type "I&Q sans erreur" (ou aussi parfaitement corrigé).

Alors que les résultats présentés jusqu'à présent ont utilisé la modulation QPSK comme schéma de modulation générique, rien ne nous empêche d'appliquer notre technique à n'importe quel type de modulation en quadrature. Il suffit de multiplier le nombre d'états associés à un point dans la constellation de signal. Autrement dit dans la couche de décision F_2 le vecteur de sortie va avoir un nombre de composantes identique au degré

de modulation. Pour la modulation QPSK, on avait 4 neurones associées, pour 8-PSK et 16-QAM nous aurons respectivement 8 et 16 éléments. La procédure mathématique n'est pas modifiée au niveau de l'algorithme de transformation des coefficients de transfert. En conséquence on s'attend à obtenir des résultats similaires à ceux présentés ci-dessus. La figure 3.10 présente les résultats pour une modulation de type 8-PSK:

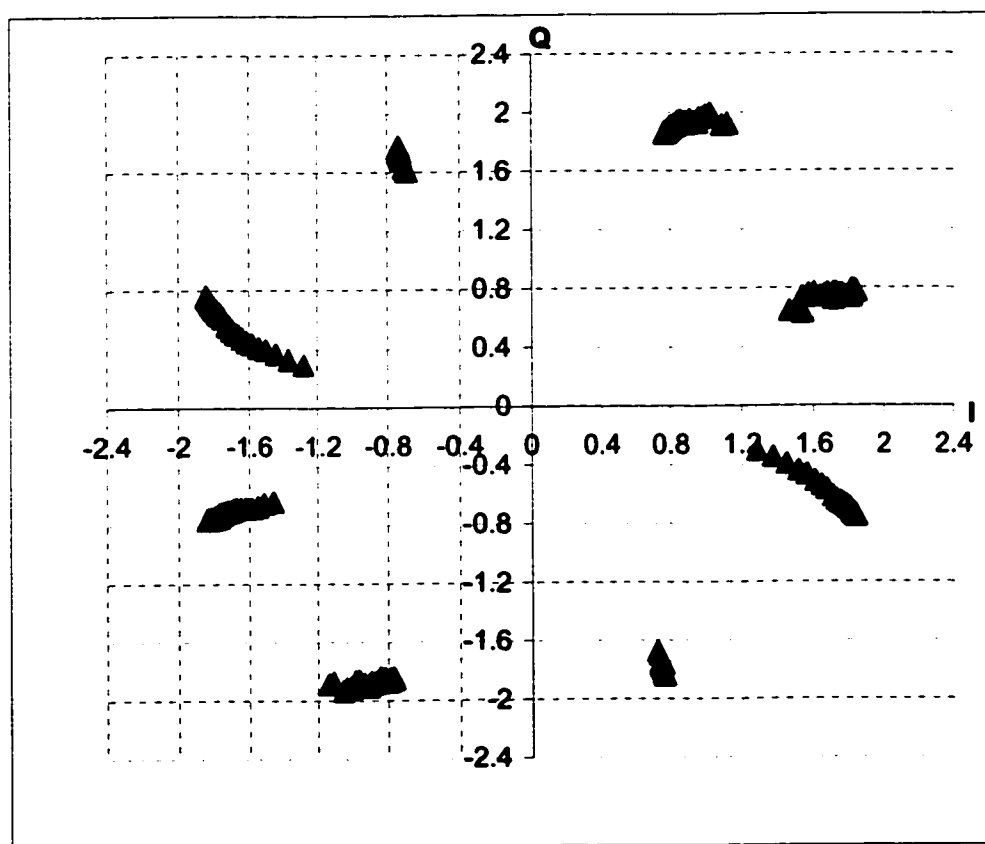


Figure 3.10 Convergence de la constellation 8-PSK

Le graphique 3.10 correspond à une erreur de phase de -12 et $+10$ degrés sur les canal 1 et 2 respectivement, et à une erreur de gain de 1.2 et 0.8. Comme d'habitude, les erreurs sont référencées par rapport au canal de référence. On voit clairement comment la constellation est modifiée pour obtenir les positions corrigées. Un résultat similaire peut

être obtenu, dans le cas du système réel où on a ajouté du bruit blanc gaussien. Pour une présentation qualitative, la figure 3.11 montre aussi le cas du système classique I&Q sans correction.

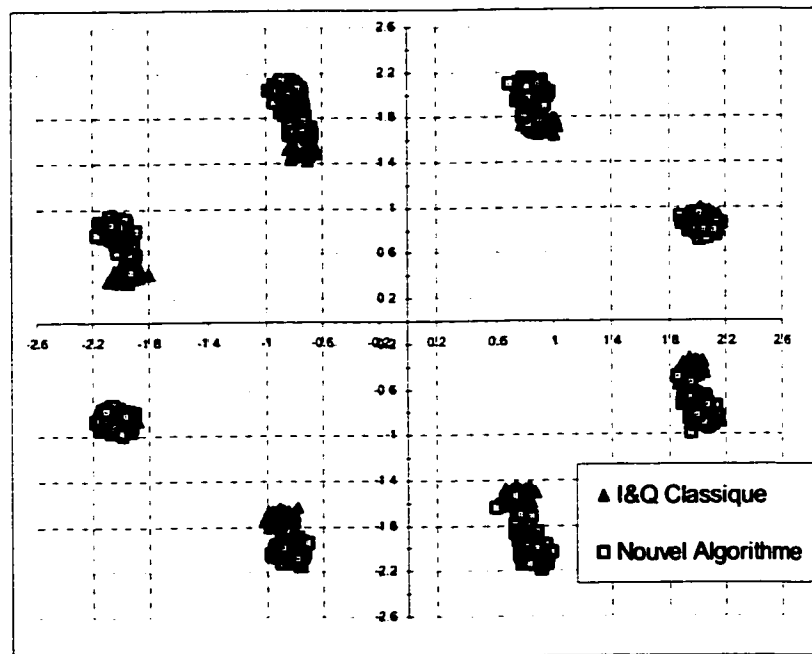


Figure 3.11 Constellation de sortie 8-PSK avec le bruit gaussien

On peut observer que le centre de gravité de chaque point de la constellation est modifié par rapport à la version non-corrigée. Suite au calcul des coefficients de transfert, la distance équivalente minimale entre deux régions de décision voisines sera augmentée en utilisant notre algorithme. Évidemment ceci a des conséquences favorables sur le taux d'erreurs binaires. Cette conclusion peut être généralisée à tout signal de type M-PSK.

Une différence importante entre les modulations de type M-QAM et M-PSK nécessite une attention spéciale. Il s'agit de l'énergie instantanée du signal modulé. Si dans le cas

de la modulation M-PSK, tous les symboles ont la même énergie, pour la variante M-QAM ce n'est pas le cas. Comme on peut le voir dans la figure 3.12, l'énergie des symboles aux "frontières" de la constellation est plus grande que celle des symboles situés dans l'intérieur de la constellation.

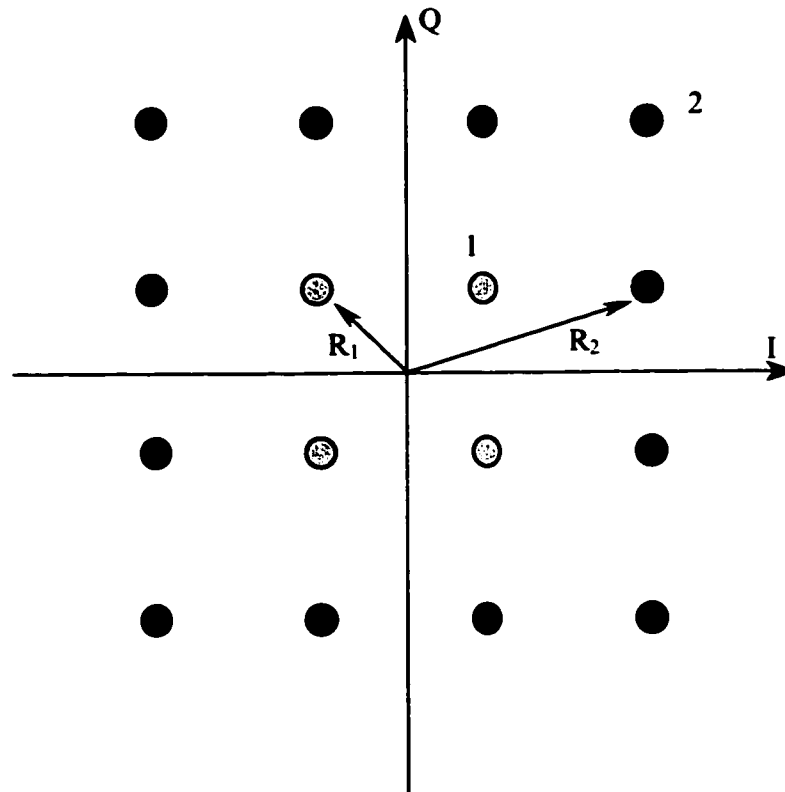


Figure 3.12 Énergie de la constellation 16-QAM

Dans le graphique 3.12 les symboles plus foncés représentent un état avec une plus grande énergie. Le principe de calcul étant très simple, l'énergie de chaque signal membre de la constellation étant proportionnelle avec le carré de la distance vers le centre du diagramme I&Q. Cette variation de l'énergie des symboles a des conséquences importantes pour l'algorithme de correction.

Dans la section 3.2 on a développé la théorie du réseau de type ART. Regardant l'équation 3.21, on a trouvé une importante propriété mathématique, le vecteur de sortie gagnant (dont l'excitation au niveau du champ de décision est maximale) est proportionnel au vecteur correspondant au niveau de la constellation. Dans le cas d'une modulation à énergie constante, il n'y a aucune difficulté pour notre traitement, en raison du fait que cette condition oblige tous les vecteurs de la constellation à être indépendants (non co-linéaires). En contrepartie, la modulation M-QAM présente des vecteurs co-linéaires, donc il peut y avoir proportionnalité entre deux membres différents de la constellation (les vecteurs d'ordre 1 et 2 dans la figure 3.12). Les résultats de notre algorithme appliqué directement seront erronés si on ne considère pas cette propriété. Autrement dit, si on ne fait pas attention, notre réseau va détecter la proportionnalité entre les deux vecteurs et il va considérer qu'ils appartiennent à une même classe. Le résultat n'est pas étonnant, il est équivalent à la détection de caractères des tailles différentes. Pour un tel réseau utilisé pour la reconnaissance des caractères, on peut dire que la lettre "G" est identique avec "g". La conséquence immédiate est que la constellation idéale de type QAM sera mal interprétée par notre algorithme.

La solution est de modifier l'interaction entre le niveau pré-traitement et le niveau de décision. On a précisé dans la section 3.3 qu'on doit à tout prix garder la linéarité du traitement afin de préserver les caractéristiques linéaires des erreurs. Dans ce cas, une des modifications apportée à l'algorithme initial était de diviser le vecteur d'entrée par une valeur maximale. Notre réseau calcule la valeur maximale du signal à l'entrée du démodulateur. Celle-ci reste quasiment constante, tant qu'aucune modification n'est apportée au circuit physique. Dans ce cas, on va utiliser cette valeur comme un indice de l'énergie du signal. Pratiquement, nous divisons la constellation en deux constellations adjacentes, formées pour une partie par les membres de l'intérieur de la constellation QAM idéale, et pour la seconde partie par les membres de la frontière de la constellation. On observe que, à l'intérieur de chaque "sous-constellation", on n'a plus de vecteurs proportionnels.

L'algorithme va calculer la différence entre l'énergie du signal d'entrée et une valeur de référence, correspondant à un signal de faible énergie. Dans le cas de la réponse positive à ce test, le champ de décision va recevoir un signal d'inhibition pour les états de "haute énergie" (les cercles foncés) et va sélectionner le gagnant seulement entre les états de faible énergie. Avec ce raisonnement simple, sans une grande augmentation de la complexité de l'algorithme on peut aisément résoudre le problème de partition dans le cas des signaux M-QAM

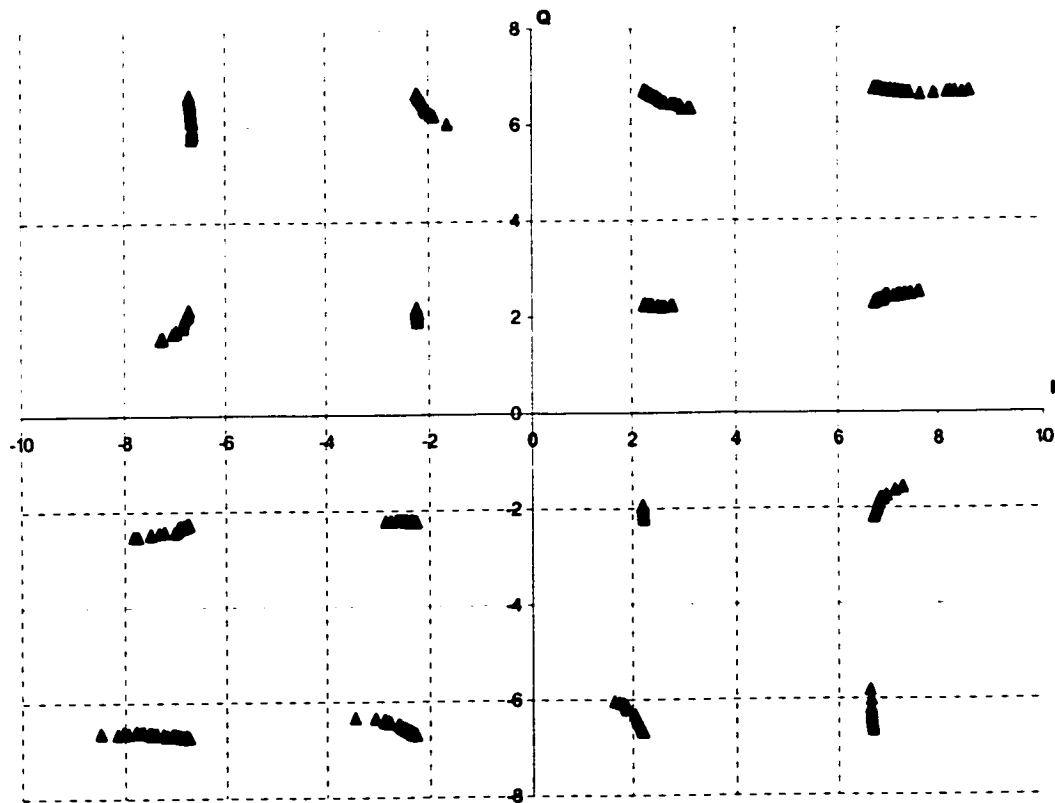


Figure 3.13 Convergence de l'algorithme pour une constellation 16-QAM

Il est ainsi évident que notre algorithme peut facilement résoudre n'importe quel problème de démodulation en présence des erreurs si on a une modulation de type I&Q. Si on est capable de transformer notre problème dans un problème de partition dans un

espace cartésien, le réseau de neurones va simplement calculer les coefficients de transfert correspondants aux positions dans la constellation. De plus l'algorithme étant indépendant de la taille du problème, on peut embarquer sur le même processeur numérique le code pour chaque type de modulation. Il s'agit dans ce cas d'un "soft-radio" dont l'unité RF est devenue un circuit générique, pouvant être utilisé dans plusieurs systèmes.

Après avoir présenté en détail, les fondements mathématiques et les performances du nouvel algorithme, nous pouvons faire le point sur les éléments originaux introduits par cette nouvelle approche.

En premier lieu, notre technique de compensation adaptative ne montre pas les limitations des méthodes de calibration à base des standards, en éliminant ainsi le besoin d'équipements additionnels et/ou le temps de calcul dédié à cette fin.

En deuxième, l'algorithme de résonance adaptative possède une complexité proportionnelle avec le type de modulation utilisée, étant strictement défini par celle-ci. Cet avantage nous a permis de l'utiliser indépendamment du format de modulation, en devenant ainsi très flexible aussi pour la conception du récepteur comme pour le matériel utilisé (il peut être embarquer sur une architecture générique de type DSP)

En troisième, le taux de transmission n'influence pas le fonctionnement de l'algorithme et nous pouvons envisager utiliser cette technique pour une version décimée de la trame de données ou pour la version intégrale. Il est simplement le choix du concepteur du système de faire le compromis entre divers paramètres. Dans le même contexte, nous devons souligner que pour toutes les modulations présentées (n-PSK, n-QAM) la structure du réseau des neurones est identique (voir figure 3.3) incluant 3 couches de traitement. Le délai de traitement devient ainsi 3 unités d'horloge de symbole dans une version intégrale et 3 fois le facteur de décimation dans une version décimée. Nous prouvons ainsi une architecture parallèle, très efficace au niveau vitesse de traitement.

CHAPITRE 4 MESURES ET FABRICATION DE CIRCUITS MMIC

On a vu dans les chapitres précédents, différents aspects concernant les avantages de l'architecture de réception directe par rapport aux démodulateurs IF super-hétérodyne et aussi, certaines techniques pour palier aux points faibles de ce type d'architecture. Néanmoins, les techniques de circuits ne peuvent pas surpasser les avantages apportés par l'intégration des fonctions sur une seule puce. Dans cette optique, ne pas avoir de composantes électroniques au niveau de la fréquence intermédiaire a un avantage important au niveau de la facilité d'intégration, et ainsi du coût de la solution homodyne.

Traditionnellement, les circuits micro-ondes, spécialement les circuits en ondes millimétriques sont des systèmes coûteux. Réduire la complexité des circuits amène la diminution du coût du circuit et, dans ce cas, le récepteur direct reste un candidat idéal pour l'implémentation des systèmes de communication sans fils. Même si notre technique de calibration est une solution pour augmenter la qualité des récepteurs directs et même si elle peut être embarquée sur des circuits génériques à faible coût (processeurs de signal), il n'en reste pas moins que notre architecture de tête réceptrice doit être capable de garder ces avantages pour les circuits en ondes millimétriques. En regardant la figure 3.1, on observe que dans notre cas, l'intégralité des fonctions RF peut être fabriquée sur le même circuit intégré, utilisant de manière uniforme la même technologie en fonction de la plage de fréquences qu'on a besoin de couvrir (bipolaire en silicium pour fréquences inférieures aux 12 GHz, ou GaAs, SiGe et InP pour les fréquences élevées).

Les résultats présentés dans les chapitres précédents ont tenu compte d'une réalisation générique de la jonction équivalente six-ports, sans considération pour une réalisation particulière. La modélisation de la jonction, fondée sur les mesures du circuit de la figure 3.4 n'impose pas des contraintes au niveau des choix du concepteur. Toutefois, l'architecture présentée dans la figure 3.1 suggère une solution complètement "active".

Autrement dit, la jonction est composée seulement des composantes actives intégrées sur la même puce. Au début de ces travaux de recherche il n'y avait aucune possibilité de fabrication des circuits monolithiques pour les ondes millimétriques au laboratoire Polygrames. Un accès aux fonderies commerciales reste très coûteux pour le milieu universitaire et il présente un cycle de retour (conception-fabrication-mesures) relativement long. Dans ces conditions, on a décidé de développer notre propre unité de fabrication, de mesures et de modélisation pour les circuits monolithiques en ondes millimétriques. La plus grosse partie de notre travail a été dédiée à ce développement, ce en partenariat avec le laboratoire de micro-électronique d'Université de Sherbrooke (GMS), avec le but précis d'obtenir des circuits pouvant être utilisés comme sous-systèmes dans notre architecture directe.

Ce chapitre va présenter les différentes techniques de mesure développées pour la caractérisation de circuits MMIC. On verra à la fois les aspects théoriques et aussi les aspects portant sur les détails pratiques de l'implémentation du système de mesure. Aussi une section du chapitre est dédiée à la technologie de fabrication de circuits MMIC en GaAs, dont les résultats sont présentés dans le chapitre suivant. L'idée de notre approche est d'obtenir un environnement intégré pour la fabrication des composantes, les mesures, l'extraction de modèles à partir de l'étape de mesures et bien sûr la conception des sous-systèmes à l'aide des outils CAO utilisant ces modèles. Un paramètre important de ce cycle de développement étant sa durée, on a essayé d'automatiser au maximum les étapes et améliorer les procédures de fabrication pour avoir un temps de réponse ("turn-around") relativement faible. En raison de la nouveauté des processus de fabrication pour les laboratoires Polygrames et GMS, nous avons essayé de développer une expertise spécifique à nos équipements en fonction des limites matérielles et budgétaires existantes. Les sections suivantes vont clairement mettre en évidence les travaux accomplis et vont présenter les détails de notre système intégré de fabrication, de mesure et de caractérisation de composantes MMIC.

4.1 Mesures des Dispositifs en Hautes Fréquences.

Comme le nom le suggère, les circuits MMIC fonctionnent à des fréquences relativement élevées. La qualité des circuits fabriqués dépend évidemment, de la précision des mesures effectuées pour extraire les modèles utilisés dans la conception des circuits. En résumé, plus les modèles sont proches de la réalité, moins la tâche du concepteur est lourde.

Pour la procédure d'extraction des modèles, nous avons développé plusieurs techniques de mesure pour la caractérisation de circuit, dont les fondements mathématiques seront présentés dans cette section. Les techniques ont été implantées au moyen du logiciel ICCAP™ (de la compagnie Agilent Technologies). L'utilisation de ce logiciel en conjonction avec la station de mesures sous-pointes, permet d'obtenir un environnement de mesures complètement automatisé.

Contrairement aux circuits de basses fréquences, qui peuvent être caractérisés sans difficultés avec des mesures de type I-V en courant continu et des oscilloscopes pour les caractéristiques AC, les circuits MMIC nécessitent des mesures complexes. Même si l'analyseur de réseau, l'appareil utilisé depuis les années '70 pour caractériser les circuits micro-ondes, peut mesurer avec précision les paramètres de répartition S du circuit étudié, certaines corrections doivent être apportées afin d'extraire l'information recherchée.

La figure 4.1 montre la sonde de mesure utilisée avec la station sous-pointes. On parle ici d'un type de connexion GSG ("ground-signal-ground"), adaptée aux mesures de types coplanaires. Ce type de mesures apporte deux avantages majeurs au niveau de la précision. En premier lieu, la transition entre le guide coaxiale et le guide coplanaire permet d'atteindre une très large bande en gardant un très faible facteur de pertes d'insertion et de réflexion. Le deuxième avantage consiste dans la présence du plan de masse au même niveau que les lignes du signal. Dans ce cas, contrairement à l'environnement de type micro-ruban, on peut fabriquer facilement des standards de

calibration comme le court circuit. Dans le cas de la version micro-ruban, pour une connexion en court circuit on doit percer de trous dans le substrat ("via-holes"). Les inductances parasites associées avec le trou, additionnées aux effets résistifs dépendants de la fréquence vont largement diminuer la précision de mesure.

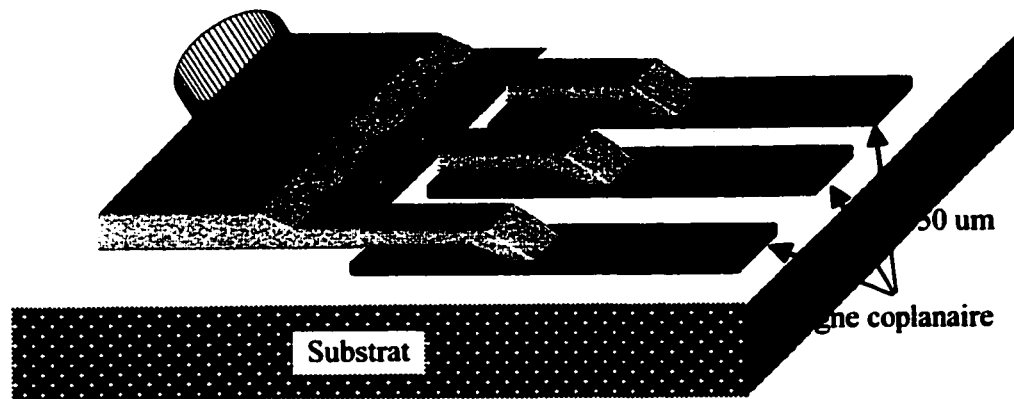


Figure 4.1 Sonde de mesure de type GSG

On note que pour les sondes de mesures sous pointes, il n'est pas possible d'avoir une connexion directe entre les deux ports de mesures. Dans ce cas, la technique de calibration doit inclure une procédure spéciale pour le changement de plan de mesure. Cette contrainte est compatible avec les dimensions des circuits à mesurer. La constante de la sonde de mesure (distance entre deux doigts consécutifs) est de 150 microns, alors la valeur de la dimension transversale du système de mesure est supérieure aux dimensions des dispositifs MMIC (un transistor PHEMT occupe une surface équivalente de $\sim 40 \times 100$ microns et les capacités peuvent avoir des dimensions de seulement 20×20 microns). Dans ce cas, il est évident que le plan de mesure de la sonde est différent du plan associé au dispositif sous test (DST). D'une façon qualitative, on présente dans la figure 4.2 le principe de mesure utilisée pour l'extraction de modèles des composants MMIC. Le plan réel de mesure est associé avec la connexion physique entre les sondes et le substrat. Pour la caractérisation, les valeurs mesurées doivent correspondre au plan effectif de mesure. En conséquence, on doit diviser l'opération de mesure en deux

étapes. La première étape correspond à la calibration par rapport aux sondes, la seconde étape est la procédure de changement du plan de mesure.

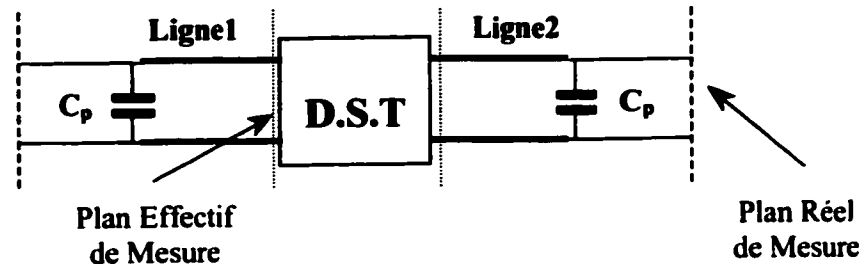


Figure 4.2 Schéma simplifié de mesure.

La calibration au niveau des sondes utilise la technique standard de calibration d'un analyseur de réseau. Si la fréquence de caractérisation (ou la plage de fréquence) est inférieure à 12-15 GHz, on utilise la calibration de type SOLT ("short-open-load-thru"). Pour couvrir l'intégralité de la bande de mesure (~40 GHz) on doit utiliser une technique plus précise comme LRM ("line-reflect-match").

Dans les deux cas, les standards de haute qualité sont utilisés sur un substrat en alumine. En particulier, on utilise les standards fournis par le fabricant des sondes (Picoprobe). Si les composantes à mesurer sont fabriquées sur un autre type de substrat (en occurrence GaAs), dans la technique de changement de plan de mesure on doit tenir compte du saut de permittivité relative du matériel (Alumine ~ 9.6, GaAs~12.9). Cette correction est représentée dans la figure 4.2 par l'effet capacitif C_p . Pour tous les circuits, la taille de plots de connexions est maintenue constante, donc on peut facilement extraire la valeur de la capacité C_p . Parfois, on monte les circuits à caractériser sur des substrats en alumine et cette correction n'est plus nécessaire.

On ne veut pas insister sur la première étape. Elle est spécifique à n'importe quelle technique de mesure en micro-ondes et le lecteur peut trouver les détails dans les guides d'application de l'analyseur de réseau. Cette étape est obligatoire quel que soit le type de

mesure (coplanaire ou coaxiale), étant le moyen d'obtenir les termes de calibration pour la correction de l'analyseur vectoriel de réseau. La technique SOLT comme celle de type LRM détermine tous les 12 termes d'erreur pour la correction.

Nous allons continuer avec la présentation des différentes techniques de changement du plan de mesure ("de-embedding"). La première technique est une puissante méthode de calibration. Il s'agit de la méthode multi-lignes développée comme une généralisation de la méthode TRL pour la calibration des analyseurs vectoriels de réseau [46,82,83,85]. Dans notre cas, on a besoin seulement de la procédure pour l'extraction de la constante de propagation sur les lignes 1 et 2 (en figure 4.2).

2.3.6 Méthode Multi-lignes.

Considérons le schéma de la figure 4.2. Si on veut calculer les paramètres de répartition S , dans le plan du dispositif (DST), il faut trouver les matrices S des diports associées aux circuits entre les points de mesure effective et le plan du DST. Le système est analysé de façon optimale à l'aide de paramètres du type chaîne. Nous avons choisi la définition de telle manière, que la matrice de la connexion "cascade" du diport A connecté en série, à la gauche du diport B, donne simplement le produit AB pour représenter le circuit en cascade.

En tenant compte de la symétrie du système entre les ports 2 et 1 de l'analyseur de réseau, on peut trouver la matrice de la connexion en cascade mais dans la direction opposée (de droite à gauche). Appelons cette connexion opposée la matrice chaîne renversée :

$$\bar{A} = \tilde{A}^{-1} = \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} A^{-1} \begin{bmatrix} 0 & 1 \\ 0 & 1 \end{bmatrix}, \quad 4.1$$

L'équation 4.1 représente la liaison apportée par la symétrie entre la matrice chaîne opposée et la matrice directe.

Si on considère la matrice de paramètres chaîne mesurée pour un standard "i", suite à notre notation le résultat est :

$$M^i = XT^i\overline{Y} \quad , \quad 4.2$$

M est composée par les valeurs mesurées et X et Y sont les matrices inconnues, représentant les termes d'erreurs et elles doivent être déterminées.

On souligne que le trait au-dessus de la matrice représente la matrice opposée (de droite à gauche), correspondant au deuxième port. De cette manière les matrices X et Y sont inter-changeables au niveau du développement mathématique.

Dans 4.2 T^i est la vraie matrice chaîne du standard. Si le standard (notre cas) est une ligne de transmission idéale et la transition sonde-ligne est parfaite, T^i devient :

$$L^i = \begin{bmatrix} e^{-\gamma l_i} & 0 \\ 0 & e^{-\gamma l_i} \end{bmatrix} = \begin{bmatrix} E_1^i & 0 \\ 0 & E_2^i \end{bmatrix} \quad , \quad 4.3$$

où γ est la constante de propagation et l_i est la longueur de la ligne "i". En pratique, la ligne utilisée comme standard n'est pas idéale. Dans ce cas on peut représenter sa matrice chaîne par :

$$T^i = (I + \delta^{1i})L^i \left(\overline{I + \delta^{2i}} \right) \quad , \quad I = \text{matrice unitaire.} \quad 4.4$$

dont les matrices δ^i modélisent les faibles imperfections dues aux connexions. On a considéré ces valeurs comme les termes d'erreur associés aux ports 1 et 2 et on a gardé l'hypothèse qu'elles sont faibles.

Si on mesure une paire de lignes de longueurs différentes, les équations de type 4.2 nous donnent :

$$M^{ij}X = XT^{ij}, \quad 4.5$$

Où on a défini

$$M^{ij} \equiv M^j(M^i)^{-1} \quad 4.6$$

$$T^{ij} \equiv T^j(T^i)^{-1} \quad 4.7$$

Considérons en première instance que les termes d'erreurs δ sont nuls. Dans ce cas T^{ij} devient simplement :

$$L^{ij} = L^j(L^i)^{-1} = \begin{bmatrix} \frac{E_1^j}{E_1^i} & 0 \\ 0 & \frac{E_2^j}{E_2^i} \end{bmatrix} = \begin{bmatrix} E_1^{ij} & 0 \\ 0 & E_2^{ij} \end{bmatrix} = \begin{bmatrix} e^{-\gamma(l_1-l_i)} & 0 \\ 0 & e^{+\gamma(l_1-l_i)} \end{bmatrix}, \quad 4.8$$

Il est clair que dans ce cas, résoudre l'équation 4.5 revient à résoudre un problème de valeurs propres. Les termes diagonaux E^{ij} de la matrice T^{ij} sont les valeurs propres, et les colonnes de la matrice X sont les vecteurs propres de la matrice M^{ij} .

Si, contrairement à notre hypothèse simplificatrice, les connexions entre les sondes et les lignes ne sont pas parfaites, alors T^{ij} n'est pas une matrice diagonale et le problème devient plus complexe parce que les valeurs propres de la matrice M^{ij} ne suffiront plus pour le calcul des termes d'erreurs. En pratique, on peut considérer que les termes δ sont faibles et la solution est obtenue par la résolution d'un problème de perturbation des valeurs propres.

En pratique, on extrait toujours les valeurs propres et les vecteurs propres de la matrice M^{ij} . Ces valeurs sont liées aux valeurs correspondantes de la matrice T^{ij} . Si V^{ij} et R^{ij} sont les vecteurs propres, respectivement les valeurs propres de la matrice T^{ij} , on obtient :

$$\mathbf{T}^{ij} \mathbf{V}^{ij} = \mathbf{V}^{ij} \mathbf{R}^{ij}, \quad 4.9$$

donc,

$$\mathbf{M}^{ij} \mathbf{U}^{ij} = \mathbf{U}^{ij} \mathbf{R}^{ij}, \quad 4.10$$

$$\mathbf{U}^{ij} = \mathbf{X} \mathbf{V}^{ij}, \quad 4.11$$

Autrement dit $\mathbf{M}^{ij}$ et $\mathbf{T}^{ij}$ ont les même valeurs propres et l'équation 4.11 montre la correspondance entre leurs vecteurs propres. Cette propriété nous permet rapidement d'extraire l'information désirée en effectuant de simples mesures des lignes de transmission de différentes longueurs. Les valeurs propres de la matrice mesurée $\mathbf{T}^{ij}$ sont calculées par la formule 4.12:

$$\lambda_1^{ij}, \lambda_2^{ij} = \frac{1}{2} \left[(\mathbf{T}_{11}^{ij} + \mathbf{T}_{22}^{ij}) \pm \sqrt{(\mathbf{T}_{11}^{ij} - \mathbf{T}_{22}^{ij})^2 + 4\mathbf{T}_{12}^{ij}\mathbf{T}_{21}^{ij}} \right], \quad 4.12$$

Dans ce cas, la valeur estimée de la constante de propagation est :

$$\gamma^{ij} = \frac{\ln(\lambda^{ij})}{l_i - l_j}, \quad 4.13$$

Les mesures de plusieurs lignes nous permettent de diminuer l'erreur de mesure sur la valeur estimée de la constante de propagation. Une analyse statistique des erreurs peut être trouvée dans [82], mais dans notre système de mesure, nous utilisons seulement 2 ou 3 paires. Les résultats sont très précis dans ce cas. La méthode de calibration multi-lignes permet aussi l'extraction de termes d'erreur pour les connexions, mais au long de notre travail, on a préféré la variante simplificatrice, dont la correction est obtenue par la présence d'une capacité parallèle (voir la figure 4.2). Avec ces paramètres, la constante de propagation et la valeur de la capacité des plots, on peut changer le plan de mesure effective par un simple calcul matriciel :

$$A_{DST} = C_p^{-1} L_2^{-1} A_{PROBE} L_1^{-1} C_p^{-1}, \quad 4.14$$

Dans l'équation 4.14, A_{DST} et A_{PROBE} sont les matrices chaîne du dispositif mesuré dans le plan effectif, respectivement le plan des sondes. C_p et L_i sont les matrices chaîne de la capacité et les lignes déterminées analytiquement par notre procédure.

La figure 4.3 montre les résultats de la mesure d'une ligne de transmission, autre que les standards utilisés.

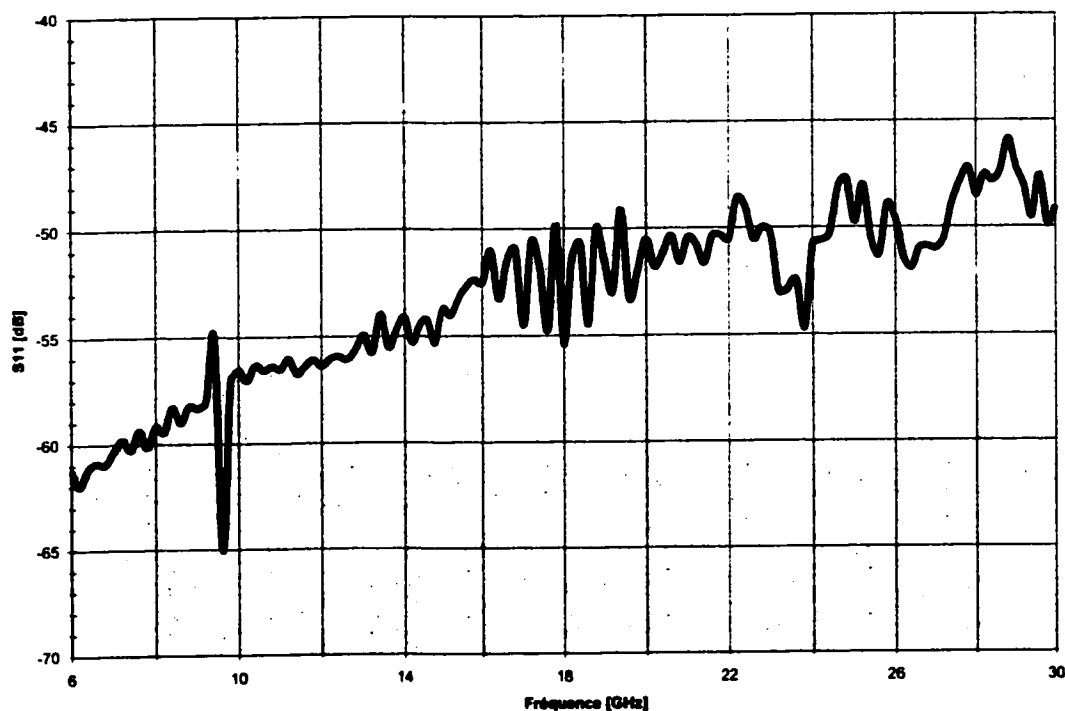


Figure 4.3 Ligne de transmission en réflexion après la calibration multi-lignes

Comme on peut remarquer, les résultats sont très bons, assurant une excellente précision pour les mesures des dispositifs à caractériser.

Évidemment la précision de la méthode dépend de la fréquence de travail, parce que la sensibilité de l'appareil est moindre pour les très hautes fréquences. Il s'agit d'un détail

très important. La qualité des mesures dépend de la puissance du signal à mesurer. Pour une plus grande précision, on a besoin d'un maximum de puissance de la source RF. Malheureusement, cette condition est incompatible avec la caractérisation de composantes actives. Les transistors sont des dispositifs à gain. La condition de mesure en petit signal qui doit être respectée en tout temps impose des niveaux de puissance faible (< -10 dbm) à l'entrée du dispositif. Si on effectue la calibration à un niveau plus grand pour augmenter notre précision de mesure, on va devoir considérer les erreurs introduites par la répétabilité des atténuateurs variables de l'analyseur. Dans ce cas un compromis doit être fait.

2.3.7 Calibration en Domaine Temporel.

Une autre méthode efficace pour le déplacement du plan de mesure tire son avantage de la capacité de mesures temporelles de l'analyseur de réseaux. L'analyseur de réseau vectoriel utilisé pour nos mesures, de type 8510C (de Hewlett-Packard) possède une option qui permet d'analyser la réponse en temps d'un circuit quelconque. Même si le fonctionnement avec cette option n'est pas similaire à celui d'un analyseur de transitions (l'outil de choix pour analyser la réponse en temps d'un circuit), nous pouvons utiliser ce mode de fonctionnement de façon optimale pour un cas spécial de calibration.

Reprenons le problème à résoudre. Pour un dispositif à mesurer DST, les valeurs mesurées correspondent à :

$$A_{\text{GLOBALE}} = A_{\text{CONN1}} * A_{\text{LIGNE1}} * A_{\text{DST}} * A_{\text{LIGNE2}} * A_{\text{CONN2}} , \quad 4.15$$

dont la notation correspond aux matrices chaîne des connecteurs, lignes etc. Si on détermine les valeurs de $A_{\text{CONN}i}$ et $A_{\text{LIGNE}i}$ par l'inversion de l'équation 4.15 on obtient la valeurs de paramètres de type chaîne du dispositif sous test (voir équation 4.14).

De même que dans la procédure de la section précédente, on va mesurer certains standards et avec un traitement mathématique on va obtenir les valeurs des matrices inconnues.

Dans la figure 4.4 on présente les standards à mesurer. Il s'agit d'une simple ligne de

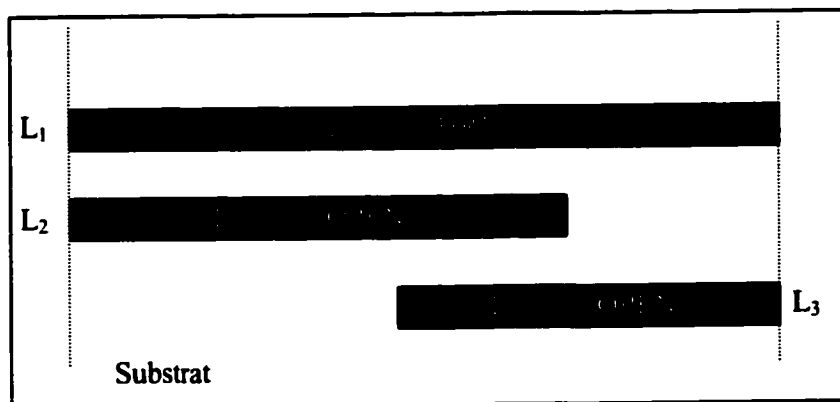


Figure 4.4 Standards de calibration

type "THRU", et deux circuits ouverts de longueurs différentes. Comme nous allons le voir, la qualité de standards de type circuit ouvert ("open") n'est pas critique en ce qui concerne la précision de l'amplitude de coefficient de réflexion et sa phase. On a besoin seulement d'un circuit hautement réflectif. Le critère de base dans le choix de standard est plutôt la simplicité de fabrication et dans le cas d'une structure micro-ruban, le circuit ouvert est clairement la plus simple structure. Pour les circuits de type guide coplanaires quelqu'un peut aussi choisir un court circuit.

L'analyseur de réseau est un instrument fonctionnant dans le domaine fréquentiel. La réponse du circuit dans le domaine temporel est calculée par l'intermédiaire de la transformée Fourier inverse de la réponse mesurée. Notre instrument possède une fonction qui permet de sélectionner une certaine fenêtre dans la réponse temporelle et effectuer la transformée Fourier directe seulement sur cette partie du signal. Cette opération de fenêtrage nous donne la possibilité de sélectionner seulement la réflexion

au niveau d'un seul connecteur (correspondant à la connexion sonde-substrat). Ayant cette possibilité, avec un nombre relativement restreint de standards, on peut extraire toute l'information nécessaire.

Regardons la figure 4.5. Le graphique montre de façon qualitative le principe de notre technique. En appliquant une fenêtre au niveau temporel on peut sélectionner seulement la réflexion correspondant au connecteur 1. Par la suite, on convertit le résultat dans le domaine fréquentiel et on extrait l'information.

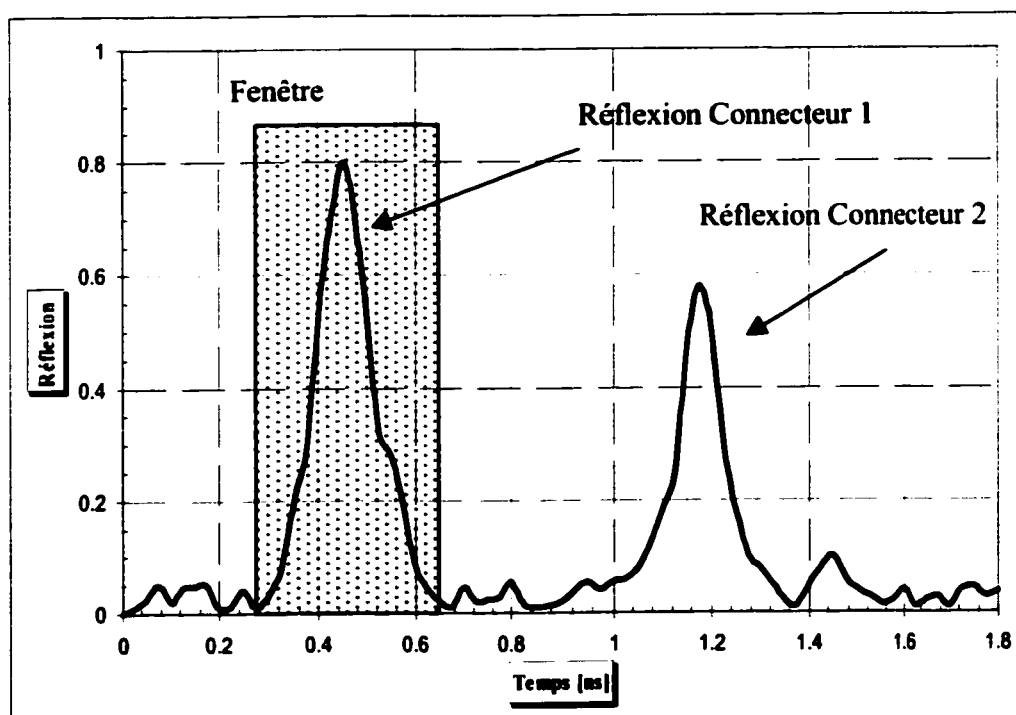


Figure 4.5 Le principe de fenêtrage.

Il faut dire que la figure 4.5 ne représente pas une mesure réelle. En pratique les transitions ne donnent pas des réflexions aussi importantes, mais le raisonnement est équivalent. En choisissant d'analyser une transition particulière, on considère que le

signal en dehors de la fenêtre sélectionnée est nul. Le signal résultant peut être manipulé dans le domaine spectral comme un ensemble de paramètres S habituels.

Nous présentons maintenant le traitement mathématique associé à cette méthode en suivant pas à pas de l'algorithme de calibration.

Regardons le schéma 4.6, où on retrouve le flux de signal pour la connexion du standard L1 (ligne "THRU"). Comme la représentation des matrices inconnues des transitions doit tenir compte de leur comportement passif et réciproque, les coefficients de transfert direct et inverse sont identiques ($A_{21}=A_{12}=A$).

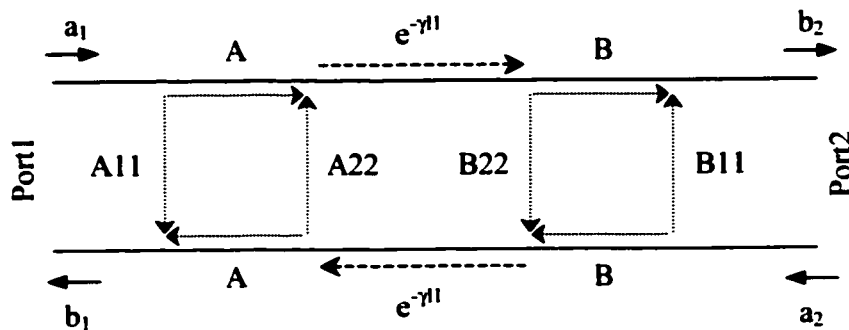


Figure 4.6 Cheminement du signal pour la connexion "THRU".

PAS 1 :

La ligne "THRU" est mesurée. Dans le domaine temporel on sélectionne seulement la réflexion au niveau de la transition 1. Dans ce cas :

$$S_{11A} = A_{11} , \quad 4.16$$

PAS 2 :

Dans la même connexion on sélectionne seulement la réflexion au niveau de la deuxième transition. On y retrouve :

$$S_{11B} = A_{21} \cdot \exp(-\gamma L_1) \cdot B_{22} \cdot \exp(-\gamma L_1) \cdot A_{12} = A^2 \cdot B_{22} \cdot \exp(-2\gamma L_1) , \quad 4.17$$

PAS 3 :

Dans la même connexion on applique le même procédé au niveau de S_{22} . Avec la fenêtre centrée sur la deuxième transition on obtient :

$$S_{22B} = B_{11} , \quad 4.18$$

PAS 4 :

La fenêtre est centrée maintenant sur la transition du port 1 :

$$S_{22A} = B^2 * A_{22} * \exp(-2\gamma L_1) , \quad 4.19$$

PAS 5 :

La dernière mesure sur la ligne "THRU" est en transmission :

$$S_{21L1} = S_{12L1} = A * B * \exp(-\gamma L_1) , \quad 4.20$$

Après les premières 5 étapes, on n'a pas suffisamment d'information pour trouver la valeur de la constante de propagation " γ " et on va mesurer les standards restants. Dans ce cas le port 1 de l'analyseur sera connecté au standard L2 et le port 2 au standard L3, respectivement. Pour une meilleure compréhension du procédé, on représente dans le schéma 4.7 le cheminement du signal correspondant.

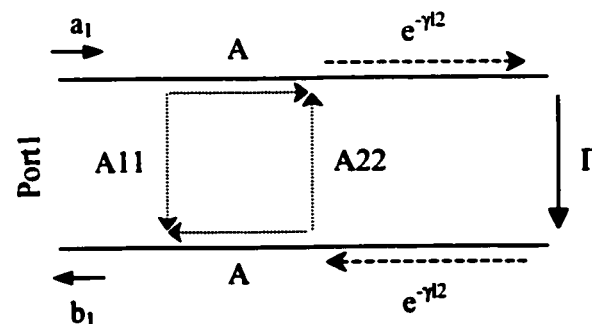


Figure 4.7 Cheminement du signal pour la connexion "OPEN".

PAS 6 :

On mesure le standard L2 ("OPEN") et la fenêtre est centrée sur la réflexion du circuit ouvert. On obtient :

$$S_{11L2} = A^2 * \Gamma * \exp(-2\gamma L_2) , \quad 4.21$$

PAS 7 :

Le même type de mesure dans le cas de la réflexion du circuit ouvert L3 au niveau de port 2 :

$$S_{22L3} = B^2 \cdot \Gamma \cdot \exp(-2\gamma L3) , \quad 4.22$$

Finalement il nous reste seulement à tourner le substrat de calibration et mesurer L3 au niveau de port 1 et L2 au niveau de port 2 :

PAS 8 :

$$S_{11L3} = A^2 \cdot \Gamma \cdot \exp(-2\gamma L3) , \quad 4.23$$

$$S_{22L2} = B^2 \cdot \Gamma \cdot \exp(-2\gamma L2) , \quad 4.24$$

Toutes ces mesures étant prises, on peut déterminer analytiquement les termes de correction, en l'occurrence les paramètres A, B, γ , Γ .

$$\gamma = \frac{1}{2 \cdot (L2 - L3)} \cdot \ln\left(\frac{1}{2} \cdot \left(\frac{S_{11L3}}{S_{11L2}} + \frac{S_{22L3}}{S_{22L2}}\right)\right) , \quad 4.25$$

$$\Gamma = \frac{\sqrt{S_{11L2} \cdot S_{22L2}}}{S_{21L1}} \cdot \exp(\gamma \cdot (2 \cdot L2 - L1)) , \quad 4.26$$

$$A_{21} = A_{12} = A = \sqrt{\frac{S_{11L2}}{\Gamma}} \cdot \exp(\gamma \cdot L2) , \quad 4.27$$

$$B_{21} = B_{12} = B = \sqrt{\frac{S_{22L2}}{\Gamma}} \cdot \exp(\gamma \cdot L2) , \quad 4.28$$

$$A_{22} = \frac{S_{22L1}}{B^2} \cdot \exp(2 \cdot \gamma \cdot L1) , \quad 4.29$$

$$B_{22} = \frac{S_{11L1}}{A^2} \cdot \exp(2 \cdot \gamma \cdot L1) , \quad 4.30$$

Cette méthode de calibration offre certains avantages au niveau de la précision et de la facilité d'implantation. Au premier, le nombre de standards utilisés est de seulement 3, et de plus ces standards peuvent être très facilement fabriqués sur n'importe quelle type de

ligne de transmission planaire. En deuxième lieu, la précision ne dépend pas d'une bonne connaissance des standards. Il est clair, suivant le développement mathématique présenté que les seules hypothèses nécessaires sont la répétabilité des connexions et le fait que les circuits réfléchis soient identiques. La première hypothèse est valable pour tout type de calibration et la deuxième est facilement respectée en pratique. Ces avantages nous ont convaincu d'implanter cette technique sous l'environnement de modélisation ICCAP™. En fait, la procédure est largement programmable, il suffit seulement de prendre les mesures associées aux trois connexions et les paramètres sont retrouvés au niveau du logiciel.

Cependant, soulignons que certains inconvénients limitent la plage d'applications de cette technique. Pour une bonne précision de la technique, la longueur physique de la ligne "THRU" doit être relativement longue pour permettre l'isolation de la réflexion de chaque connecteur. Cette condition est obligatoire pour éviter le chevauchement du signal réfléchi au port 1 et celui réfléchi au port 2. En pratique cette condition signifie que les dimensions des standards sont relativement grandes. Si on veut limiter le coût de fabrication des standards, (les standards doivent être fabriqués sur le même substrat que les circuits à mesurer) on préférera la méthode multi-lignes.

Un deuxième point faible provient du système de mesure et il n'est pas spécifique à cette technique de calibration. La qualité de la technique temporelle dépend de la précision numérique de la transformée de Fourier. Pour atteindre les niveaux satisfaisants, on doit utiliser une caractéristique de type passe-bas pour la réponse en temps ce qui oblige de choisir des fréquences multiples de 45 MHz à partir de cette valeur. Aussi, la fréquence maximale doit être très élevée (au moins 26.5 GHz) pour assurer une bonne résolution de la transformée de Fourier. Ces conditions sont parfois en contradiction avec les techniques de modélisation du logiciel ICCAP™. Il faut se souvenir que le but final de nos mesures reste de caractériser les composants MMIC. Une condition du logiciel nous oblige de garder la même plage de fréquence pour la calibration que pour la

procédure de modélisation. Dans ce cas, une plage de fréquences excessivement large augmente de façon exponentielle le temps de mesures, ce qui rend nos mesures impraticables.

Néanmoins, nous avons utilisé la technique d'analyse dans le domaine temporel pour obtenir l'information de départ pour les valeurs de capacités équivalentes simulant la connexion sonde-ligne. Pour les circuits mesurés en monture micro-ruban, la technique a donné des résultats similaires à ceux obtenus avec la technique multi-lignes. Dans ce cas, pour les circuits MHMIC sur alumine réalisés à PolyGrames, cette technique est facile à utiliser et elle peut être implantée comme une procédure de calibration habituelle de l'analyseur de réseaux, sans faire une calibration préalable au niveau des connexions.

2.3.8 Changement de Plan de Mesure pour les Structures Non-uniformes.

Les nouvelles composantes actives micro-ondes ont des dimensions de plus en plus petites. En fonction du critère de conception du transistor utilisé par les ingénieurs (faible figure de bruit, dissipation de chaleur) le dessin ("layout") physique est relativement non-uniforme. Il devient alors difficile pour l'opérateur du système de mesure de restreindre la structure des éléments parasites à une simple capacité et une ligne de transmission uniforme. Le même problème se pose dans le cas des composantes encapsulées, dont les caractéristiques du boîtier influencent de façon importante les performances globales de la composante.

Pour traiter ce type des structures, Cho et Burk [29] ont conçu une méthode pratique de calibration, qui a été utilisée parfois dans nos procédés de mesures et caractérisation. Il s'agit d'une technique que peut être généralisée facilement pour une classe très large de mesures sous-pointes. La figure 4.8 montre quatre structures de connexion qui doivent être disponibles pour la calibration. Même si les structures représentent une connexion pour un transistor bipolaire en monture micro-ruban, il ne s'agit pas du seul cas possible,

car d'autres types de connexion (soit en dessin coplanaire soit ligne à fente etc.) plus ou moins uniformes peuvent être considérés de la même manière.

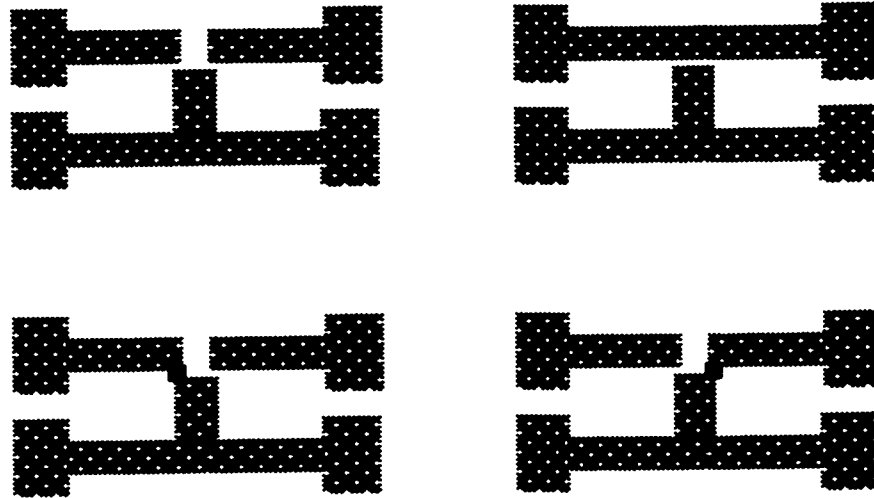


Figure 4.8 Dessin des structures de calibration.

La procédure utilise les paramètres Y ou Z pour représenter les éléments parasites, selon le type de connexion (série et parallèle). Un modèle électrique associé avec la monture du dispositif DST est présenté dans la figure 4.9. G_i représentent les éléments en connexion parallèle dus aux contacts sonde-substrat et Z_i les éléments série. Une telle représentation est possible aussi pour la modélisation des boîtiers. Si on dispose des boîtiers avec des connexions en circuit ouvert, court circuit et connexion directe, on peut obtenir un solide algorithme de caractérisation des éléments parasites.

La technique suit trois étapes, correspondant aux connexions consécutives des trois "standards". Pour la connexion "OPEN" on a:

$$G_1 = Y_{11OPEN} + Y_{12OPEN} ,$$

$$G_2 = Y_{22OPEN} + Y_{12OPEN} ,$$

$$G_3 = -Y_{12OPEN}$$

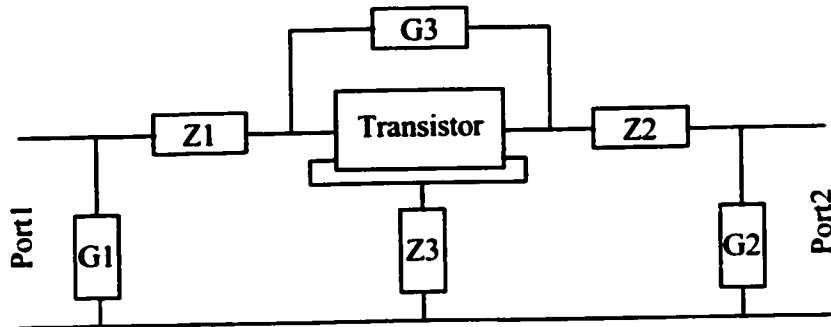


Figure 4.9 Modèle électrique de la monture.

En trouvant les valeurs G_i on peut soustraire les valeurs G_1 et G_2 par une simple procédure :

$$\begin{aligned} Y_{11} &= Y_{11m} - G_1, \\ Y_{22} &= Y_{22m} - G_2, \end{aligned} \quad 4.32$$

Dans l'équation 4.32 Y_{xxm} est la valeur mesurée au niveau de connexions sonde-substrat. Ensuite on doit extraire les éléments série Z_i . Pour cela, il suffit de mesurer les paramètres Y en connexion "THRU" et "SHORT". La figure 4.10 montre les schémas électriques pour l'intégralité de la procédure.

Les valeurs des éléments série peuvent être calculées par :

$$\begin{aligned} Z_1 &= \frac{1/Y_{12thru} + 1/Y_{11short1} - 1/Y_{22short2}}{2} \\ Z_2 &= \frac{1/Y_{12thru} - 1/Y_{11short1} + 1/Y_{22short2}}{2}, \\ Z_1 &= \frac{-1/Y_{12thru} + 1/Y_{11short1} + 1/Y_{22short2}}{2} \end{aligned} \quad 4.33$$

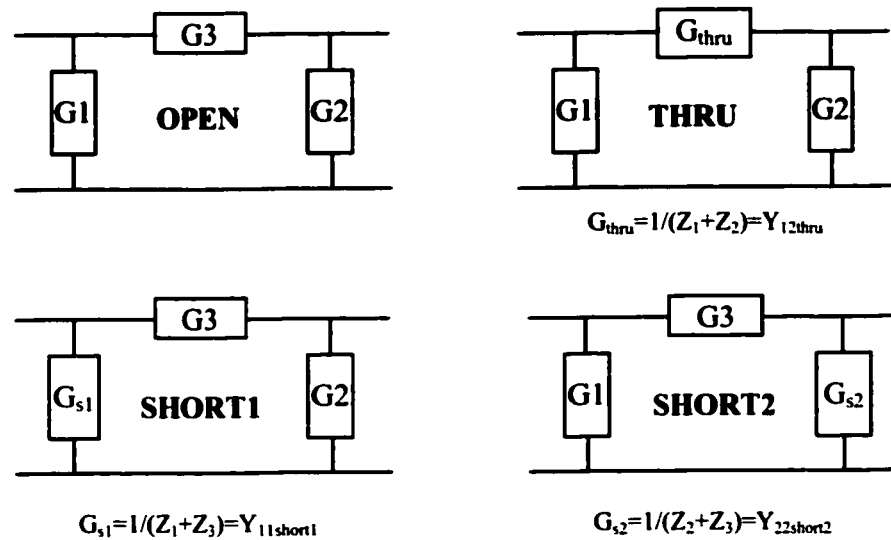


Figure 4.10 Schéma électrique des connexions.

Avec les valeurs des éléments "parasites" série déterminées, on peut trouver les valeurs corrigées des paramètres Z au niveau du plan du dispositif sous test :

$$\begin{aligned}
 Z_{11s} &= Z_{11} - Z_1 - Z_3 & Z_{12s} &= Z_{12} - Z_3 \\
 Z_{21s} &= Z_{21} - Z_3 & Z_{22s} &= Z_{22} - Z_2 - Z_3
 \end{aligned}
 \tag{4.34}$$

Le dernier pas de la calibration reste l'extraction de la valeur G_3 pour obtenir les valeurs corrigées des paramètres Y du transistor :

$$\begin{aligned}
 Y_{11} &= Y_{11s} - G_3 & Y_{21} &= Y_{21s} + G_3 \\
 Y_{12} &= Y_{12s} + G_3 & Y_{22} &= Y_{22s} - G_3
 \end{aligned}
 \tag{4.35}$$

La technique de Cho et Burk est facilement intégrable dans un processus de mesure automatisé. Elle permet une adaptabilité accrue au niveau du dessin des composantes et éventuellement elle peut servir dans la caractérisation des boîtiers pour les composantes ou des techniques de connexion comme les systèmes "flip-chip". Aussi, dans le cas de

composantes MMIC, elle peut nous apporter des informations supplémentaires sur les capacités parasites du dessin, information nécessaire pour l'optimisation de la fréquence de coupure du transistor.

Toutes les mesures effectuées pour la modélisation ont fait l'usage intensif d'une des trois méthodes présentées ou d'une combinaison de la méthode multi-lignes et des structures associées avec la méthode Cho. Avant chaque mesure de dispositif, la calibration a été testée pour valider les processus de correction. La figure 4.11 montre l'amplitude du coefficient de réflexion d'une structure de capacité MMIC mais avec les deux plans métalliques en court-circuit. Comme on peut observer le plan de mesure est exactement au niveau de la composante.

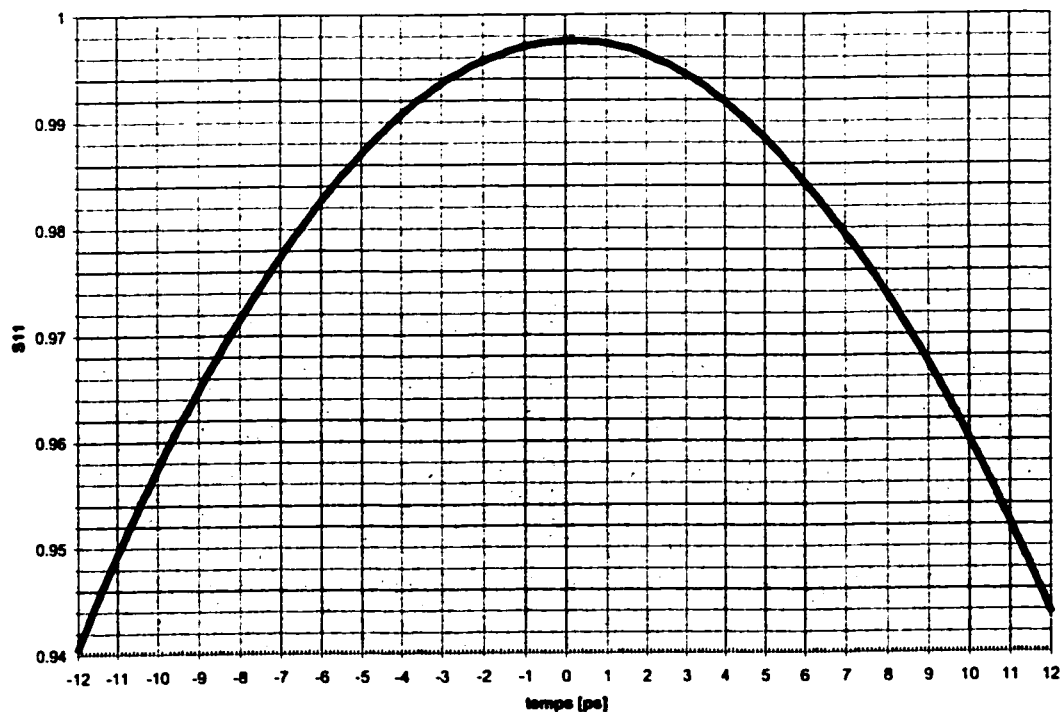


Figure 4.11 Vérification de la correction.

Le dessin est représentatif pour les mesures sous-pointes dans le sens où il représente le plot de connexion, une ligne d'accès de 40 μm en longueur et la structure non-uniforme

du court-circuit. Le graphique présente les valeurs après la correction et le changement du plan de mesure. Puisque les structures à modéliser avaient un dessin identique au niveau des connexions on voit qu'on a obtenu une calibration de qualité

4.2 Fabrication des Circuits MMIC

Un grand effort a été consacré au développement d'une unité de fabrication de composantes MMIC. Nous allons consacrer cette section aux procédés développés, avec les détails spécifiques et une vue générale sur les composantes MMIC utilisées à grande échelle.

Au début de notre travail, le laboratoire PolyGrames possédait une facilité de fabrication de circuit hybrides ("MHMIC"), qui nous permettait d'intégrer sur un substrat d'Alumine les structures planaires de lignes de transmission et des résistances déposées en Ti. Pour obtenir un circuit complet on ajoutait des composantes actives discrètes et capacités (puce ou en boîtier). En fonction du type de circuit les connexions électriques sont réalisées par simple collage, par fil soudé ("wire-bond"), par ruban etc. Malgré la robustesse de cette technique et son faible coût, elle montrait des limitations en ce qui concerne nos objectifs. La résolution maximale est d'environ 25,4 μm (1 mils) dû au processus photographique qui utilise une imprimante graphique habituelle comme source pour la fabrication de masque. Le manque de capacités intégrées et les "parasites" des structures d'interconnexions rendait le processus moins flexible pour le concepteur. Mais, plus important, c'était l'impossibilité d'avoir une reproductibilité accrue, étant donné que les transistors utilisés sont des composantes discrètes.

Dans ces conditions nous avons commencé à travailler pour développer un vrai procédé MMIC capable de fabriquer des circuits qui fonctionnent à des hautes fréquences. Pour obtenir cela on a besoin des deux conditions "technologiques". La première condition concerne le transistor lui-même. Pour un fonctionnement en ondes millimétriques on doit utiliser des semiconducteurs à haute mobilité. En principe, les

transistors de type PHEMT à base de GaAs ou InP sont les meilleurs candidats pour cette application et représentent un quasi-standard en industrie. Nous avons décidé de travailler avec ce type de matériaux, cependant comme l'épitaxie de la structure pseudo-morphique demande des outils extrêmement coûteux et une expertise spécifique, les circuits MMIC que nous avons fabriqués l'ont été sur des gaufres dopées disponibles commercialement.

En second lieu, il s'agissait d'avoir une précision suffisante au niveau du masque du dessin pour obtenir les structures à grilles très courtes (~200 nm). Dans ce cas on a développé un partenariat avec le laboratoire GSM de l'Université de Sherbrooke où on peut faire de la lithographie par faisceau d'électrons ("e-beam lithography"). Pour augmenter la vitesse de fabrication des circuits d'essai, notre laboratoire a acquis en 1999 un système de photo-plotter à rayon laser, capable d'une résolution de 2 microns. Dans ce cas, une grande partie des étapes de fabrication se font à l'aide d'une photolithographie classique avec les masques produits par le photo-plotter, la lithographie par faisceau d'électrons n'étant utilisée que dans les étapes qui requièrent une grande précision.

Les figures 4.12-4.19 présentent les principales étapes dans la fabrication d'un circuit MMIC complet.



Figure 4.12 Structure PHEMT sur GaAs, épitaxie par jet moléculaire

La structure active est déposée sur un substrat semi-isolant de GaAs d'une épaisseur de 600 μm et une permittivité relative 12,9. La première étape est la définition des zones actives par la gravure des couches dopées (fig. 4.13). La gravure est réalisée par une solution acide $\text{H}_3\text{PO}_4/\text{H}_2\text{O}_2/\text{H}_2\text{O}$ (5:2:200) dans des bains à température contrôlée (20°C)

pour une durée d'environ 130 secondes. Cette étape nous permet d'isoler les transistors et les résistances enterrées.

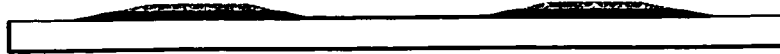


Figure 4.13 Gravure MESA (Définition Zones Actives)

Chaque zone active est ensuite utilisée pour la fabrication des transistors. Dans la deuxième étape, les contacts ohmiques sont fabriqués par une évaporation successive des différents métaux (Ni:Ge:Au:Ni:Au 10/30/50/10/50 nm) suivi par le recuit de 1 min. aux environs de 365°C dans le four à recuit rapide, pour stabiliser la structure (fig. 4.14).

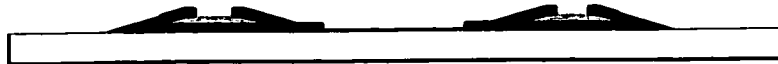


Figure 4.14 Contacts Ohmiques

Une gravure sélective est réalisée pour enlever la couche fortement dopée (une solution acide citrique/citrate de potassium et peroxyde (5:5:2) et la formation du contact Schottky au niveau de la grille.



Figure 4.15 Définition de la grille par lithographie à faisceau d'électrons

Une écriture directe (fig. 4.15) dans un système à deux couches de résine à l'aide du microscope électronique nous permet l'obtention d'une grille de type champignon (pour diminuer la résistance série).



Figure 4.16 Première couche de métal d'interconnexions

Par la suite, une couche métallique (Cr/Au 10/500nm) est déposée par évaporation suivie par le soulèvement. La mince couche de Cr est nécessaire pour l'adhérence de la couche d'or au substrat cristallin. Cette couche représente le premier niveau d'interconnexions qui est utilisé pour la définition des lignes de transmission et l'établissement des connexions électriques. Pour obtenir une bonne qualité du processus de soulèvement un procédé utilisant une résine positive a été mis au point. Ceci est nécessaire en raison de l'épaisseur relativement grande du métal évaporé (fig. 4.16).



Figure 4.17 Résistance NiCr déposée par évaporation

Le procédé permet l'intégration des résistances déposées en NiCr. La photolithographie est utilisée pour la définition du patron des résistances. Ce type de résistance présente certains avantages par rapport aux variantes enterrées (utilisant le canal conducteur). Elles sont stables par rapport à l'intensité du champ électrique (linéaires), peuvent atteindre des valeurs faibles (~5-10 ohms) et surtout elles présentent un coefficient de variation par rapport à la température largement inférieur. En pratique, elles peuvent être ajustées après la fabrication des masques par un procédé laser. La fabrication nécessite un niveau de masque supplémentaire et l'évaporation d'une fine couche de NiCr (80-120 nm) en fonction de la valeur de la résistance par surface désirée. Comme une alternative d'amélioration future, le dépôt de la couche résistive pourrait se faire par pulvérisation cathodique ("sputtering"). Aussi une résistivité accrue de l'alliage NiCr peut nous

permettre une valeur d'épaisseur un peu plus élevée pour un meilleur contact avec la couche de métal.

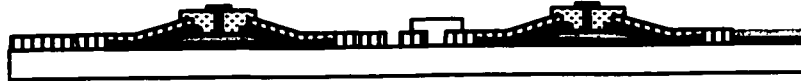


Figure 4.18 Passivation Si_3N_4

La déposition du nitrure de silicium par un procédé au plasma (réacteur PECVD) sert comme moyen de passivation pour augmenter la fiabilité des circuits et en même temps elle définit les zones du diélectrique pour les capacités de type MIM ("métal-isolant-métal"). Les valeurs habituelles d'épaisseur sont d'environ 120-180 nm pour la couche de nitrure. Le contrôle de l'épaisseur est nécessaire pour l'obtention d'une bonne tolérance au niveau de la valeur de la capacité surfacique. À l'aide d'un ellipsomètre on peut mesurer l'indice de réfraction (qualité du dépôt) et l'épaisseur.



Figure 4.19 Métal 2 interconnexions type pont-à-air

Une deuxième couche de métal nous permet la réalisation du plan métallique supérieur des capacités MIM et les interconnexions de deuxième niveau comme les ponts-à-air (figure 4.19). En fait, la structure métallique du deuxième niveau est constituée pratiquement de trois masques consécutifs : le masque du métal 2 et sa mise en œuvre identique avec le procédé pour le métal 1 (CrAu), le masques des piliers de ponts et le masque pour les ponts-à-air. Ce dernier sert aussi de masque pour le plot de connexions externes. La définition de cette couche est faite en deux étapes, mais avec un seul masque de niveau. Une couche initiale (métal intermédiaire), déposée par évaporation

(CrAu 5/25 nm) ou pulvérisation cathodique (AuPd pour 3 min) et une deuxième réalisée par électroplaquage. Avec cette technique on peut augmenter l'épaisseur équivalente des lignes métallique donc une meilleure capacité de transport du courant et plus faibles pertes dans les lignes de transmission. Notre procédé a été développé avec le cuivre pour diminuer le coût global de circuit, mais on peut facilement le modifier pour l'or si le problème d'oxydation de cuivre s'avère trop important. Avec une telle technique on obtient pour des échantillons de test une épaisseur cumulative d'environ 2-2,5 μm de métal, suffisante pour la plupart des applications. Dans le futur, pour des applications nécessitant un courant plus fort on peut augmenter cette épaisseur vers 4-6 μm .

La difficulté de fabrication des trous métallisés nous a conduit à préférer les structures coplanaires aux structures micro-ruban. Les dispositifs micro-ruban sont mieux connus et les modèles sont largement disponibles dans les outils CAO couramment utilisés. Pour cette raison, dans l'industrie les techniques micro-ruban dominent le marché actuel. Cependant, pour réaliser un circuit micro-ruban deux étapes supplémentaires sont nécessaires : l'aminçissement du substrat aux valeurs des 50-100 μm et la fabrication des trous métallisés par attaque chimique suivie de l'électroplaquage. La première étape rend les gaufre très fragile avec une grande possibilité de cassure, ce qui entraîne une augmentation du coût global de fabrication. Cette technique introduit aussi des variations de performance en fonction de la précision du contrôle d'épaisseur. La deuxième étape, l'alignement des trous, réalisée du côté opposé de la structure planaire pose parfois des contraintes importantes au niveau de l'espace utilisé. On a donc opté pour l'option des structures CPW pour leur simplicité de fabrication et leur faible coût. Soulignons que dans le milieu universitaire [11] on considère de plus en plus le dessin en structure coplanaire en raison d'une diminution de l'espace utilisé et d'une meilleure isolation électrique entre les différents étages du circuit.

La mise en œuvre de chaque étape de ce procédé a nécessité des grands efforts au niveau du développement des techniques de fabrication.

La plupart des "recettes" de fabrication ont été mises aux points spécifiquement pour les outils disponibles dans le laboratoire



Figure 4.20 Système de fabrication de masques à balayage électro-optique

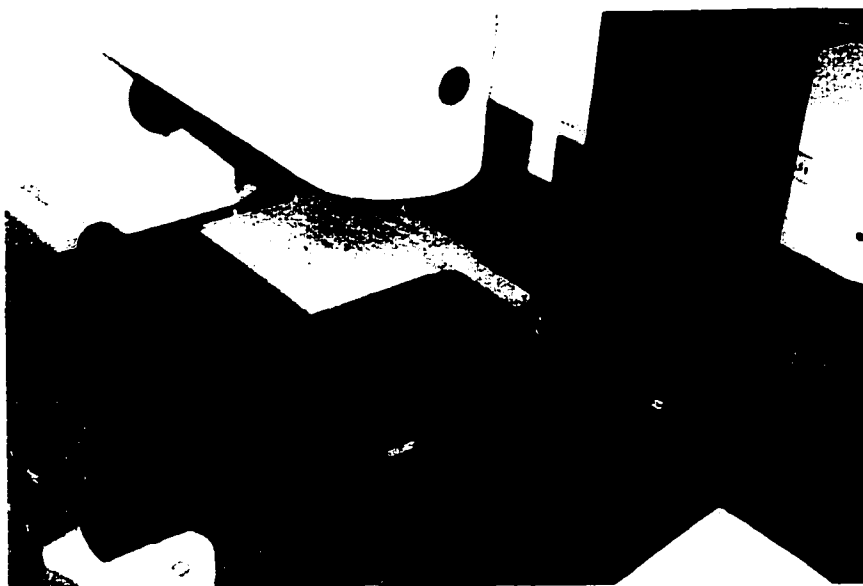


Figure 4.21 Détail du porte-échantillon.

GSM de Sherbrooke. Une autre partie du travail a été consacrée au développement de la technique de fabrication de masques pour la photolithographie à l'aide du photoplotter. Les masques fabriqués à PolyGrames nous permettent de diminuer la durée du cycle de fabrication. Les figures 4.20 et 4.21 montrent le système à laser utilisé pour la fabrication de masques.

Le principe de fonctionnement du photoplotter est relativement simple. On contrôle le balayage (en x et y) d'un flux lumineux laser (longueur d'onde ~ 457 nm, bleu) au moyen de deux modulateurs acousto-optiques. Le contrôle de l'angle de déflexion est réalisé par un synthétiseur de fréquence, et il permet l'exposition d'une surface rectangulaire ($\sim 205/205$ mm). Un porte-échantillon peut se déplacer par rapport au centre d'exposition, ce qui permet de réaliser des masques d'une grandeur maximale de 15 cm. Le système est contrôlé par ordinateur (voir figure 4.20) et permet une résolution maximale de $2\text{ }\mu\text{m}$ entre deux structures et des lignes d'une largeur minimale de $\sim 4\text{ }\mu\text{m}$.

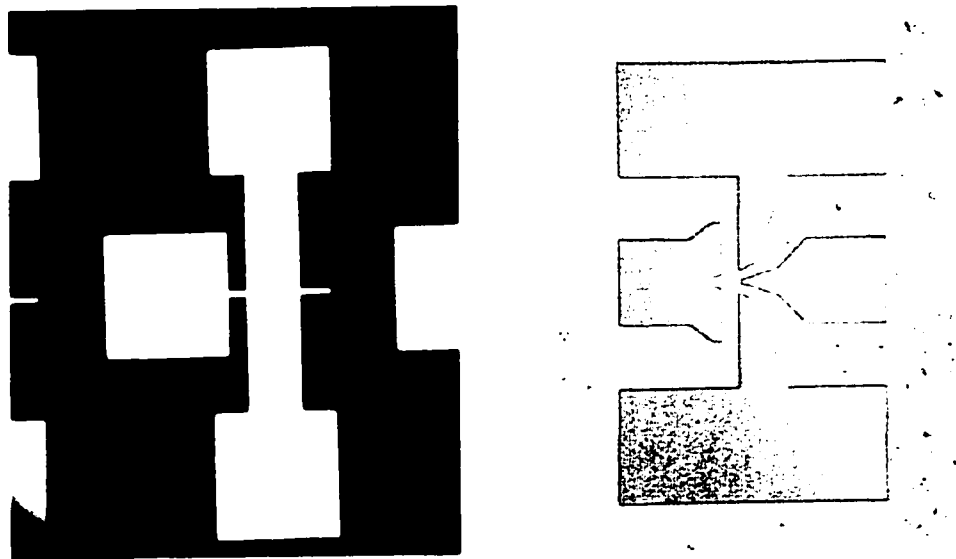


Figure 4.22 Masque pour lithographie

Dans la figure 4.22, on voit deux portions de masque fabriqué dans notre laboratoire. Un a servi comme structure de test et l'autre est un détail exact du masque de niveau métal 1 d'un transistor PHEMT en architecture "T".

Les procédés de fabrication des masques ont été mis au point dans le laboratoire PolyGrames pour des structures de 2,54 x 2,54 cm en chrome. Nous pouvons fabriquer aussi des dimensions plus grandes, mais au cours du notre travail nous avons utilisé seulement cette taille de masque.

CHAPITRE 5 MODÉLISATION DES COMPOSANTES MMIC

Après avoir vu les détails de la procédure de mesure et de fabrication pour les composantes micro-ondes, nous allons continuer la présentation de nos travaux avec le développement des algorithmes d'extraction de modèles. Indépendamment de la technique de fabrication, le besoin du concepteur n'est pas seulement d'avoir un dispositif performant mais aussi de concevoir un dispositif dont le comportement électrique est bien modélisé, dans un format compatible avec les outils de CAO. Pour obtenir de bons modèles deux conditions doivent être réunies. Premièrement, obtenir une précision de mesure suffisante, ceci passe par des techniques de calibration performantes. Deuxièmement, il reste à résoudre le problème de l'extraction des paramètres. Pour atteindre ce double objectif, un système de mesure intégré a été mis au point. L'ensemble des algorithmes de modélisation présentés dans ce chapitre ont été développés sur une plate-forme contrôlée par le logiciel ICCAP™. Le logiciel a un double rôle : il sert en tant que centre de contrôle pour les équipements de mesure (il gère le plan de mesure et assure le synchronisme entre les différents appareils) et il possède aussi un puissant outil de modélisation et de simulation, puisqu'il est intégré au système CAO de la compagnie Agilent (en fonction de la tâche spécifique il utilise le simulateur MDS ou ADS).

La figure 5.1 présente une vue d'ensemble du système de mesure sous pointes, dédié aux composantes MMIC. Une station de mesure manuelle de type CASCADE 9000™ est connectée à l'analyseur de réseau HP8510C et à un système automatisé de polarisation DC (HP4142B). Un microscope optique est relié à un système vidéo, permettant à l'utilisateur un bon contrôle de l'opération de la station. Le logiciel ICCAP embarqué sur une station de travail prend le contrôle externe des appareils via GPIB. Comme le logiciel possède deux types de moyens d'automatisation (en mode MACRO par des commandes internes utilisant un langage similaire au BASIC ou externe par des routines

en langage C ajoutées au fichier exécutable), on obtient une procédure d'extraction sans intervention de l'utilisateur.



Figure 5.1 Système de mesure sous-pointes.

En procédant ainsi, on peut faire l'usage de certains types de modèles à base de données, modèles qui nécessitent un grand nombre de mesures (le lecteur va trouver les détails dans la section dédiée à la modélisation du transistor).

La majorité des algorithmes développés au long de nos travaux permettent d'extraire de façon automatique, les paramètres du modèle ou, dans certains cas (pour les transistors), directement le modèle sous un format compatible avec le logiciel ADS. Pour les

inductances, pour lesquelles il n'existe pas de modèle analytique, on doit aussi regarder les résultats d'une analyse numérique du champ électromagnétique ("full-wave") à l'aide du logiciel MOMENTUM. Dans ce cas, l'optimisation finale est plus facile.

Il faut souligner que les efforts ont été focalisés plutôt sur la caractérisation du transistor et nous avons construit un algorithme complet d'extraction. Comme la fabrication du transistor suppose des contraintes plus élevées que celles des composantes passives, notre effort initial a été dirigé, tout au début, vers la fabrication de composantes passives. Il faut dire que, on y trouve de plus en plus des librairies commerciales pour la simulation de composantes MMIC passives, comme un ajout au noyau des outils CAO. N'ayant pas eu à notre disposition une librairie comme "COPLAN[®]", il a fallu développer nos propres modèles.

5.1 Capacités Métal-Isolant-Métal (MIM).

Composante souvent utilisée dans les circuits intégrés, les capacités MIM sont rencontrées dans diverses applications comme : l'adaptation, le filtrage, les coupleurs directionnels, couplage AC etc. Dans notre cas, des capacités à base de nitrure de silicium (voir chapitre précédent) ont été fabriquées. La constante diélectrique d'environ 6,8 permet une capacité surfacique entre 380 et 420 F/mm² en fonction de l'épaisseur de la couche déposée (140-160nm). Cette valeur nous permet d'obtenir des valeurs entre 0.18 fF (~20x20 um) et 20 pF (~200x200 um). Des valeurs plus grandes sont évitées par les concepteurs en raison de la grande surface occupée.

Dans la figure 5.2 on voit une structure MIM (en connexion série) observé sous microscope électronique. Habituellement, le Métal 1 forme le premier plan métallique de la capacité et le Métal 2 se pose sur la couche de diélectrique (format "sandwich"). Parfois (ce n'est pas le cas sur notre figure) même la deuxième couche métallique est réalisée par électroplaquage pour diminuer les pertes ohmiques. Dans notre cas, on voit que le pont-à-air, qui sert de connexion série est plaqué (sa rugosité est plus prononcée que celle des couches évaporées).

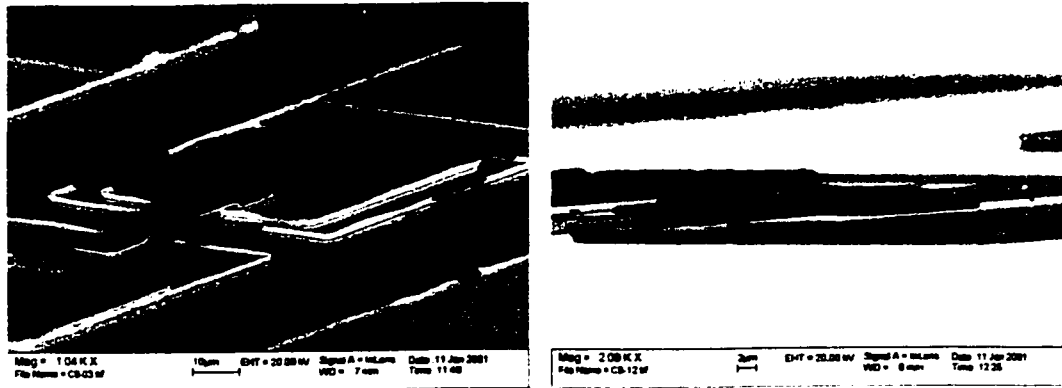


Figure 5.2 Capacité MIM en vue de haut et vue latérale.

Comme c'est une composante très utilisée, une grande attention a été portée pour trouver un modèle adéquat et on peut trouver dans la littérature plusieurs approches pour la simulation des capacités MIM. Habituellement on préfère les modèles relativement simples surtout dans le cas des capacités faibles (petite surface). Dans ce cas les dimensions physiques sont largement inférieures à la longueur d'onde même pour les ondes millimétriques et les phénomènes distribués sont peu importants. Le schéma 5.3 montre deux variantes que nous avons adoptées pour notre modèle.

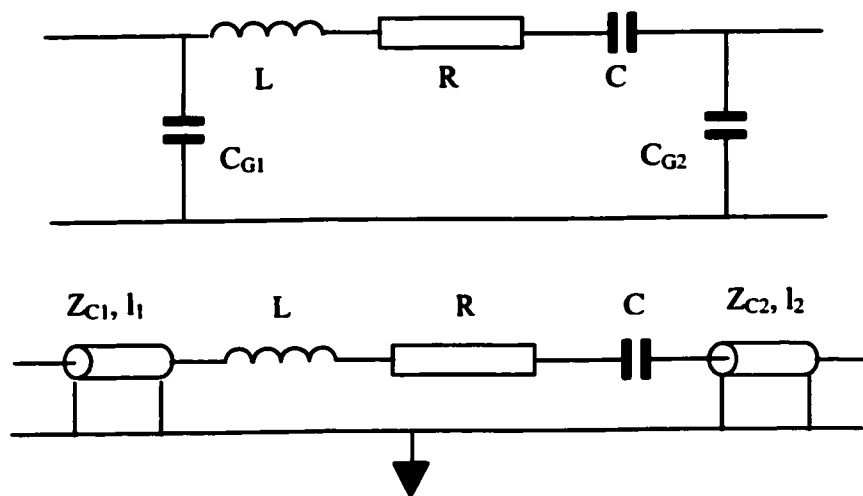


Figure 5.3 Modèle de la capacité

Dans les deux cas, l'inductance et la résistance représentent l'influence du pont-à-air associé à la deuxième couche métallique du MIM. Pour la première variante, deux capacités parasites sont ajoutées pour compléter le modèle coplanaire. Normalement la résistance doit avoir une valeur dépendante de la fréquence en raison des pertes ohmiques variables avec la fréquence. En pratique une modélisation de ce genre ne s'est pas avérée nécessaire et le modèle fonctionne bien avec une valeur constante. La deuxième variante tiens compte un peu plus des phénomènes distribués pour une meilleure correspondance entre les valeurs mesurées et simulées de la phase de paramètres S. La figure 5.3 présente la variante de connexion série. Dans le cas d'une connexion vers le plan de masse, un port va obligatoirement disparaître, ainsi que les éléments associés avec ce deuxième port.

Pour la fabrication et l'extraction du modèle, on a conçu un masque des capacités de plusieurs dimensions et de plusieurs formes. Le tableau 5.1 présente les mesures de la capacité pour un échantillon de test. Les valeurs sont extraites dans une routine ICCAP conçue dans ce but. La capacité est mesurée en basse fréquence (45 MHz-2GHz) et une procédure d'optimisation est mise en œuvre pour le calcul des autres paramètres. L'optimisation est poursuivie sur une bande étendue des fréquences (1-30 GHz). Le code définit les dimensions (longueur x largeur) exprimées en microns.

CS20x20	CS40x40	CS60x60	CS80x80	CS100x100	CS120x120	CODE
0.183862	0.674107	1.50607	2.61789	4.10775	5.80411	Valeur [pF]
CS20x20	CS30x30	CS50x50	CS70x70	CS90x90	CS110x110	Code
0.177483	xxx	1.07036	2.07744	3.35274	4.97887	xx = défaut
CS40x50	CS30x40	CS30x50	CS40x60	CS40x80	CS40x70	Code
0.824715	0.527075	0.624931	1.01898	1.33839	1.16923	Valeur [pF]

Tableau 5.1 Masque de capacités mesurées.

Sur un échantillon il y a 4 masques des capacités identiques. Pour les capacités avec les dimensions identiques on calcule la moyenne sur l'ensemble et les valeurs pour chaque moyenne sont considérées dans un système surdéterminé pour obtenir la capacité surfacique. Une bonne corrélation a été trouvée entre les résultats sur différents échantillons et on peut estimer une variation d'environ 2% pour les capacités fabriquées en même temps et meilleure que 5% entre deux échantillons fabriqués en différé.

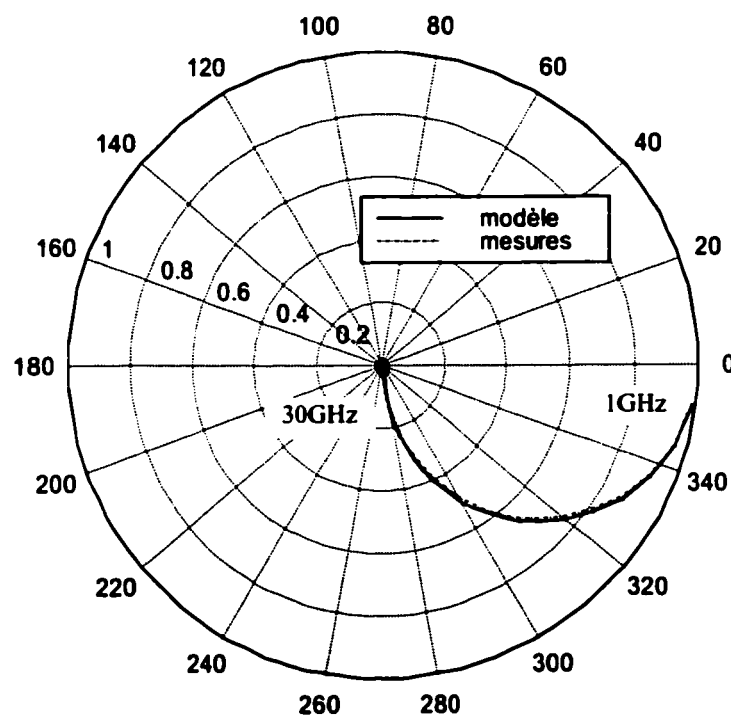


Figure 5.4 Comparaison entre les valeurs mesurées et modélisées.

La figure 5.4 montre la comparaison entre le coefficient de réflexion mesuré et la valeur simulée à l'aide de notre modèle. Une excellente correspondance a été trouvée. Pour toutes les capacités la valeur d'inductance série parasite est maintenue constante et elle est d'environ 10 pH. La résistance série a une valeur de bases fréquence d'environ 0,12 Ω . Comme on peut l'observer, le coefficient de réflexion montre un excellent comportement en fréquence. La fréquence de résonance est largement au-delà de la

fréquence supérieure de notre système de mesure. Ceci montre que nous en obtenons ainsi un bon facteur de qualité pour les capacités MIM fabriquées. La tension de claquage est supérieure à 12 V (la valeur maximale mesurable avec nos dispositifs), ce qui est suffisant pour notre procédé, en considérant que les dispositifs pHEMT fonctionnent à des tensions faibles (1,2-5 V).

5.2 Les Inductances Spirales.

Les inductances intégrées, une autre composante nécessaire aux applications de filtrage ou d'adaptation, sont souvent rencontrées dans les circuits MMIC. Du point de vue réalisation elles sont relativement simples à fabriquer, utilisant les niveaux de masque métal 1, métal 2, piliers du pont et pont-à-air. La figure 5.5 montre une version surélevée pour un meilleur comportement en fréquence.

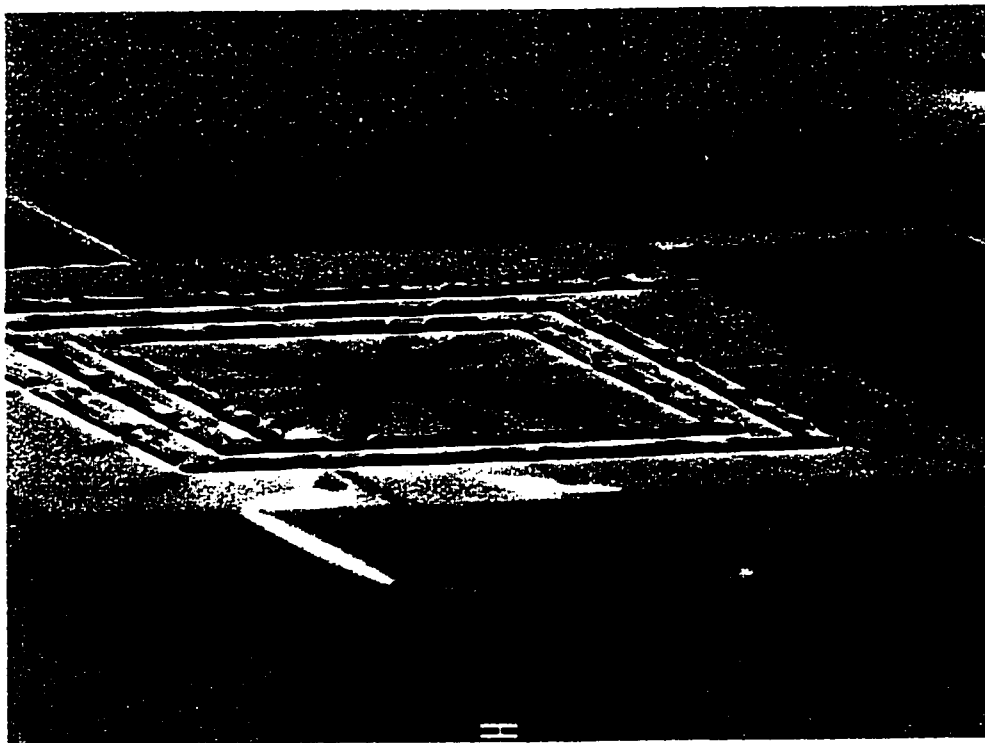


Figure 5.5 Inductance surélevée

Contrairement aux capacités MIM les modèles pour les inductances sont moins précis et largement dépendant de la plage de fréquence. Les simulateurs comme ADS utilisent un modèle de synthèse (en fonction de paramètres physiques) mais il est applicable seulement en version micro-ruban et il est valable seulement jusqu'à 4 GHz. De plus, la plage de valeurs des dimensions est spécifique aux circuits hybrides donc la précision sur des variantes MMIC (avec des dimensions plus faibles) est très approximative.

Pour une réalisation intégrée le concepteur doit respecter certaines contraintes. En principe, les inductances remplacent les circuits résonnants, qui avec leurs lignes de transmission sont trop "gourmandes" en espace (or l'espace est le paramètre principal de calcul du prix d'un circuit MMIC). Cette possibilité de diminution de la taille doit respecter les conditions électriques DC. En dehors de leur rôle dans les applications de filtrage, les inductances doivent aussi être utilisées comme chemin d'alimentation DC pour le courant de polarisation. Par rapport au niveau de puissance RF ces valeurs de courant peuvent se montrer critiques. Pour une analyse qualitative, l'épaisseur cumulée de notre combinaison métal² et pont-à-air est d'environ 1,5 microns. Pour une telle valeur le courant critique est d'environ 4 mA/um. Dans ce cas une inductance ayant une largeur de ligne de 12 microns peut transporter un courant maximal de 48 mA. Parfois cette valeur est trop faible et le concepteur doit augmenter la largeur de la ligne; mais cette opération peut abaisser significativement la fréquence de résonance de l'inductance en raison de la capacité de couplage entre les lignes parallèles.

Couramment, en raison de ces difficultés de modélisation, les concepteurs [11,70] préfèrent utiliser une base de données utilisant des valeurs mesurées. On construit et mesure une large classe de dispositifs avec des diamètres, les largeurs de ligne, des nombres de tours et des espacements entre les lignes différents. Cette approche est aussi validée par le fait que les ports d'accès ont des positions fixes en pratique. Si on considère (voir la figure 5.5) que le port numéro 1 est en position Est, le port 2 doit être en position SUD, NORD ou OUEST. Pour une valeur fixe de la largeur de la ligne

conductrice et de l'espacement entre deux lignes adjacentes, on peut avoir des valeurs discrètes d'inductances en fonction de nombre de tours ($1\frac{1}{4}$, $1\frac{1}{2}$, $1\frac{3}{4}$, $2\frac{1}{4}$ etc.). Pour obtenir une variation quasi-continue de la valeur de l'inductance, on préfère modifier le diamètre intérieur du premier tour. Avec un ensemble d'inductances mesurées on a un choix pour la modélisation : soit on prend un fichier correspondant aux paramètres S de chaque dispositif et le simulateur va interpoler entre les valeurs mesurées, soit on construit un modèle électrique équivalent basé sur des éléments concentrés et on considère seulement les valeurs de ces éléments. La deuxième approche est plus performante au niveau de la mémoire utilisée et elle permet une orientation rapide pour le concepteur au niveau du choix des valeurs. La première approche permet d'obtenir une meilleure précision absolue (valeurs mesurées!), et elle donne un temps de simulation très court.

Notre algorithme de modélisation utilise la deuxième variante avec un circuit équivalent présenté dans la figure 5.6.

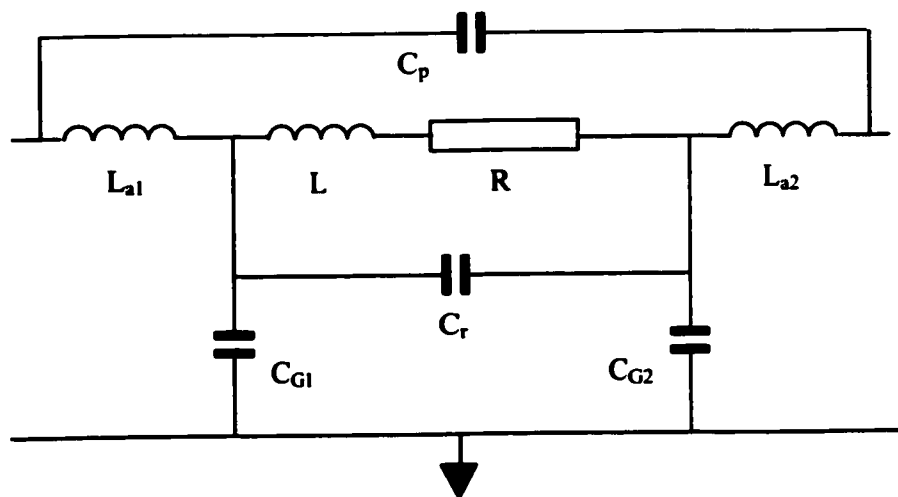


Figure 5.6 Modèle électrique de l'inductance

En dehors des valeurs L et R qui représente l'inductance de basses fréquence et la résistance série, le circuit équivalent possède d'autres composantes pour faciliter la

modélisation en haute fréquence. Les deux inductances série, modélisent l'accès à la structure en spirale réalisée par un pont-à-air d'un côté et une ligne mince pour l'autre port. La capacité C_p est la capacité parasite causée par le dessin physique de la structure (voir figure 5.7) et les capacités C_G et C_r représentent les phénomènes de couplage par rapport au plan de masse, respectivement entre les lignes adjacentes. La figure 5.7 montre un masque pour une inductance spirale d'un diamètre extérieur de 224 microns à $3\frac{1}{2}$ tours. La ligne centrale a une largeur de $12\text{ }\mu\text{m}$ et un espacement de $8\text{ }\mu\text{m}$ entre deux lignes adjacentes.

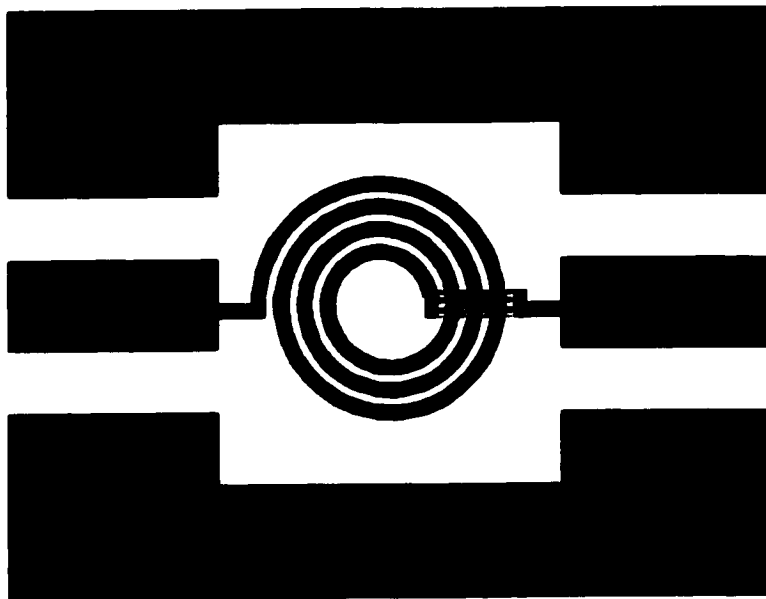


Figure 5.7 Masque pour une inductance spirale.

Du point de vue de la simulation, le circuit de la figure 5.6 est paramétré, et pour chaque valeur d'inductance ont fait correspondre une valeur pour chaque élément du circuit. Une approche plus poussée sera de trouver une courbe qui relie les valeurs de chaque élément électrique en fonction des dimensions pour être capable de faire une optimisation avec une telle composante. Jusqu'à date, à notre connaissance, aucune caractérisation de ce type n'est disponible pour une large plage de fréquence. Avec le modèle présenté on a

comparé les résultats de mesures et de simulations. Pour une connexion série, comme celle de la figure 5.7 on voit dans la figure 5.8 les résultats pour le coefficient de réflexion. Dans le tableau 5.2 on trouve les valeurs de chaque élément du modèle.

L [nH]	L _{a1} [nH]	L _{a2} [nH]	C _p [fF]	C _r [fF]	C _{G1} [fF]	C _{G2} [fF]	R[Ω]
1.41	0.25	0.37	13.5	12.4	33.2	22.6	0.6

Tableau 5.2 Paramètres électriques du modèle de l'inductance

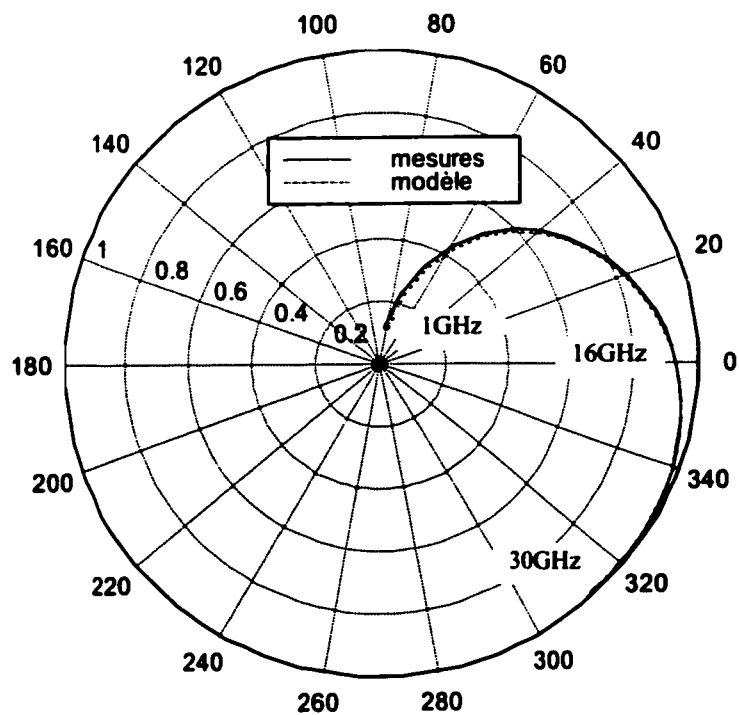


Figure 5.8 Coefficient de réflexion mesuré et modélisé de l'inductance série.

L'inductance a été mesurée entre 1 et 30 GHz. La valeur de basse fréquence de l'inductance est relativement grande (2.23 nH), ce qui est une indication que la fréquence

de résonance n'est pas trop élevée. Les capacités parasites de couplage et les pertes de la ligne centrale modifient la trajectoire de la courbe par rapport à la courbe idéale qui décrit un cercle. On y trouve une fréquence de résonance à environ 16 GHz. Les valeurs plus faibles d'inductance peuvent avoir une fréquence de résonance plus élevée (~40-60 GHz). Au-delà de ces fréquences, les effets des ponts-à-air de connexions deviennent comparables à ceux de l'inductance centrale et il est difficile d'obtenir une bonne reproductibilité.

5.3 Résistances Déposées

Un des procédés développé dans le laboratoire GRM a été dédié au dépôt des couches résistives. Outre ce dépôt spécifique, deux solutions inhérentes à la technologie de fabrication des transistors peuvent aussi être utilisées.

- On peut utiliser la couche épitaxiale du canal, dont la résistance spécifique est relativement élevée. Cette solution est utilisée seulement dans les applications où la précision n'est pas critique.
- On peut aussi fabriquer une résistance, en utilisant la couche fortement dopée, utilisée pour la fabrication des contacts ohmiques ("Cap Layer"). Par rapport aux résistances de canal ($480\Omega/\text{carré}$ pour notre type de substrat) celles-ci ont une plus faible résistance spécifique ($90\Omega/\text{carré}$). Dans ce cas même les résistances avec des petites valeurs peuvent être fabriquées.
- Certains manufacturiers préfèrent même ajouter dans leur structure épitaxiale, une couche spécialement dédiée à la fabrication des résistances enterrées. Cependant, les résistances utilisant une couche de semiconducteur présentent des variations importantes à cause de la sensibilité de la dynamique de déposition à température élevée, et du nombre de porteurs dans la couche déposée. En plus de ces désavantages, ces couches présentent aussi des phénomènes non-linéaires en fonction de la puissance du signal.

La solution pour une résistance facilement intégrable et possédant aussi des bonnes performances électriques est la déposition de minces couches métalliques résistives. L'alliage de NiCr est le standard actuel dans l'industrie et il est facilement déposé par évaporation ou pulvérisation cathodique.

La figure 5.9 montre une structure de test, réalisée au GSM pour mesurer et caractériser le procédé de fabrication. On voit une structure de mesure Kelvin (technique de mesure en 4 points) connectée à la couche résistive pour extraire l'information sur la valeur de la résistance. La valeur de la résistivité par carré est contrôlée par l'épaisseur de la couche déposée, en l'occurrence le temps de dépôt dans l'évaporateur. La valeur ciblée au début a été de $20\Omega/\text{carré}$ pour faciliter la fabrication des faibles résistances, mais en fonction de la valeur finale désirée par le concepteur on peut modifier le temps de déposition.

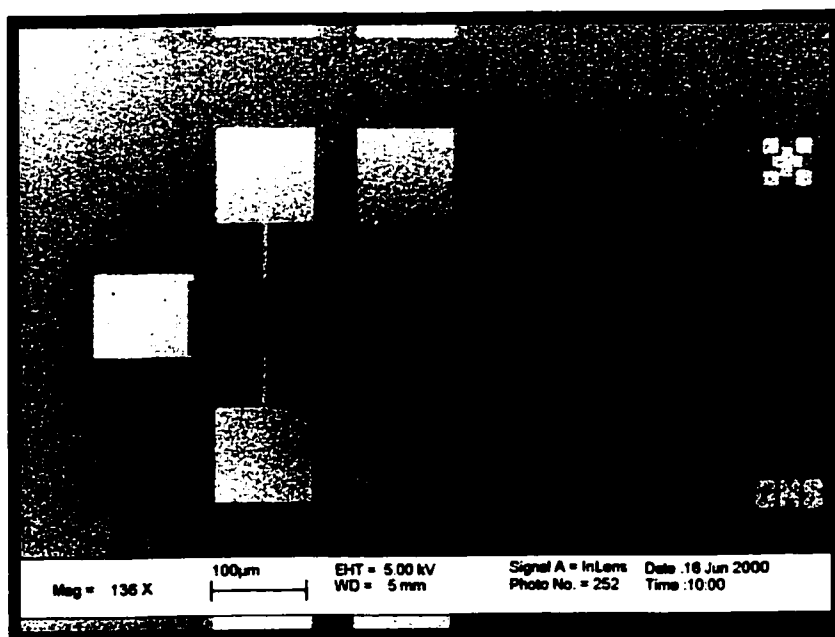


Figure 5.9 Structure test Kelvin pour résistance NiCr.

Comme pour les autres structures, on a fabriqué un masque avec plusieurs structures pour mesurer la résistance spécifique. Les longueurs de résistances sont intentionnellement différentes pour l'obtention des valeurs de résistance variables. La figure 5.10 montre les résultats du test utilisé pour calibrer le temps d'évaporation. Comme on peut le voir, le procédé a une excellente tolérance ($\sim 2\%$). Pour le courant critique un test non-destructif a indiqué une valeur d'environ 2mA/mm. Il est conseillé d'accepter une certaine marge de fiabilité et le concepteur ne va donc pas utiliser plus de 80% de la capacité en courant, de la résistance.

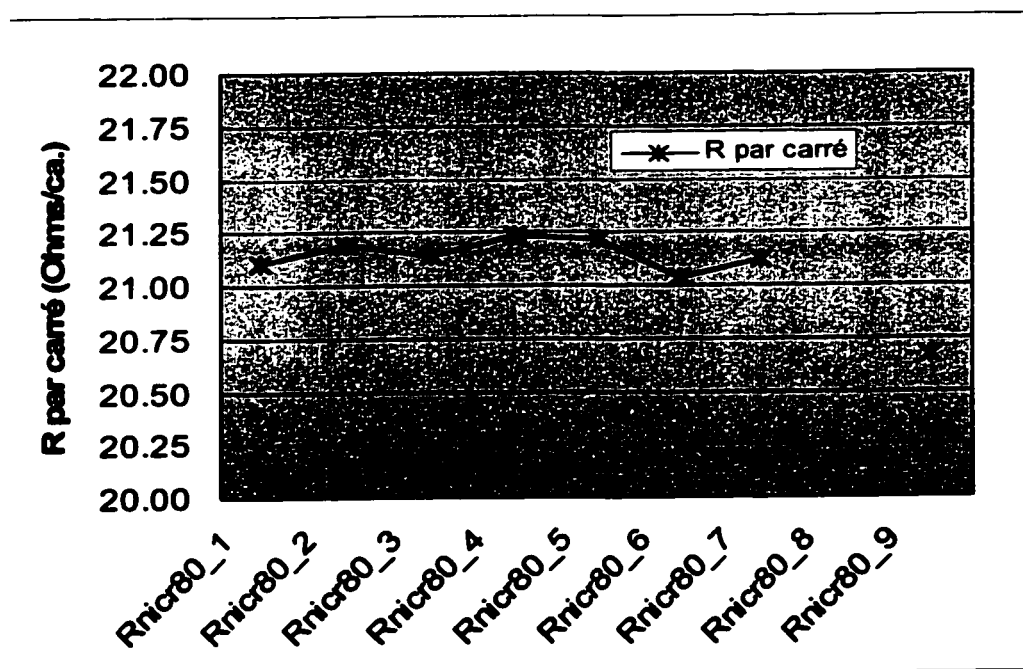


Figure 5.10 L'extraction de la résistance spécifique ($\Omega/\text{carré}$).

Les résistances déposées, ayant des dimensions très faibles, ne présentent pas une caractéristique distribuée. En principe elles sont très proches de la composante électrique idéale. Pour ces raisons, on veut d'habitude utiliser seulement une simple résistance (dépendante de la largeur et la longueur physique) comme la représentation électrique.

Dans le cas d'une simulation en ondes millimétriques, on peut utiliser le modèle fourni par le logiciel ADS, en spécifiant seulement les détails de notre procédé (épaisseur de la couche de métallisation, résistance spécifique et dimensions).

Finalement, soulignons que pour un circuit complet, on a besoin aussi d'autres composantes passives, comme les coudes, les jonctions en "T" et en "croix" etc. Pour ce type de composantes, il existe des librairies commerciales, mais pour s'assurer d'une excellente corrélation entre les mesures et les modèles, on doit procéder à des simulations numériques. Indépendamment de la méthode utilisée, on peut fabriquer les composantes mentionnées sans difficulté, en utilisant les procédés de fabrication MMIC décrits dans le chapitre précédent.

5.4 Le Transistor PHEMT

La composante clé des circuits intégrés est bien sûr le transistor. Son comportement "non-linéaire" nous permet la réalisation de fonctions de circuit complexes. Le besoin d'une composante pouvant fonctionner dans le domaine des ondes millimétriques impose certains contraintes au niveau de la réalisation physique du dispositif. Pour obtenir des fréquences de travail supérieures à 30 GHz on doit avoir un temps de transit des électrons d'environ 1-2 ps. Cette valeur peut être obtenue à deux conditions : une longueur de transit courte (< 200 nm) et une mobilité de porteurs très grande. Ces deux conditions ont imposé comme unique choix le transistor à haute mobilité d'électrons ("HEMT").

La figure 5.11 montre la structure transversale du transistor PHEMT qu'on a fabriqué à Sherbrooke. Même si ce n'est pas une structure conventionnelle, puisqu'elle présente une double hétérostructure, elle est similaire aux autres types de transistors HEMT.

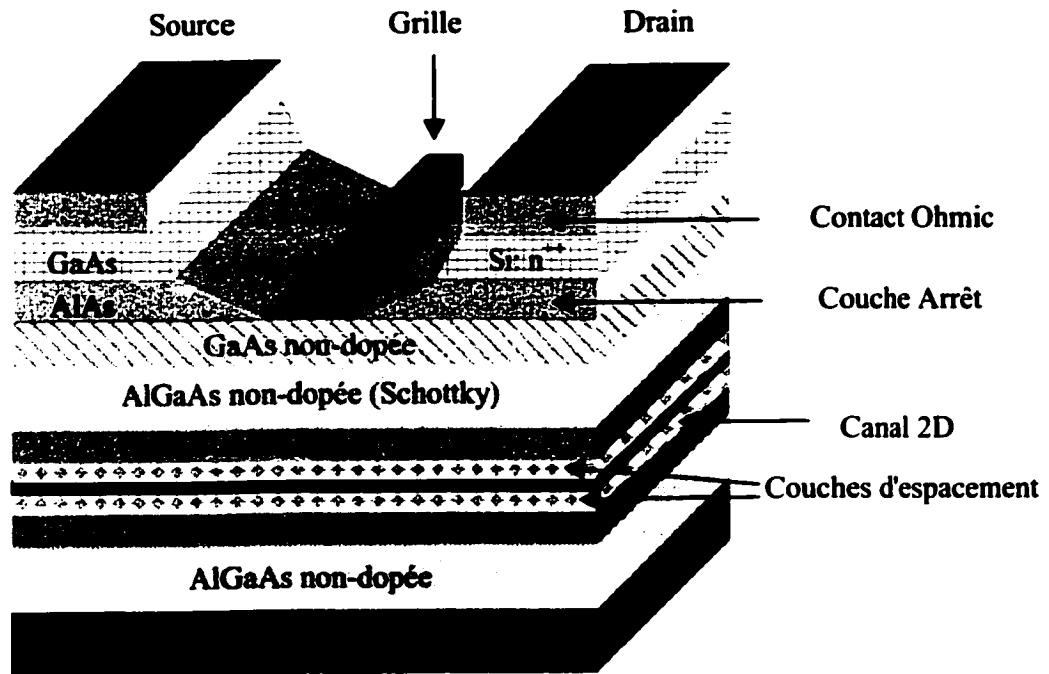


Figure 5.11 Structure transversale du transistor PHEMT

Le principe fondamental du transistor à haute mobilité d'électrons est la présence d'une différence entre les niveaux de la bande interdite entre deux matériaux avec une constante cristalline similaire (ex. GaAs et AlGaAs). Cette différence va générer la présence d'un gaz d'électrons bidimensionnel. Ces électrons sont fournis par la couche dopée et ils sont confinés à l'interface entre les deux couches. L'avantage fondamental par rapport au transistor MESFET consiste dans le mécanisme de contrôle du flux d'électrons. Pour le MESFET la tension de grille contrôle la profondeur de la zone d'appauvrissement, tandis que pour le PHEMT la tension de grille contrôle la densité de porteurs confinés dans le canal de conduction. Alors, la conduction prend place dans une zone non-dopée, donc les électrons du canal n'interagissent pas avec les ions donneurs (ils ne sont pas freinés). En séparant ainsi les électrons de leurs donneurs on augmente de façon importante la valeur de la mobilité de volume.

Plusieurs variantes ont été développées pour améliorer les hétérostructures. La première version du HEMT utilisait une structure à constante cristalline adaptée ("lattice matched") de type AlGaAs/GaAs, mais la faible différence de la bande de conduction (seulement 0.2eV) donnait un faible facteur de confinement pour les porteurs dans le canal. Au fur à mesure, on a amené des améliorations aux transistors, et quatre grandes catégories de PHEMT peuvent être rencontrées sur le marché aujourd'hui :

- PHEMT sur substrat de GaAs avec une hétérostructure plus complexe de type $\text{Al}_x\text{Ga}_{1-x}\text{As}/\text{In}_y\text{Ga}_{1-y}\text{As}$. En utilisant le In on augmente la différence entre les bandes de conduction et on obtient aussi une meilleure mobilité. La structure utilisée pour la fabrication de notre transistor fait partie de cette catégorie.
- PHEMT sur substrat de InP avec une hétérostructure de type $\text{Al}_{0.48}\text{In}_{0.52}\text{As}/\text{In}_x\text{Ga}_{1-x}\text{As}$ ($0.53 < x < 0.8$). En utilisant ce substrat on obtient de meilleures valeurs pour la mobilité ($\sim 10400 \text{ cm}^2/\text{Vs}$ à 300 degrés K)
- PHEMT de type "lattice-matched" LM-HEMT sur substrat de InP. La structure est fixe : $\text{Al}_{0.48}\text{In}_{0.52}\text{As}/\text{In}_{0.53}\text{Ga}_{0.47}\text{As}$.
- HEMT Métamorphique (MM-HEMT), qui correspond à une structure de type $\text{Al}_{1-y}\text{In}_y\text{As}/\text{In}_x\text{Ga}_{1-x}\text{As}$ ($0.3 < x < 0.5$).

Évidemment, la qualité de la composante finale dépend de la qualité de la hétérostructure et le dépôt par épitaxie des différentes couches doit être contrôlé rigoureusement. Comme il n'y avait pas la possibilité de faire l'épitaxie dans nos laboratoires on a utilisé une hétérostructure prédéfinie d'un fabricant français (PICOGIGA).

En principe deux types d'architecture de transistor peuvent être envisagés : la structure en "T" et celle de type multi-doigts (voir figure 5.12).

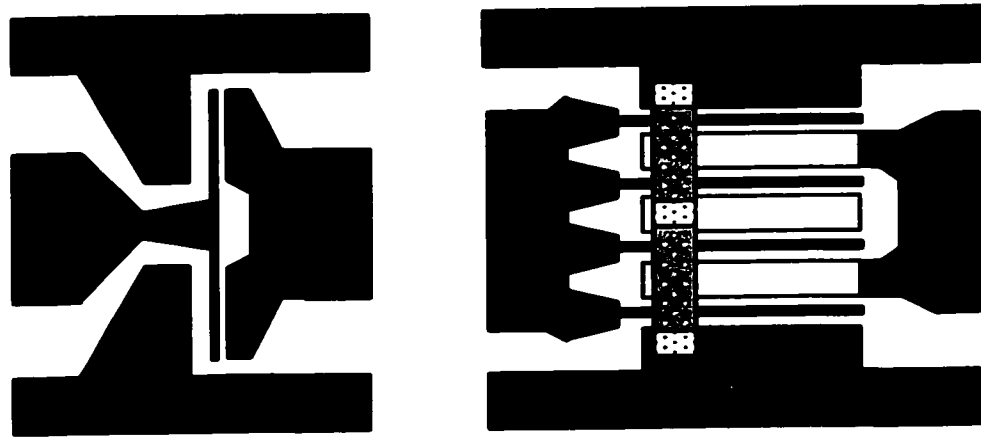


Figure 5.12 L'architecture en "T" (gauche) et "multi-doigts" (droite).

La structure de type multi-doigts simule une connexion de type parallèle entre plusieurs transistors. Elle donne la possibilité d'augmenter le courant équivalent DC du transistor et elle est surtout utilisée pour les applications pour lesquelles le gain est important. La structure en "T" est l'architecture de choix pour application faible bruit et elle permet l'obtention d'une meilleure fréquence de coupure. La raison pour cette dernière propriété est que pour les deux structures, le produit bande passante-gain est identique si on compare deux transistors avec la même périphérie (même largeur de grille équivalente). Dans le deuxième cas la présence du pont-à-air va augmenter la valeur de la capacité parasite grille-source qui influence directement la fréquence de coupure.

Dans le cadre de nos travaux, le développement des procédés de fabrication et de modélisation, a été fait en utilisant la variante en "T" du transistor. En parallèle avec les travaux de fabrication, on a bâti sous l'environnement ICCAP plusieurs algorithmes pour la caractérisation du transistor. Les algorithmes ont été testés dans une procédure complexe d'extraction de modèles pour un transistor commercial (Fujitsu FHR20X),

transistor qui présente une caractéristique de matériel et une architecture semblable à la structure développée à UDS.

5.4.1 Modèle Petit Signal.

La première procédure d'extraction présentée dans ce mémoire fournit l'information au niveau d'un modèle équivalent en petit signal. Ce type de modèle est extrêmement précis et il peut être utilisé directement pour la conception des amplificateurs faible bruit ou de petit signal. Notre intérêt portait principalement pour l'extraction des valeurs pour les composantes électriques extrinsèques, une condition obligatoire pour l'obtention d'un modèle de qualité en mode "grand signal".

Le modèle électrique considéré est illustré dans le schéma 5.12. Le format en "TI" du transistor intrinsèque reste la manière usuelle pour modéliser le circuit dans les simulateurs commerciaux. De plus, ce format est compatible avec les modèles de grand signal et les résultats sont directement exportables dans le modèle de haut niveau. Deux approches d'extraction ont été utilisées. La première, relativement récente est une combinaison des plusieurs méthodes d'optimisation ([16],[39],[60]) et elle nous permet l'obtention de l'intégralité des paramètres électriques du modèle. La deuxième approche, qui a ses origines dans les premières techniques de caractérisation ([9],[15],[37]) tient compte de l'aspect symétrique d'un transistor à effet du champ dans le cas d'une connexion diode ("cold FET", $V_{DS}=0$ V) et c'est une procédure applicable exclusivement pour l'extraction des composantes parasites de type inductif et résistif.

Au modèle de la figure 5.13, certains auteurs ajoutent parfois des capacités associées au boîtier, entre les ports externes (grille, source, drain). Dans notre cas, il ne s'agit pas d'une composante encapsulée, mais la procédure d'extraction qui sera présentée en détail est facilement modifiable pour accepter une architecture étendue du type mentionné.

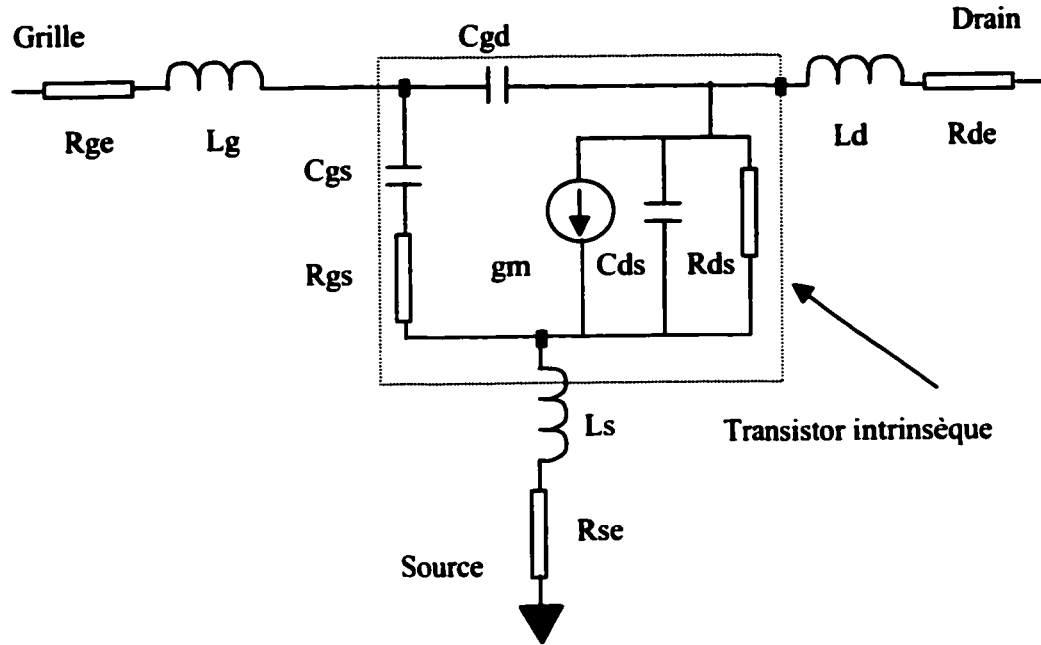


Figure 5.13 Modèle électrique du transistor en mode "petit signal".

Pour le calcul des composantes extrinsèques on a construit un algorithme automatisé qui est une version adaptée de la technique présentée en [39]. En principe, cette technique suit l'approche développée pour la première fois en [73] et elle se base sur le concept de l'indépendance des paramètres intrinsèques par rapport à la fréquence.

Si nous considérons la matrice Y du transistor intrinsèque (figure 5.12), on peut écrire :

$$Y_{int} = \begin{bmatrix} \frac{j\omega C_{gs}}{1 + j\omega C_{gs} R_{gs}} & -j\omega C_{gd} \\ \frac{gm \cdot \exp(-j\omega\tau)}{1 + j\omega C_{gs} R_{gs}} - j\omega C_{gd} & gds + j\omega(C_{gd} + C_{ds}) \end{bmatrix}, \quad 5.1$$

où " τ " est le temps de transit associé avec la source de courant commandé, " gm " est la transconductance et " gds " est la conductance de sortie.

Pour les éléments extrinsèques on peut déduire la matrice Z suivante :

$$Z_{ext} = \begin{bmatrix} R_{gc} + R_{sc} + j\omega(L_g + L_s) & R_{sc} + j\omega L_s \\ R_{sc} + j\omega L_s & R_{dc} + R_{sc} + j\omega(L_d + L_s) \end{bmatrix}, \quad 5.2$$

Dans ce cas on peut calculer la matrice complète du dispositif sous test :

$$Z_t = Z_{ext} + Y_{int}^{-1}, \quad 5.3$$

La matrice Z_t peut être déterminée par les mesures de paramètres S au niveau du transistor. La procédure débute à partir de la matrice complète mesurée, auquel on soustrait Z_{ext} pour obtenir la matrice Z_{int} des éléments intrinsèques du modèle. En connaissant la matrice Y_{int} par inversion on va pouvoir calculer analytiquement les paramètres électriques :

$$d(\omega) = \frac{\text{Re}(Y_{11}(\omega) + Y_{12}(\omega))}{\text{Im}(Y_{11}(\omega) + Y_{12}(\omega))}, \quad c(\omega) = (Y_{21}(\omega) - Y_{12}(\omega))(1 + jd(\omega)), \quad 5.4$$

$$C_{gs} = \frac{(1 + d^2(\omega))}{\omega} \text{Im}(Y_{11}(\omega) + Y_{12}(\omega)), \quad 5.5$$

$$C_{ds} = \frac{\text{Im}(Y_{22}(\omega) + Y_{12}(\omega))}{\omega}, \quad 5.6$$

$$R_{gs} = \frac{d^2(\omega)}{(1 + d^2(\omega))\text{Re}(Y_{11}(\omega) + Y_{12}(\omega))}, \quad 5.7$$

$$C_{gd} = \frac{-\text{Im}(Y_{12}(\omega))}{\omega}, \quad 5.8$$

$$gm = \sqrt{c^2(\omega)}, \quad 5.9$$

$$\tau = -\frac{1}{\omega} \arctan\left(\frac{\text{Im}(c(\omega))}{\text{Re}(c(\omega))}\right), \quad 5.10$$

$$g_{ds} = \text{Re}(Y_{22}(\omega) + Y_{12}(\omega)) \quad , \quad 5.11$$

Dans 5.5-5.11 Y_{ij} sont les coefficients de la matrice Y_i . On observe que les paramètres intrinsèques dépendent de la fréquence et des termes de la matrice Z_{ext} . Donc on peut représenter chaque paramètre du circuit intrinsèque comme une fonction du type:

$$C_{gs} = C_{gs}(\omega, Z_{ext}) \quad , \quad C_{gd} = C_{gd}(\omega, Z_{ext}) \quad , \quad \text{etc.} \quad 5.12$$

En attribuant des valeurs aux paramètres extrinsèques on peut représenter ces fonctions en fonction de la fréquence. Si les valeurs attribuées sont différentes par rapport aux valeurs réelles, chaque paramètre va présenter une variation en fonction de la fréquence, ce qui contredit l'hypothèse d'invariance. À partir de cette considération une procédure d'optimisations sur les valeurs de paramètres extrinsèques est mise en place avec comme cible la variation de chaque paramètre du modèle intrinsèque. Dans ce cas on peut trouver une fonction coût de type [50,113]:

$$C(X_e) = \sum_{j=1}^6 \sum_{i=1}^{NP} w_j |f_j(\omega_i, X_e) - m_j| / |m_j| + \sum_{p=1}^2 \sum_{q=1}^2 \sum_{i=1}^{NP} |S_{pq}^M(\omega_i) - S_{pq}^C(\omega_i, X_e)| / |S_{pq}^C(\omega_i, X_e)| \quad , \quad 5.13$$

Dans l'équation 5.13 X_e est le vecteur des paramètres extrinsèques, f_j est la fonction associée à un paramètre intrinsèque ($j=1\dots 6$), m_j la valeur moyenne de la fonction sur l'ensemble de points de mesure en fréquence (NP), S^M , S^C sont les paramètres S mesurés respectivement calculés. W_j sont des facteurs de poids.

Di Martino et al. [39] considèrent aussi le coût sur une plage de points de polarisation. Malgré la généralité de l'expression on a trouvé que la procédure d'optimisation en considérant la valeur moyenne est difficile en raison de l'instabilité de la dérivée de premier ordre. Par conséquent, on a adopté le critère de la variance (erreur quadratique)

similaire à la version initiale [113]. Dans ce cas les deux termes d'erreur du membre droite de l'équation 5.13 deviennent :

$$T_1 = \frac{1}{NP-1} \sum_{j=1}^6 \sum_{i=1}^{NP} w_j |f_j(\omega_i, X_c) - m_j|^2, \quad 5.14$$

$$T_2 = \sum_{p=1}^2 \sum_{q=1}^2 \sum_{i=1}^{NP} w_{pq} |S_{pq}^M(\omega_i) - S_{pq}^C(\omega_i, X_c)|^2, \quad 5.15$$

L'avantage de l'expression en termes quadratiques est la disponibilité des routines d'optimisation puissante dans le logiciel ICCAP. Comme tous les problèmes d'optimisation, les valeurs de départ peuvent influencer la solution finale. En [107] les auteurs ont conçu une routine particulière, en utilisant une technique spéciale de annelage ("annealing"). Cette technique évite les minimaux locaux, mais en revanche, elle demande un temps d'exécution très long. Pour diminuer la complexité de l'algorithme on a utilisé une procédure de type Levenberg-Marquart.

Les figures 5.14 et 5.15 montrent les résultats de notre procédure pour le transistor Fujitsu FHR20X, avec la tension de drain $V_{ds}=2V$ et la tension de grille $V_{gs}=-0.2V$. Comme on peut voir sur une large gamme de fréquence les paramètres intrinsèques restent presque constants, en conformité avec l'hypothèse considérée.

En plus de cette observation, la valeur de la transconductance "gm" et le temps de transit "τ" sont compatibles avec les valeurs extraites par une procédure classique de modélisation grand signal (modèle Curtice [6,35, 36]), procédure qui sera analysée dans un paragraphe ultérieur.

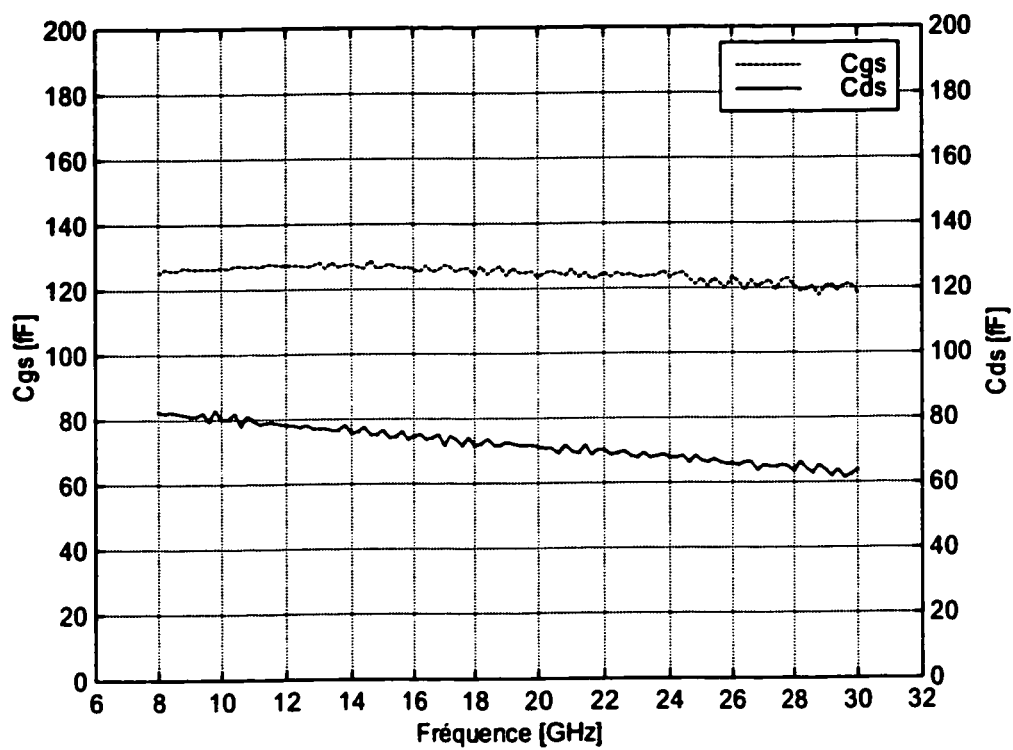


Figure 5.14 Extraction de capacités intrinsèques C_{gs} et C_{ds}

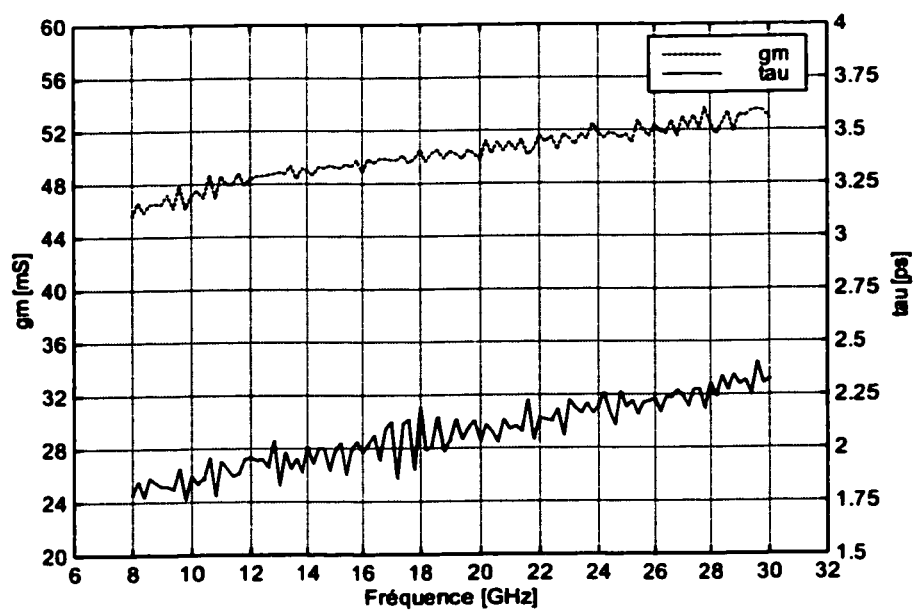


Figure 5.15 Transconductance et le temps de transit.

L'algorithme d'extraction s'est avéré robuste sur toute la plage de variation de points de polarisation, à condition qu'on "garde" le transistor dans la partie linéaire de la caractéristique. Dans ces conditions il est évident que la procédure de modélisation donne d'excellents résultats pour la caractérisation petit signal. La figure 5.16 montre la comparaison entre les valeurs simulées et mesurées des coefficients de réflexion à l'entrée et à la sortie du dispositif sous test (S_{11} et S_{22}).

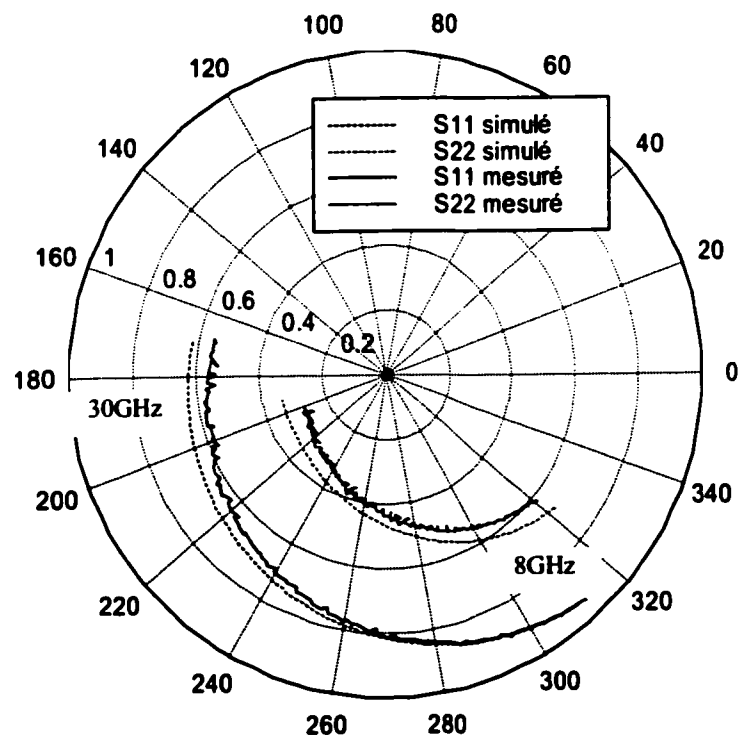


Figure 5.16 Coefficients de réflexion S_{11} et S_{22} (mesurés et simulés).

Une bonne corrélation a été trouvée aussi dans le cas des paramètres de transfert.

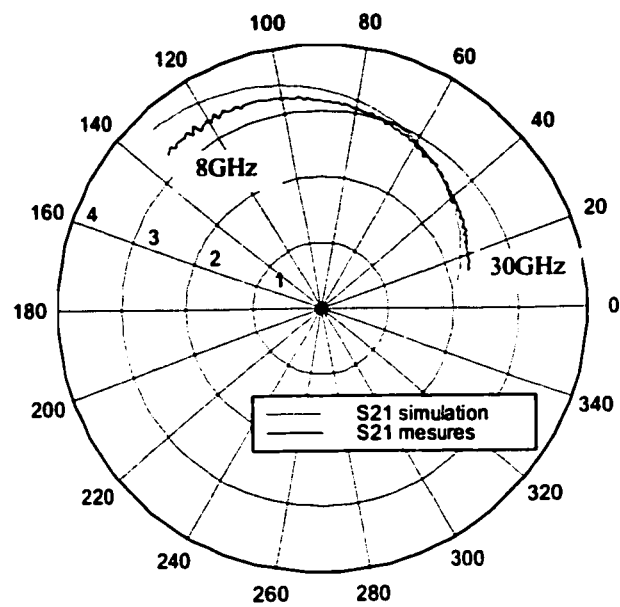


Figure 5.17 Paramètres de transfert S21 (mesurés et simulés).

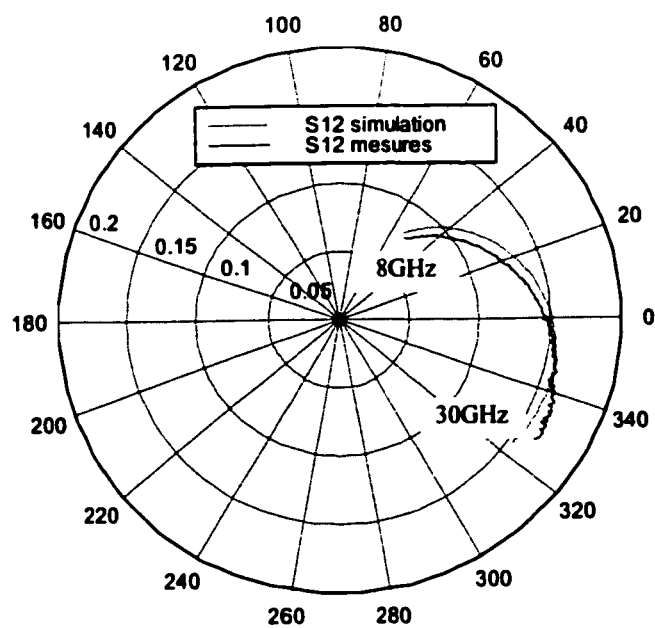


Figure 5.18 Paramètres de transfert S12 (mesurés et simulés).

Le modèle de petit signal peut être utilisé avec succès pour la conception des amplificateurs en classe A. Aussi, avec une mesure de la figure de bruit il est idéal pour la fabrication des amplificateurs faible bruit (LNA). Notre principal objectif a été de trouver une solution robuste pour l'extraction précise des inductances et des résistances parasites. Ces valeurs sont indispensables pour les modèles généraux en grand signal. Le tableau 5.3 indique les valeurs trouvées par la procédure d'optimisation présentée.

Lg [pH]	Rge [Ω]	Ld [pH]	Rde [Ω]	Ls [pH]	Rs [Ω]
189.3	7.8	190.8	2.6	5.3	1.98

Tableau 5.3 Valeurs des composantes électriques parasites.

Les valeurs trouvées sont en concordance avec le dessin du circuit et la connexion par fil soudé, le transistor étant monté sur une base de céramique. On s'est connecté à la grille et au drain en utilisant un fil d'environ 250 microns (~ 10 mils). La source a été connectée par 4 fils (deux de chaque côté, voir figure 5.11 gauche). En conséquence, son inductance est très faible. La résistance de source est un peu plus faible que la résistance de drain (en raison de la plus grande surface métallique) et elle présente une valeur compatible avec des dispositifs similaires. En raison de sa longueur de 0,18 microns la résistance de grille est la plus grande.

La deuxième méthode utilisée pour la validation des éléments extrinsèques part aussi du modèle de la figure 5.13. Si on considère le transistor polarisé avec une tension $V_{DS}=0V$, le circuit devient un circuit réciproque dont la matrice Z est donnée par les relations 5.16-5.18 [18,37] pour une tension de grille inférieure à ~ 0.6 V. Dans ces équations R_{CH} est la résistance du canal, R_{ge} est la résistance de la grille et C est la capacité totale grille-source et grille-drain.

$$Z_{11} = R_{se} + R_{CH}/3 + R_{ge} + j\omega(L_s + L_g) - 1/j\omega C, \quad 5.16$$

$$Z_{12} = Z_{21} = R_{se} + R_{CH}/2 + j\omega L_s, \quad 5.17$$

$$Z_{22} = R_{se} + R_{de} + R_{CH} + j\omega(L_s + L_d), \quad 5.18$$

Si la tension de grille dépasse le seuil d'ouverture de la jonction Schottky, Anholt [7] a montré que les équations se modifient et elles deviennent :

$$Z_{11} = R_{se} + \alpha_G R_{CH} + R_{ge} + R_{GG} + j\omega(L_s + L_g), \quad 5.19$$

$$Z_{12} = Z_{21} = R_{se} + \alpha R_{CH} + j\omega L_s, \quad 5.20$$

$$Z_{22} = R_{se} + R_{de} + 2\alpha R_{CH} + j\omega(L_s + L_d), \quad 5.21$$

$R_{GG} = nkT/I_G$ est la résistance de la diode de grille et α , α_G sont des coefficients qui modélisent la dépendance de Z avec le courant de grille.

À partir de ces expressions, on peut déduire facilement les inductances en considérant la partie imaginaire des équations 5.19-5.21. Pour l'extraction des résistances, on observe qu'il y a plus d'inconnues que d'équations. Plusieurs solutions sont indiquées dans la littérature. La majorité de ces solutions sont obtenues en faisant des hypothèses quant à la variation des paramètres α et α_G . Aussi R_{CH} peut être déduit à partir de mesures de conductivité sur les matériaux. Dans notre cas, nous avons préféré adopter la méthode Yang-Long [132] qui est fondée seulement sur l'hypothèse qu'une jonction grille-source peut être modélisée par une structure distribuée de diodes identiques. Dans ce cas, on peut écrire le courant de grille comme [132]:

$$I_G = W l_g J_s \exp((V_G - V_s)/nV_T) F, \quad 5.22$$

$$F = \frac{1}{l_g} \int_0^{l_g} \exp(-V'_d x / (l_g n V_T)) dx = (1 - \exp(-u)) / u, \quad u = V'_d / (n V_T) \quad 5.23$$

V_T = tension thermique, n = facteur d'idéalité, l_g = largeur de la grille

À partir de l'équation 5.23 on mesure I_G en deux points de polarisation différents et on peut écrire :

$$V_{G2} - V_{G1} = V_{S2} - V_{S1} + R_s(I_{G2} - I_{G1}) + nV_T \ln(F_1/F_2) , \quad 5.24$$

Comme dans l'équation 5.24, on connaît tous les termes ou du moins on peut les mesurer, la résistance de source peut être calculée directement :

$$R_s = (V_{G2} - V_{G1} + nV_T \ln(F_1/F_2)) / (I_{G2} - I_{G1}) , \quad 5.25$$

En échangeant le rôle de la source et du drain on peut trouver de la même façon la résistance de drain R_d . En connaissant les deux valeurs R_s et R_d , on obtient facilement les valeurs de R_{CH} , α et α_G à partir des équations 5.19-5.21.

Deux observations concernant la précision de cette méthode nécessitent une attention particulière. On doit absolument s'assurer que la tension drain-source est nulle. Cette condition est difficile à respecter, dans le cas d'un courant de grille important. Les sources de polarisation de type HP4142, permettent sous la commande du logiciel ICCAP, de synchroniser le courant du drain et de la source au même niveau mais de sens inverse. En conséquence, on évite le problème de la source flottante.

Pour les transistors PHEMT, la résistance de canal peut être relativement élevée. L'expression simplificatrice dans les formules 5.16 et 5.19 n'est pas très précise dans ce cas, et la partie imaginaire dépend aussi de la résistance du canal et de la capacité totale de la structure. En conséquence, les valeurs des inductances extraites peuvent être erronées. Une procédure d'optimisation finale comme celle présentée auparavant est fortement recommandable pour l'obtention des bons résultats.

5.4.2 Modèle Grand Signal.

L'objectif final pour un travail de caractérisation d'une composante active est l'obtention d'un modèle qui rend compte d'une ou de plusieurs fonctions de transfert entre les tensions et courants dans les ports d'accès de la composante. Un tel modèle, qui ne comporte aucune restriction en termes de niveau du signal, s'appelle modèle large signal. Il existe plusieurs approches pour la caractérisation des transistors à effet de champs et en particulier les PHEMTs. Sans faire une présentation exhaustive, on peut distinguer trois grandes classes de modèles :

- **Modèles physiques.** Si on applique les équations de Schroedinger et de Maxwell en considérant la description matérielle de chaque couche et leur dimensions géométriques, on peut obtenir une liaison de type fonctionnelle entre les inconnues désirées : tensions et courants aux bornes de transistor.
- **Modèles analytiques ou semi-empiriques.** À partir d'observations physiques et de mesures spécialisées on peut extraire un modèle électrique en utilisant de composantes localisées dont la valeur est décrite par de fonctions analytiques. Le type de fonction a été choisi parmi celles qui donnent une meilleure approximation par rapport aux données mesurées.
- **Modèles à base des mesures.** En mesurant le dispositif sur une large plage des points de fonctionnement (domaine exhaustif de fonctionnement) on y construit un modèle de haut niveau dont la composante est représentée par sa réponse aux excitations.

Théoriquement la première classe et potentiellement la dernière donne la meilleure précision. Pratiquement, à quelques exceptions, la deuxième classe représente le quasi-monopole pour les modèles utilisés aujourd'hui en pratique. Les raisons pour une telle réalité sont multiples, mais ils ont comme principale source le compromis entre précision et les résultats pratiques. Si on adopte une approche physique, on est obligé de

résoudre un système d'équations non-linéaire exceptionnellement complexe. Habituellement, on a besoin d'une procédure très laborieuse en utilisant des méthodes de calcul numériques et statistiques comme les techniques de type Monte-Carlo. Le temps de calcul est long, et en plus on ne peut pas extraire des coefficients statistiques de variation de composantes. En pratique les taux de dopage et l'épaisseur des couches varient en fonction de la tolérance des procédés de fabrication. Une telle variation n'est pas quantifiable dans un modèle physique sous forme des données statistiques (moyenne, variance etc.). En conséquence il est presque impossible d'estimer le rendement de fabrication ("yield") pour un certain circuit (ex. amplificateur). En pratique ce type de modèles est exclusivement utilisé pour l'analyse de composante par le fabricant et le développement de nouvelles composantes.

La troisième classe à priori peut obtenir la précision maximale. Malheureusement, les mesures sur une plage très large de tensions et de fréquences vont rapidement augmenter la taille de la base de données ce qui rend impossible leur utilisation dans des ordinateurs peu performants. Mais, l'avantage de la précision a imposé quelques variantes qui utilisent une taille mémoire raisonnable. Il y a deux types de modèles très performants appartenant à cette classe : le modèle de type ROOT et les modèles à base de réseaux de neurones. Les derniers ne sont pas encore disponibles sous les simulateurs commerciaux, donc ils ont moins de popularité parmi les concepteurs. En revanche, le modèle de type ROOT est un standard actuel en industrie et nous l'avons utilisé comme une des architectures de modélisation de base pour un algorithme de caractérisation complètement automatisé.

La deuxième classe qui est celle utilisée depuis l'apparition des transistors, permet le transfert du modèle facilement, en prenant seulement les paramètres associés à chaque famille de fonctions caractéristiques. En fonction de l'expression mathématique utilisée pour la modélisation d'un phénomène particulier, on reconnaît une multitude de modèles: Curtice, Chalmers, Materka, TOM etc.

Pour notre technique d'extraction, on a choisi les modèles de ROOT et HEMT1 (une variante du modèle d'Angelov [5,6] proposée par la compagnie Agilent). Les résultats de modélisation sur le même transistor FHR20X seront présentés avec une brève description du fondement du modèle et les particularités d'implantation qu'on a choisies.

Si on considère le modèle standard présenté par ROOT [105,106] (voir figure 5.19) on peut voir que les éléments du transistor intrinsèque doivent être considérés variables en fonction du point de polarisation.

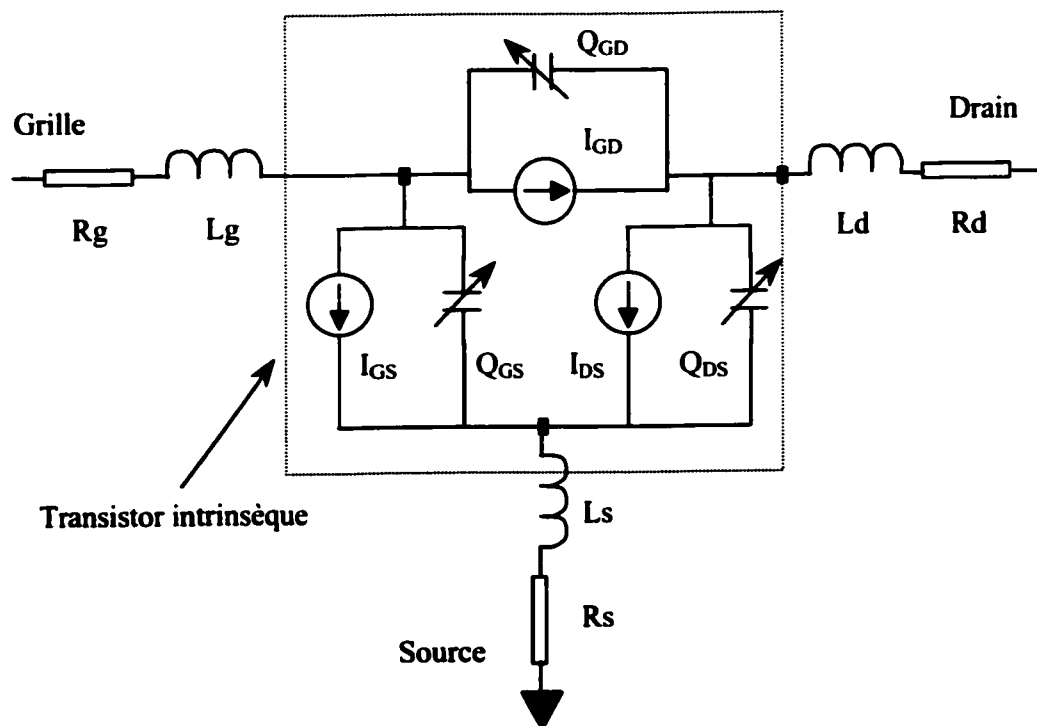


Figure 5.19 Modèle du PHEMT en mode grand signal.

Les simulateurs non-linéaire commerciaux utilisent dans leur grande majorité, la description en paramètres "Y" pour les interconnexions, dont les tensions aux nœuds sont les inconnues et les courants les variables. En variant les courants pour satisfaire

l'équation de Kirchoff aux nœuds, on y trouve le vecteur des inconnues. Dans ce cas, on peut écrire pour le transistor vu comme diport que :

$$\rho_{F_1} = \frac{\text{Imag}(Y_{11}^M(V_{GS}, V_{DS}, \omega))}{\omega} \hat{v}_{GS} + \frac{\text{Imag}(Y_{12}^M(V_{GS}, V_{DS}, \omega))}{\omega} \hat{v}_{DS}, \quad 5.26$$

$$\rho_{F_2} = \frac{\text{Imag}(Y_{21}^M(V_{GS}, V_{DS}, \omega))}{\omega} \hat{v}_{GS} + \frac{\text{Imag}(Y_{22}^M(V_{GS}, V_{DS}, \omega))}{\omega} \hat{v}_{DS}, \quad 5.27$$

$$\rho_{F_3} = \text{Real}(Y_{21}^M(V_{GS}, V_{DS}, \omega)) \hat{v}_{GS} + \text{Real}(Y_{22}^M(V_{GS}, V_{DS}, \omega)) \hat{v}_{DS}, \quad 5.28$$

Dans les équations 5.26-5.28, on a considéré le vecteur $d\vec{V} = \rho_{GS} dV_{GS} + \rho_{DS} dV_{DS}$, le vecteur unitaire de la différentielle d'ordre 1 pour les fonctions F_i (en deux variable V_{GS} et V_{DS}). Les paramètres Y_{ij}^M sont extraits d'une mesure de paramètres "S" dans le point de polarisation (V_{GS}, V_{DS}).

Deux observations majeures sont importantes pour les mesures 5.26-5.28. Les tensions sont les valeurs DC mesurées au niveau du dispositif et les paramètres S associés sont mesurés au niveau du transistor intrinsèque. Autrement dit, on doit absolument connaître les éléments extrinsèques (inductances et résistances dans la figure 5.17) pour calculer les valeurs Y_{ij}^M . Dans ce cas, la technique présentée dans la section précédente est une solution pour obtenir un bon modèle non-linéaire.

La deuxième observation est liée au comportement en fréquence. Même si dans 5.26-5.28 on a exprimé une variance en fréquence pour mettre en évidence la mesure avec un analyseur de réseaux, on ne doit pas avoir une variation des fonctions F_i avec la fréquence, les mesures étant faites à une seule fréquence. Cette hypothèse peut être vérifiée en effectuant les mesures à deux fréquences différentes et en comparant la différence entre les résultats obtenus. Un manque de concordance peut suggérer deux causes potentielles (en supposant des mesures correctes) : une inconsistance dans le

modèle (on doit ajouter des composantes) ou une vraie variation avec la fréquence due au phénomène non-linéaire comme la capture d'électrons ("trapping effect").

Si on fait des mesures sur un grand échantillon de paires (V_{GS}, V_{DS}) et si dans le plan cartésien (V_{GS}, V_{DS}) , on effectue une intégrale de contour pour un champ vectoriel conservatif on trouve une valeur nulle:

$$\oint \vec{F}_i(V_{GS}, V_{DS}) d\vec{V} = 0, \quad 5.29$$

N'importe quel champ vectoriel conservatif peut être exprimé comme un champ de gradient : $\nabla \phi_i = \vec{F}_i$. En regardant les équations 5.26-5.28 on observe qu'il représente soit une capacité, soit un courant. Mais la capacité est la dérivée de premier ordre d'une charge électrique par rapport à la tension et le courant est aussi la dérivée de la charge par rapport au temps. Dans ce cas on sait du principe de conservation de charge et la continuité de courant que les champs de gradient résultants sont conservatifs. En l'occurrence les potentiels ϕ_i peuvent être considérés comme les fonctions de haut niveau qui décrivent le dispositif. En mesurant dans un ensemble de point on peut interpoler les résultats pour obtenir la réponse de circuit pour une excitation donnée. Plus fine est la résolution entre deux points adjacents dans l'espace cartésien (V_{GS}, V_{DS}) , plus on se rapproche de la condition 5.29.

La procédure d'extraction du modèle est dans ce cas très claire. À partir des mesures à plusieurs points de polarisation on calcule pour chaque point les valeurs des fonctions ϕ_i . Le modèle ROOT (présenté ci-dessus) est intégré dans le logiciel ICCAP. On a seulement changé la procédure d'extraction de paramètres extrinsèques en utilisant la technique présentée dans la section 5.4.1. Pour notre transistor, pour une excellente précision le système de mesure demande entre 5 et 7 heures. De plus, cette technique nous a aussi permis de vérifier les performances de la méthode d'extraction des

paramètres extrinsèques. En variant, dans la simulation “grand-signal”, les valeurs des inductances et résistances parasites, on compare la réponse mesurée du transistor à celle obtenue par la simulation. La meilleure correspondance a été trouvée en utilisant les valeurs des paramètres extrinsèques obtenus par la procédure d'optimisation en petit signal.

Voici dans les figures 5.20-5.22 une comparaison entre les valeurs simulées et mesurées. Les valeurs correspondent à un transistor différent de celui utilisé pour l'extraction du modèle, mais de même type (si quelqu'un utilise le même transistor la correspondance est naturellement parfaite).

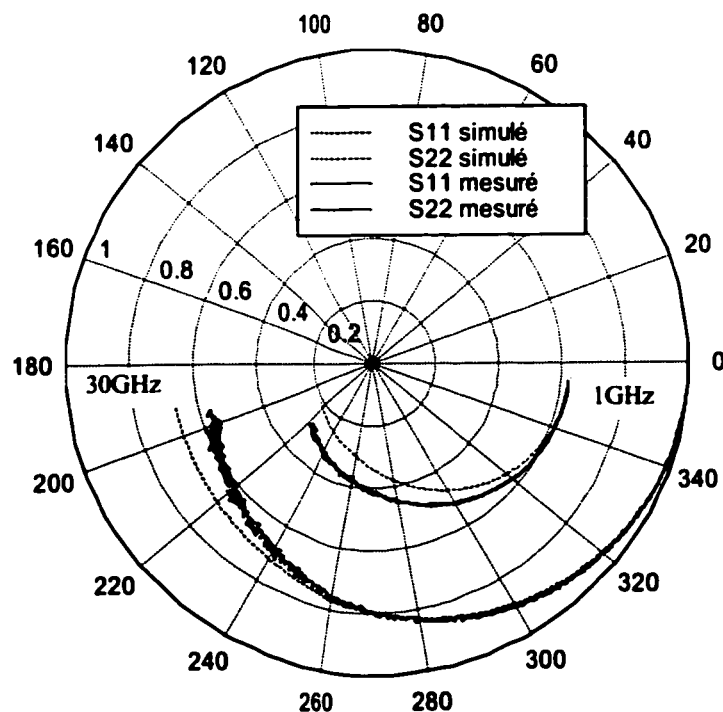


Figure 5.20 Coefficients de réflexion S_{11} et S_{22} (mesurés et modèle Root).

Pour une meilleure approximation en ondes millimétriques on doit considérer des composantes distribuées extrinsèques dans l'architecture du modèle.

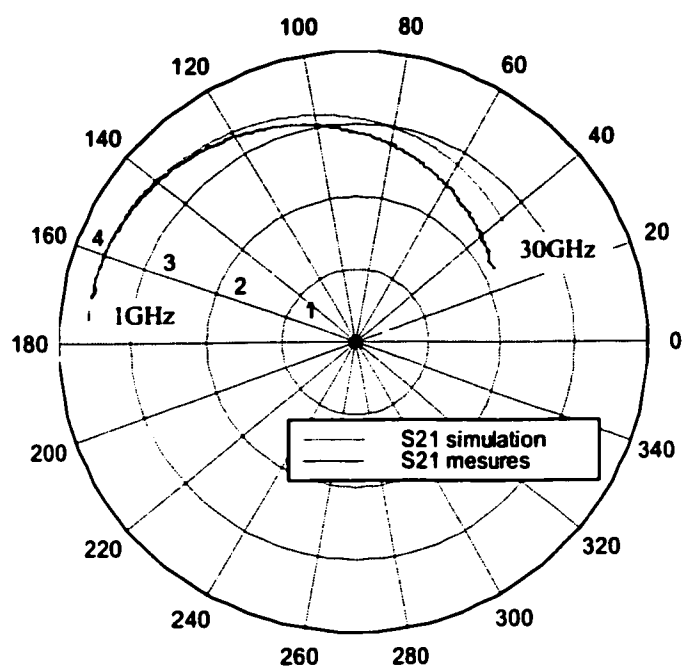


Figure 5.21 Paramètres de transfert S21 (mesurés et modèle ROOT).

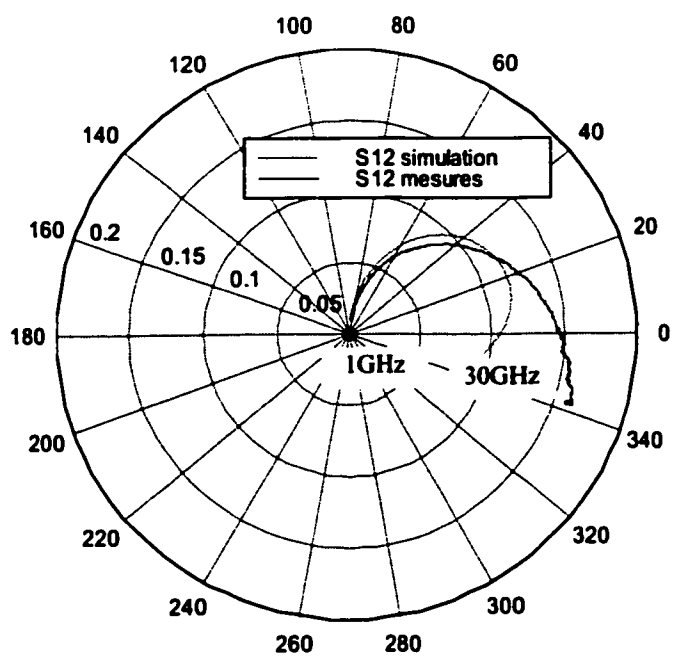


Figure 5.22 Paramètres de transfert S12 (mesurés et modèle ROOT).

Une deuxième approche de modélisation utilisée dans nos travaux est le modèle analytique de type Curtice ([35,36,87]). Ce modèle considère un circuit électrique équivalent (voir figure 5.23) qui reflète la structure physique des courants et charges dans la région active du dispositif. Comme on l'avait indiqué, une telle approche analytique, dont les expressions sont sous forme d'équations mathématiques permet aux fabricants de suivre l'aspect statistique de la fabrication (dispersion des paramètres électriques), en caractérisant la variation de chaque paramètre par rapport au procédé de fabrication.

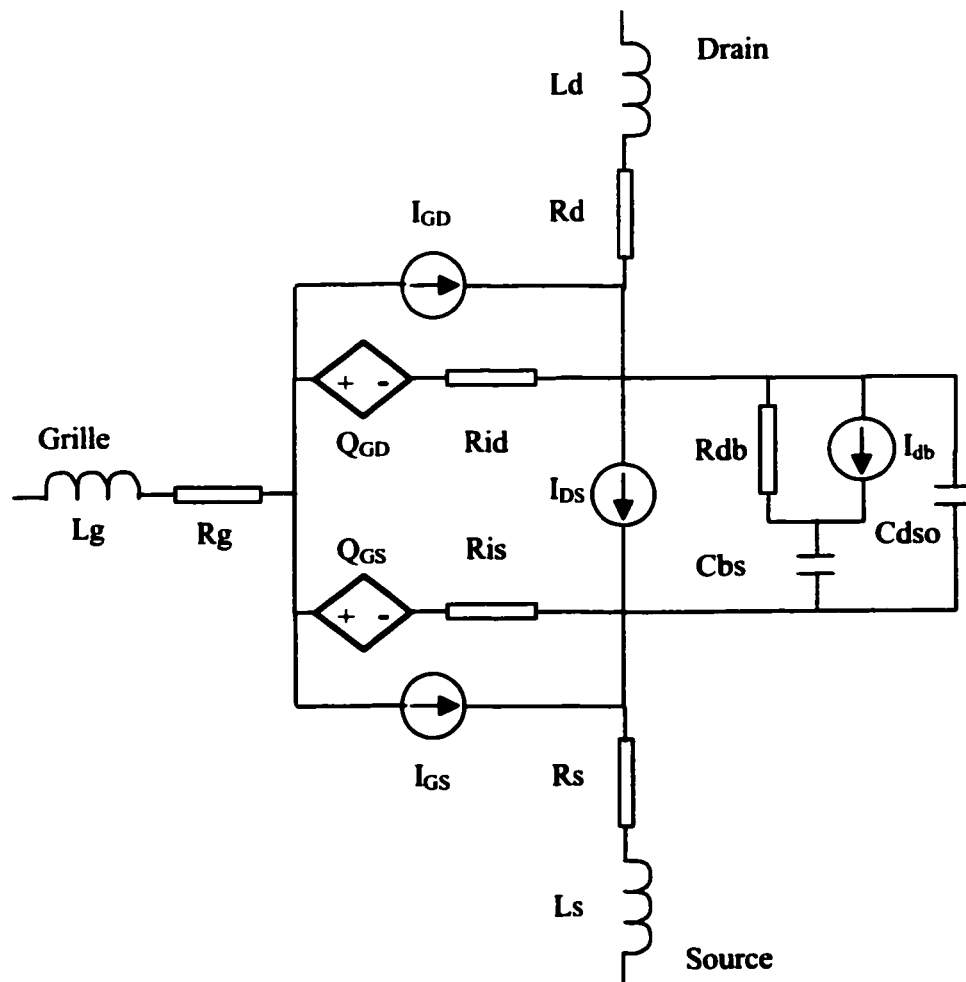


Figure 5.23 Modèle EEHEMT1.

L'évolution de ce modèle, par rapport à sa première version, a vu l'addition d'un certain nombre de paramètres supplémentaires pour une meilleure correspondance entre les valeurs mesurées et les simulations, et aussi la modification de fonctions mathématiques pour la description de chaque paramètre. Cette modification a été apportée surtout pour améliorer les méthodes de simulations, les fonctions de ce modèle étant continues jusqu'à la dérivée de deuxième ordre.

La plus importante modification est la variation de la transconductance et de la conductance de sortie en fonction de la fréquence. Les premières versions extraient ces valeurs à partir de mesures DC. La plupart de transistors présentent une dispersion de ces valeurs en hautes fréquences. Ce phénomène a nécessité l'introduction d'une source supplémentaire de courant (I_{db}) et les paramètres R_{db} et C_{bs} .

Une autre modification a été introduite pour le comportement thermique du transistor. Dans ce cas, la modification du point de polarisation en raison de la puissance dissipée ("self-heating") est modélisée pour les 3 paramètres de conduction:

$$I_{ds} = \frac{I'_{ds}}{1 + \frac{P_{diss}}{PEFF}}, \quad 5.30$$

$$g_m = \frac{g'_m}{\left[1 + \frac{P_{diss}}{PEFF}\right]^2}, \quad 5.31$$

$$g_{ds} = \frac{g'_{ds} - \frac{I'^2_{ds}}{PEFF}}{\left[1 + \frac{P_{diss}}{PEFF}\right]^2}, \quad P_{diss} = I'_{ds} V_{ds} \quad 5.32$$

$PEFF$ est le paramètre de modèle et il est déterminé par des mesures en température ou des mesures pulsées de type isothermiques.

Le plus gros avantage de ce modèle est une excellente correspondance de la transconductance mesurée avec les valeurs prédites par les simulations. Dans la figure 5.24 on observe la comparaison entre les valeurs du modèle et celles mesurées.

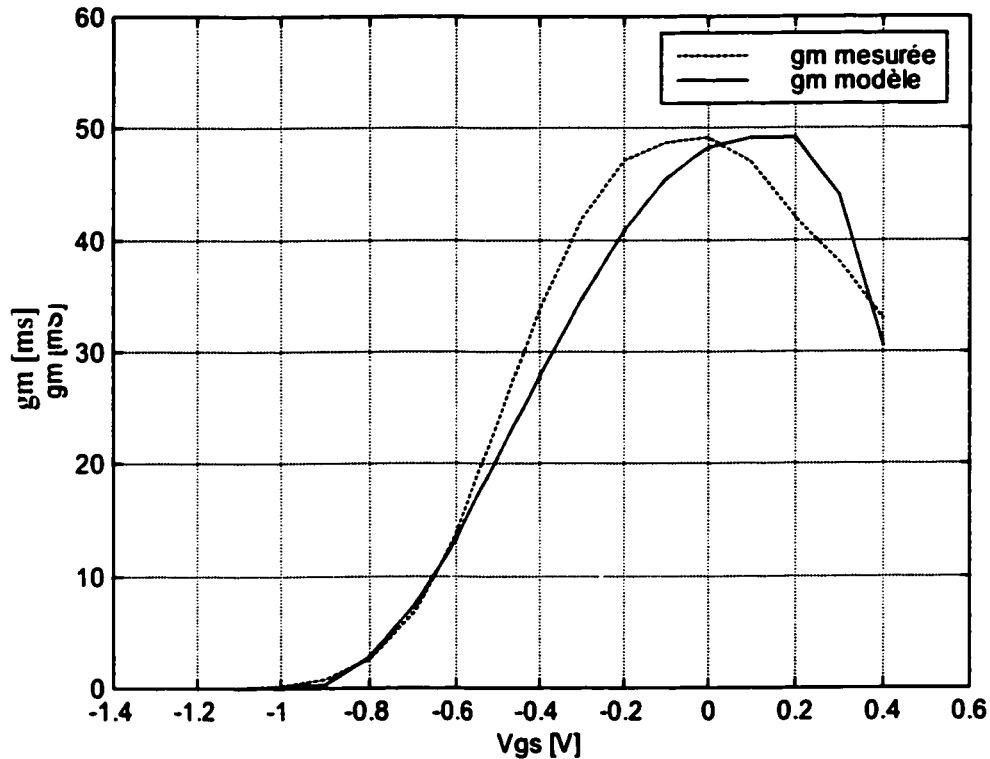


Figure 5.24 Transconductance en mesure DC.

La différence pour les plus grosses valeurs a deux causes. D'une part, le transistor mesuré avait une différence de ~ 0.1 V dans la tension de seuil par rapport à celui utilisée pour l'extraction du modèle et d'autre part, les mesures sont faites en DC sans un contrôle de la température. Le transistor utilisé pour l'extraction a été utilisé plus longtemps en mesures, dont il a fonctionné à une température un peu plus élevée. Le bon comportement en transmission donne de bons résultats pour la modélisation de paramètres S21 et S12 sur une très grande plage de fréquence (voir figures 5.25-5.26 pour les courbes représentatives).

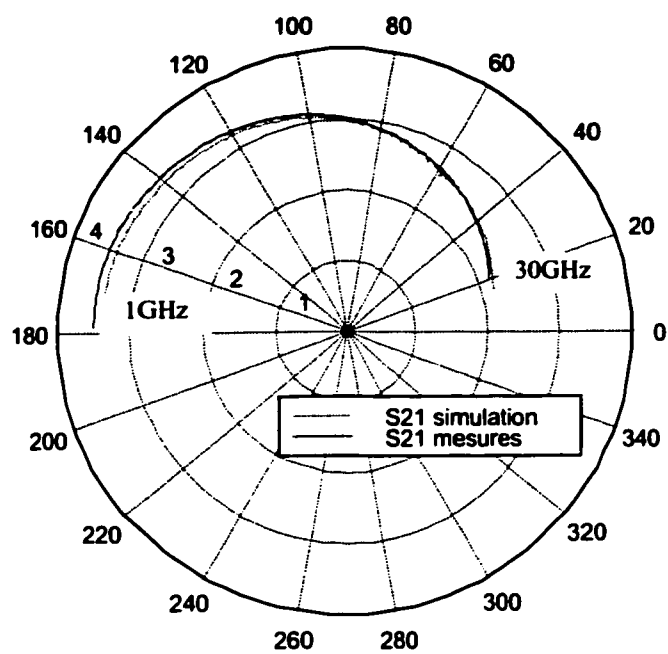


Figure 5.25 S21 modèle EEHEMT1

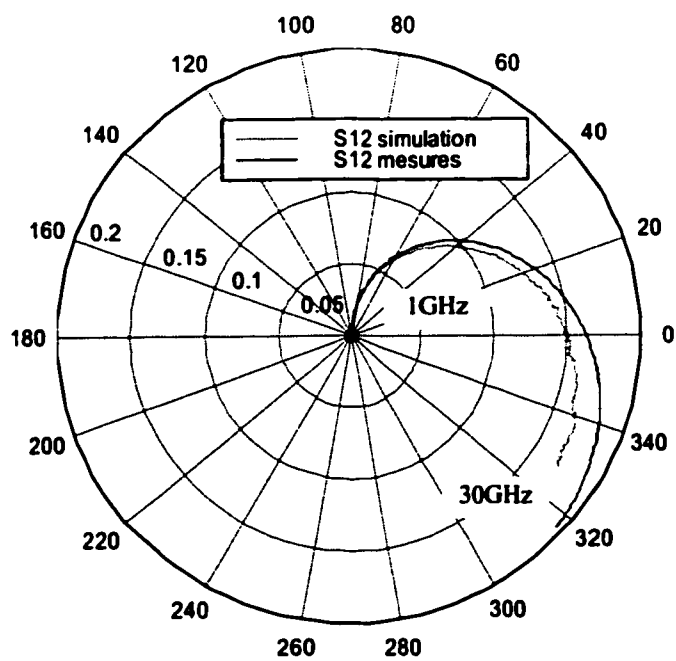


Figure 5.26 S12 modèle EEHEMT1.

Avant terminer notre présentation de ce chapitre, nous devons expliquer la méthodologie utilisée pour la validation des modèles. La comparaison entre les résultats expérimentaux et celles de la modélisation nous ont permis mettre en évidence les performances de chaque type de modèle. Certaines hypothèses ont été utilisées:

Dans le cas du modèle en petit signal (voir les figures 5.14-5.18) les courbes simulées sont déterminées en utilisant les valeurs moyennes des paramètres électriques. Une valeur moyenne est déterminée pour chaque paramètre, en considérant la variation par rapport à la fréquence. Dans ces cas, les courbes de validations nous donnent une expression quantitative sur la précision du procédé d'extraction.

Dans le cas du modèle ROOT (grand signal), en tenant compte du fait qu'il s'agit d'une méthode de modélisation à base de mesures, la comparaison est faite entre les caractéristiques simulées par le logiciel ADS (en utilisant le modèle extrait) et les mesures sur un autre transistor similaire. Si nous considérons les mesures du même dispositif utilisé pour l'extraction du modèle, la variation montrerait seulement les performances des fonctions d'interpolation et non pas la capacité du modèle de caractériser le comportement électrique du dispositif. Pour diminuer les effets de la variation physique d'une composante à l'autre, nous utilisons d'habitude un dispositif provenant du même lot de fabrication. Ainsi, nous limitons l'influence des paramètres technologique dans les résultats.

Cette technique peut permettre éventuellement de suivre les performances du procédé technologique, en comparant les valeurs simulées des deux modèles extraits à partir des composantes provenant de lots différents.

Pour rester dans le cadre de cette hypothèse, nous avons considéré le même type de comparaison pour le modèle analytique de type Curtice (voir les graphiques 5.24-5.26).

CHAPITRE 6 RÉSULTATS FINAUX ET CONCLUSIONS

On a vu dans le chapitre 3 comment l'architecture POLYCOM est mise en valeur par une procédure de correction d'erreurs analytique. Aussi, dans l'optique de la fabrication d'un récepteur intégré, en parallèle à la mise en place des installations de l'UDS, nous avons conçu un système complet de mesure, dont les détails de conception et les bases théoriques sont expliqués dans les chapitres IV et V. Nous allons maintenant présenter les résultats liés à la fabrication des transistors, et par la suite suggérer des améliorations autant au niveau des caractéristiques électriques que des procédés de fabrication développés.

Pour la validation de notre concept on a poursuivi notre travail par la conception d'une composante clé de l'architecture de la figure 3.1 : le mélangeur. Les premiers résultats sur les récepteurs directs [74-79] ont été validés sur des structures utilisant la jonction six-ports conventionnelle [41-46] ou la structure de la figure 3.4. Comme on l'a déjà expliqué, ces structures n'offrent aucune isolation naturelle entre le port RF et l'oscillateur local. Dans le cas d'un récepteur direct le transfert de la puissance de l'oscillateur local vers l'antenne de réception va générer un signal parasite qui va entraîner des perturbations pour les usagers utilisant la même bande. Les standards de communications imposent des limites sur ce signal, et l'amplificateur de faible bruit n'offre qu'une isolation limitée par rapport à celle demandée (l'isolation $S_{12} \sim -40$ dB comparé à -70 dB demandé typiquement par les standards). Dans notre cas, la présence d'un mélangeur peut apporter une isolation importante qui s'additionne à celle du transfert inverse de l'amplificateur faible bruit. La deuxième partie de ce chapitre va présenter les résultats relatifs la conception d'un mélangeur pour l'architecture POLYCOM.

La troisième partie de ce chapitre présentera une transition vers une ligne différentielle. Dans le chapitre 3, on a souligné que les écarts DC représentent un désavantage majeur

pour l'architecture directe et que le concepteur doit considérer une solution pour diminuer les effets néfastes de ce phénomène. Cette transition, conçue dans les plages de fréquences du système LMDS, doit être utilisée en conjonction avec les mélangeurs différentiels. L'oscillateur et l'amplificateur LNA peuvent être quant à eux conçus dans une technologie unipolaire (référéncée par rapport au plan de masse).

Nous allons conclure ce chapitre par quelques recommandations pour des projets futurs et nous allons faire le point sur l'ensemble du travail accompli dans le cadre de ce doctorat.

6.1 Caractéristiques du Transistor PHEMT.

Les installations de fabrication et de caractérisation étant en place (voir les chapitres IV et V), on a fabriqué une série de transistors en configuration "T" (voir figure 5.11). Les figures 6.1 et 6.2 montrent les images prises par le microscope électronique sur la configuration de grille.

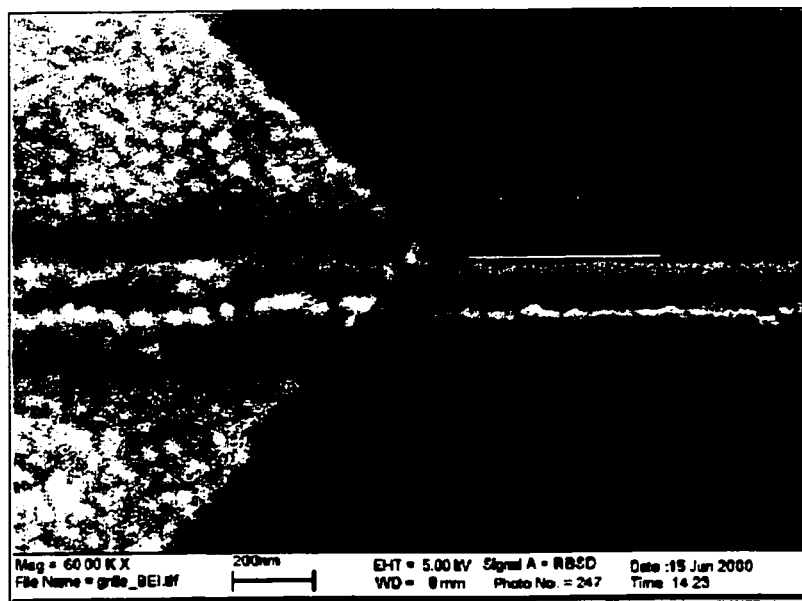


Figure 6.1 Grille du transistor PHEMT (~200 nm).

Comme on peut observer, on peut obtenir facilement des dimensions compatibles (longueur de la grille ~ 200 nm) avec la plage de fréquences de travail. La dimension transversale du transistor est de 100 microns et on s'attendait à des résultats similaires au niveau de courant de drain avec le transistor Fujitsu.

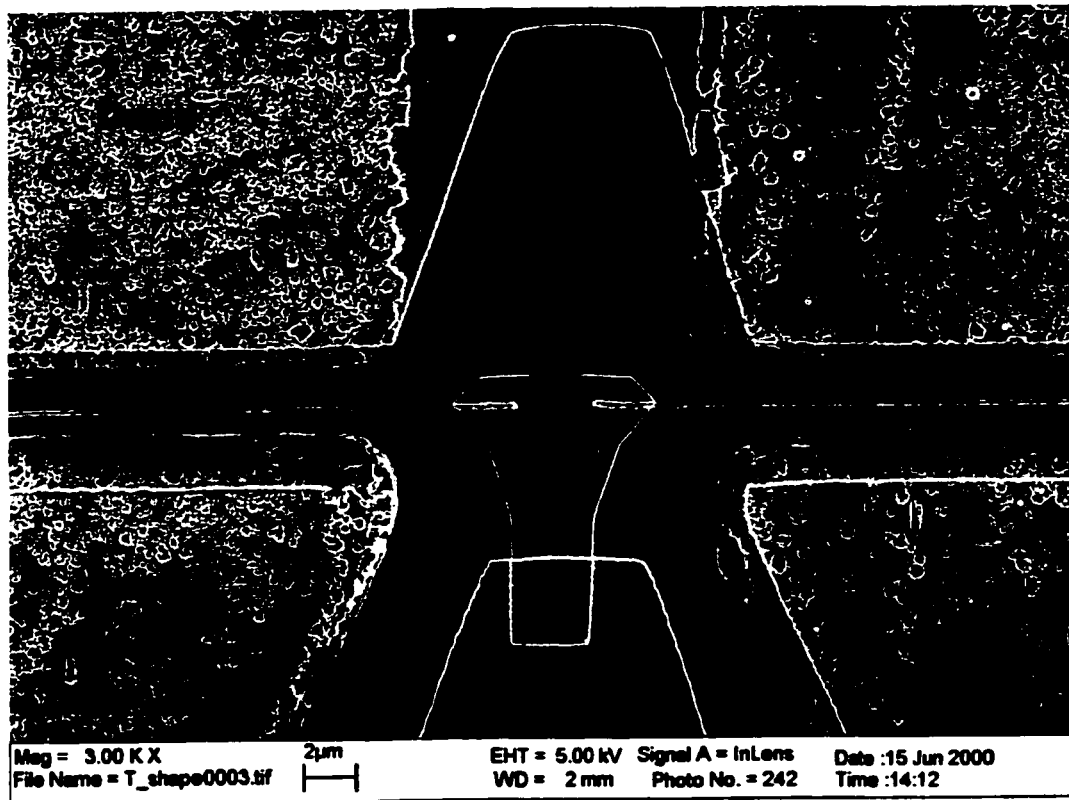


Figure 6.2 Transistor PHEMT (Vue du haut).

La texture non-uniforme de l'image 6.2 est due au recuit du contact ohmique. Le traitement en température, modifie l'aspect lisse de la surface métallique déposée par évaporation.

Sur la figure 6.3 on observe aussi un détail sur la connexion de grille, qui doit être faite sur une structure spéciale pour une bonne transition entre le substrat isolant de GaAs et la structure active. Avec une analyse plus attentive de la figure 6.2, on observe la différence de niveau. Pour franchir cette marche, on doit utiliser soit une structure de connexion plus épaisse, soit un pont-à-air. En raison de la dimension réduite de la grille, il est difficile de fabriquer avec une bonne tolérance le pont et pour notre première version nous avons préféré la variante plus robuste.

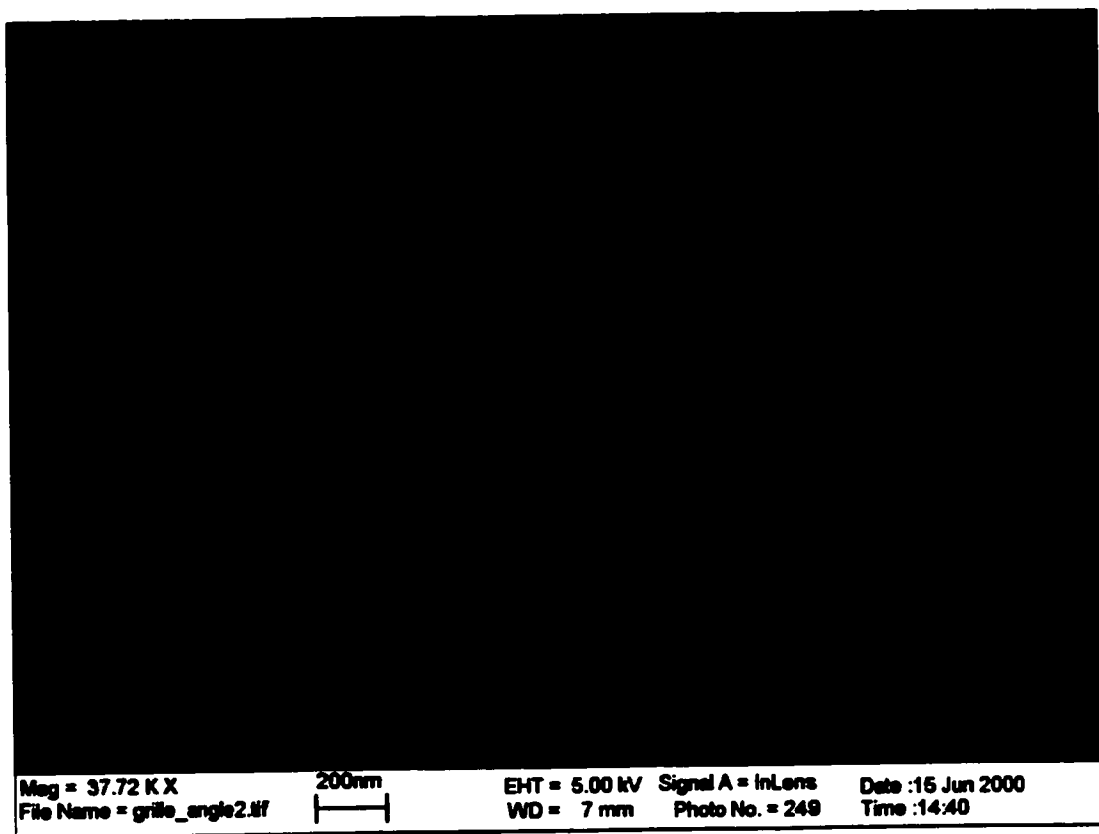


Figure 6.3 Détail de la grille

Les transistors fabriqués au laboratoire GRM ont été mesurés ensuite à PolyGrames pour l'extraction des paramètres électriques.

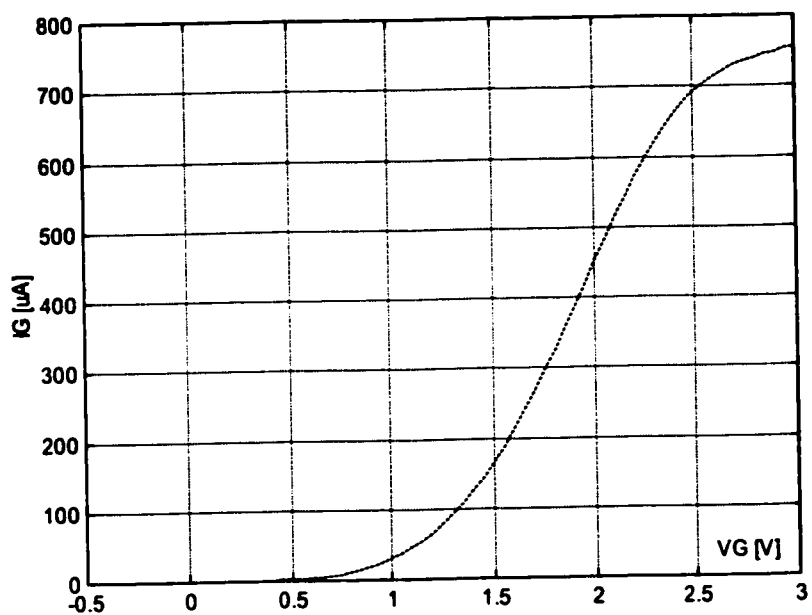


Figure 6.4 Courant de la grille en connexion diode.

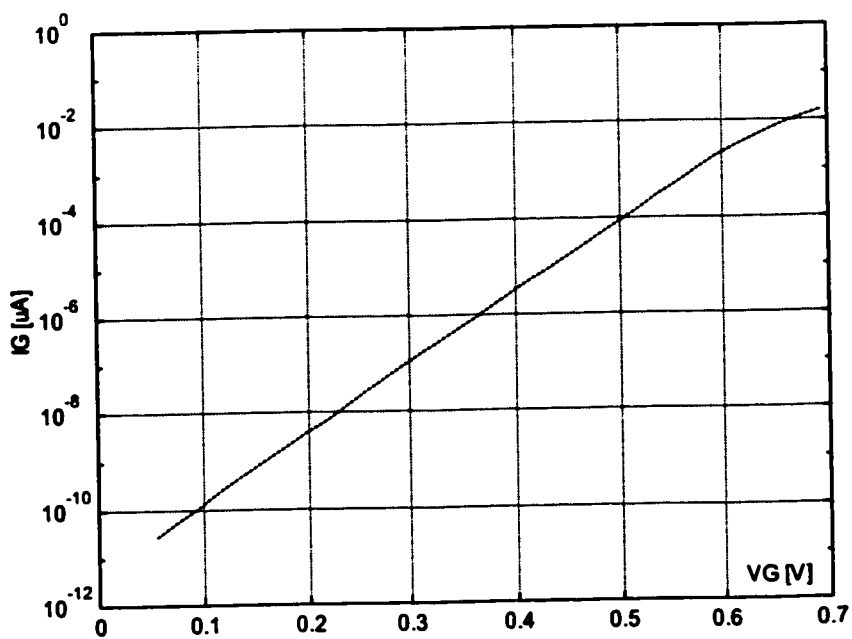


Figure 6.5 Courant de grille en représentation logarithmique

Le premier test a été la caractéristique courant-tension en connexion diode, qui nous a permis valider la présence du contact Schottky au niveau de la grille. En conséquence, le drain et la source ont été mises au potentiel zéro.

La figure 6.4 montre la variation de courant du dispositif et on peut observer la forme classique du courant d'une jonction. Une portion de cette caractéristique est montrée dans figure 6.5, à l'échelle logarithmique pour vérifier la linéarité de la dépendance exponentielle par rapport à la tension.

Un autre test montre dans la figure 6.6 la caractéristique de sortie I_D - V_D et on y retrouve des valeurs de courant compatibles avec la taille du transistor et les spécifications du matériel.

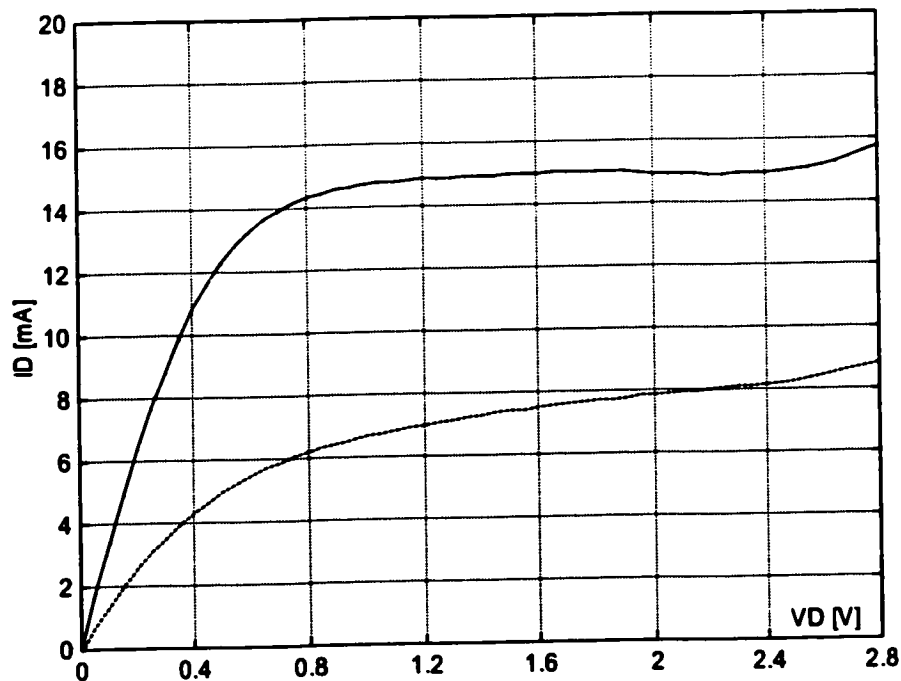


Figure 6.6 Caractéristique de sortie I_D - V_D pour deux tensions de grille.

Bien que les résultats soient très encourageants, quelques paramètres de fabrication doivent être modifiés pour l'obtention de meilleures performances. L'analyse du courant

de grille en connexion diode et la caractéristique de sortie montrent un très faible courant de grille pour la jonction en conduction directe. En plus, la valeur de tension de grille 'seuil' pour contrôler efficacement le courant de drain est relativement élevée. Ceci est probablement dû au fait que la gravure au niveau de la couche du niveau Schottky est incomplète, ce qui se traduit en pratique par un "faible" contact Schottky (tension de seuil très grande). Nous allons revenir sur cet aspect en analysant les résultats de caractérisation en fréquence.

Le dernier test en courant continu a été la mesure du courant de fuite. Les dimensions d'un transistor pour ondes millimétrique peuvent induire une valeur non-négligeable de ce courant de grille. En améliorant les paramètres de fabrication, le concepteur doit minimiser l'effet de la tension de drain sur le courant de grille (effet de contre-réaction). Le graphique 6.7 montre les valeurs de courant de grille en fonction de la tension de drain.

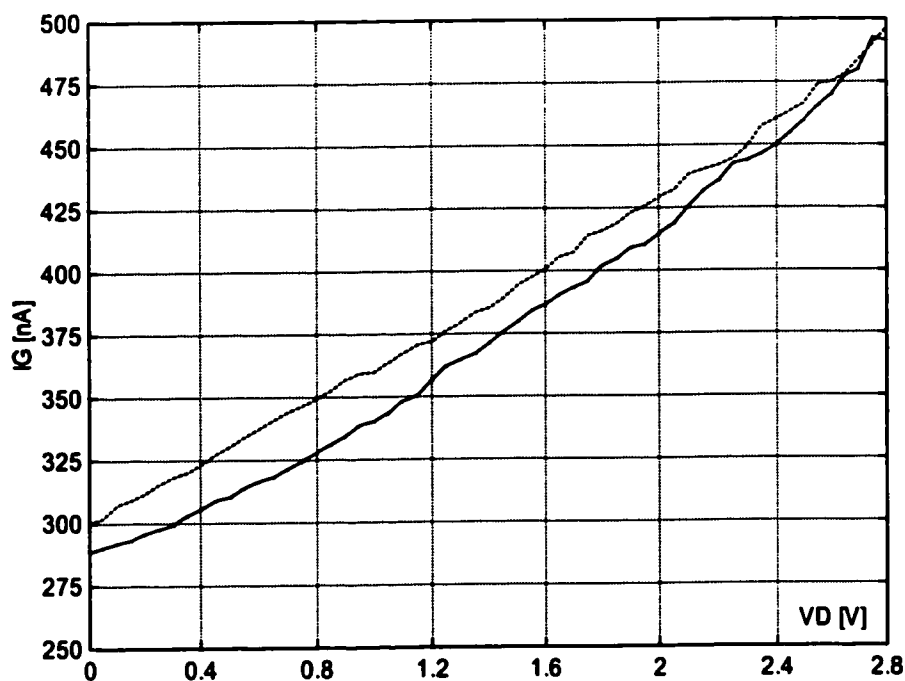


Figure 6.7 Courant de fuite.

Les valeurs sont largement négligeables dans notre cas, mais dans le processus d'amélioration du contact Schottky, on doit mesurer ce courant qui reste un paramètre à vérifier.

Soulignons que cette première variante fabriquée utilisait une grille en section rectangulaire pour simplifier la fabrication. En revanche, cette approche va induire une grande résistance de grille en raison de la section réduite de la grille. Pour des raisons liées au procédé de soulèvement, il est difficile de fabriquer une grille avec une largeur de section de 200 nm et avec une hauteur plus grande. Cette résistance parasite de grille va influencer directement le gain du transistor, son facteur de bruit et surtout sa fréquence de coupure.

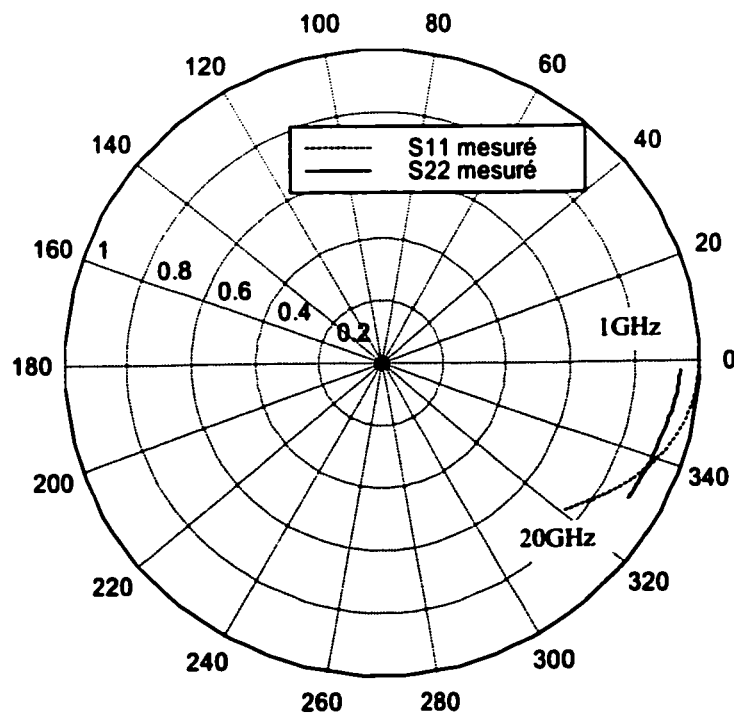


Figure 6.8 Transistor PHEMT en réflexion.

Nous avons effectué les mesures en paramètre S du transistor. La figure 6.8 montre les coefficients de réflexion et on peut observer que la résistance de grille est élevée (~24-26 ohms)

En raison de cette caractéristique notre transistor est stable; on a mesuré alors la variation de gain maximale disponible, G_{max} , en fonction de la fréquence (figure 6.9). On peut déduire à partir du graphique 6.9 la fréquence maximale d'oscillation, un facteur important pour un dispositif en ondes millimétriques.

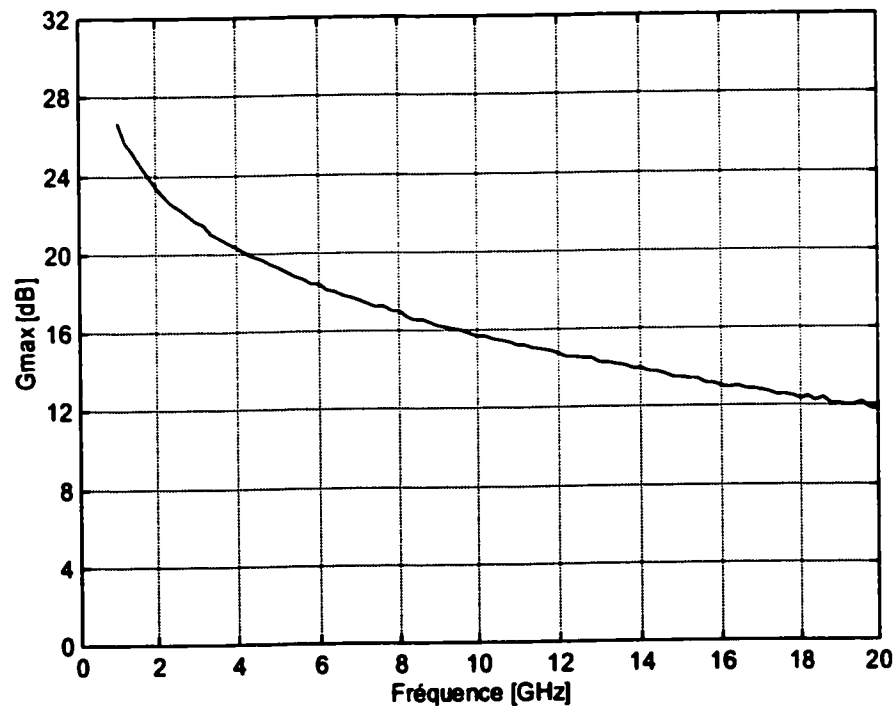


Figure 6.9 Gain Maximal pour le transistor

Si on fait l'extrapolation de cette courbe on trouve une fréquence maximale d'oscillation de ~60 GHz. Si on considère une variante dont on diminue la résistance de grille vers 3-4 ohms on obtient une fréquence $f_{max} \sim 120$ GHz et une fréquence de coupure f_c

d'environ 90 GHz, des valeurs standard pour le type de matériel qu'on a utilisé et une dimension de grille de 200 nm.

Ces premiers résultats, nous ont convaincu de la possibilité d'obtenir de bonnes performances pour un dispositif avec cette géométrie, et nous avons travaillé pour le développement du procédé de fabrication d'une grille champignon qui va avoir une résistance de grille très faible. Le procédé développé au laboratoire GRM utilise deux couches superposées de résine de sensibilité différentes. La lithographie par faisceau d'électrons se fait en deux temps: une première passe avec une forte dose électronique et une deuxième passe avec une dose plus faible.

Plusieurs tests ont été menés pour optimiser ce procédé, et les profils obtenus confirment la validité de cette méthode (voir figure 6.10).

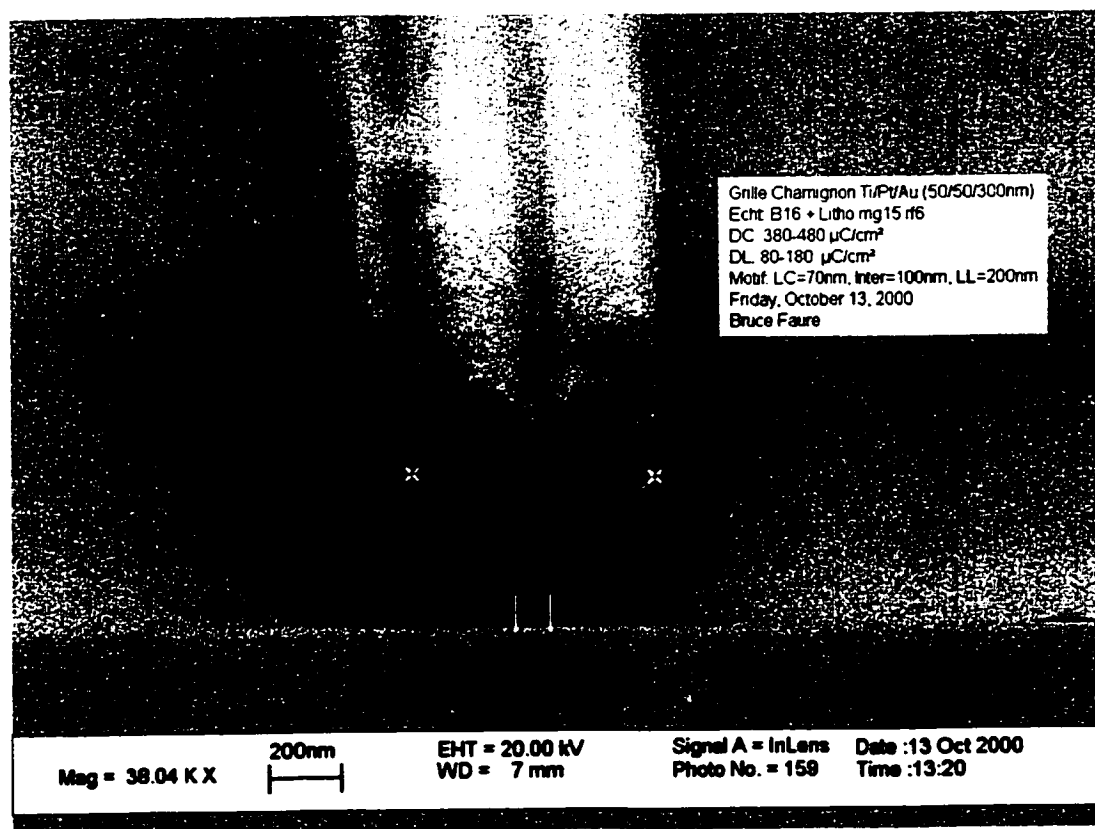


Figure 6.10 Profil de la grille champignon.

À la date d'écriture de ce manuscrit, le personnel du GRM poursuivait le travail d'amélioration de la procédure de gravure de la couche fortement dopée (temps d'immersion, procédure de nettoyage etc.) pour fabriquer des contacts Schottky performants ce qui permettra le contrôle adéquate de la tension de seuil par les paramètres technologiques.

6.2 Mélangeur à Transistor PHEMT.

Ayant présentée le processus de modélisation dans le chapitre précédent, la conception du mélangeur a utilisé le modèle du transistor Fujitsu, modèle extrait par le système de caractérisation. Ce modèle est inclus dans un fichier compatible avec le format du logiciel ADS[®].

Le but principal de notre approche a été de fournir une isolation suffisante entre l'oscillateur local et le port RF. En théorie, concevoir le mélangeur avec des transistors à effet de champs permet d'obtenir un gain positif, ce qui peut améliorer la sensibilité du récepteur. Cependant dans notre cas, on a voulu préserver deux paramètres importants, soit la linéarité et l'isolation, au détriment du gain. En conséquence on a accepté une certaine perte (~ 3 dB) mais celle-ci reste inférieure à celle des mélangeurs à diodes. Pour obtenir une isolation accrue, on a imposé l'utilisation d'un mélangeur sous-harmonique. La raison est très simple. Considérant le cas d'un récepteur homodyne il n'y a aucune façon d'augmenter l'isolation entre l'oscillateur local et le port RF parce que leur fréquence sont identiques et on peut uniquement obtenir l'isolation introduite par le mélangeur lui même (~ 15 dB). En revanche, pour une variante sous-harmonique, les valeurs de la fréquence de l'oscillateur local et du signal RF étant différentes (dans notre cas f_{LO} est la moitié de f_{RF}) on peut utiliser une structure simple de réjection. Pour la valeur de mélange (le deuxième harmonique) sa valeur est déjà largement atténuée par la composante.

La conception a utilisé une architecture proposée par Dubuc [40] mais nous l'avons modifié de façon importante (dans [40] il s'agit d'une analyse du mélangeur sur la fondamentale) pour l'utilisation en fonctionnement sous-harmonique.

Dans la figure 6.11 on voit la structure simplifiée du schéma du mélangeur. Le signal RF est appliqué dans la grille du transistor et l'oscillateur local fournit le signal pour le pompage de la source. Pratiquement, la grille contrôle la transconductance qui dans ce cas est une fonction non-linéaire de la puissance du signal de l'oscillateur local. Deux structures d'adaptation sont ajoutées dans la grille et la source pour obtenir les pertes minimales. Pour la conception de ces structures on a utilisé la méthode d'équilibrage harmonique ("harmonic balance") en conjonction avec une approche de type "load-pull" pour l'optimisation des pertes d'insertion.

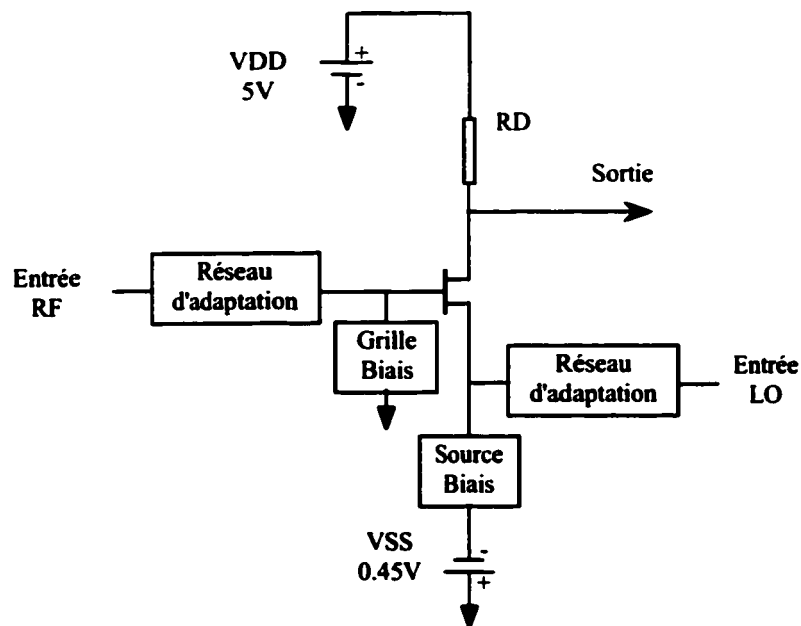


Figure 6.11 Mélangeur à transistor PHEMT.

L'analyse de fonctionnement de ce type de mélangeur peut se faire en regardant le rôle de chaque composante du schéma 6.11. Le circuit, pour être compatible avec la structure

utilisée dans la partie de modélisation à été monté sur un substrat de céramique (Alumine, $\epsilon_r = 9.9$) avec une épaisseur de 127 microns (5 mils) et le dessin utilise des lignes micro-ruban. La plage de fréquences de travail (13.5-14.5 GHz pour LO et 27-29 GHz pour RF) ne nécessite pas une épaisseur aussi faible (254 microns suffiraient) mais, nous avons utilisé cette valeur en raison de la compatibilité avec notre système de mesure sous-pointes. En pratique pour un récepteur intégré, on va utiliser le substrat de GaAs et certaines solutions adoptées pour la version hybride seront remplacées par des composantes localisées.

Les réseaux de polarisation, de source et de grille ont un double rôle. D'une part, ils induisent évidemment le découplage entre le circuit d'alimentation DC et les ports des signaux des hautes fréquences, en établissant le point de polarisation du transistor. La valeur de tensions de polarisations a été établie pour permettre l'excursion de courant sur toute la plage possible (à partir de la tension de seuil vers $IDSS$).

D'autre part, il consiste en l'élimination de la composante "parasite" (le signal sur la fréquence fondamentale du LO pour le signal RF et la fréquence de la porteuse pour la source). Dans notre version les circuits ont fait usage des lignes de transmission en circuit ouvert.

Pour une version intégrée, les mêmes circuits peuvent être réalisés par une solution très élégante et simple, suggérée dans la figure 6.12.

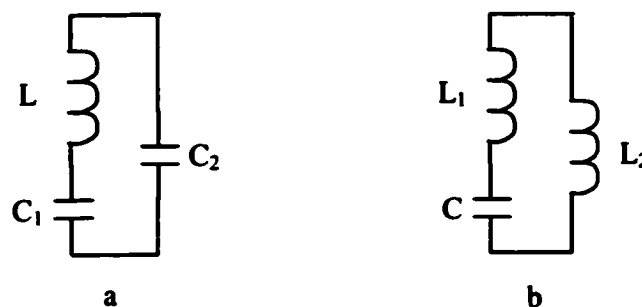


Figure 6.12 Réseaux biais pour la grille (a) et source (b).

Dans le cas de la variante (a), le circuit résonant série est un virtuel court-circuit à la fréquence de l'oscillateur local (ex. 14 GHz). Pour le signal RF (28GHz) son impédance est de type inductif et la capacité C_2 va assurer la résonance parallèle (impédance infinie). De façon analogue, pour la version (b), le circuit série est en court-circuit pour le RF et se comporte de façon capacitive pour LO. Avec l'inductance L_2 on assure la résonance parallèle.

Dans le cas d'un récepteur direct, on doit tenir compte du fait que la connexion externe doit être réalisée par un couplage DC. Si on regarde la structure du schéma 6.11, on observe que le potentiel du drain n'est pas nul. En conséquence, la résistance R_D influence aussi le point de polarisation du transistor est aussi le gain, étant impédance de charge pour la sortie. En pratique, le mélangeur est suivi par un étage d'amplification en bande de base. Dans ce cas on peut avoir 3 solutions possibles pour une meilleure architecture :

- On fait un compromis entre les tensions de polarisation du drain et de la source pour avoir un potentiel de sortie nul. Cette variante va avoir un écart DC donné par une variation des paramètres du transistor.
- On utilise un circuit pour le changement de niveau, comme un étage source commune ou émetteur commun ("level-shifting"). Aussi on présente le même écart DC, mais on peut polariser le circuit mélangeur avec de valeurs standard.
- On utilise une architecture différentielle (voir figure 6.18). Dans ce cas le niveau DC présent n'influence pas le transfert d'information, mais on doit veiller à ce qu'il ne limite pas la bande dynamique du signal. D'habitude l'excursion du signal à la sortie n'est pas plus grande que quelques centaines de mV, donc le potentiel existant ne pénalise pas trop la conception.

Pour notre circuit de test, on n'a pas utilisé un amplificateur après le mélangeur, donc on a testé le circuit avec un oscilloscope en sortie. Dans ce cas, le circuit de changement de

niveau est incorporé dans l'appareil de test. Cette configuration nous a obligé construire ainsi, une structure avec de lignes micro-ruban jouant le rôle de sonde de mesure. La raison est très simple. Pour un bon fonctionnement du mélangeur (pertes minimales) le circuit d'entrée de l'oscilloscope doit ne pas influencer l'impédance de charge optimale calculée par la méthode d'équilibrage harmonique. Le comportement en fréquence de cette sonde est présenté dans la figure 6.13 et on observe clairement que la sonde se comporte comme un circuit ouvert dans les plages de fréquences désignées (13.5-14.5 GHz respectivement 27-29 GHz).

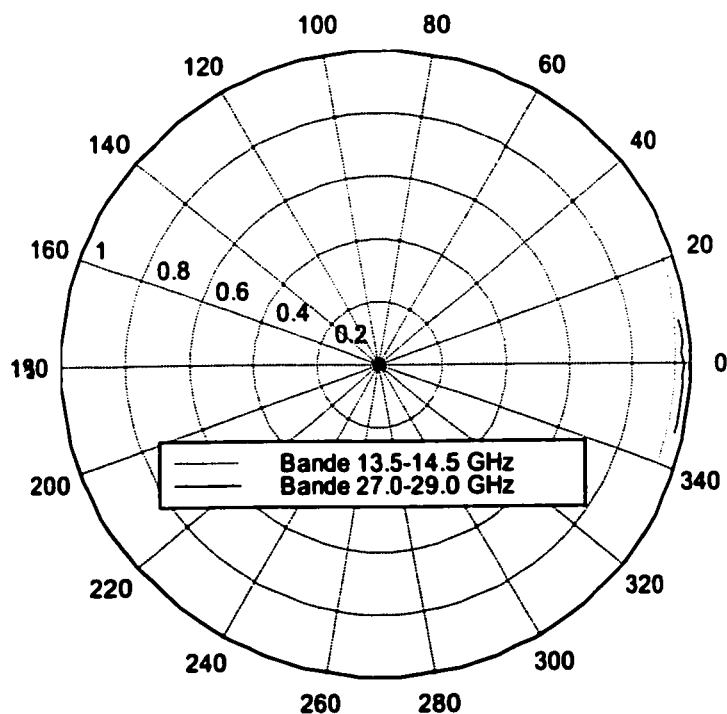


Figure 6.13 Impédance de la sonde.

En revanche le circuit de la sonde est un court circuit pour les basses fréquences (<200MHz). Parfois on peut satisfaire des conditions d'oscillation, surtout pour les transistors pHEMT qui présentent un haut gain. En conséquence, on doit surveiller le comportement et éventuellement ajouter une résistance de ballast dans la source.

L'optimisation à l'aide de la méthode d'équilibrage harmonique nous a permis de vérifier notre modèle extrait, d'une façon indirecte. Le contrôle du signal RF de la source non-linéaire peut être modélisé comme une série des sources de courant, chacune utilisant une fréquence de type $f_{RF} \pm mf_{LO}$. Notre mélangeur utilise la combinaison entre la deuxième harmonique du signal et le signal RF. Le produit d'intérêt est le signal dans la bande de base. Le minimum de pertes de conversion a été établi par de simulations dans le cas d'une terminaison purement résistive, avec une valeur $R_D \sim 400$ ohms. Le résultat est en parfaite concordance avec la modélisation de petit signal du chapitre précédent, où on a trouvé une valeur similaire pour la résistance de sortie du transistor. Dans ce cas la résistance de charge satisfait la condition de maximum de transfert de puissance.

La figure 6.14 montre la variation de la tension de grille en fonction du niveau de signal LO. Cette courbe nous a permis de choisir la valeur optimale de la puissance de pompage.

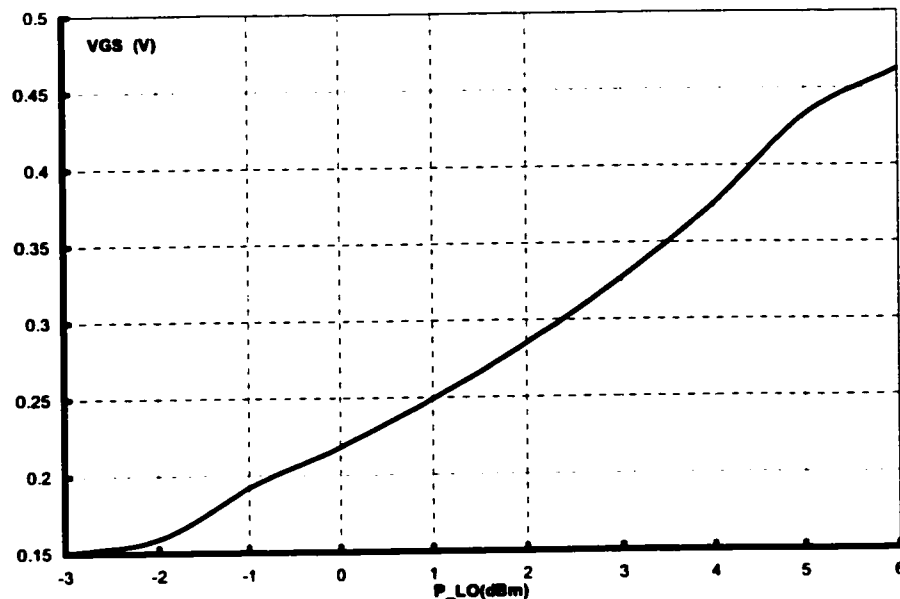


Figure 6.14 Plage de variation de la tension de grille

Comme on s'y attendait, plus on augmente le niveau de la puissance, plus la variation est grande. Pour le point de polarisation choisi, la valeur optimale se trouve entre 3 et 4 dBm, en raison de la tension de seuil de notre composante. Une variation plus grande va diminuer le gain du transistor. Ce phénomène a été validé par les mesures des pertes d'insertion. Dans le graphique suivant, on obtient le maximum de conversion (minimum de pertes d'insertion) pour un niveau près de 4 dBm.

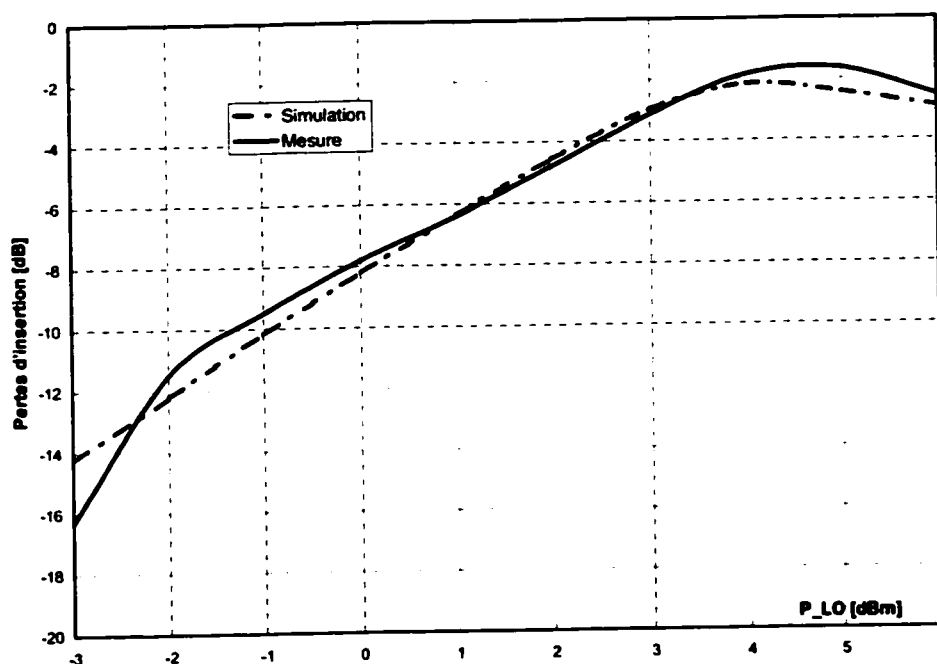


Figure 6.15 Pertes d'insertion

En réalité, les mesures sont un peu plus optimistes en raison de la résistance de charge. On a mesuré la tension de sortie et on a calculé la puissance avec la valeur optimale de 400 ohms. La valeur réelle était un peu plus grande en raison de la variation du procédé. On estime une diminution d'environ 1 dB et une perte globale d'environ 3 dB. La valeur n'est pas trop importante tant que le signal de sortie est suffisamment grand pour être amplifié. Dans notre cas, on obtient un signal d'environ 58 mV en amplitude pour un niveau RF de -20 dBm. Pour cette composante on a choisi un bon compromis pour la bande dynamique, en considérant une architecture dont le signal RF varie à l'entrée du

mélangeur entre -30 et -10 dBm. Si on essaie d'augmenter la sensibilité, on devrait utiliser un transistor avec un gain plus important, celui utilisé étant optimisé plutôt pour les applications faible bruit.

La figure 6.16 montre les pertes par réflexion simulées du mélangeur. Une mesure précise de ce paramètre nous était impossible car nous ne pouvions fabriquer avec précision une connexion répétable entre la sonde en structure coplanaire et la structure micro-ruban.

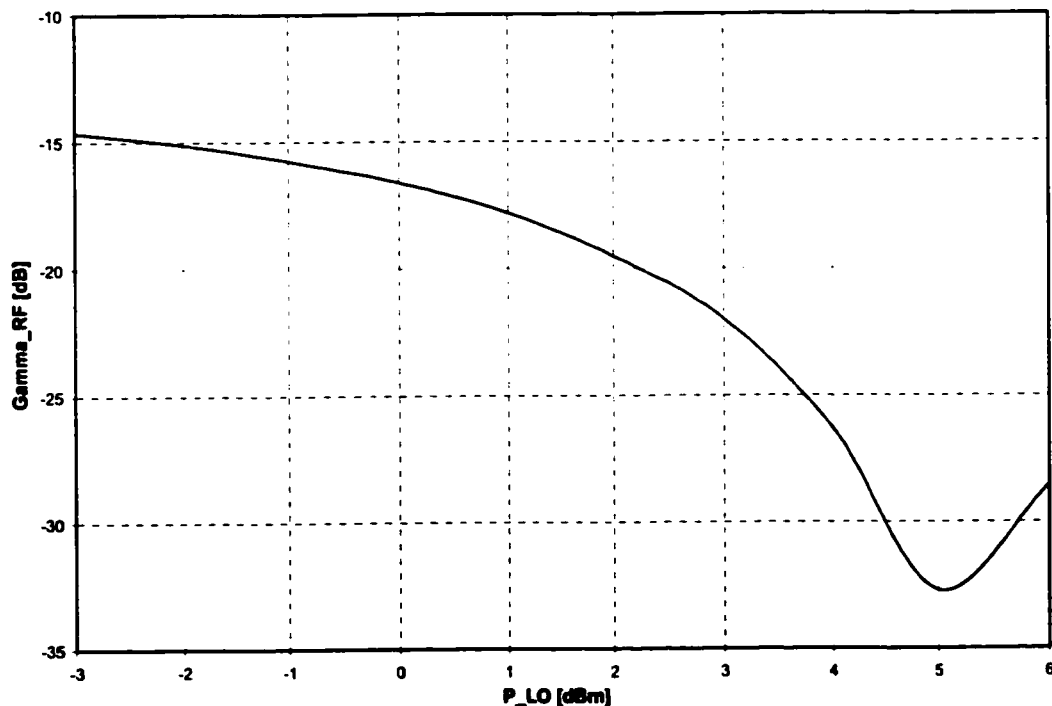


Figure 6.16 Pertes par réflexion du mélangeur.

La présence des trous métallisés nous empêche d'avoir une précision suffisante. Pour une puissance de 3 dBm, qui assure les pertes minimales de conversion, la valeur de perte par réflexion pour le signal RF que nous avons mesurée était meilleure que -12 dBm, dans la bande passante du mélangeur.

L'isolation calculée du notre dispositif est présentée dans la figure 6.17 pour une puissance du signal de l'oscillateur local optimale de 3 dBm. On s'attendait à un bon résultat pour ce type de mélangeur. La fondamentale de l'oscillateur local est largement atténuée par une structure de rejection dont l'influence sur le signal RF est minimale, en raison de la différence importante de fréquence.

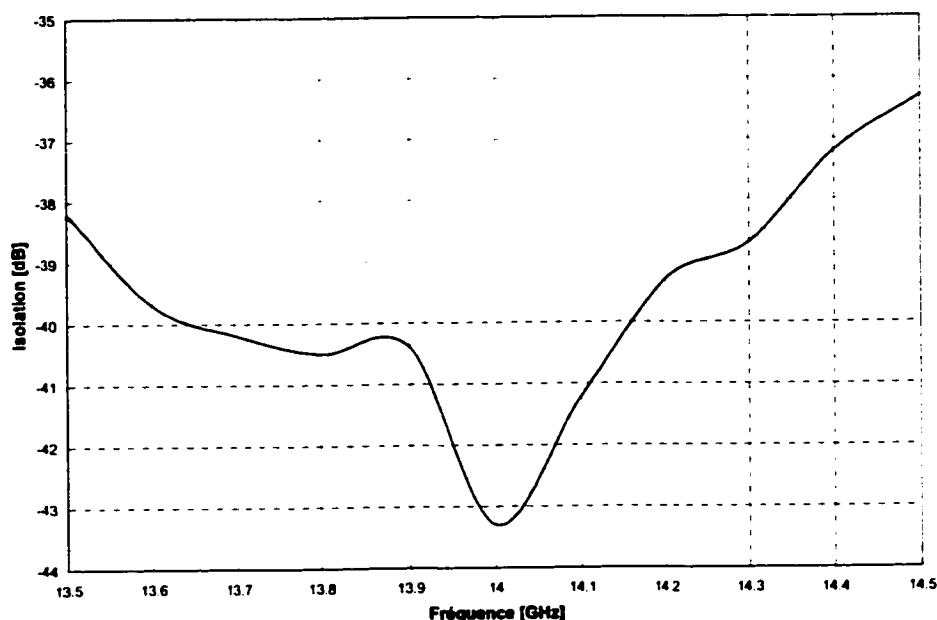


Figure 6.17 L'isolation LO-RF.

Les résultats présentés ci-dessus montre la viabilité de cette architecture de mélangeur, ainsi que ses avantages au niveau de fonctionnement d'un récepteur direct. Bien que le circuit physique ait été réalisé en technologie MHMIC, la structure est identique pour une version intégrée avec un gain important au niveau d'espace utilisée. En fait cette architecture tire ses avantages de la technologie MMIC qui permet de réaliser des éléments concentrés avec des tolérances contrôlées.

Un autre aspect qui mérite d'être souligné est la possibilité d'inclure dans le circuit mélangeur un étage tampon ("buffer") qui sert aussi comme un circuit d'amplification et

surtout comme un changeur de niveau d'impédance (la sortie est vue comme une source de tension).

6.3 Transition CPWFG à CPS.

Dans le cadre de l'optimisation de notre mélangeur pour le fonctionnement du récepteur direct, on doit minimiser les écarts DC à l'entrée du système d'échantillonnage (convertisseurs analogique-numériques). Comme après l'étage de conversion de fréquence, il y a seulement des amplificateurs vidéo et un filtre passe-bas, l'unique source importante d'écarts DC est la connexion entre le mélangeur et l'amplificateur vidéo. À défaut d'un circuit plus complexe de correction de la composante continue parasite, l'utilisation des structures différentielles reste la solution la plus appropriée pour éliminer cet effet. La figure 6.18 montre de façon schématique cette connexion.

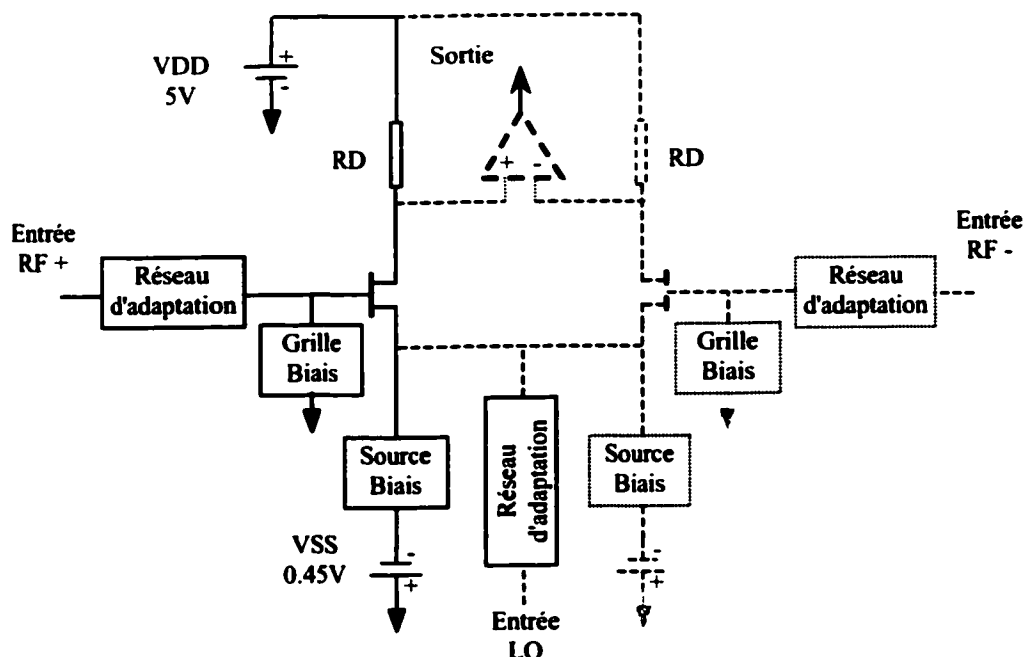


Figure 6.18 Mélangeur différentiel.

On peut considérer la structure différentielle soit pour le signal RF (figure ci-dessus) soit pour le signal LO.

Indépendamment de notre choix, nous avons besoin d'un transfert entre une structure de transmission unipolaire (ex. ligne coplanaire) et une structure bipolaire ou balancée (ex. rubans coplanaires CPS). Nous avons choisi un type spécial de symétriseur ("BALUN") pour obtenir de bonnes propriétés pour notre application. Les symétriseurs en "Y" [66,67] forment une classe spéciale de dispositif à large bande. Leur principal avantage est le petit gabarit et dans notre plage de fréquence (28 GHz), ils peuvent être intégrés à moindre coût. La figure 6.19 montre la structure réalisée (a) et son modèle électrique équivalent (b).

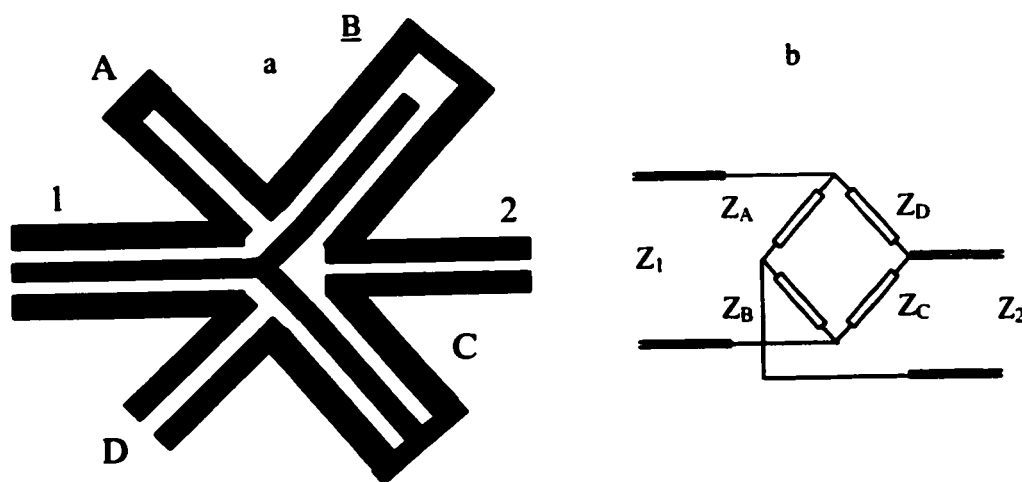


Figure 6.19 Transition CPWFG vers CPS – dessin (a) et modèle électrique (b)

Si on considère le modèle électrique de la figure 6.y (b), l'impédance d'entrée peut être trouvée comme étant :

$$Z_{in} = \frac{1}{Z'_A + Z'_B + 1} \left[Z'_A Z'_B + \frac{(Z'_A + Z'_B)(Z'_A + 1)(Z'_B + 1)}{(Z'_A + Z'_B + 2)} \right], \quad 6.1$$

si on considère que $Z_A=Z_C=Z_3$ et $Z_B=Z_D=Z_4$ et toutes les impédances sont normalisées par rapport à la valeur Z_2 ($Z_n = Z_n/Z_2$).

Dans ce cas la condition d'adaptation parfaite ($Z_{in}=1$) peut se réduire à :

$$Z_3'Z_4' = 1 \text{ ou } \Gamma_3 = -\Gamma_4, \text{ dans le cas des coefficients de réflexion} \quad 6.2$$

L'expression 6.2 montre qu'un court-circuit et un circuit ouvert peuvent satisfaire la condition d'adaptation parfaite. La qualité de l'adaptation, respectivement les pertes d'insertion de cette structure sont préservées tant que les imperfections du court-circuit et circuit ouvert réels sont négligeables. En théorie, tant qu'on garde la même longueur électrique pour les circuits, on obtient une excellente adaptation, indépendant de la fréquence. En pratique, la bande est limitée par l'influence des bras de réactances ("stubs") et une bonne approximation est gardée seulement si la longueur électrique est inférieure à $\lambda_G/8$.

Pour calculer la longueur des courts-circuits et circuits ouverts, en respectant les conditions des équations 6.1 et 6.2, nous avons utilisé des formules de synthèse à partir d'une approche quasi-stationnaire, modifiée pour inclure la dispersion et la variation de l'impédance caractéristique avec la fréquence. Ces formules sont déduites à l'aide de la théorie de transformations conformes pour les deux types de géométrie (coplanaire avec plan de masse fini et ruban coplanaire) Étant donné que les logiciels de simulation n'ont pas de modèles valables pour ce type des lignes, il a fallu utiliser une méthode de calcul numérique du champ électromagnétique (la méthode de moments)

Les mesures ont validé notre conception. Dans les figures 6.20 et 6.21, on peut observer la différence entre les valeurs simulées et celles mesurées pour les paramètres S de la structure. Pour les mesures nous avons utilisé une structure 'bout-à-bout' ("back-to-back") en raison du fait que la structure CPS ne peut pas être mesurée avec des probes coplanaire. Pour cette raison, dans la figure 6.20, les valeurs S11 et S22 proviennent en

réalité du coefficient de réflexion du côté CPWFG de la transition. Les faibles différences entre les deux courbes sont explicables par la variation de fabrication entre les ponts-à-air des deux transitions.

Nous avons utilisé un circuit hybride réalisé au laboratoire PolyGrames et les ponts-à-air sont réalisés par des fils soudés. C'est un procédé qui n'est pas très reproductible. Indépendamment de cette variance, les pertes par réflexion obtenus sont minimales (inférieurs à -20 dB dans la bande passante). D'un point de vue plus général, ce type de transition a montré une bande passante relativement large.

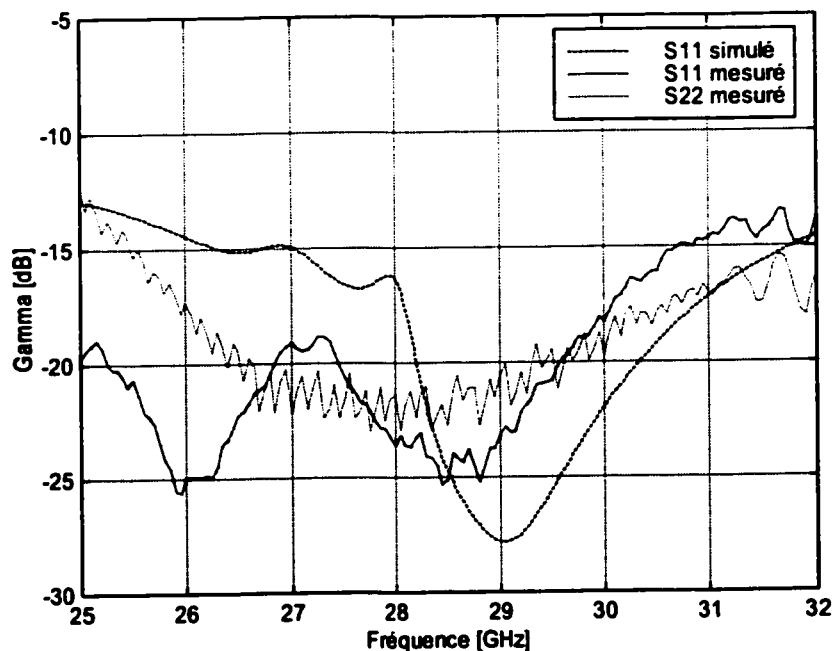


Figure 6.20 Pertes par réflexion de la transition CPWFG-CPS.

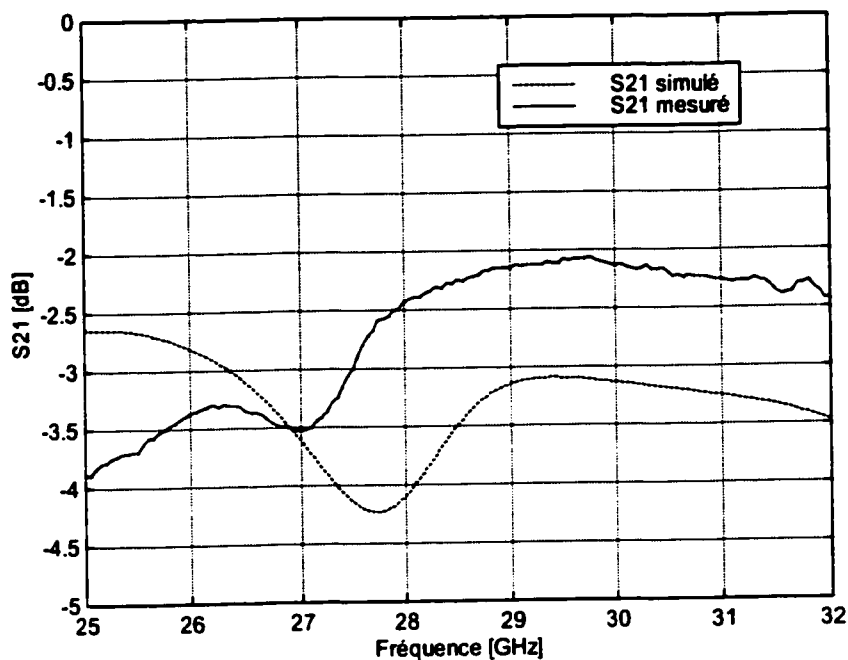


Figure 6.21 Pertes d'insertion pour la transition CPWFG-CPS-CPWFG

Pour le cas des pertes d'insertion, on doit faire une remarque importante. La structure simulée est un peu plus longue par rapport à celle fabriquée, ce qui fait que le résultat de simulation obtenu montre une plus grande atténuation par rapport à la valeur mesurée. La raison de cette limitation dans la simulation provient du mode de calcul de logiciel MOMENTUM. Pour ce type de structure, on n'a pas la possibilité de faire une calibration au niveau des ports d'entrée. Dans ce cas, au niveau du port excité (structure coplanaire) on peut retrouver le mode impair de propagation. Le logiciel de simulation fonctionne seulement pour le cas d'une transmission unimodale. La présence des deux modes (paire et impair) perturbe la solution et le logiciel affiche alors un terme d'erreur de calibration. La solution est de prolonger la structure pour que ce mode soit atténué par la présence du ponts-à-air. Finalement les valeurs mesurées montrent une perte d'insertion de seulement 1,25 dB (2,5 dB bout au bout) à 28 GHz. Cette valeur peut probablement être réduite dans le cas d'un procédé de fabrication multicouches (comme

procédés MMIC), qui permet la réalisation de ponts-à-air avec une meilleure précision (ex. circuits intégrés).

6.4 Technologie Homodyne. Évaluation et Synthèse

Nous avons décrit, dans les chapitres précédents les procédés (expérimentaux et analytiques) qui nous ont conduit vers une approche intégrée au niveau de la conception d'un système de récepteur homodyne. Les composantes obtenues respectent les critères énoncés au chapitre 2 (voir page 36) dans le but de mettre en place un récepteur direct utilisant l'architecture Polycom (figure 3.1)

Au niveau de performances physiques de la technologie PHEMT, nous avons vu que la fréquence maximale d'oscillation de notre transistor (environ 100 GHz) est en ligne avec la technologie similaire existante. En conséquence, nous pouvons utiliser les dispositifs avec succès dans la bande de fréquences des micro-ondes et ondes millimétriques basses (2-40 GHz).

Évidemment, le choix d'utilisation d'une technologie de type PHEMT (GaAs ou InP), HBT ou SiGe est surtout influencer des facteurs de coût et de fiabilité. Il n'en reste pas moins que nos résultats au niveau de gain sont suffisants pour assurer un bon fonctionnement des amplificateurs et mélangeurs à base de cette technologie. De plus, les procédures d'extraction de modèles sont facilement utilisables aussi pour d'autres technologies de circuits intégrés, sans modifications majeures (plage de polarisation etc.)

Par rapport aux récepteurs homodyne de type six-ports, nous avons introduit un mélangeur sous-harmonique qui nous permet intégrer entièrement la partie analogique d'ensemble transmetteur - récepteur. En utilisant les transistors PHEMT, nous pouvons

réaliser aussi l'amplificateur faible bruit (LNA) et l'amplificateur de puissance (étage final dans l'architecture de transmission) dans la même technologie.

Les résultats expérimentaux sur les pertes de conversion et la puissance de l'oscillateur local -3 dB et 5 dBm respectivement montrent des meilleures performances que les mélangeurs résistifs à base de pont à diodes (-10 à -12 dB des pertes pour la version sous-harmonique).

Le fonctionnement sous-harmonique nous a permis d'améliorer aussi l'isolation entre le signal de l'oscillateur locale et le port d'entrée. Cette amélioration est un avantage significatif dans le cas d'un récepteur direct, pour éviter le rayonnement du signal parasite dans la bande passante de la communication. Aussi, la caractéristique linéaire de gain de conversion nous assure d'une variation minimale de phase dans la plage de fréquences, variation qui est une des sources d'erreurs analysées dans le chapitre 2 et 3.

En ajoutant une structure BALUN, notre circuit peut être réalisé dans une architecture différentielle, ayant comme principal avantage la diminution significative de la composante continue parasite. Dans ce cas, un de principaux effets dangereux dans la technologie homodyne est largement atténué, préservant ainsi la bande dynamique du récepteur. Cette structure a des dimensions relatives petites et elle est une bonne solution pour une version intégrée de récepteur homodyne, étant aussi compatible avec les nouvelles techniques de réception fondées sur les antennes planaires.

Une réalisation complètement intégrée de la tête réceptrice (front-end) nécessite aussi d'autres éléments comme les capacités, les inductances et les résistances déposés. Nous avons réalisé physiquement ces composantes et nous avons développé les modèles électriques, utilisable pour la conception des circuits. Les résultats sont comparables avec les technologies utilisées en pratique pour la fabrication des circuits pour les récepteurs actuels (soit dans l'architecture homodyne ou super-hétérodyne). Les facteurs de qualité mesurés (pour les capacités et les inductances) permettent aussi un fonctionnement dans le domaine des ondes millimétrique (voir figures 5.4 et 5.8).

Sur le plan analytique, l'algorithme adaptatif proposée augmente de façon significative les performances du système, mettant en place une technique de compensation robuste et flexible. La possibilité d'utiliser la même approche indépendamment du type de modulation utilisé nous a permis de faire une distinction entre la procédure de compensation et la modulation. De plus, nous avons montré que l'algorithme peut aussi fonctionner avec une séquence décimée des symboles, pouvant être utilisé dans le cas des systèmes à haut débit. Une structure avec échantillonnage direct qui n'utilise pas une approche homodyne sera limitée par la vitesse des convertisseurs analogique-numérique (voir les performances des analyseurs de spectre utilisant la technique de transformation Fourier).

Les résultats du chapitre 3 sur la performance de notre algorithme par rapport à un système de type I&Q ont démontré la supériorité de la nouvelle architecture homodyne. Nous avons montré les tests sur la capacité de correction et nous avons vu que même des erreurs de 15 degrés et 25% respectivement pour la phase et le gain différentiel, sont éliminées par le réseau des neurones. Dans ce contexte, l'analyse de notre mélangeur nous indique qu'on pourra utiliser nos composantes avec succès.

Les procédés de fabrication utilisés font usage des techniques de lithographie, soulèvement et déposition par évaporation simple et qui garantissent un faible coût de fabrication. En ajoutant la structure directe qui élimine le besoin de l'étage de la fréquence intermédiaire, nous observons rapidement le potentiel de notre approche.

6.5 Conclusions et Travaux Futurs

Le projet présenté tout au long de cette thèse s'inscrit dans la série de développements apportés à l'architecture directe de récepteurs, la topologie favorite ces derniers temps dans les recherches pour l'amélioration de têtes réceptrice intégrées.

La principale nouveauté présentée dans le cadre de ce travail reste le développement de l'algorithme de démodulation adaptatif, une méthode extrêmement puissante pour la correction d'imperfections physiques et qui rend la tête réceptrice indépendante du système de modulation numérique utilisé. Cet algorithme à base de réseaux de neurones est complètement différent des approches de calibration proposées et implantées jusqu'à maintenant, approches qui demandaient l'utilisation de composantes supplémentaires ou une correction préalable du circuit.

Une validation de la théorie conçue a été présentée et le lecteur a pu voir la facilité d'implémentation d'une telle procédure. Cette méthode a mis en évidence la spécificité de notre application (récepteur homodyne) par rapport aux architectures basées sur la jonction six-ports, dispositif qui du point de vue chronologique a introduit le concept d'architecture directe au Centre de Recherche PolyGrames. Cette spécificité nous a conduit à éliminer le canal de référence et à des modifications au niveau des éléments mélangeurs (les détecteurs de puissance ont été remplacés par des mélangeurs dans l'architecture Polycom).

L'analyse globale de l'architecture directe nous a permis de voir comment les désavantages de cette topologie peuvent être contournés et nous pensons que le gain important au niveau du coût de fabrication permettra à ce type de récepteurs de prendre place, de plus en plus, parmi les radios commerciales.

Un deuxième volet de notre recherche a porté sur le développement d'un procédé de fabrication MMIC dans nos laboratoires et en même temps, sur la mise en place d'outils de caractérisation des composantes qui seront utilisées dans le futur pour la conception des circuits. À notre connaissance c'est la première fois, dans le milieu universitaire canadien, qu'on présente de tels résultats de caractérisation sur une technologie capable de fonctionner dans les bandes millimétriques.

Nous avons vu comment a été bâti le système de mesure automatique sous-pointes, autour du logiciel ICCAP. Les algorithmes de calibration et d'extraction des paramètres présentés dans le cadre de ce thèse ont été validés sur des produits commerciaux (ex. transistor faible bruit Fujitsu FHR20X). Des modèles de composantes passives et actives ont été validés par le système de mesure, modèles qui sont compatibles avec des simulateurs commerciaux (ex. ADS™).

En amont des résultats montrés ici, un effort considérable a été fait pour mettre au point les procédures de fabrication dans les deux laboratoires concernés (PolyGrames et GRM). Étant donné qu'au début de ces travaux certains appareils nécessaires à la fabrication ou aux expériences n'étaient pas disponibles, nous avons mis au point tous les détails de fabrication de masques (niveau d'exposition de laser, temps de développement de la résine etc.) et les recettes pour les différentes étapes de fabrication des circuits MMIC : la gravure MESA, le soulèvement pour les couches métalliques épaisses, l'électroplaquage, les doses d'exposition pour la lithographie par faisceau d'électrons, etc. Ce travail a été nécessaire en raison de l'impossibilité d'accéder rapidement et à des coûts modérés aux services des fonderies commerciales. Dans de telles conditions, notre travail, et celui du personnel du GRM a mis à la disposition des deux laboratoires une facilité de fabrication avec un temps de réponse ("turn-around") très court (~ deux semaines pour un circuit complet). Même s'il reste beaucoup des progrès à faire pour l'amélioration des performances des composantes fabriquées, il n'en reste pas moins qu'une expertise de premier ordre a été acquise et il existe maintenant la possibilité de faire des recherches plus poussées, tant au niveau d'une fréquence de travail plus élevée qu'au niveau de l'amélioration de caractéristiques.

Nous avons vu comment le système de mesure et de caractérisation a permis l'extraction du modèle du transistor. L'étape de la conception d'un mélangeur utilisable dans notre architecture a alors suivi. On voit alors comment le cercle de développement peut être fermé avec l'analyse et la mise en pratique des nouveaux concepts.

Notre projet montre aussi quelles avenues peuvent être explorées pour une amélioration significative des résultats. Les premières sont évidentes, elles consistent en l'amélioration du procédé de fabrication des transistors PHEMT (l'étape de "RECESS") et en la conception et l'intégration de chaque sous-système dans l'architecture PolyCom.

De plus, on peut suggérer de faire une étude analytique d'une procédure de synchronisation du signal d'oscillateur local, à l'aide d'un algorithme intégré à notre méthode de correction d'erreurs. Une idée pour aborder ce problème pourrait être la suivante : la rotation de la constellation (effet de la différence de phase instantanée entre la porteuse RF et le signal de l'oscillateur local) est facilement détectable et corrigée au niveau du traitement de signal. Cette technique possède l'avantage de réduire les coûts dans le cas d'une architecture de synchronisation très complexe. L'idée fondamentale est que dans l'espace cartésien de la constellation, n'importe quelle rotation est une opération linéaire et qu'elle peut être détectée par une technique similaire à celle utilisée dans la synchronisation du code dans le cas des communications de type DS-SS ("Direct Sequence Spread Spectrum") utilisé pour les systèmes CDMA.

Enfin, un sujet qui mérite une investigation particulière est la modélisation au niveau de comportement ("behavioral models") pour les transistors PHEMT. Dans les références [110, 111] certains aspects au niveau d'une caractérisation de type "boîte-noire" sont étudiés. La description d'un tel modèle demande aussi des mesures non-linéaires (effectuées sous un nouveau système de mesure de AGILENT™) Par rapport aux techniques classiques basés sur l'utilisation des paramètres S, le système de mesure non-linéaires nous permet d'avoir une approche de type 'forme d'onde' au lieu d'essayer d'améliorer la corrélation entre les valeurs mesurées et la fonction analytique décrivant certain paramètres du modèle. Avec des méthodes de modélisation dans lesquelles les dispositifs sont décrits dans le domaine fréquentiel (disponibles dans les simulateurs MDS/ADS) on pourrait décrire la réponse d'un transistor plutôt en termes des formes

d'ondes. Une telle approche évitera les limites des modèles électriques basés sur l'utilisation des composantes concentrées et de plus, elle peut s'avérer mieux adaptée dans le domaine des ondes millimétriques où les phénomènes distribués jouent un rôle important.

BIBLIOGRAPHIE

- [1] ABE, M., SASHO, N., BRANKOVIC, V., KRUPCEVIC, D. (2000). Direct conversion receiver based on six-port technology. 2000 European Conference on Wireless Technology, 5-6 Oct., Conference Proceedings, 139-142.
- [2] ABIDI, A.A. (1995). Direct-Conversion Radio Transceivers for Digital Communications. IEEE Journal of Solid-State Circuits, Vol.30, No.12, 1399-1410.
- [3] AKOS, D.M., TSUI, J.B.Y. (1996). Design and implementation of a direct digitization GPS receiver front end. IEEE Transaction on Microwave Theory and Techniques, Vol. 44, No. 12, 2334-2339.
- [4] AN, H. (1993). Broadband Active and Passive Microstrip Antennas. Thèse de doctorat, Université Catholique de Leuven, Belgique.
- [5] ANGELOV, I., BENGTTSSON, L., GARCIA, M. (1996). Extension of the Chalmers Nonlinear HEMT and MESFET Model. IEEE Transactions on Microwave Theory and Techniques, Vol. 44, No. 10, 1664-1674.
- [6] ANGELOV, I., ZIRATH, H., RORSMAN, N. (1992). A new empirical nonlinear model for HEMT and MESFET devices. IEEE Transactions on Microwave Theory and Techniques, Vol. 40, No. 12, 2258-2266.
- [7] ANHOLT, R. (1995). Electrical and Thermal Characterization of MESFETs, HEMTs and HBTs. Artech House, Boston.
- [8] ANHOLT, R. Swirhun S. (1991). Equivalent-circuit parameter extraction for cold GaAs MESFET's. IEEE Transactions on Microwave Theory and Techniques, Vol. 39, No. 7, 1243-1247.
- [9] ANHOLT, R. Swirhun S. (1991). Measurement and analysis of GaAs MESFET parasitic capacitances. IEEE Transactions on Microwave Theory and Techniques, Vol. 39, No. 7, 1247-1251.
- [10] ARNOLD, E., GOLIO, M., MILLER, M., BECKWITH, B. (1990). Direct extraction of GaAs MESFET intrinsic element and parasitic inductance values. 1990 IEEE MTT-S International Microwave Symposium Digest. Vol.1, 359-362.

- [11] BAEYENS, Y., (1997). Monolithic Microwave Integrated Circuits Using GaAs and InP based Heterojunction Field-Effect Transistors. Thèse de doctorat, Université Catholique de Leuven, Belgique.
- [12] BARTA, G.S., JPNES, K.E., HERRICK, C.G., STRID, E.W. (1986). Surface-mounted GaAs active splitter and attenuator MMIC's used in a 1-10 GHz leveling loop. IEEE Transactions on Microwave Theory and Techniques, Vol. 34, No. 12, 1569-1575.
- [13] BASU, S., HAYDEN, L. (1997). An SOLR calibration for accurate measurement of orthogonal on-wafer DUTs. 1997 IEEE MTT-S International Microwave Symposium Digest, Vol. 3, 1335-1338.
- [14] BENGTTSSON, L., GARCIA, M., ANGELOV, I. (1995). An extraction Program for nonlinear transistor model parameters for HEMTS and MESFETs. Microwave Journal, Vol. 37, No.1, 146-153.
- [15] BERROTH, M., BOSCH, R. (1991). High-frequency equivalent circuit of GaAs Fet's for arge signal applications. IEEE Transactions on Microwave Theory and Techniques, Vol. 39, No. 2, 224-229.
- [16] BILBRO, G.L., STEER, M.B., TREW, R.J., CHANG, C-R., SKAGGS, S.G. (1990). Extraction of the parameters of equivalent circuits of microwave transistors using tree annealing. IEEE Transactions on Microwave Theory and Techniques, Vol. 38, No. 11, 1711-1718.
- [17] CARLIN, H.J. (1977). A new approach to gain-bandwidth problems. IEEE Transactions on Circuits and Systems, Vol. 24, No. 4, 170-175.
- [18] CARLIN, H.J., KOMIAK, J.J. (1979). A new method of broad-band equalization applied to microwave amplifiers. IEEE Transactions on Microwave Theory and Techniques, Vol. 27, No. 2, 93-99.
- [19] CARLIN, H.J., YARMAN, B.S. (1983). The double matching problem: analytic and real frequency solutions. IEEE Transactions on Circuits and Systems, Vol. 30, No. 1, 15-28.

- [20] CARPENTER G.A., GROSSBERG, S. (1987). A massively Parallel Architecture for a self-organizing Neural Pattern recognition machine. Computer Vision, Graphics and Image Processing, vol. 37, 54-115.
- [21] CARPENTER G.A., GROSSBERG, S. (1987). ART 2:self-organization of stable category recognition codes for analog input patterns. Applied Optics Vol. 26, No. 23, 4919-4930.
- [22] CARPENTER G.A., GROSSBERG, S. (1988). The ART of adaptive pattern recognition by a self-organizing neural network. IEEE Computer, 77-88.
- [23] CARPENTER G.A., GROSSBERG, S. (1990). Adaptive resonance theory: neural network architectures for self-organizing pattern recognition. Intl. Conference on Parallel Processing in Neural Systems and Computers, 383-389.
- [24] CARPENTER G.A., GROSSBERG, S. ROSEN, D.B.(1991). ART 2-A: an adaptive resonance algorithm for rapid category learning and recognition. Neural Networks, Vol. 4, 493-495.
- [25] CAVERS, J.K., LIAO, M.W. (1993). Adaptive Compensation for Imbalance and Offset Losses in Direct conversion Transceivers. IEEE Transactions on Vehicular Technology, Vol.42, No.4, 581-588.
- [26] CHANG, K. (1991). Handbook of Microwave and Optical Components. John Wiley, New York.
- [27] CHANG, K. (1994). Microwave Solid-State Circuits and Applications. John Wiley, New York.
- [28] CHANG, K. (2000). RF and Microwave Wireless Systems. Wiley, New York.
- [29] CHO, H., BURK, D. (1991). A three-step method for the de-embedding of high frequency S-parameter measurements. IEEE Transactions on Electron Devices, Vol. 38, No. 6 1371-1375.
- [30] CHOU, S.Y., ANTONIADIS, D.A. (1987). Relationship between measured and intrinsic transconductances of FET's. IEEE Transactions on Electron Devices, Vol. 34, No.2, 448-450.

- [31] CHOW, Y.L., FENG, N.N., FANG, D.G. (1999). A simple method for ohmic loss in conductors with cross-section dimensions on the order of skin depth. Microwave and Optical technology Letters, Vol. 20, No.5, 302-304.
- [32] COLEF, G., KARMEI, P.R., ETTEBERG, M. (1990). New in-situ calibration of diode detectors used in six-port network analyzers. IEEE Transaction on Instrumentation and Measurement, Vol. 39, No.1, 201-204.
- [33] COSTA, J.C., MILLER, M., GOLIO, M., NORRIS, G. (1992). Fast, accurate, on-wafer extraction of parasitic resistance and inductances in GaAs MESFETs and HEMTs. 1992 IEEE MTT-S International Microwave Symposium Digest. Vol.2, 1011-1014.
- [34] COUPEZ, J.P., PERICHON, R.A. (1988). A novel microwave transmission phase-shifter. 18th European Microwave Conference, Stockholm, Sweden, 1023-1027.
- [35] CURTICE, W.R. (1980). A MESFET model for use in the design of GaAs integrated circuits. IEEE Transactions on Microwave Theory and Techniques, Vol. 28, No. 5, 448-456.
- [36] CURTICE, W.R., ETTEBERG, M. (1985). A nonlinear GaAs FET model for use in the design of output circuits for power amplifiers. IEEE Transactions on Microwave Theory and Techniques, Vol. 33, No. 12, 1383-1393.
- [37] DAMBRINE, G., CAPPY, A., HELIODORE, F., PLAYEZ, E. (1988). A new method for determining the FET small-signal equivalent circuit. IEEE Transactions on Microwave Theory and Techniques, Vol. 36, No. 7, 1151-1159.
- [38] DEGROOT, D.C.; WALKER, D.K.; MARKS, R.B. (1996). Impedance mismatch effects on propagation constant measurements. Electrical Performance of Electronic Packaging Conference Proceedings, Napa, CA, USA; 28-30 Oct., 141-143.
- [39] Di MARTINO, A., PISA, S., TOMMASINO, P., TRIFILETTI, A. (1999). A new algorithm to extract the nonlinear model of MESFETs and HEMTs, Microwave and Optical technology Letters, Vol. 20, No.5, 295-297.

- [40] DUBUC, D., PARRA, T., GRAFFEUIL, J. (2000). Original topology of GaAs-PHEMT mixer. 30th European Microwave Conference-Paris 2000, Conference Proceedings, Vol. 1, 165-168.
- [41] ENGEN, G.F. (1977). An improved circuit for implementing the six-port technique of microwave measurements IEEE Transaction on Microwave Theory and Techniques, Vol. 25, No. 12, 1080-1083.
- [42] ENGEN, G.F. (1977). The six-port reflectometer: an alternative network analyzer. IEEE Transaction on Microwave Theory and Techniques, Vol. 25, No. 12, 1075-1080.
- [43] ENGEN, G.F. (1978). Calibrating the six-port reflectometer by means of sliding terminations. IEEE Transaction on Microwave Theory and Techniques, Vol. 26, No. 12, 951-957.
- [44] ENGEN, G.F. (1978). The six-port measurement technique. A Status Report. Microwave Journal, Vol.20, Vol. 5, 20-22,24,84,87-89.
- [45] ENGEN, G.F., HOER, C.A. (1972). Application of an arbitrary 6-port junction to power-measurement problems. IEEE Transaction on Instrumentation and Measurement, Vol. 21, No.4, 470-473.
- [46] ENGEN, G.F., HOER, C.A. (1979). "Thru-Reflect-Line": an improved technique for calibrating the dual six-port automatic network analyzer. IEEE Transactions on Microwave Theory and Techniques, Vol. 27, No. 12, 987-993.
- [47] FEHER, K. (1995). Wireless Digital Communications : Modulation and Spread Spectrum Applications. Prentice-Hall, Upper Saddle River, N.J.
- [48] FEHER, K. (1997). Advanced Digital Communications : Systems and Signal Processing Techniques. Noble, Atlanta.
- [49] GHANNOUCHI, F.M., BOSISIO, R.G. (1991). A wideband millimetric wave six-port reflectometer using four diode detectors calibrated without a power ratio standard. IEEE Transaction on Instrumentation and Measurement, Vol. 40, No.6, 1043-1046.

- [50] GHAZINOUR, A., JANSEN, R.H. (1998). Robust, model-independent generation of intrinsic characteristics and multi-bias parameter extraction for MEFETS/HEMTs. 1998 IEEE MTT-S International Microwave Symposium Digest, Vol. 1, 149-152.
- [51] GOLIO, M. (1991). Microwave MESFETs and HEMTs. Artech House, Boston.
- [52] GONZALES, G. (1997). Microwave Transistor Amplifiers. Analysis and Design. Prentice-Hall, Upper Saddle River, N.J.
- [53] GRAMMER, W., YNGVESSON, K.S. (1993). Complanar waveguide transitions to slotline: design and microprobe characterization. IEEE Transactions on Microwave Theory and Techniques, Vol. 41, No. 9, 1653-1658.
- [54] HALCHIN, D., MILLER, M., GOLIO, M., TEHRANI, S. (1994). HEMT models for large signal circuit simulation. 1994 IEEE MTT-S International Microwave Symposium Digest, Vol.3, 985-988.
- [55] HESSELBARTH, J., WIEDMANN, F., HUYART, B. (1996). New structures for six-port reflectometers covering very large bandwidths. IEEE Instrumentation and Measurement Technology Conference, Brussels Belgium, Juin 4-6, 1263-1268.
- [56] HEWLETT PACKARD (1997). HP IC-CAP Modeling Software, Ver. 5.0, Manuel d'utilisateur.
- [57] HEWLETT PACKARD. (1999). IC-CAP Modeling Reference, VI. 1-3.
- [58] HINDSON, D., HUANG, X., de LÉSELEUC, M., CARON, M. (1998). Preliminary performance measurements on a 20 GHz five-port direct receiver. Fourth Ka Band Utilization Conference, Venice, Italy, Conference Proceedings, 449-455.
- [59] HODGETTS, T.E., GRIFFIN, E.J. (1983) A unified treatment of the theory of six-port reflectometer calibration using the minimum of standards. Report 83003 RSRE Malvern, Août., 1-40.
- [60] HOER, C.A., (1972). The six-port Coupler: a new approach to measuring voltage, current, power, impedance, and phase. IEEE Transaction on Instrumentation and Measurement, Vol. 21, No.4, 466-470.

- [61] HOER, C.A., ROE, K.C., MCKAY ALLRED, C. (1976). Measuring and minimizing diode detector nonlinearity. IEEE Transaction on Instrumentation and Measurement, Vol. 25, No.4, 324-328.
- [62] HOLMSTROM, R.P., BLOSS, W.L., CHI, J.Y. (1986). A gate probe method of determining parasitic resistance in MESFET's. IEEE Electron Device Letters, Vol. 7, No.7, 410-412.
- [63] HUANG, X., CARON, M., HINDSON, D. (1999). Adaptive I/Q-regeneration in 5-port junction based direct receivers. Proceedings of APCC/OECC'99 5th Asia Pacific Conference on Communications 4th Optoelectronics and Communications Conference, Vol.1, 717-720.
- [64] HUANG, X., HINDSON, D., de LÉSÉLEUC, M., CARON, M. (1999). I/Q-channel regeneration in 5-port junction based direct receivers. 1999 IEEE MTT-S International Topical Symposium on Technologies for Wireless Applications, Vancouver, BC, Canada; 21-24 Feb., Conference Proceedings, 169-173.
- [65] JANEZIC, M.D., JARGON, J.A. (1999). Complex permittivity Determination from propagation constant measurements. IEEE Microwave and Guided Letters, Vol. 9, No.2, 76-78.
- [66] JOKANOVIC, B., TRIFUNOVIC, V. (1996). Star mixer with high port-to-port isolation. Electronic Letters, Vol.32, No.24, 2251-2252.
- [67] JOKANOVIC, B., TRIFUNOVIC, V. (1997). Review of printed baluns: characteristics and applications. TELSIKS'97, Nis, Yugoslavia, 8-10 Octobre, Conference Proceedings, 287-298.
- [68] JUDAH, S.R., WRIGHT, A.S. (1990). A dual six-port automatic network analyzer incorporating a biphas-bimodulation element. IEEE Transaction on Microwave Theory and Techniques, Vol. 38, No. 3, 1083-1085.
- [69] KOMPA, G. (1998). Reliable extraction of small-signal elements of a generalized distributed FET model. 1998 IEEE MTT-S International Microwave Symposium Digest, Vol. 1, 291-294.

- [70] KUMAR, S.; SHARMA, R.; HARSH R. (1998). Microwave characterisation of GaAs based MMIC passive components. Physics of Semiconductor Devices Conference, Delhi, India; 16-20 Dec. 1997, Proceedings, Vol.2, 914-917.
- [71] KWON, Y., DEAKIN, D.S., SOVERO, E.A., HIGGINS, J.A. (1996). High performance Ka-band monolithic low-noise amplifiers using 0.2- μ m dry-recessed GaAs PHEMT's. IEEE Microwave and Guided Letters, Vol.6, No.7, 253-255.
- [72] KWON, Y., PAVLIDIS, D., BROCK, T.L., STREIT, D.C. (1993). A D-band monolithic fundamental oscillator using InP-based HEMT's. IEEE Transactions on Microwave Theory and Techniques, Vol. 41, No. 12, 2336-2344
- [73] LENK, F., DOERNER, R. (1998). New Extraction method for FET extrinsic capacitances using active bias conditions. 1998 IEEE MTT-S International Microwave Symposium Digest, Vol. 1, 279-282.
- [74] LI, J., BOSISIO, R.G., WU, K. (1994). A collision avoidance radar using six-port phase/shift discriminator (SPFD). 1994 IEEE MTT-Symposium Digest, 1553-1556.
- [75] LI, J., BOSISIO, R.G., WU, K. (1995). A new direct digital receiver performing coherent PSK reception. 1995 IEEE MTT-Symposium Digest, 1007-1010.
- [76] LI, J., BOSISIO, R.G., WU, K. (1995). Computer and measurement simulation of a new digital receiver operating directly at millimetric-wave frequencies. IEEE Transaction on Microwave Theory and Techniques, Vol. 43, No. 12, 2766-2772.
- [77] LI, J., BOSISIO, R.G., WU, K. (1995). Dual-tone calibration of six-port junction and its application to the six-port direct digital millimetric receiver. IEEE Transaction on Microwave Theory and Techniques, Vol. 44, No. 1, 93-99.
- [78] LI, J., BOSISIO, R.G., WU, K. (1995). Modeling of the six-port discriminator used in a microwave direct digital receiver. 1995 Canadian Conference on Electrical and Computer Engineering Digest, 1164-1165.
- [79] LIU, Y. (1995). Calibrating an industrial microwave six-port instrument using the artificial neural network technique. IEEE Transaction on Instrumentation and Measurement, Vol. 25, No.2, 651-656.

- [80] LOPER, R.K. (1990). A tri-phase direct conversion receiver. MILCOM 90. 1990 IEEE Military Communications Conference. Conference Vol.3, 1228-1232.
- [81] LORD, A.J. (1999). Comparing the Accuracy and repeatability of on-wafer calibration techniques to 110 GHz. Cascade Microtech Inc. Notes d'applications.
- [82] MARKS, R.B. (1991). A multiline method of network analyzer calibration. IEEE Transaction on Microwave Theory and Techniques, Vol. 39, No. 7, 1205-1215.
- [83] MARKS, R.B. (1991). Characteristic impedance determination using propagation constant measurement. IEEE Microwave and Guided Letters, Vol. 1, No.6, 141-143.
- [84] MARKS, R.B., JARGON, J.A., PAO, C.K., WEN, C.P. (1995). Microwave characterization of flip-chip MMIC interconnections. 1995 IEEE MTT-S International Symposium Digest Vol. 3, 1463-1466.
- [85] MARKS, R.B., JARGON, J.A., PAO, C.K., WEN, C.P. (1995). Microwave characterization of flip-chip MMIC components. 45th Electronic Components and Technology Conference, Las Vegas, NV, USA; 21-24 Mai, Proceedings , 343-350.
- [86] MARSHALL, C.B. (1986). Low-power integrable paging receiver architecture. IEE Proceedings, Vol. 133, Pt. F, No.5, 449-455.
- [87] MATERKA, A., KACPRZAK, T. (1985) Computer Calculation of Large-Signal GaAs FET amplifier characteristics. IEEE Transactions on Microwave Theory and Techniques, Vol. 33, No. 2, 129-135.
- [88] McDERMOTT, M.G., SWEENEY, C., BENEDEK, M., DAWE, G. (1988). Monolithic Ka band VCO using quarter micron GaAs MESFETs and integrated high-Q varactors. 1990 IEEE Microwave and Millimetric-Wave Monolithic Circuits Symposium, Conference Proceedings, 103-106.
- [89] MEHROTRA, K., MOHAM, C.K., RANKA, S. (1997). Elements of Artificial Neural Networks. MIT Press, Cambridge, Ma.
- [90] MELLOR, D.J., LINVILL, J.G. (1975). Synthesis of interstage network of prescribed gain over frequency slopes. IEEE Transactions on Microwave Theory and Techniques, Vol. 23, No. 12, 1013-1020.

- [91] MILLER, M., GOLIO, M., BECKWITH, B., ARNOLD, E. (1990) Choosing an optimum large signal model for GaAs MESFETs and HEMTs. 1990 IEEE MTT-S International Microwave Symposium Digest. Vol.3, 1279-1282.
- [92] MONDAL, J.P. (1987). An experimental verification of a simple distributed model of MIM capacitors for MMIC applications. IEEE Transactions on Microwave Theory and Techniques, Vol. 35, No. 4, 403-408.
- [93] POIRÉ, P., GHANNOUCHI, F.M. (1996). Immunity of six-port measurement techniques to the AM-PM noise of the microwave signal generator. IEEE Instrumentation and Measurement Technology Conference, Brussels Belgium, Juin 4-6, 1259-1262.
- [94] POIRÉ, P. (1999). Mesures Load-Pull Multi-Harmonique avec Forme d'Onde et Application à la Conception d'Amplificateurs en Classe F. Thèse de Doctorat, École Polytechnique de Montréal, Canada.
- [95] PROAKIS, J.G. (1995). Digital Communications. McGraw-Hill, New York.
- [96] PROAKIS, J.G., MANOLAKIS, D.G. (1996). Digital Signal Processing: Principles Algorithms and Applications. Prentice-Hall, Upper Saddle River, N.J.
- [97] PROBERT, P.J., CARROLL, J.E. (1982). Design features of multi-port reflectometers. IEEE Proceedings, Vol. 129 Pt. II, No. 5, 245-251.
- [98] RAZAVI, B. (1997). Design Considerations for Direct-Conversion Receivers. IEEE transactions on Circuits and Systems-II: Analog and Digital Signal Processing. Vol. 44, No.6, 428-435.
- [99] REEVES, G.K., HARRISON, H.B. (1982). Obtaining the specific contact resistance from transmission line model measurements. IEEE Electron Device Letters, Vol. 3, No.5, 111-113.
- [100] REEVES, G.K., HARRISON, H.B. (1995). A tri-layer transmission line model applied to alloyed ohmic contacts. Solid-State Electronics, Vol. 38, No. 5, 1115-1117.

- [101] REEVES, G.K., LEECH, P.W., HARRISON, H.B. (1995). Understanding the sheet resistance parameter of alloyed ohmic contacts using a transmission line model. Solid-State Electronics, Vol. 38, No. 4, 745-751.
- [102] RIBLET, G.P., HANSSON, E.R.B. (1984). Some properties of the matched symmetrical six-port junction. IEEE Transaction on Microwave Theory and Techniques, Vol. 32, No. 2, 164-170.
- [103] ROHDE, U.L. (1997). Microwave and Wireless Synthesizers : Theory and Design. John Wiley, New York.
- [104] ROHDE, U.L., NEWKIRK, D.P. (2000). RF/Microwave Circuit Design for Wireless Applications. John Wiley, New York.
- [105] ROOT, D.E., FAN, S. (1992). Experimental evaluation of large-signal modeling, assumptions based on vector analysis of bias-dependent S-parameter data from MESFETs and HEMTs. 1992 IEEE MTT-S International Microwave Symposium Digest, Vol. 1, 255-258.
- [106] ROOT, D.E., PIROLA, M., FAN,S., ANKLAM, W.J., COGNATA, A. (1993). Measurement-based large-signal diode modeling system for circuit and device design. IEEE Transactions on Microwave Theory and Techniques, Vol. 41, No. 12, 2211-2216.
- [107] RUTENBAR, R.B. (1989). Simulated annealing algorithms: an overview. IEEE Circuits and Devices Magazine, Vol.5, No.1, 19-26.
- [108] SADHIR, V., BAHL, I.J., WILLEMS, D.A (1994). CAD compatible accurate models of microwave passive lumped elements for MMIC applications. International Journal of Microwave and Millimetric-Wave Computer-Aided Engineering, Vol.4, No.2, 148-162.
- [109] SCHREURS, D., VERSPECHT, J., VANDERBERGHE, S., CARCHON, G., Van der ZANDEN, K., NAUWELAERS, B. (1999). Easy and accurate empirical transistor model parameter estimation from vectorial large-signal measurements. 1999 IEEE MTT-S International Microwave Symposium Digest Vol.2, 753-756.

- [110] SCHREURS, D., VISAN, T., VANDERBERGHE, S., Van MEER, H., Van der ZANDEN, K., NAUWELAERS, B., BOSISIO, R.G. (1999). IM3 suppression using a technology independent method based on vectorial large-signal measurements. 53rd ARFTG Conference Digest, 82-88.
- [111] SCHREURS, D., VISAN, T., VANDERBERGHE, S., Van MEER, H., Van der ZANDEN, K., NAUWELAERS, B., BOSISIO, R.G. (1999). Power amplifier linearisation using a straightforward technology independent method based on vectorial large-signal measurements. 29th European Microwave Conference 99. Incorporating MIOP '99. Conference Proceedings Vol.2, 76-79.
- [112] SHIH, F.Y., MOH, J. (1990). Improved adaptive resonance theory. Proceedings of SPIE Vol. 1382 Intelligent Robots and Computer Vision IX: Neural, Biological and 3-D Methods, 26-36.
- [113] SHIRAKAWA, K., OIKAWA, H., SHIMURA, T., KAWASAKI, Y., OHASHI, Y., SAITO.T. (1995). An approach to determining an equivalent circuit for HEMT's. IEEE Transactions on Microwave Theory and Techniques, Vol. 43, No. 3, 499-503.
- [114] SOARES, R.A., GOUZIEN, P., LEGAUD, P., FOLLOT, G. (1989). A unified mathematical approach to two-port calibration techniques and some applications. IEEE Transactions on Microwave Theory and Techniques, Vol. 37, No. 11, 1669-1674.
- [115] SOLOMON, M.N., McCLAY, C.P., CRONSON H.M., CHU, A., WEITZMAN P.S. (1993). A monolithic six-port module for built-in-test applications. 1993 IEEE Instrumentation and Measurement Technology Conference, Irvine, CA, USA; 18-20 Mai, 123-126.
- [116] STATZ, H., NEWMAN, P., SMITH, I.W., PUCCEL, R.A., HAUS, H.A. (1987). GaAs FET device and circuit simulation in SPICE. IEEE Transactions on Electron Devices, Vol. 34, No.2, 160-169.

- [117] STAUDINGER, J., MILLER, M., BECKWITH, B., HALCHIN, D. (1991). An accurate HEMT large signal model usable in SPICE simulators. 1991 IEEE MTT-S International Microwave Symposium Digest. Vol.1, 99-102.
- [118] TRIPATHI, V.K. (1975). Asymetric coupled transmission lines in an inhomogenous medium. IEEE Transactions on Microwave Theory and Techniques, Vol. 23, No. 9, 734-739.
- [119] VAI, M-K., PRASAD, S., LI, N.C., KAI, F. (1989). Modeling of microwave semiconductor devices using simulated annealing optimization. IEEE Transactions on Electron Devices, Vol. 36, No.4, 761-762.
- [120] VENDELIN, G., PAVIO, A.M., ROHDE, U.L. (1990). Microwave Circuit Design Using Linear and Nonlinear Techniques. John Wiley, New York.
- [121] VISAN, T., BEAUVAIS J., XU, Y., JACQUET P., BOSISIO R.G. (2000). Advances in Integrated Six Port Direct Digital Millimetre Wave Receivers. 24th Workshop on Compound Semiconductor Devices and Integrated Circuits WOCSDICE 2000 Mai 29-Juin 02, 2000 Aegean Sea, Greece, Conference Proceedings.
- [122] VISAN, T., BOSISIO R.G. (1999). Six-Port Technology (SPT) in RF Microwave and Optical Sensors. DYCON'99 Août 5-7 1999 Ottawa, Canada, Conference Proceedings.
- [123] VISAN, T., BOSISIO, R.G., BEAUVAIS, J. (2000). A new approach to calibrate the six-port direct digital millimetric receivers (DDMR). Progress in Electromagnetics Research Symposium, Juillet 5-14 2000 Cambridge, Ma. USA, Conference Proceedings, 880-881.
- [124] VISAN, T., BOSISIO, R.G., BEAUVAIS, J. (2000). New phase & gain imbalance correction algorithm for six-port based direct digital millimetric receivers. Microwave and Optical technology Letters, Vol. 27, No.6, 432-438.
- [125] WALKER, N.G., CARROLL, J.E., (1984). Simultaneous phase and amplitude measurements on optical signals using a multiport junction. Electronics Letters, Vol. 20, November 8th, No.23, 981-983.

- [126] WANG, H., DOW, G.S., ALLEN, B.R., TON, T.-N., TAN, K.L., CHANG, K.W., CHEN T.-H., BERENZ, J., LIN, T.S. (1992). High-performance W-band monolithic pseudomorphic InGaAs HEMT LNA's and design/analysis methodology. IEEE Transactions on Microwave Theory and Techniques, Vol. 40, No. 3, 417-428.
- [127] WEYDMAN, M.P. (1977). A semiautomated six port for measuring millimetric-wave power and complex reflection coefficient. IEEE Transaction on Microwave Theory and Techniques, Vol. 25, No. 12, 1083-1085.
- [128] WILLIAMS, D.F., BELQUIN, J.-M., DAMBRINE, G., FENTON, R. (1996). On-wafer measurement at millimeter wave frequencies. 1996 IEEE MTT-S International Symposium Digest, 1683-1686.
- [129] WILLIAMS, D.F., MARKS, R.B. (1991). Transmission line capacitance measurement. IEEE Microwave and Guided Letters, Vol. 1, No.9, 243-245.
- [130] WILLIAMS, D.F., MARKS, R.B. (1993). Accurate transmission line characterization. IEEE Microwave and Guided Letters, Vol. 3, No.8, 247-249.
- [131] Xu, Y., Bosisio, R.G. (1999). On the real-time calibration of six-port receivers (SPRs). Microwave and Optical technology Letters, Vol. 20, No.5, 318-322.
- [132] YANG, L., LONG, S.I. (1986). New method to measure the source and drain resistance of the GaAs MESFET. IEEE Electron Device Letters, Vol. 7, No.2, 75-77.
- [133] YARMAN, B.S., CARLIN, H.J. (1982). A simplified "real frequency" technique applied to broad-band multistage microwave amplifiers. IEEE Transactions on Microwave Theory and Techniques, Vol. 30, No. 12, 2216-2222.
- [134] YOUNG, G.P., SCANLAN, O.S. (1981). Matching network design studies for microwave transistors amplifiers. IEEE Transactions on Microwave Theory and Techniques, Vol. 29, No. 10, 1027-1035.
- [135] ZHAOWU, C., BINCHUN, X. (1987). Linearization of diode detector characteristics. 1987 IEEE MTT-S Digest, Vol. J-3, 265-267.