



Modèle hydrodynamique de transistor MOSFET et méthodes numériques, pour l'émission et la détection d'onde électromagnétique THz.

Mirijason Richard Razafindrakoto

► To cite this version:

Mirijason Richard Razafindrakoto. Modèle hydrodynamique de transistor MOSFET et méthodes numériques, pour l'émission et la détection d'onde électromagnétique THz.. Autre [cond-mat.other]. Université Montpellier, 2017. Français. NNT : 2017MONT035 . tel-01923358

HAL Id: tel-01923358

<https://theses.hal.science/tel-01923358>

Submitted on 15 Nov 2018

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

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

THÈSE

Pour obtenir le grade de
Docteur

Délivré par l'Université de Montpellier

**Préparée au sein de l'école doctorale Information, Structures
et Systèmes (I2S)
Et de l'unité de recherche Laboratoire Charles Coulomb**

Spécialité : Physique

Présentée par Mirijason Richard Razafindrakoto

**Modèle hydrodynamique de transistors MOSFET
et méthodes numériques pour l'émission et la
détection d'onde électromagnétique THz**

Soutenue le 31 Mars 2017 devant le jury composé de

Didier FELBACQ, Professeur des Universités,
Laboratoire Charles Coulomb

Directeur de thèse

Gérald BASTARD, Directeur de Recherche,
Laboratoire Pierre Aigrain

Rapporteur

Bruno LOMBARD, Chargé de Recherche,
Laboratoire de Mécanique et d'Acoustique
Emmanuel ABRAHAM, Directeur de Recherche,

Rapporteur

Laboratoire Ondes et Matière d'Aquitaine
Dominique COQUILLAT, Directrice de Recherche,
Laboratoire Charles Coulomb

Examineur

Examinatrice

Emmanuel KLING, Ingénieur,
SAFRAN Defense & Electronics

Invité



Résumé

Du fait de ses propriétés intéressantes, le domaine de fréquence térahertz (THz) du spectre électromagnétique peut avoir de nombreuses applications technologiques, de l'imagerie à la spectroscopie en passant par les télécommunications. Toutefois, les contraintes technologiques empêchant l'émission et la détection efficaces de ces ondes par des systèmes conventionnels ont valu à cette partie du spectre électromagnétique le nom de gap THz. Au cours des deux dernières décennies, plusieurs solutions novatrices sont apparues. Parmi elles, l'utilisation de transistors à effet de champ s'est imposée comme une solution originale, bon marché, avec un fort potentiel d'intégration. Le mécanisme identifié fait intervenir l'interaction entre les ondes THz et des ondes de courant (dites ondes plasma) dans le canal du transistor. Le canal du transistor agit tel une cavité pour ces ondes plasma. Le dispositif peut alors se comporter de manière résonante ou non-résonante en fonction de divers paramètres. Dans ce manuscrit, nous étudions numériquement ces différents régimes à l'aide de modèles hydrodynamiques. Les modèles utilisés élargissent les phénomènes pris en compte dans de précédentes études théoriques. Les résultats portent sur la détection d'ondes THz par des transistors et dans une moindre mesure sur leur émission. Dans le régime non-résonant, nous étudions dans quelle mesure la plage de linéarité de détection peut être étendue. Dans le régime résonant, nous montrons l'existence de nouvelles fréquences de résonance, permettant d'élargir le spectre d'intérêt de ces détecteurs.

Abstract

Due to its interesting properties, the electromagnetic THz frequency range may lead to numerous technological applications, ranging from imaging to spectroscopy or even communications. However, technological constraints prevented the efficient emission and detection of such waves with conventional electronics, leading to the idea of the terahertz gap. In the last decades, multiple novel solutions to resolve this gap have been proposed. Amongst these, one may find the use of simple field effect transistors as the most promising one. Their production benefits from currently available CMOS technology thus drastically decreasing the fabrication cost of such a device while allowing it to be easily integrated within electronic circuits. The mechanism behind the emission and detection is the interaction between THz electromagnetic radiations and current oscillations, that is plasma waves, in the transistor's channel. This channel forms a cavity for plasma oscillations, hence, the device may act either resonantly or non-resonantly, depending on various parameters. This thesis deals with the numerical simulation of the transistor in different regimes using hydrodynamical models. These models account for multiple phenomena that have been considered in previous theoretical studies. Some theoretical results on both the emission and detection of THz radiation are presented. In the non-resonant case, we study how one can increase the linear regime of detection. In the resonant case, we show the existence of unexpected resonance frequencies, enlarging the detection spectrum of such detectors.

Table des matières

Introduction	9
I Théorie hydrodynamique du transistor à effet de champ	13
1 Le modèle Dyakonov-Shur	15
1.1 Le MOSFET	16
1.1.1 Éléments de base de la théorie des bandes pour la jonction PN .	16
1.1.2 Le dopage, élément clé vers la jonction PN	21
1.1.3 Une application technologique : le MOSFET	22
1.2 Hypothèse hydrodynamiques et équations de Dyakonov-Shur	23
1.3 Instabilité de Dyakonov-Shur	25
1.4 Introduction de la résistivité : effet de pincement	27
1.5 Détection non-linéaire - Conversion AC/DC d'un signal THz incident .	30
1.6 Régimes de détection : résonants et non-résonants	33
1.6.1 Régime résonant	33
1.6.2 Régime non résonant	35
2 Résolution de quelques limitations du modèle de Dyakonov et Shur	37
2.1 Tension négative : transistor sous le seuil	37
2.1.1 Modèle tension-charge sous le seuil	37
2.1.2 Détection THz sous le seuil	38
2.2 Modèle unifié de la charge	40
II Simulation numérique du modèle Dyakonov-Shur	45
3 Instabilité de Dyakonov-Shur, saturation et ressaut hydraulique	47
3.1 Régime d'oscillation aux temps longs	47
3.1.1 Limites de la méthode perturbative	47
3.1.2 Etude de la discontinuité	50
3.2 Amortissement par résistivité et diminution du gain	52
3.2.1 Modes perturbatifs et pulsation complexe dans l'espace de Fourier	53
3.2.2 Résolution directe FDTD du gain	56

3.3	Introduction des collisions électron-électron : canal visqueux	58
3.3.1	Origines microscopique de la viscosité et contribution mathématiques	58
3.3.2	Application de la méthode perturbative : contribution stationnaire	61
3.3.3	Application de la méthode perturbative : disparition de la discontinuité	66
4	Détection térahertz non perturbative	69
4.1	Résonances sub-harmoniques	69
4.1.1	Résolution FDTD du système Dyakonov-Shur à forte puissance	69
4.1.2	Développement perturbatif général	71
4.1.3	Mélange de fréquence comme source de signal sous-harmonique .	73
4.1.4	Comparaison directe des deux approches	74
4.2	Décalage du spectre vers le rouge	76
4.3	Extension de la plage de linéarité	80
III	Analyse numérique	83
5	Algorithmes numériques	85
5.1	Méthode de discrétisation temporelle	85
5.2	Méthodes de discrétisation spatiale	87
5.2.1	La méthode des différences finies	87
5.2.2	La méthode des éléments finis	94
5.2.3	Les méthodes spectrales	94
	Conclusion	99
A		101
A.1	Inversion de la relation ρ/U_{GS} unifiée	101

Remerciements

Je tiens à remercier ma fiancée, Christelle, et mon fils, Joshua, pour m'avoir accompagné et soutenu tout au long de ma thèse. Dans les moments les plus intenses de ces trois dernières années, j'ai pu compter sur leur aide et leur amour afin de fournir le travail nécessaire à l'aboutissement de ce manuscrit. Je tiens également à remercier mes parents et ma famille qui m'ont montré à quel point ils sont fiers de mes accomplissements. Je remercie également mes amis qui m'ont soutenu grâce à leurs encouragements et leur compassion. Je remercie aussi Florian Bigourdan, jeune chercheur avec qui j'ai beaucoup collaboré durant cette thèse, ainsi que les différentes personnes avec qui j'ai partagé un bureau au cours de ces trois dernières années, que ce fut pour quelques semaines ou quelques années.

Enfin, et non des moindres, je remercie mon directeur de thèse, Didier Felbacq, pour l'opportunité qu'il m'a présentée, ainsi que de m'avoir accompagné tout au long de ma thèse, me guidant quand j'en ai eu besoin.

Introduction

L'étude des ondes électromagnétiques, et de leurs interactions avec la matière, ont permis la création de nouvelles technologies aux applications diverses et variées. D'abord conjecturées par J.C.Maxwell en 1865, à la suite de sa découverte de la relation entre électricité et magnétisme quelques années auparavant, leur existence physique fût mise en évidence par H.Hertz près d'une vingtaine d'années plus tard. Le spectre électromagnétique, qui désigne l'ensemble des fréquences auxquelles les ondes électromagnétiques peuvent osciller, peut être divisé en deux grandes catégories : le régime basses fréquences en-dessous du térahertz, et le régime hautes fréquences au-dessus du térahertz. Le régime basse fréquences comprend notamment les ondes radios, qui furent le type d'onde mis en évidence par H.Hertz dans son expérience, mais aussi les micro-ondes et l'infrarouge par exemple. L'émission et la détection de ce type d'onde sont très bien expliquées par la théorie classique, par opposition à quantique, de la matière et des ondes électromagnétiques. Selon cette théorie, les ondes électromagnétiques sont les oscillations d'un champ qui existe en tout point de l'espace, le champ électromagnétique. Ces ondes sont émises et détectées par la mise en mouvement de particules ou tout autre objet électriquement chargé, et en particulier par la variation de courant dans les matériaux conducteurs. De fait, le phénomène d'émission de rayonnement électromagnétique classique nécessite d'avoir accès à des sources de courants oscillants à la fréquence d'émission désirée. Malheureusement, la fréquence de coupure limitée des composants électroniques prévient l'utilisation de ce phénomène au delà de quelques centaines de GHz, fréquence au-delà de laquelle on ne dispose pas de composants capables de générer un courant alternatif suffisant.

A l'opposé du spectre électromagnétique, vers les hautes fréquences, on a notamment les rayons gammas, rayons les plus énergétiques connus du spectre électromagnétique, mais également les rayons X ou encore ultra-violets. Ce type de rayonnement peut être émis et détecté en utilisant les propriétés quantiques de la matière. La mécanique quantique décrit le rayonnement électromagnétique comme étant composé de quanta que l'on appelle photon. Ces photons interagissent de façon discrète avec la matière, comme si les ondes électromagnétiques étaient composées de particules individuelles. Chaque photon possède une énergie proportionnelle à la fréquence de l'onde électromagnétique à laquelle il est associé. Ce quantum d'énergie ne peut qu'être absorbé ou émis dans son intégralité par l'absorption ou l'émission d'un photon individuel. Ce phénomène, allié à la quantification des états d'énergie électronique des atomes, explique les raies d'absorption et d'émission de ces derniers. En effet, l'existence de ces raies s'explique par le fait que les électrons d'un atome ne peuvent posséder que des valeurs d'énergie spécifique autour de leur noyaux. C'est en passant d'un état d'énergie à un autre qu'un électron émet ou absorbe un photon. En excitant continuellement les atomes d'un milieu, ces derniers, se désexcitant naturellement après un certain temps, émettent des photons distribués en fréquence selon un spectre spécifique dépendant de l'atome choisi. A l'origine postulés par une tentative désespérée de Planck d'expliquer le rayonnement du corps noir, l'existence de ces photons est aujourd'hui acceptée et supportée par de nombreuses expériences comme l'effet photo-électrique ou la diffusion Compton. Cependant, l'efficacité de cette méthode d'émission dépend de la fréquence et

donc de l'énergie du photon émis et n'est efficace que pour des fréquences typiquement au-dessus du domaine térahertz.

Ainsi, les deux méthodes d'émission et de détection décrites ici ont laissé pendant un certain temps une plage du domaine électromagnétique inaccessible aux applications technologiques du fait de leur inefficacité respective. Cette plage du spectre électromagnétique fut baptisée le gap térahertz. Cependant, l'utilisation de cette région du spectre électromagnétique peut avoir de nombreuses applications. En effet, l'accès à différents domaines du spectre électromagnétique a permis l'émergence de nouvelles technologies exploitant les interactions différentes de ces ondes avec la matière. Ainsi, la transparence de l'atmosphère aux ondes radios a permis l'utilisation de ces ondes pour la communication terrestre longue distance. Les photons micro-ondes correspondent à des transitions énergétiques de l'eau, permettant le chauffage efficace de matériaux contenant de l'eau. Les transitions discrètes alliées au phénomène d'émission stimulée de rayonnement par les atomes ont permis l'invention du laser aujourd'hui utilisé dans de nombreux domaines allant du médical à la gravure en passant par la lecture de données. Les rayons X permettent quant à eux l'analyse cristallographique de certains solides ou encore l'imagerie médicale du fait que les tissus mous du corps y sont transparents. Cependant, la haute énergie des photons X les classe dans le domaine des rayonnements ionisants, dangereux pour la santé car pouvant rendre chimiquement actifs des éléments normalement inertes ou détruire certaines molécules en séparant leurs constituants. Certains matériaux ont de fortes bandes de transmission dans le domaine térahertz. Ainsi, l'accès au gap térahertz pourrait permettre de voir à travers ces matériaux tels que des poussières ou certains tissus. L'utilisation du rayonnement térahertz pourrait donc permettre l'accroissement de la sécurité des lieux publics en permettant la détection d'armes ou autre objets dissimulés tout en étant sûrs pour la santé. En effet, les photons térahertz ont des énergies bien plus faibles que les rayons X et sont non-ionisants. Par ailleurs, les transitions énergétiques de beaucoup de molécules se trouvent dans le gap térahertz. Ainsi, la spectroscopie térahertz pourrait permettre l'identification de ces molécules dont le spectre d'émission, ou indifféremment d'absorption, constitue une forme d'empreinte digitale de ces molécules. Enfin, l'accès aux fréquences térahertz pourrait permettre l'ouverture de nouvelles fréquences de télécommunications. Le fort potentiel d'application d'utilisation du domaine térahertz a donc stimulé l'intérêt porté au gap térahertz et à la résolution du problème du contrôle du rayonnement électromagnétique térahertz.

Afin de résoudre le problème de l'émission et de la détection térahertz, plusieurs mécanismes et dispositifs ont été développés avec plus ou moins de succès. Pour l'émission d'ondes térahertz, on peut par exemple utiliser des lampes à incandescence, un synchrotron, ou encore des lasers à cascade quantique. Pour la détection, on dispose par exemple de micro-bolomètres. Le mécanisme que nous allons exposer propose une résolution simple et élégante du problème du gap térahertz avec l'utilisation d'un simple transistor à effet de champ. Ce dispositif présent dans tout appareil électronique moderne, dont la fonction initiale est d'agir comme un simple interrupteur à commande électrique, possède l'avantage d'être compact, des milliards d'entre eux tiennent dans

un processeur, et a un faible coût de production. De plus, leur technologie de fabrication permet leur intégration directe dans un circuit intégré, souvent conçu à partir de ces mêmes transistors.

En 1993, M.Dyakonov et M.Shur publient un papier dans lequel ils prédisent que sous certaines conditions, un simple transistor à effet de champ puisse émettre un rayonnement électromagnétique lorsque soumis à une charge constante. La fréquence du rayonnement est modulable à la fois par les caractéristiques physiques du transistor définies lors de sa fabrication, mais également par la charge appliquée à ce dernier permettant d'utiliser un dispositif unique à différentes fréquences. Pour un transistor de longueur sub-micronique, ils montrent que la plage de variabilité fréquentielle couvre le gap térahertz, proposant là une résolution au problème de l'émission térahertz. Plus tard, en 1996, ils expliquent comment leur modèle du comportement d'un transistor peut rendre compte de la détection de telles ondes.

Dans ce manuscrit, nous allons étudier le principe de fonctionnement du mécanisme qu'ils ont proposé. Nous allons donc commencer par introduire le modèle hydrodynamique qu'ils ont utilisé pour prédire l'émission d'onde térahertz. Puis voir quelles sont les limitations de leur modèle initial et les résolutions possibles. Enfin, nous modéliserons le système Dyakonov-Shur numériquement afin de sonder des régimes du modèle qui n'ont pas encore été étudiés. Nous montrerons l'apparition de nouveaux phénomènes dont l'approche promulguée par M.Dyakonov et M.Shur ne peut rendre compte. Enfin nous allons étudier le schéma numérique utilisé pour faire ces nouvelles prédictions.

Première partie

Théorie hydrodynamique du transistor à effet de champ

Chapitre 1

Le modèle Dyakonov-Shur

Le transistor MOSFET¹ (voir figure 1.1) est un dispositif électronique à trois pôles, la source, la grille et le drain. C'est un interrupteur dont le caractère passant ou bloquant entre la source et le drain est contrôlé par le potentiel de la grille. Il constitue, avec les autres types de transistor, l'élément de base de l'électronique logique car une information binaire, 0 ou 1, peut être encodée dans son état. Son fonctionnement repose sur celui de la jonction P-N (voir figure 1.4). Dans ce premier chapitre, nous allons commencer par introduire des notions de base de physique de la matière condensée afin de comprendre l'origine microscopique du fonctionnement de ce dispositif. Nous allons donc nous intéresser au comportement de la matière dans les milieux ordonnés, solides cristallins, et plus particulièrement au comportement des électrons dans ces milieux. C'est la compréhension des propriétés de conduction électrique de différents matériaux grâce à la mécanique quantique qui a permis l'invention des ordinateurs, téléphones portables et autres dispositifs électroniques composés de processeurs et autres composants à base de transistors. Nous allons par la suite introduire le modèle de Dyakonov et Shur[1], ainsi que les hypothèses permettant sa formulation mathématique et montrer comment ce modèle propose de résoudre le problème de l'émission et de la détection d'ondes électromagnétiques THz.

1. Metal Oxide Semiconductor Field Effect Transistor : Transistor à effet de champ à grille métal oxide.

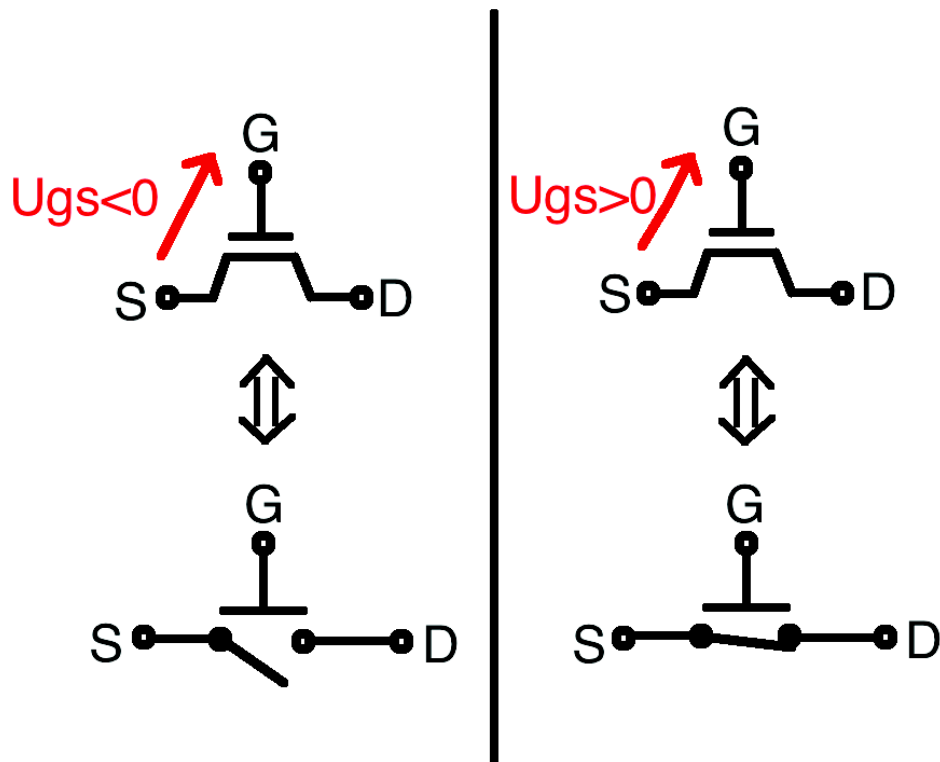


FIGURE 1.1 – Schéma de fonctionnement d'un transistor. En haut, la représentation électronique d'un transistor. En bas, l'effet du potentiel U_{GS} sur l'interrupteur équivalent entre la source et le drain. Lorsque la tension U_{GS} est positive, l'interrupteur source-drain est fermé et le courant peut passer.

1.1 Le MOSFET

1.1.1 Éléments de base de la théorie des bandes pour la jonction PN

La jonction PN est une jonction, autrement dit un contact, entre deux matériaux semi-conducteurs, l'un dont le dopage est dit P et l'autre dit N. Le caractère P, pour positif, ou N, pour négatif, définit le type de charges mobiles dans le matériau permettant de conduire un courant électrique à travers celui-ci. Dans le premier cas, certains atomes du matériau considéré, comme du silicium pur par exemple, sont remplacés par un atome situé sur une colonne à gauche de celui-ci sur le tableau périodique, par exemple du bore ou du gallium. Ces atomes possédant un électron en moins sur leur couche externe diminuent le nombre total d'électrons dans le semi-conducteur. Les

électrons possédant une charge électrique négative, la diminution de leur nombre est à l'origine de l'appellation de dopage positif. Au contraire dans le second cas, ce sont des atomes d'une colonne à droite de l'atome initial qui remplacent ce dernier, comme le phosphore. Ces atomes ajoutent des électrons supplémentaires dans le semi-conducteurs, d'où l'appellation de dopage négatif.

C'est la théorie des bandes[2], issue de la mécanique quantique, qui permet de comprendre le fonctionnement de ces matériaux. Lorsqu'un électron est confiné, dans une boîte, par une force, autour d'un atome ou autre, la mécanique quantique nous dit que son énergie ne peut prendre qu'un ensemble de valeurs discrètes. La transition d'un électron d'un niveau d'énergie à un autre se fait par l'émission ou l'absorption d'un photon qui possède toute la différence d'énergie entre les deux niveaux. Ce photon est alors associé à une fréquence unique par la relation de Planck :

$$E_i - E_f = \hbar\omega$$

où E_i est l'énergie du niveau initial de l'électron avant la transition, E_f son niveau d'énergie final, $\hbar$ la constante de Planck réduite et ω la pulsation du photon émis. C'est ce phénomène qui permet d'expliquer les raies d'absorption et d'émission des atomes. Lorsque plusieurs atomes identiques, et donc possédant les mêmes niveaux d'énergie, sont proches les uns des autres, leurs interactions entraînent une modification des niveaux d'énergie du système global, ce que l'on appelle la levée de dégénérescence (voir figure 1.2). Ainsi, lorsque N systèmes identiques sont en interaction, au lieu d'avoir N copies de chaque niveau d'énergie du système unique isolé, on obtient N niveaux d'énergie distincts à partir de chaque niveau de l'atome isolé. Dans la limite d'un nombre macroscopique d'atomes et d'électrons, par exemple lors de la formation d'un solide cristallin, le grand nombre de niveaux d'énergie semblables fait que l'on ne distingue plus le caractère discret des états d'énergie. Toute une bande continue d'énergie est alors associée à chaque niveau d'énergie du système isolé initial (voir figure 1.2).

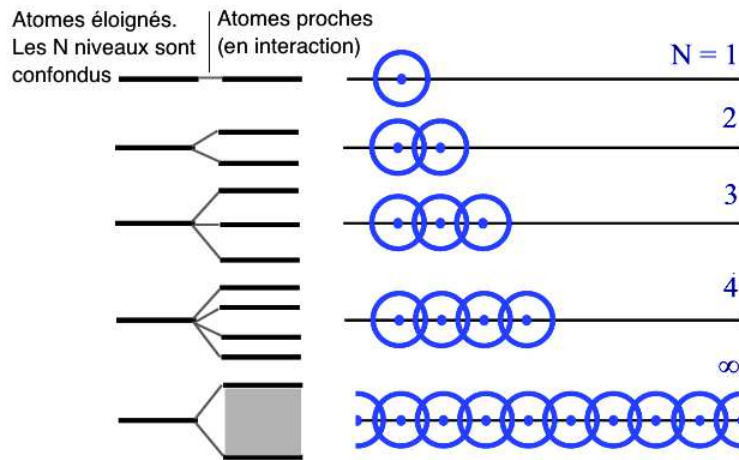


FIGURE 1.2 – Levée de dégénérescence des niveaux d'énergie d'atomes identiques en interaction (proches les uns des autres). Chaque ligne noire épaisse représente un niveau d'énergie ou état accessible aux électrons. Pour des atomes identiques isolés, tous les atomes possèdent les mêmes niveaux d'énergie. Lorsque plusieurs atomes identiques sont proches (en interaction) les N niveaux d'énergie identiques se mélangent et donnent lieu à N niveaux collectifs d'énergies différentes. Dans la limite macroscopique ($N \rightarrow +\infty$), une bande d'énergie est associée à chaque niveau initial.

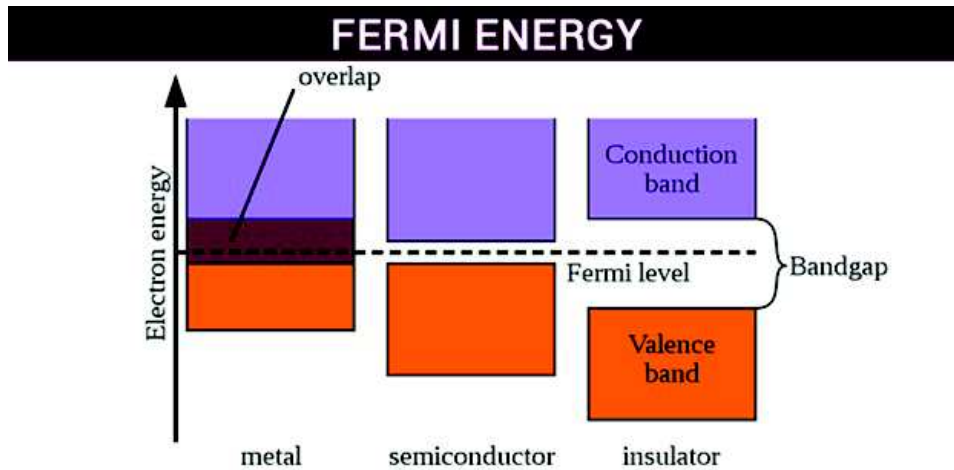


FIGURE 1.3 – Position du niveau de Fermi dans différents matériaux. A gauche, le niveau de Fermi se trouve dans une bande d'énergie et le matériau est conducteur. A droite, le niveau de Fermi est loin des bandes de conduction et de valence et le matériau est isolant. Au milieu, le niveau de Fermi est suffisamment proche des bandes de conduction et/ou de valence pour permettre la transition spontanée de certains électrons de valence dans la bande de conduction grâce à l'énergie thermique environnante.

Les différents types de matériaux (conducteurs, isolants et semi-conducteurs) s'expliquent alors par le principe d'exclusion de Pauli. Les électrons étant des fermions, deux électrons identiques ne peuvent pas occuper le même état quantique et les électrons d'un solide remplissent les niveaux d'énergie disponibles un à un, du moins énergétique au plus énergétique. Le plus haut niveau d'énergie occupé à température nulle (0 K) définit ce que l'on appelle le niveau de Fermi (voir figure 1.3). Lorsque ce niveau de Fermi se trouve dans une bande d'énergie autorisée, le matériau est conducteur.

Lorsqu'une différence de potentiel électrostatique est appliquée à un tel matériau, la conservation de l'énergie de l'électron entre deux points A et B de sa trajectoire s'écrit :

$$T_A + V_A = T_B + V_B$$

où T_M est l'énergie cinétique de l'électron au point M et V_M son énergie électrostatique en ce même point. On en déduit la variation d'énergie cinétique de l'électron entre ces deux points : $\Delta E = T_B - T_A = V_A - V_B$. Le potentiel électrostatique variant de façon continue dans un milieu homogène, cette variation d'énergie doit se faire continûment le long de la trajectoire de l'électron. Or, cela n'est possible que pour les électrons proches du niveau de Fermi. En effet, suffisamment au-dessus du niveau de Fermi, il n'y a aucun électron à accélérer à température finie, que le niveau considéré soit dans une bande d'énergie ou non. Au contraire, suffisamment en dessous du niveau de Fermi, tous les états sont déjà occupés par des électrons. Ces derniers sont alors bloqués à leur niveaux d'énergie respectifs par la présence des autres électrons. Afin de pouvoir accélérer un de ces électrons, il faut lui apporter instantanément une énergie lui permettant au moins d'atteindre le niveau de Fermi où des niveaux d'énergie commencent à être libres. Enfin, autour du niveau de Fermi, la moitié des niveaux sont occupés à température non-nulle. Si ce niveau se trouve dans une bande d'énergie, alors un électron autour de ce niveau peut passer continûment de l'état d'énergie E_A à E_B en passant par les états inoccupés autour du niveau de Fermi.

Lorsqu'au contraire le niveau de Fermi se trouve entre deux bandes d'énergie dont l'écart, appelé gap, est grand devant l'énergie thermique environnante, le matériau est isolant. On appelle alors bande de conduction la première bande inoccupée au-dessus du niveau de Fermi et bande de valence la dernière bande pleine juste en-dessous de celui-ci. Afin de faire circuler un courant dans un tel matériau, il est nécessaire de fournir aux électrons les plus hauts de la bande de valence une énergie au moins égale au gap afin que ces derniers passent dans la bande de conduction. En effet, la bande de valence, étant pleine, donne lieu à un courant macroscopique nul car à chaque état de vecteur d'onde $\vec{k}$ occupé correspond un état occupé de vecteur d'onde $-\vec{k}$ dont les contributions au courant s'annulent par symétrie. Ainsi, seuls les électrons de la bande de conduction (et l'absence d'électron dans la bande de valence) peuvent contribuer au courant macroscopique. Dans les isolants, la bande de conduction est vide car l'énergie thermique n'est pas suffisante afin d'exciter spontanément un électron de valence vers la bande de conduction par définition. Néanmoins, si le champ électrique appliqué au matériau est suffisamment intense, celui-ci peut potentiellement permettre la transition d'un électron de la bande de valence vers la bande de conduction.

Le raisonnement qualitatif permettant cette conclusion est le suivant : si la variation d'énergie électrostatique d'un électron d'un site atomique au voisin est supérieure ou de l'ordre du gap, l'énergie électrostatique apportée à l'électron est suffisante pour que celui-ci passe de la bande de valence sur l'atome initial à la bande de conduction sur l'atome voisin. Un calcul d'ordre de grandeur montre alors que cela correspondrait à un champ électrique de l'ordre de 10^{11} V.m^{-1} dans du dioxyde de silicium, un isolant très utilisé en électronique. La rigidité diélectrique, champ électrique à partir duquel un matériau isolant devient conducteur, réelle de ce matériau n'est cependant que de quelques 10^8 V.m^{-1} [3]. La considération de l'effet tunnel permettrait de revoir l'estimation précédente à la baisse. En effet, l'effet tunnel permet aux électrons, et de manière générale à tout système quantique, de passer de l'autre côté d'une barrière de potentiel sans avoir à fournir l'énergie nécessaire pour atteindre le haut de la barrière. Ainsi, pour un champ électrique légèrement inférieur à 10^{11} V.m^{-1} , l'électron serait classiquement bloqué autour de son noyau atomique car cet électron se trouverait dans le gap en allant vers les atomes voisins. Cependant, l'énergie électrostatique cumulée sur quelques sites atomiques successifs lui permettrait d'atteindre la bande de conduction loin de son atome de départ. Un électron peut alors passer instantanément de la bande de valence autour de son atome initial vers la bande de conduction loin de celui-ci sans passer par les positions classiquement interdites grâce à l'effet tunnel. Ce phénomène de conduction dans les isolants est appelé claquage. Celui-ci est généralement destructeur car les nombreuses collisions des électrons de conduction, très énergétiques du fait du fort apport d'énergie électrostatique, avec le réseau cristallin endommagent ce dernier et peuvent finir par compromettre son intégrité.

Enfin, les matériaux tels que le niveau de Fermi se trouve dans le gap mais à une distance inférieure à ou de l'ordre de l'énergie thermique environnante des bandes autorisées sont dits semi-conducteurs. Dans ces milieux, l'énergie thermique permet la transition spontanée de quelques électrons de valence vers la bande de conduction. Cette transition laisse des états inoccupés dans la bande de valence que l'on appelle des trous. Ces trous sont ce que l'on appelle des quasi-particules car hormis le fait qu'ils ne correspondent pas à de vraies particules que l'on pourrait extraire du matériau, possèdent toutes les autres propriétés associées aux particules. Leur considération en tant que tel permet une grande simplification de la modélisation des semi-conducteurs. En effet, au lieu de considérer une bande de conduction presque pleine avec un grand nombre d'électrons, on la remplace par une bande de conduction pleine, et une densité de trous dont le comportement est facile à modéliser. Le comportement de la bande pleine est trivial car sa contribution au courant est nulle comme dans le cas des isolants. Le seul effet de la bande de valence est désormais dû à sa densité de charge électrique négative. En effet, la remplir d'électron signifie introduire une charge de fond négative. La matière étant électriquement neutre, cette charge doit être compensée par une charge électrique positive attribuée aux trous. La transition d'un électron de valence vers la bande de conduction correspond à la création d'une paire électron-trou, tandis que la chute d'un électron de conduction dans la bande de valence correspond à son annihilation.

1.1.2 Le dopage, élément clé vers la jonction PN

Pour améliorer la conduction d'un semi-conducteur, on peut le doper[4]. Cette opération consiste à remplacer certains atomes composant le cristal par un autre atome possédant un électron en plus (dopage N qui augmente le niveau de Fermi et le nombre d'électrons) ou en moins (dopage P qui baisse le niveau de Fermi et augmente le nombre de trous). En réalisant une jonction PN, les trous naturellement plus nombreux côté P que N et les électrons plus nombreux côté N que P diffusent spontanément à travers l'interface et se recombinent. En conséquence, il se forme ce que l'on appelle une zone de charge d'espace à la jonction où la recombinaison des paires électron-trou laisse la charge de fond de la bande de valence non compensée. Un champ électrique se forme alors, prévenant une diffusion supplémentaire des porteurs de charges libres (trous et électrons). Lorsque le milieu dopé P est connecté à la charge positive d'une batterie et le milieu dopé N à sa charge négative, le champ électrique induit par la batterie pousse les électrons et les trous vers la jonction où ceux-ci se recombinent. Dans ce cas, on parle de polarisation directe (voir figure 1.4). Les trous étant chargés positivement et les électrons négativement, leur mouvement contribuent tous deux à un courant positif orienté de la zone P vers la zone N. En polarisation indirecte, le champ électrique renforce le champ existant à l'équilibre thermodynamique de la jonction et le courant ne passe pas car aucun porteur n'est présent pour assurer la conduction autour de la jonction. La jonction PN se comporte alors comme une diode.

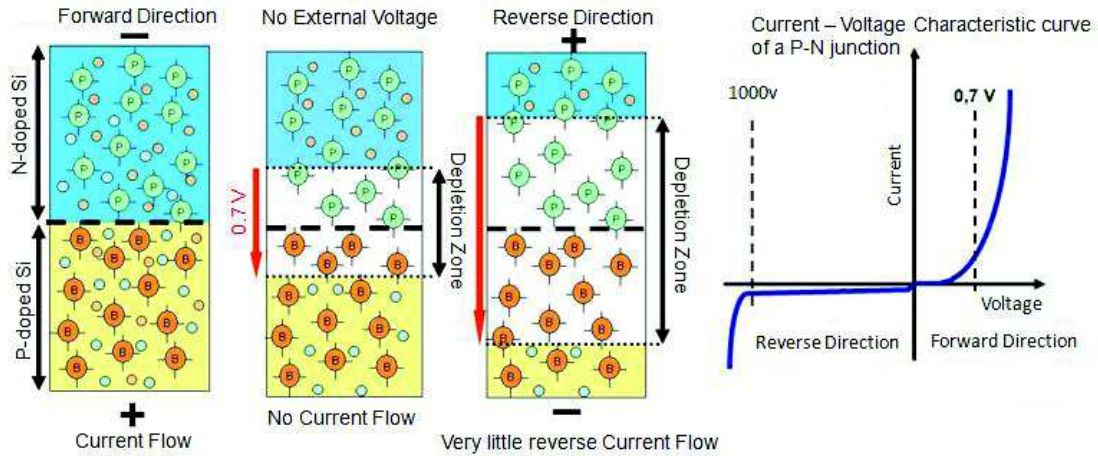


FIGURE 1.4 – Schéma de fonctionnement de la jonction P-N. Les trois schéma à gauche représentent la jonction PN de la gauche vers la droite en polarisation directe, en l'absence de polarisation externe, et en polarisation indirecte. Tout à droite, une courbe caractéristique typique d'une jonction PN est représentée. On voit bien l'effet de la diode qui laisse facilement passer le courant dans un sens (à partir de 0,7 V) et difficilement dans l'autre (très peu de courant avant d'atteindre la tension de claquage à -1000 V).

1.1.3 Une application technologique : le MOSFET

Comme mentionné au début, le MOSFET est, comme d'autres transistors, un dispositif électronique dont l'intérêt principal est de servir d'interrupteur. Grâce à cette fonction, il constitue l'élément de base de l'électronique moderne. En effet, les deux états passant et bloquant permettent d'encoder une information binaire et de réaliser les différents types d'opération nécessaires à la logique binaire. Ainsi, on le trouve dans tous les processeurs d'ordinateur, téléphone et autres dispositifs électroniques. On peut également l'utiliser pour réaliser des amplificateurs opérationnels, dispositifs servant à amplifier un signal d'entrée. Enfin, le MOSFET repose sur la technologie CMOS dont l'amélioration des techniques de fabrication a permis leur miniaturisation et leur production en masse, réduisant les coûts de fabrication de ces dispositifs.

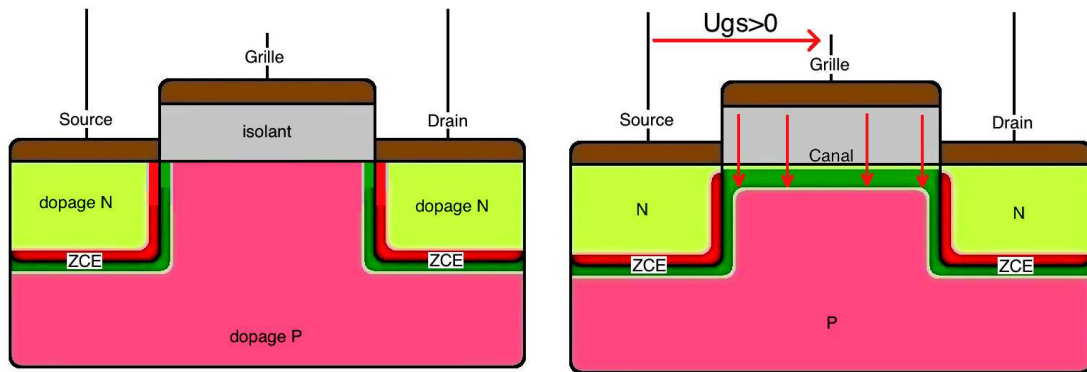


FIGURE 1.5 – Schéma d'un MOSFET N-P-N. Pour le transistor P-N-P, il suffit d'inverser les mentions électron et trou ainsi que dopage P et dopage N. A gauche, le transistor est bloquant, la zone de charge d'espace (ZCE) empêche les électrons de passer de la source vers le drain. A droite, l'application d'un potentiel suffisant entre la source et la grille génère un champ électrique qui repousse les trous loin de l'isolant et attire les électrons de la source et du drain. Un canal de conduction se forme alors à l'interface substrat/isolant et assure la conduction du courant entre la source et le drain. Le transistor est alors passant.

La composition interne d'un MOSFET est schématisée sur la figure 1.5. Le substrat du MOSFET y est schématisé en rose. Il est constitué d'un semi-conducteur dopé P ou N suivant le type de porteurs que l'on désire dans l'état passant. En vert pâle, sont schématisées les deux zones externes du MOSFET, la source et le drain, entre lesquelles on souhaite obtenir le caractère passant ou bloquant. Celles-ci sont dopées de manière complémentaire au substrat, donc N si le substrat est dopé P, P sinon, et définissent le type de porteurs dans l'état passant. Des liaisons métalliques en marron assurent les contacts avec le reste du circuit électrique. A l'équilibre, l'état du transistor entre la source et le drain est naturellement bloquant. En effet, la jonction entre les deux zones externes est assurée par le substrat, dont le dopage ne permet pas le passage du

courant. L'ensemble source-substrat-drain est en effet constitué de deux jonctions, NP puis PN (ou inversement), agissant comme des diodes en sens inverse. Afin d'obtenir l'effet d'interrupteur, un contact métallique appelé grille est ajoutée entre la source et le drain, espacée du substrat par une couche d'isolant, en gris sur le schéma, permettant d'éviter le court-circuit du substrat. Un champ électrique généré par le potentiel de la grille suffisamment intense permet de repousser les porteurs de courant du substrat et d'attirer ceux de la source et du drain. Une zone d'inversion des porteurs de charge se crée à l'interface entre le substrat et l'isolant et assure la conductivité du courant entre la source et le drain. Cette zone d'inversion constitue le canal de conduction du MOSFET. Dans la partie suivante, on s'intéressera au comportement de ce canal dans des conditions particulières de fonctionnement, au travers du modèle Dyakonov-Shur.

1.2 Hypothèse hydrodynamiques et équations de Dyakonov-Shur

Dans le but d'étudier la génération et la détection d'onde électromagnétique THz, M.Dyakonov et M.Shur proposèrent un modèle hydrodynamique du canal de conduction d'un MOSFET. Ainsi, l'ensemble des électrons du canal est considéré comme un gaz classique. Cette approche est justifiée si la densité d'électrons est suffisante, auquel cas on peut représenter les différentes grandeurs macroscopiques mesurables du système par différents champs. Du fait du fort confinement électrostatique, le gaz d'électrons est restreint à l'interface bidimensionnelle entre le substrat et l'isolant le séparant de la grille. Dans une approche simplificatrice, leur modèle restreint également le mouvement des électrons selon une seule dimension spatiale donnée par la direction source-drain. Ainsi, le modèle Dyakonov-Shur modélise le canal d'un MOSFET par un champ de densité électronique ρ et un champ de vitesse électronique v à une dimension spatiale x évoluant dans le temps t . Leurs équations d'évolutions sont données par les lois de conservations de la charge (1.1) et de la quantité de mouvement ou équation de Navier-Stokes (1.2) :

$$\frac{\partial}{\partial t}\rho + \frac{\partial}{\partial x}(\rho v) = 0 \quad (1.1)$$

$$\frac{\partial}{\partial t}v + v\frac{\partial}{\partial x}v = \frac{f}{m^*} \quad (1.2)$$

où f est la résultante des forces appliquées en un point du fluide et m^* la masse effective des électrons dans le canal. Leur modèle initial ne considère que l'effet de l'interaction électromagnétique. Les électrons possédant une charge électrique non-nulle, ceux-ci interagissent avec le champ électromagnétique. Dans une approche simplificatrice, ce modèle ne considère cependant pas l'intégralité des interactions électromagnétiques mais une approche quasi-statique que l'on justifiera *a posteriori*. Par ailleurs, le problème étant essentiellement unidimensionnel, il néglige également la contribution du champ

magnétique. Dans ce cas, la force f est simplement donnée par la force électrostatique $q\vec{E}$ où q est la charge de l'électron, incluant sa convention de signe, et $\vec{E}$ le champ électrique à évaluer dans le canal, solution de l'équation de Maxwell-Gauss : $\vec{\nabla} \cdot (\epsilon \vec{E}) = q\rho(t, x) \delta(z)$. Du fait de l'approche quasi-statique, le champ électrique peut être donné par le gradient d'un potentiel scalaire $\vec{E} = -\vec{\nabla}\phi$. Comme le mouvement des électrons selon la direction $(0y)$ ainsi que la dépendance des champs selon cette direction sont négligés, l'équation de Maxwell-Gauss est à considérer dans le plan (Oxz) figure 1.6.

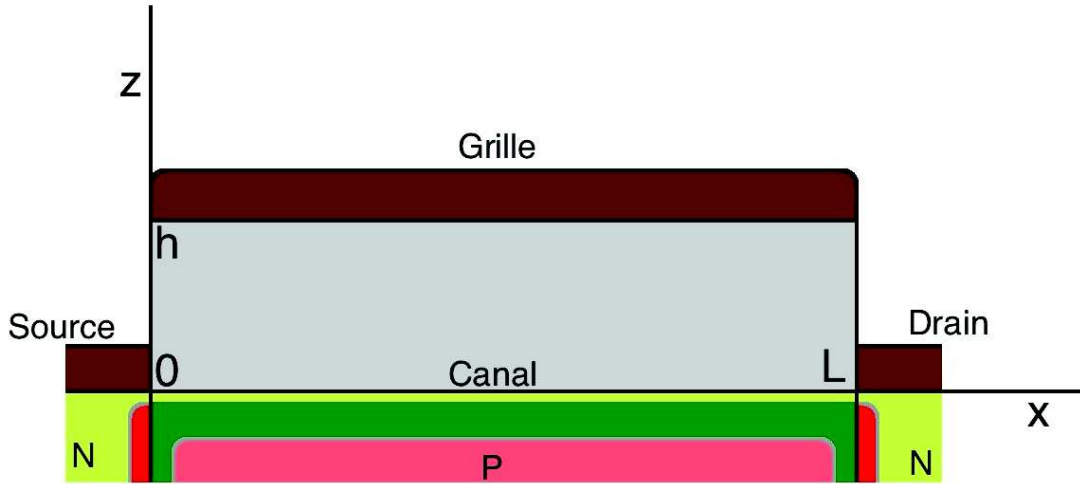


FIGURE 1.6 – Système de coordonnées utilisé. La dépendance selon la troisième dimension y (entrant conventionnellement dans la feuille) est supposée nulle.

Lorsque la longueur du canal L est très grande devant la distance le séparant de la grille h , on peut évaluer le potentiel électrostatique à l'aide de l'approximation du canal graduel (1.3) :

$$q\rho = CU \quad (1.3)$$

Cette approximation modélise le système grille-canal par une capacité C , liant la différence de potentiel locale $U(t, x) = \phi(t, x, 0) - \phi(t, x, h)$ entre la grille et le canal, à la densité d'électrons dans le canal $\rho(t, x)$. On aboutit finalement au système Dyakonov-Shur (1.4) :

$$\begin{cases} \frac{\partial}{\partial t}\rho + \frac{\partial}{\partial x}(\rho v) = 0 \\ \frac{\partial}{\partial t}v + \alpha \frac{\partial}{\partial x}\rho + v \frac{\partial}{\partial x}v = 0 \end{cases} \quad (1.4)$$

où $\alpha = \frac{q^2}{m^*C}$ est une constante qui peut être arbitrairement fixée à 1 par une redéfinition de l'unité de densité, ce que l'on fera désormais. Ce système, initialement proposé pour

l'émission spontanée d'ondes électromagnétiques, sera modifié par la suite pour rendre compte des régimes de détection de ces ondes.

1.3 Instabilité de Dyakonov-Shur

Pour trouver une solution unique au système Dyakonov-Shur (1.4), il est nécessaire de spécifier deux conditions aux limites reflétant deux contraintes physiques imposées. M.Dyakonov et M.Shur proposèrent que les conditions considérées soient :

$$\begin{cases} \text{un courant de drain constant} & : & q\rho(t, L) v(t, L) = i_D \\ \text{une tension grille-source constante} & : & U(t, 0) = U_{GS} \end{cases} \quad (1.5)$$

qui, du fait de l'approximation du canal graduel (1.3), implique une densité constante $\rho(t, 0) = \rho_{GS}$ au niveau de la source. Suivant la valeur du courant de drain i_D choisie, M.Dyakonov et M.Shur ont montré que bien que les conditions aux bords soient statiques, la solution statique dans le canal pouvait être instable et que les champs de densité et de vitesse pouvaient se mettre à osciller spontanément. Le canal du transistor étant considéré comme un gaz de particules chargé, celui-ci se comporte comme un plasma. Les oscillations du canal sont donc qualifiées d'ondes plasma. La formation spontanée d'ondes plasma devrait se produire à des fréquences bien précises, multiples impairs d'une fréquence fondamentale se situant dans le gap THz pour un MOSFET de longueur inférieure au micromètre. On peut désormais justifier l'approximation quasi-statique. En effet, une onde électromagnétique THz est associée à une longueur d'onde d'environ $300\mu\text{m}$, largement supérieur à la longueur du canal considérée ($\lesssim 1\mu\text{m}$) qui est le critère de justification de cette approximation.

Ces ondes plasma sont associées à une variation du courant dans le canal du transistor, entraînant une variation de son moment dipolaire. Celui-ci devrait donc être une source de rayonnement électromagnétique aux fréquences d'oscillation du transistor. Cette oscillation source de rayonnement électromagnétique générée par des contraintes constantes est similaire au son émis par une flûte. Dans cette dernière, un utilisateur souffle de manière constante dans le bec d'une flûte qui, via des instabilités hydrodynamiques fournies par la présence d'un biseau, génère une oscillation périodique de l'air, le son.

Afin de montrer que leur système d'équation peut être instable, M.Dyakonov et M.Shur proposèrent une approche perturbative. Pour cela, ils commencèrent par donner une solution explicite statique de leur système d'équation qui est simplement donnée par un champ de densité ρ_s et un champ de vitesse v_s constants dans le canal. Par la condition de normalisation ($\alpha = 1$), on obtient $\rho_s = \frac{qU_{GS}}{m^*}$ qui est homogène à une vitesse au carré, dimension dont on explique la signification plus loin.

Ils ajoutèrent par la suite à chaque champ ρ_s et v_s une petite contribution ρ_p et v_p , tel que la somme soit toujours solution de leur système. Ainsi, le champ de densité total

s'écrit par exemple : $\rho(t, x) = \rho_s(x) + \rho_p(t, x)$. On aboutit alors au système suivant :

$$\begin{cases} \frac{\partial}{\partial t} \rho_p + \frac{\partial}{\partial x} (v_s \rho_p + \rho_s v_p + \rho_p v_p) = 0 \\ \frac{\partial}{\partial t} v_p + \frac{\partial}{\partial x} \left(\rho_p + v_s v_p + \frac{1}{2} v_p^2 \right) = 0 \end{cases} \quad (1.6)$$

muni des nouvelles conditions aux bords :

$$\begin{aligned} \rho_p(t, 0) &= 0 \\ v_s(L) \rho_p(t, L) + \rho_s(L) v_p(t, L) + \rho_p(t, L) v_p(t, L) &= 0 \end{aligned}$$

Comme les amplitudes des perturbations ρ_p et v_p sont supposées faibles, les termes faisant apparaître un produit d'au moins deux de ces termes peuvent être négligés en première approximation. Etant donné que ρ_s et v_s sont constants, le système d'équations (1.6) peut être réécrit sous la forme d'une équation d'onde linéaire à coefficients constants dont la solution s'écrit comme la somme d'ondes se propageant à la vitesse

$s = \sqrt{\rho_s} = \sqrt{\frac{qU_{GS}}{m^*}}$ par rapport au gaz d'électron. La densité stationnaire ρ_s est donc le carré de la vitesse de propagation typique des ondes plasma. Lorsque le courant de drain imposé est non-nul, le gaz d'électron dérive à la vitesse v_s , définie positive dans le sens source-drain. Cette vitesse de dérive induit la relation de dispersion $\omega = k(s + v_s)$ pour les ondes allant de la source vers le drain et $\omega = k(s - v_s)$ pour les ondes allant du drain vers la source. En injectant une solution de la forme $\exp(-i(\omega t - kx))$ dans le système (1.6), on aboutit à la solution (1.7) suivante :

$$\begin{cases} \rho_p(t, x) = \Re \left(\frac{sv_p(0, 0)}{2} \left(\exp \left(i \frac{\omega x}{v_s + s} \right) - \exp \left(i \frac{\omega x}{v_s - s} \right) \right) \exp(-i\omega t) \right) \\ v_p(t, x) = \Re \left(\frac{v_p(0, 0)}{2} \left(\exp \left(i \frac{\omega x}{v_s + s} \right) + \exp \left(i \frac{\omega x}{v_s - s} \right) \right) \exp(-i\omega t) \right) \end{cases} \quad (1.7)$$

où $\omega = \frac{s^2 - v_s^2}{2Ls} \pi (2n + 1) + i \frac{s^2 - v_s^2}{2Ls} \ln \left(\frac{s + v_s}{s - v_s} \right)$, $\forall n \in \mathbb{Z}$.

On appelle g le gain perturbatif du canal, correspondant à la partie imaginaire de ω . D'après les conventions utilisées, un gain négatif correspond à une décroissance exponentielle de la contribution perturbative. Dans ce cas, une perturbation de l'état statique disparaît à long terme et on parle alors d'un régime statique stable. Si au contraire le gain est positif, alors une perturbation de la solution statique croît exponentiellement avec le temps. Cette perturbation finit alors par devenir en un temps fini la contribution dominante des champs de densité et de vitesse. On parle alors de régime statique instable. Dans ce régime, cette instabilité excite spontanément les résonances du canal qui se comporte comme une cavité. Pour un canal d'une longueur comprise entre 100 nm et 1 μm , les fréquences de résonances de celui-ci sont situées dans le THz et des ondes plasma devraient s'établir dans le MOSFET à ces fréquences.

Ces ondes plasma étant associées à une variation du moment dipolaire électrique, le MOSFET devrait spontanément émettre un rayonnement électromagnétique THz dans le régime instable. Malheureusement, une telle émission spontanée n'a été observé qu'à température cryogénique (4K) [5]. L'absence d'oscillation spontanée et donc d'émission d'onde électromagnétique à plus haute température, notamment ambiante (300K), est attribuée à la contribution de forces dissipatives dans le canal venant amortir les ondes plasma dans le canal et ainsi diminuer la valeur du gain g . Une forme de dissipation non négligeable est donnée par la résistivité au courant du canal, contribution que l'on va introduire par la suite.

1.4 Introduction de la résistivité : effet de pincement

Selon le modèle de Drude, la résistivité d'un matériau conducteur est due à la réinitialisation de la vitesse des électrons suite à un choc avec un élément non-conducteur. Parmi les particules sources de collisions résistives, on compte les impuretés, présence d'éléments dans le réseau cristallin différents de l'élément constitutif du milieu, et les phonons, quasi-particules issues de l'application des lois de la mécanique quantique aux oscillations du réseau cristallin. A la suite d'une collision, le vecteur vitesse de l'électron est modifié aléatoirement de manière isotrope. Lorsqu'une force $\vec{F}$ constante est appliquée sur les électrons, ces derniers sont accélérés entre chaque collision. Ainsi, la vitesse d'un électron entre deux collisions aux instants t_1 et t_2 est donnée par :

$$\vec{v}(t) = \vec{v}(t_1) + \frac{\vec{F}}{m^*} (t - t_1)$$

Comme les collisions sont isotropes, la vitesse initiale moyenne $\langle \vec{v}(t_1) \rangle$ des électrons est nulle. La vitesse moyenne des électrons dans le matériau est donc uniquement spécifiée par le second terme et est alors proportionnelle à la force appliquée. C'est de cette valeur moyenne que l'on déduit la loi d'Ohm locale $\vec{j} = \sigma \vec{E}$ où σ est la conductivité locale du matériau. Ainsi, on retrouve bien que sans force extérieure appliquée, il n'y a pas de courant macroscopique.

Dans le modèle Dyakonov-Shur, cette résistivité peut être prise en compte en introduisant une force résistive $f = -\frac{m^*v}{\tau}$ où τ est le temps de libre parcours moyen des électrons ou temps de relaxation de l'impulsion. Le système Dyakonov-Shur résistif (1.8) s'écrit alors :

$$\begin{cases} \frac{\partial}{\partial t} \rho + \frac{\partial}{\partial x} (\rho v) = 0 \\ \frac{\partial}{\partial t} v + \frac{\partial}{\partial x} \rho + v \frac{\partial}{\partial x} v + \frac{v}{\tau} = 0 \end{cases} \quad (1.8)$$

Pour maintenir un courant stationnaire dans le canal, il devient alors nécessaire d'appliquer une force exerçant un travail sur le fluide. Ce travail est fourni par la force

électrostatique et correspond à la différence de potentiel entre la source et le drain. On en déduit que les champs de densité et de vitesse stationnaire ne sont plus constants en fonction de la position dans le canal mais solutions du système algébrique (1.9) suivant :

$$\begin{cases} \rho_s(x) v_s(x) = \rho_s(0) v_s(0) \\ \frac{1}{2\rho_s(0)} \rho_s^3(x) + \left(\frac{v_s(0)}{\tau} x - \frac{1}{2} \rho_s(0) - v_s^2(0) \right) \rho_s(x) + \rho_s(0) v_s^2(0) = 0 \end{cases} \quad (1.9)$$

La deuxième équation du système (1.9) est une équation polynômiale de degré trois sur son inconnue $\rho_s(x)$, ce qui implique trois couples de solutions (ρ_s, v_s) Fig. 1.7. La formule de Cardan permet d'obtenir explicitement les trois couples de solutions par radicaux, dont l'expression longue, peu esthétique et peu éclairante sur la nature des solutions ne sera pas explicitement donnée ici.

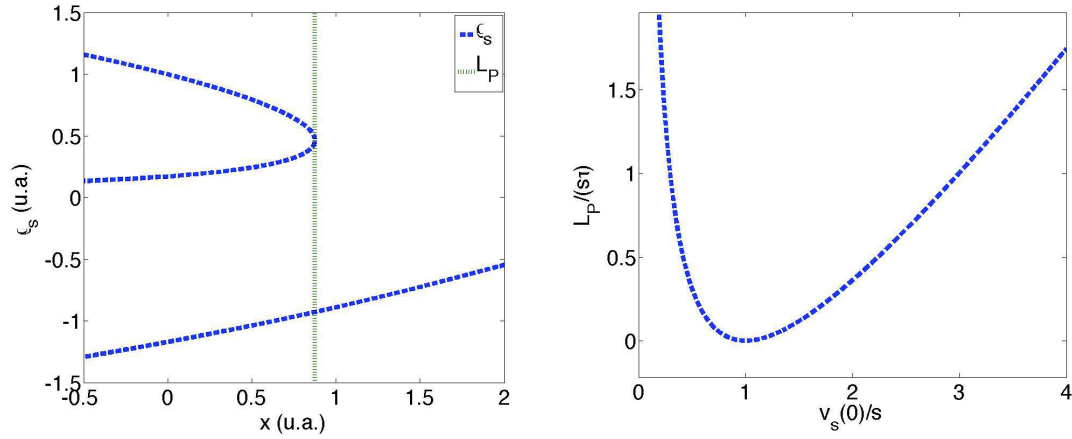


FIGURE 1.7 – Solutions du système (1.9) pour $v_s(0) > 0$ à gauche et longueur de pincement en fonction de la vitesse de dérive à la source à droite. La solution négative existe pour toute longueur de canal mais n'est pas physique. Les deux solutions restantes existent jusqu'en $x = L_P$ et sont donc possible si et seulement si $L_P > L$.

Des trois solutions, une est réelle, continue et différentiable pour toute position x dans $\mathbb{R}$ mais est associée à une densité strictement négative qui n'est donc pas physique. La solution à considérer figure donc parmi les deux solutions restantes. Celles-ci sont positives, réelles, continues et différentiables si et seulement si $x \in]-\infty; L_P]$. Ainsi, la solution stationnaire n'est définie que si la longueur du canal L est inférieure à une grandeur L_P que l'on va chercher à interpréter.

Lorsque le courant de drain est important, le potentiel U_{DS} associé peut devenir supérieur au potentiel U_{GS} imposé. En conséquence, le potentiel ϕ proche du drain peut devenir supérieur au potentiel ϕ_G de la grille. Dans ce cas, le potentiel de la grille n'est plus suffisant pour assurer l'existence du canal sur toute sa longueur. En particulier, celui-ci disparaît du côté du drain (voir figure 1.8). On dit que le canal se

pince. La longueur L_P correspond donc à la longueur de pincement, distance au-delà de laquelle le canal n'est plus défini. Cette distance est donnée, dans le cadre du modèle Dyakonov-Shur, par :

$$L_P = \frac{s\tau}{2} \left(2 \frac{v_s(0)}{s} - 3 \left(\frac{v_s(0)}{s} \right)^{\frac{1}{3}} + \frac{s}{v_s(0)} \right) \quad (1.10)$$

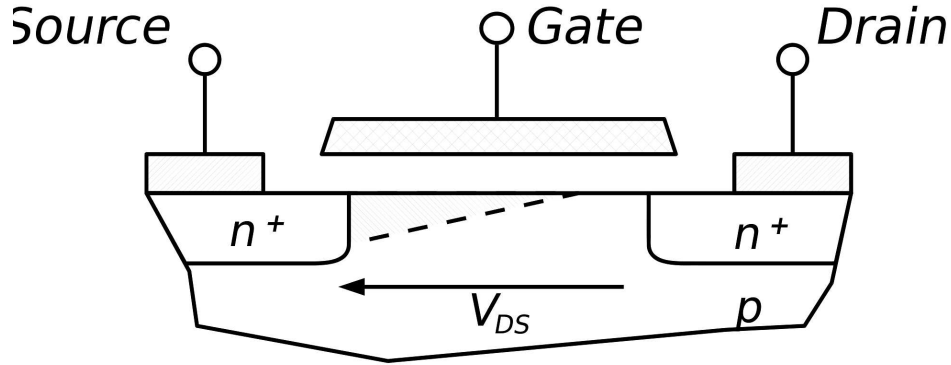


FIGURE 1.8 – Origine du pincement selon le modèle usuel. La distance entre la ligne pointillée et la ligne horizontale juste au-dessus d'elle représente la densité d'électron dans le canal. A cause du fort potentiel V_{DS} , le potentiel de drain passe en dessous du seuil et le canal se pince.

Dans ce cadre, on peut donner une autre interprétation à l'origine du pincement du canal. En l'absence de résistivité, on sait que le système Dyakonov-Shur permet la propagation de petites perturbations autour de la solution stationnaire à la vitesse locale s . En présence de résistivité, ces ondes se déplacent désormais à la vitesse locale $c = s\sqrt{\frac{U}{U_{GS}}}$ où U est le potentiel local dans le canal. La vitesse s représente désormais l'ordre de grandeur caractéristique de la vitesse des ondes plasma. Lorsque l'on fait varier le courant de drain i_D d'une valeur nulle à une valeur non-nulle de façon adiabatique, on peut voir la modification progressive du courant de drain comme source d'ondes plasma. Ces ondes, émises au niveau du drain, se propagent vers la source et modifient la solution stationnaire sur leur passage. Lorsque le courant de drain imposé devient trop important, il arrive un moment où la vitesse de dérive du gaz d'électrons v devient égale à la vitesse locale des ondes plasma. A cet instant, l'onde de modification, générée par la variation du courant de drain, n'est plus capable de remonter le courant dans le canal. Le courant ne pouvant plus augmenter d'avantage, celui-ci sature à la valeur i_P . On retrouve là le comportement typique des transistors MOSFET (voir figure 1.9).

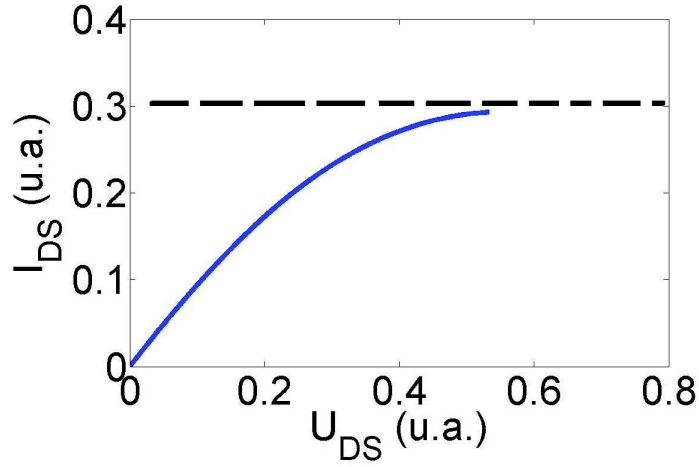


FIGURE 1.9 – Courbe caractéristique courant-tension donnée par le modèle Dyakonov-Shur en unités adimensionnées. On retrouve la forme typique de cette courbe jusqu’à la saturation. A faible potentiel drain-source, le courant est linéaire avec le potentiel. Au delà du potentiel drain-source de saturation (de l’ordre de 0,5 ici), le courant sature et tend vers une asymptote horizontale (en noire).

1.5 Détection non-linéaire - Conversion AC/DC d’un signal THz incident

Le modèle Dyakonov-Shur fut initialement formulé dans le but de résoudre le problème de l’émission d’ondes électromagnétiques THz. Cependant, il est également adapté pour la détection de telles ondes[6]. En effet, en régime de détection, le modèle Dyakonov-Shur prédit la rectification basse fréquence du potentiel drain-source suite à une excitation THz. Cette conversion des hautes fréquences vers les basses fréquences constitue le point fort du mécanisme de Dyakonov-Shur pour la détection d’ondes électromagnétiques THz. D’un point de vue classique, l’absorption d’une onde électromagnétique dans un conducteur est due à la mise en mouvement de ses charges libres. Il s’établit alors un courant oscillant à la fréquence de l’onde incidente dans le conducteur. Pour des fréquences THz, ce courant oscille bien trop rapidement pour être détecté par des dispositifs électroniques. En effet, ces derniers ne répondant pas instantanément à une excitation donnée possède une fréquence de coupure de l’ordre de l’inverse de leur temps de réponse. Au-delà de cette fréquence de coupure (pouvant atteindre quelques centaines de gigahertz dans les transistors à base de graphène [7]), le dispositif n’a pas le temps de répondre et celui-ci ne détecte rien. Par le mécanisme de Dyakonov-Shur, l’arrivée d’une onde THz sur le canal d’un MOSFET est mise en évidence par la rectification de son potentiel drain-source moyen. Cette rectification ne nécessitant pas la commutation du transistor entre son état passant et bloquant, elle se

produit même lorsque la fréquence de l'onde incidente est supérieure à la fréquence de coupure du transistor. Le transistor reste en effet passant lors de la détection de l'onde THz incidente. La rectification du potentiel drain-source étant à fréquence nulle, elle est aisément mesurable à l'aide d'un système de mesure conventionnel. En couplant une antenne THz à un MOSFET, il est ainsi possible de détecter des ondes THz via une électronique simple et bon marché.

La rectification du potentiel drain source, suite à l'excitation du canal par une onde électromagnétique THz, est due aux interactions non-linéaires du modèle Dyakonov-Shur. Si une antenne est couplée entre la source et la grille d'un MOSFET, l'absorption d'une onde incidente par l'antenne engendre l'oscillation du potentiel grille-source. Ce potentiel étant lié à la densité électronique dans le canal, une onde plasma est générée au niveau de la source et se propage alors à travers le canal en direction du drain. Le système Dyakonov-Shur étant fortement non-linéaire, cette propagation se fait de manière non-linéaire. Notamment, en se propageant, l'onde plasma modifie la densité moyenne d'électron sur son passage. Cette densité étant liée au potentiel du canal, celui-ci est également rectifié. C'est la différence entre la source et le drain de ce potentiel rectifié qui constitue alors le signal de détection.

Pour évaluer la rectification du potentiel drain-source en fonction de l'amplitude de l'oscillation induite par l'antenne, M.Dyakonov et M.Shur ont recours à la méthode des perturbations [6]. Pour un courant de drain nul, la solution stationnaire sans excitation du potentiel de source du système (1.8) est donnée par :

$$\begin{cases} \rho_s(x) = \rho_{GS} \\ v_s(x) = 0 \end{cases} \quad (1.11)$$

Ils décomposèrent ensuite les champs de densité et de vitesse comme la série de Fourier suivante :

$$\begin{cases} \rho = \sum_{n \in \mathbb{N}} \rho_n(t, x) \\ v = \sum_{n \in \mathbb{N}} v_n(t, x) \end{cases}$$

où chaque ρ_n et v_n oscille à la pulsation $n\omega$ où ω est la pulsation de l'onde incidente sur l'antenne. Le terme ρ_0 est le signal que l'on recherche car ce dernier est indépendant du temps. Si l'excitation est supposée faible (l'amplitude d'oscillation U_A générée par l'antenne est faible devant le potentiel U_{GS} du générateur) on peut tronquer la série au terme ρ_1 . On a alors $\rho_0 - \rho_s$ et v_0 de l'ordre de U_A^2 et ρ_1 et v_1 de l'ordre de U_A . Les équations que l'on doit résoudre sont alors (en identifiant les termes oscillants à la

même fréquence) :

$$\left\{ \begin{array}{l} \frac{\partial}{\partial t} \rho_1 + \rho_{GS} \frac{\partial}{\partial x} v_1 = 0 \\ \frac{\partial}{\partial t} v_1 + \frac{\partial}{\partial x} \rho_1 + \frac{v_1}{\tau} = 0 \\ \frac{\partial}{\partial x} (\rho_0 v_0 + \langle \rho_1 v_1 \rangle) = 0 \\ \frac{\partial}{\partial x} \left(\rho_0 + \frac{1}{2} v_0^2 + \frac{1}{2} \langle v_1^2 \rangle \right) + \frac{v_0}{\tau} = 0 \end{array} \right. \quad (1.12)$$

où la notation $\langle f \rangle$ désigne la valeur moyenne temporelle de f sur une période. Les deux premières équations du système (1.12) donnent la propagation au premier ordre de l'onde plasma générée par le potentiel oscillant de l'antenne. On déduit de ces équations la relation de dispersion de ces ondes :

$$\omega^2 + i \frac{\omega}{\tau} - s^2 k^2 = 0$$

On en déduit que les solutions de ces équations sont données par la superposition d'exponentielles complexes $\exp(-i(\omega t \pm kx))$ dont les coefficients sont données par les conditions aux bords. On trouve alors :

$$\left\{ \begin{array}{l} \rho_1(t, x) = \Re \left(s^2 \frac{U_A}{U_{GS}} \exp(-i\omega t) \frac{\cos(k(x-L))}{\cos(kL)} \right) \\ v_1(t, x) = \Re \left(\frac{U_A}{U_{GS}} \frac{i\omega}{k} \exp(-i\omega t) \frac{\sin(k(x-L))}{\cos(kL)} \right) \end{array} \right.$$

où $k = \frac{\omega}{s} \sqrt{1 + \frac{i}{\omega\tau}}$ le vecteur d'onde de l'onde plasma prend une valeur complexe.

Les deux dernières équations du système (1.12) donnent la rectification des champs moyens en fonction de l'amplitude des ondes de densité dans le canal. En ne retenant que les termes d'ordre U_A^2 , le système se réécrit :

$$\left\{ \begin{array}{l} \frac{\partial}{\partial x} (\rho_{GS} v_0 + \langle \rho_1 v_1 \rangle) = 0 \\ \frac{\partial}{\partial x} \left(\rho_0 + \frac{1}{2} \langle v_1^2 \rangle \right) + \frac{v_0}{\tau} = 0 \end{array} \right.$$

D'où on déduit :

$$\left\{ \begin{array}{l} \rho_{GS} v_0(x) + \langle \rho_1(t, x) v_1(t, x) \rangle_t = \text{constant} = 0 \\ \Delta U_{DS} \propto \rho_0(L) - \rho_0(0) = \frac{\langle v_1^2(t, 0) - v_1^2(t, L) \rangle_t}{2} - \int_0^L dx \frac{v_0(x)}{\tau} \end{array} \right.$$

On trouve finalement que la rectification du potentiel drain-source est donnée par :

$$\Delta U_{DS} = \frac{U_A^2}{4U_{GS}} f\left(\omega\tau, \frac{s\tau}{L}\right) \quad (1.13)$$

où f est le facteur de résonance. Il s'agit d'une grandeur sans dimension qui dépend de deux paramètres adimensionnés fixés par les propriétés du transistor, le potentiel du générateur U_{GS} et de la fréquence de l'onde électromagnétique incidente :

- $\omega\tau$ qui détermine le caractère résonant ou non du transistor
- $\frac{s\tau}{L}$ qui est relié à la longueur d'absorption de l'onde plasma dans le canal

1.6 Régimes de détection : résonants et non-résonants

L'expression générale du facteur de résonance f est donné par :

$$f\left(\omega\tau, \frac{s\tau}{L}\right) = f(\tilde{\omega}, \tilde{s}) = \frac{(\beta - 1) \sin^2(\tilde{k}_R) + (\beta + 1) \sinh^2(\tilde{k}_I)}{\cos^2(\tilde{k}_R) + \sinh^2(\tilde{k}_I)} \quad (1.14)$$

où $\beta = \frac{2\tilde{\omega}}{\sqrt{1 + \tilde{\omega}^2}}$ et $\tilde{k} = \frac{\tilde{\omega}}{\tilde{s}} \sqrt{1 + \frac{i}{\tilde{\omega}}}$ dont $\tilde{k}_R$ et $\tilde{k}_I$ désignent les parties réelle et imaginaire respectivement [6]. Le facteur de résonance peut alors se comporter de deux façons :

- avoir une structure résonante, auquel cas celui-ci est grand autour de fréquences spécifiques et faible ailleurs.
- être large bande, auquel cas il ne possède aucun pic de résonance mais est plus ou moins indépendant de la fréquence d'excitation.

Les paramètres contrôlant le caractère résonant ou non de ce facteur sont les deux grandeurs adimensionnées dont il dépend.

1.6.1 Régime résonant

On obtient la rectification résonante lorsque ces deux paramètres sont simultanément grands devant 1. Cela correspond alors aux critères suivants :

- $\tilde{\omega} = \omega\tau \gg 1$: la fréquence de collision électron-phonon est faible devant la fréquence de l'onde électromagnétique incidente.
- $\tilde{s} = \frac{s\tau}{L} \gg 1$: la distance de propagation de l'onde plasma donnée par $s\tau$ ici est grande devant la longueur du canal L .

Dans ce cas, le facteur de résonance se réduit à :

$$f\left(\omega\tau, \frac{s\tau}{L}\right) \approx \frac{3 \sinh^2\left(\frac{L}{2s\tau}\right) + \sin^2\left(\frac{\omega L}{s}\right)}{\sinh^2\left(\frac{L}{2s\tau}\right) + \cos^2\left(\frac{\omega L}{s}\right)} \quad (1.15)$$

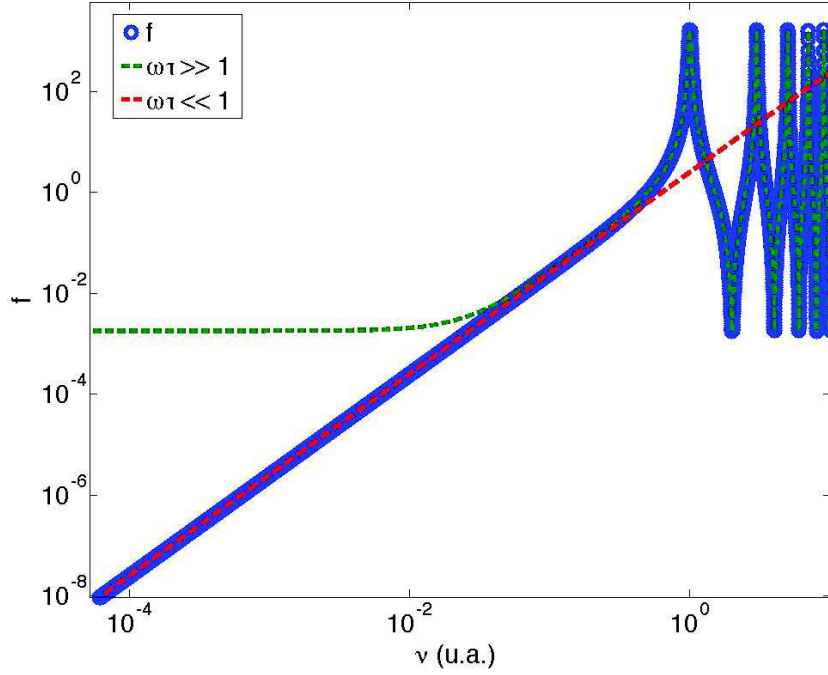


FIGURE 1.10 – Facteur de résonance f dans le cas résonant. La formule exacte est donnée par la courbe bleue. La courbe verte représente la limite de f lorsque $\omega\tau \gg 1$. Le caractère résonant est mis en évidence par l'apparition de pics de résonance, autour desquels le facteur de résonance varie de plusieurs ordres de grandeur. La courbe rouge représente la limite de f lorsque $\omega\tau \ll 1$.

Le caractère résonant est dû aux interférences entre l'onde plasma initialement générée à la source par l'antenne et ses réflexions multiples par les bords du canal. En effet, la distance de propagation étant grande devant la longueur du canal, l'onde plasma peut se propager sans absorption significative dans le canal. Celle-ci peut donc être réfléchi de nombreuses fois par les bords du canal avant d'être dissipée. Les interférences entre les multiples réflexions de l'onde plasma sont alors constructives lorsque $\frac{\omega L}{s} = \frac{\pi}{2} (1 + 2n)$ pour n entier, correspondant exactement aux fréquences propres d'oscillation du canal. Dans ce cas, l'onde plasma stationnaire qui s'établit dans le canal est de grande amplitude et les effets non-linéaires menant à la rectification sont alors importants (voir figure 1.10). Au contraire, les interférences destructives aux autres fréquences induisent une oscillation de densité de faible amplitude, conduisant à une faible rectification.

1.6.2 Régime non résonant

On obtient la limite non-résonante dans deux cas :

- lorsque $\frac{s\tau}{L} \ll 1$.
- et/ou lorsque $\omega\tau \ll 1$.

Dans le premier cas, l'absence de résonance est due à l'absence de réflexion de l'onde plasma et donc d'interférences. En effet, cette limite correspond au cas où le canal est long devant la distance de propagation de l'onde plasma. Lorsque celle-ci est générée à la source, elle est dissipée avant d'atteindre le drain. Dans ce cas, la contribution hyperbolique est dominante dans la formule (1.15) et la contribution qui dépend de la fréquence peut être négligée. On obtient alors $f \approx 3$.

Dans le second cas, $\tilde{\omega}$ tend vers 0, on obtient :

$$f(\tilde{\omega}, \tilde{s}) \approx \frac{(1 + 2\tilde{\omega}) \sinh^2\left(\frac{1}{\tilde{s}}\sqrt{\frac{\tilde{\omega}}{2}}\left(1 - \frac{\tilde{\omega}}{2}\right)\right) - (1 - 2\tilde{\omega}) \sin^2\left(\frac{1}{\tilde{s}}\sqrt{\frac{\tilde{\omega}}{2}}\left(1 + \frac{\tilde{\omega}}{2}\right)\right)}{\sinh^2\left(\frac{1}{\tilde{s}}\sqrt{\frac{\tilde{\omega}}{2}}\left(1 - \frac{\tilde{\omega}}{2}\right)\right) + \cos^2\left(\frac{1}{\tilde{s}}\sqrt{\frac{\tilde{\omega}}{2}}\left(1 + \frac{\tilde{\omega}}{2}\right)\right)} \quad (1.16)$$

Il est à noter que cette formule est différente de la référence [6]. En effet, celle-ci donne :

$$f_{DS}(\tilde{\omega}, \tilde{s}) \approx \frac{\sinh^2\left(\frac{1}{\tilde{s}}\sqrt{\frac{\tilde{\omega}}{2}}\right) - \sin^2\left(\frac{1}{\tilde{s}}\sqrt{\frac{\tilde{\omega}}{2}}\right)}{\sinh^2\left(\frac{1}{\tilde{s}}\sqrt{\frac{\tilde{\omega}}{2}}\right) + \cos^2\left(\frac{1}{\tilde{s}}\sqrt{\frac{\tilde{\omega}}{2}}\right)} \quad (1.17)$$

La différence entre les deux est simplement due à la prise en compte du terme suivant dans le développement de Taylor de β et $\tilde{k}$ en $\tilde{\omega}$. Lorsque $\tilde{s} \ll 1$, ce terme est négligeable et la formule (1.17) est suffisante. Cependant, lorsque $\tilde{s} \gg 1$, la formule (1.17) s'annule exactement à l'ordre $\tilde{\omega}$. En effet, à cet ordre, $\sinh^2\left(\frac{1}{\tilde{s}}\sqrt{\frac{\tilde{\omega}}{2}}\right) = \frac{\tilde{\omega}}{2\tilde{s}^2} = \sin^2\left(\frac{1}{\tilde{s}}\sqrt{\frac{\tilde{\omega}}{2}}\right)$.

Pour calculer le premier terme non-nul du facteur de résonance, il est nécessaire de considérer le terme d'ordre suivant, d'où la formule (1.16). Dans ce cas, on a :

$$f(\tilde{\omega}, \tilde{s}) \approx (1 + 6\tilde{s}^2) f_{DS}(\tilde{\omega}, \tilde{s})$$

Ainsi, l'évaluation du facteur de résonance par la formule asymptotique donnée en [6] est $1 + 6\tilde{s}^2$ fois inférieure à la vraie valeur asymptotique de la formule (1.14).

Dans la limite $\tilde{\omega} \ll 1$, la dissipation de l'onde plasma est très forte du fait du grand nombre de collisions électron-phonon par période d'oscillation de l'onde. La longueur d'absorption de cette onde est donnée par $s\sqrt{\frac{\tau}{\omega}}$ et est égale à sa longueur d'onde. Ainsi, même si le canal est court devant cette longueur de propagation, il n'y a pas d'interférences destructives dans le canal. En effet, le décalage de phase entre les multiples réflexions décroît avec la longueur du canal et tend vers 0 lorsque L devient très inférieur à la longueur de propagation. En conséquence, les interférences sont constructives pour toute fréquence d'excitation. Le facteur de résonance ne présente alors pas de pics de résonance (voir figure 1.11).

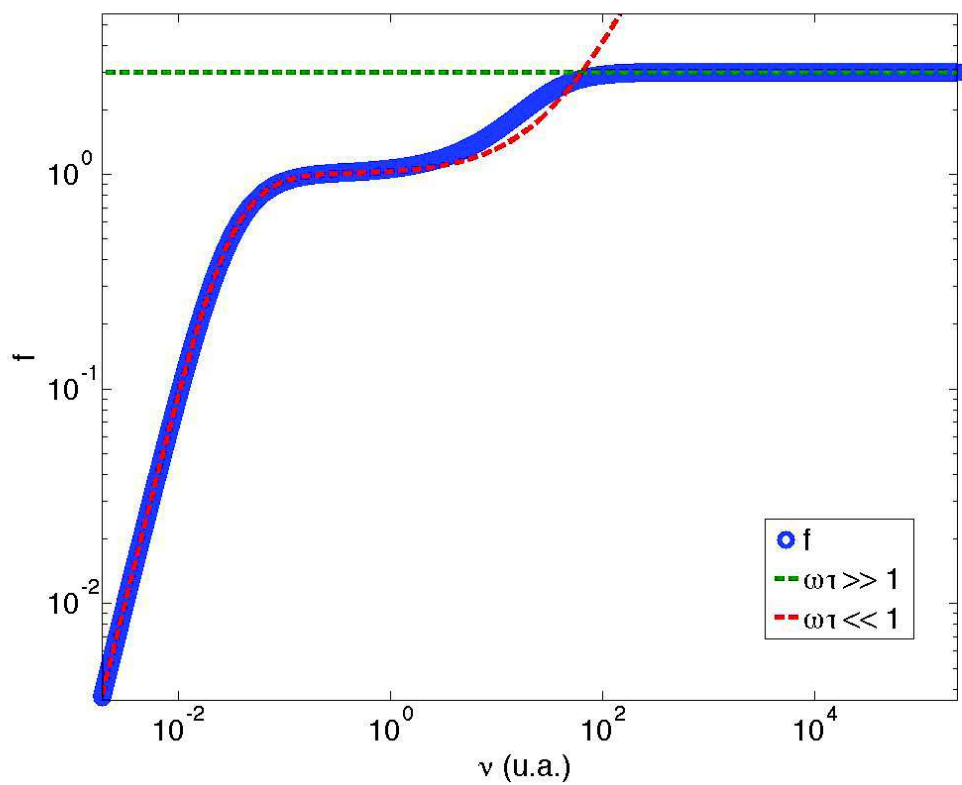


FIGURE 1.11 – Facteur de résonance dans le régime non-résonant (en bleu). Que $\omega\tau$ soit grand (limite en vert) ou faible (limite en rouge), il n’y a pas de pic de résonance car le canal est trop long ici.

Chapitre 2

Résolution de quelques limitations du modèle de Dyakonov et Shur

Le modèle Dyakonov-Shur permet de modéliser le comportement d'un transistor MOSFET lorsque le potentiel grille-source imposé est positif. En effet, l'approximation de la capacité plane relie la densité d'électron dans le canal au potentiel électrostatique grille-canal par une relation de proportionnalité. Justifiable lorsque ce potentiel est suffisamment grand, cette approximation n'est évidemment plus valable dans le cas d'un potentiel négatif car il serait associé à une densité négative. Une nouvelle relation charge/potentiel doit alors être considérée ou la résolution directe de l'équation de Poisson doit être effectuée. Dans ce chapitre, nous allons nous intéresser à deux nouveaux modèles de relation charge/potentiel. Le premier concerne la description du transistor lorsque U_{GS} est négatif et le second pour tout potentiel U_{GS} , unifiant le comportement au-dessus et en-dessous du seuil de formation du canal.

2.1 Tension négative : transistor sous le seuil

2.1.1 Modèle tension-charge sous le seuil

Lorsque le potentiel de grille est sous le seuil de formation du canal, les électrons du canal ne sont plus confinés à l'interface substrat/isolant. Un modèle de capacité doit alors prendre en compte l'extension spatiale du canal dans le substrat. Dans ce cas, la densité d'électron du canal est exponentiellement décroissante avec le potentiel de grille[8]. On obtient alors la nouvelle relation potentiel/densité suivante :

$$\rho = \rho_0 \exp\left(\frac{U}{U_0}\right) \quad (2.1)$$

où ρ_0 est la densité électronique à la tension de référence U_0 . En injectant la relation (2.1) dans le système (1.4), et en prenant la résistivité en compte, on aboutit au système

suivant :

$$\begin{cases} \frac{\partial}{\partial t} \rho + \frac{\partial}{\partial x} (\rho v) = 0 \\ \frac{\partial}{\partial t} v + \frac{qU_0}{m^*} \frac{\partial \rho}{\partial x} + v \frac{\partial}{\partial x} v + \frac{v}{\tau} = 0 \end{cases} \quad (2.2)$$

avec les conditions aux bords suivantes :

$$\begin{cases} \rho(t, 0) = \rho_0 \exp\left(\frac{U_{GS}}{U_0}\right) \\ q\rho(t, L)v(t, L) = i_D \end{cases}$$

2.1.2 Détection THz sous le seuil

La théorie de la détection THz au-dessus de la tension seuil du canal, discutée en partie 1.5, prédit une rectification du potentiel drain-source donnée par l'équation (1.13) que l'on rappelle :

$$\Delta U_{DS} = \frac{U_A^2}{4U_{DS}} f(\tilde{\omega}, \tilde{s})$$

Dans le cas de la détection résonante, autour d'une fréquence de résonance, cette équation se réduit à :

$$\Delta U_{DS} \approx \frac{q\tau^2}{m^*L^2} U_A^2$$

Ainsi, on voit que dans ce cas, la rectification est indépendante du potentiel grille-source imposé par le générateur. Au contraire, loin des fréquences de résonance, on obtient :

$$\Delta U_{DS} \approx \frac{3U_A^2}{4U_{GS}}$$

On remarque donc que la rectification est d'autant plus importante que le potentiel du générateur est faible. De même, dans la limite $\omega\tau \ll 1$, la détection devient non-résonante et on obtient une relation similaire entre le potentiel de rectification et le potentiel du générateur :

$$\Delta U_{DS} \approx \frac{U_A^2}{4U_{GS}} \quad (2.3)$$

La résistivité relativement importante des MOSFET actuels restreint les régimes de détection expérimentale aux cas non-résonants. Ainsi, il est intéressant de travailler avec un potentiel U_{GS} proche de zéro. Dans la limite U_{GS} nul, on obtient la divergence de la rectification d'après la formule (1.13). Cette prédiction entre en contradiction directe avec les résultats expérimentaux, qui montrent une rectification finie à potentiel nul relativement au seuil de formation du canal [8], ainsi qu'avec l'intuition physique qu'aucun phénomène réel n'est divergent.

Lorsque le potentiel U_{GS} tend vers 0, l'approximation du canal graduel, permettant l'établissement de la formule (2.3), n'est plus correcte. En effet, on sait déjà que cette

approximation ne peut pas être valable pour des potentiels négatifs. Il existe donc une zone de transition autour de U_{GS} nul permettant le passage de la relation (1.3) à la relation (2.1). Un modèle valable à la fois en-dessous du seuil de formation du canal et au-dessus de ce dernier est ce que l'on appelle un modèle unifié et on en décrira certains dans la section 2.2 suivante.

La théorie de la détection non-résonante d'onde électromagnétique THz sous le seuil du transistor est décrite par la référence [8]. En utilisant la loi d'ohm locale, l'équation de Navier-Stokes se réduit à :

$$\rho v = -\frac{qU_0\tau}{m^*} \frac{\partial}{\partial x} \rho \quad (2.4)$$

Cette approximation suppose $\omega\tau$ très faible devant 1 ainsi qu'une contribution advective négligeable. Dans ce cas, on peut injecter la relation (2.4) dans l'équation de continuité pour aboutir à l'équation de la chaleur :

$$\frac{\partial}{\partial t} \rho = D \Delta \rho \quad (2.5)$$

Le coefficient de diffusion D est donné par :

$$D = \frac{qU_0\tau}{m^*}$$

La résolution de l'équation (2.5) est académique. La seule difficulté vient des conditions aux bords. La condition imposée coté drain est simple :

$$\left. \frac{\partial}{\partial x} \rho \right|_L = -\frac{j_L}{D}$$

où j_L est le courant imposé au niveau du drain. Coté source en revanche, la condition au bord est donnée par :

$$\rho(t, 0) = \rho_0 \exp\left(\frac{U_{GS}}{U_0}\right) \exp\left(\frac{U_A}{U_0} \cos(\omega t)\right) \quad (2.6)$$

La décomposition de la condition (2.6) en série de Fourier peut se faire à l'aide des fonctions spéciales de Bessel. Si on ne s'intéresse qu'au cas de faible excitation, on peut se contenter du développement limité de (2.6) à l'ordre 2 en U_A . Pour un canal long devant la longueur d'absorption de l'onde plasma générée dans le canal, on obtient la rectification drain-source suivante :

$$\Delta U_{DS} = \frac{U_A^2}{4U_0} \quad (2.7)$$

La rectification non-résonante sous le seuil du transistor est donc indépendante du potentiel grille-source appliqué et est finie. Dans la référence [8], les auteurs ont en réalité également considéré la contribution d'un courant de fuite j_F entre la grille et le canal. Un tel phénomène se traduit par l'apparition d'un terme source dans l'équation

de conservation. Ce terme est positif si les électrons entrent dans le canal par la grille en traversant la couche isolante les séparant, et négatif s'ils sortent du canal. Ce courant de fuite est une conséquence de l'effet tunnel que l'on a mentionné en partie 1.1.1. Bien qu'un isolant soit situé entre la grille et le canal du transistor, lorsque celui-ci est suffisamment fin, les électrons peuvent passer du canal à la grille et vice-versa par effet tunnel. Ce courant de fuite est dû au fait que la fonction d'onde des électrons décroît exponentiellement à travers l'isolant. Cependant, lorsque ce dernier est suffisamment fin, cette fonction d'onde peut avoir une amplitude suffisamment importante de l'autre côté de l'isolant. Les électrons peuvent alors traverser l'isolant avec une probabilité non-négligeable.

Lorsqu'un tel courant de fuite est considéré, la rectification du potentiel drain-source en est affectée. On obtient alors :

$$\Delta U_{DS} = \frac{U_A^2}{4U_0} \left(\frac{1}{1 + \zeta \exp\left(-\frac{U_{GS}}{U_0}\right)} - \frac{1}{\left(1 + \zeta \exp\left(-\frac{U_{GS}}{U_0}\right)\right)^2 |\cosh(kL)|^2} \right)$$

où $\zeta = \frac{j_L L^2 m^*}{2C\tau U_0^2}$ est la contribution du courant de fuite et $k = \sqrt{\frac{i\omega m^*}{qU_0\tau}}$ est le vecteur d'onde complexe de l'onde plasma. Du fait du courant de fuite, on obtient la décroissance exponentielle de la rectification avec U_{GS} tendant vers $-\infty$. On constate donc qu'il est effectivement plus intéressant d'utiliser un potentiel grille-source proche du seuil du transistor afin de maximiser le signal de rectification. Un modèle unifié s'impose donc.

2.2 Modèle unifié de la charge

Pour pouvoir décrire le fonctionnement d'un MOSFET à la fois en-dessous et au-dessus du seuil de formation du canal, il est nécessaire d'avoir recours à un modèle unifié de sa charge. Un tel modèle est donné par la référence [9], où l'on a la relation suivante entre le potentiel grille-source U_{GS} et la densité d'électron dans le canal ρ :

$$U_{GS} = V_T + \frac{q}{C} (\rho - \rho_T) + \eta V_{th} \ln \left(\frac{\rho}{\rho_T} \right) \quad (2.8)$$

Les différents paramètres de l'équation (2.8) représentent les grandeurs physiques suivante :

- V_T : tension seuil de formation du canal
- q : charge de l'électron
- C : capacité du système grille-canal
- ρ_T : densité d'électron à la tension seuil
- η : facteur d'idéalité du transistor
- V_{th} : tension thermique du transistor

La relation (2.8) peut être inversée pour obtenir la densité d'électron dans le canal en fonction du potentiel appliqué à l'aide de la fonction W_0 de Lambert (calcul en annexe A.1). On obtient alors :

$$\rho = \frac{C\eta V_{th}}{q} W_0 \left(\frac{q\rho_T}{C\eta V_{th}} \exp \left(\frac{U_{GS} - V_T + \frac{q\rho_T}{C}}{\eta V_{th}} \right) \right) \quad (2.9)$$

Dans la limite $U_{GS} - V_T$ fortement négatif, on obtient grâce au développement limité de W_0 autour de 0 la formule suivante :

$$\rho \approx \underbrace{\rho_T \exp \left(\frac{q\rho_T}{C\eta V_{th}} \right)}_{\rho_0} \exp \left(\frac{U_{GS} - V_T}{\eta V_{th}} \right)$$

Au contraire, dans la limite $U_{GS} - V_T$ fortement positif, on obtient grâce au développement asymptotique de W_0 à l'infini :

$$\rho \approx \frac{C}{q} (U_{GS} - V_T)$$

Le modèle unifié (2.8) tend donc bien vers les deux comportements asymptotiques attendus en-dessous et au-dessus du seuil V_T .

L'utilisation de ce modèle est nécessaire lorsque la tension dans le canal passe alternativement en-dessous et au-dessus de V_T . Notamment, en régime de détection, il est nécessaire lorsque l'amplitude du potentiel U_A générée par l'antenne est supérieure à la tension de commutation $U_{GS} - V_T$. On qualifiera ce régime de régime de forte intensité. La formule (2.8) sera considérée dans la modélisation numérique pour le problème de la détection à forte intensité.

La fonction W_0 de Lambert est une fonction spéciale, dont l'évaluation n'est pas implémentée systématiquement dans tous les logiciels de calculs. Ainsi, il est parfois préféré au modèle (2.8) un autre modèle unifié [10] plus simple, où on a la relation suivante :

$$\rho = \rho_0 \ln \left(1 + \exp \left(\frac{U_{GS} - V_T}{\eta V_{th}} \right) \right) \quad (2.10)$$

L'équation (2.10) possède les mêmes développements asymptotiques en 0 et à l'infini que l'équation (2.8) au premier ordre. Les deux modèles sont néanmoins légèrement différents, notamment lorsque U_{GS} est autour de la tension de seuil V_T . Ceux-ci sont représentés en figure 2.1 ainsi qu'avec leurs développements asymptotiques.

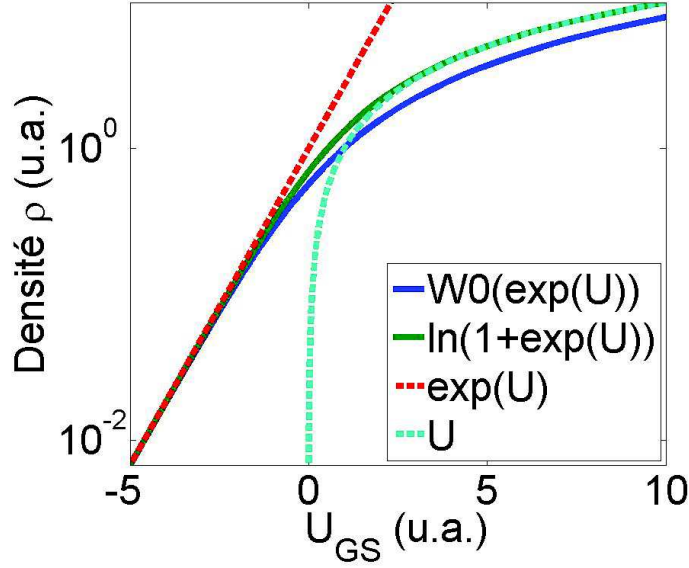


FIGURE 2.1 – Comparaison des deux modèles unifiés. En $+\infty$, le premier modèle (2.8) tend vers $\rho \approx U - \ln(U)$. L'écart relatif avec le second modèle (2.10) décroît donc vers 0.

En régime de faible excitation, on peut obtenir facilement la rectification ΔU_{DS} en linéarisant la relation densité-tension autour de la tension statique U_{GS} . En effet, de cette manière, on retrouve une relation linéaire entre la densité du canal et la tension dans celui-ci. On peut donc appliquer la théorie Dyakonov-Shur au-dessus du seuil de formation du canal en remplaçant U_{GS} par une tension de référence :

$$U_{ref}(U_{GS}) = \frac{\rho(U_{GS})}{\frac{d}{dU}\rho|_{U_{GS}}}$$

En effet, la linéarisation de la relation densité-potentiel revient à une redéfinition du paramètre α du système (1.4). Cela modifie donc la vitesse de référence de déplacement des ondes plasma s , et donc, du fait de sa relation avec la tension U_{GS} , cela modifie le potentiel effectif perçu dans le canal. Cette tension de référence est donnée par :

$$U_{ref} = \eta V_{th} \left\{ 1 + W_0 \left[\frac{qn_T}{C\eta V_{th}} \exp \left(\frac{qn_T}{C\eta V_{th}} + \frac{U_{GS} - V_T}{\eta V_{th}} \right) \right] \right\} \quad (2.11)$$

pour le premier modèle unifié et :

$$U_{ref} = U_0 \left(1 + \exp \left(-\frac{U_{GS}}{U_0} \right) \right) \ln \left(1 + \exp \left(\frac{U_{GS}}{U_0} \right) \right) \quad (2.12)$$

pour le second. En remplaçant U_{GS} dans l'équation (2.3) par la relation (2.12), on retrouve bien le résultat donné en [8] dans le cas d'un courant de fuite nul. Les inverses

des tensions de référence données par les équations (2.11) et (2.12) sont représentées en fonction du potentiel U_{GS} sur la figure 2.2. Le potentiel de rectification U_{DS} , proportionnel à l'inverse de la tension de référence, est toujours fini, quelque soit U_{GS} . A forte tension négative, la rectification est constante et on retrouve le résultat sous le seuil en l'absence de courant de fuite. A forte tension positive, la rectification est inversement proportionnelle au potentiel U_{GS} . Ainsi, le modèle unifié permet de modéliser le comportement d'un MOSFET soumis à une faible excitation, que celui-ci soit en-dessous ou au-dessus du seuil de formation du canal. Notamment, le comportement à U_{GS} proche du seuil est le plus intéressant. En effet, c'est autour du seuil que la rectification non-résonante est maximale, du fait de la décroissance exponentielle du potentiel de rectification sous le seuil en présence d'un courant de fuite.

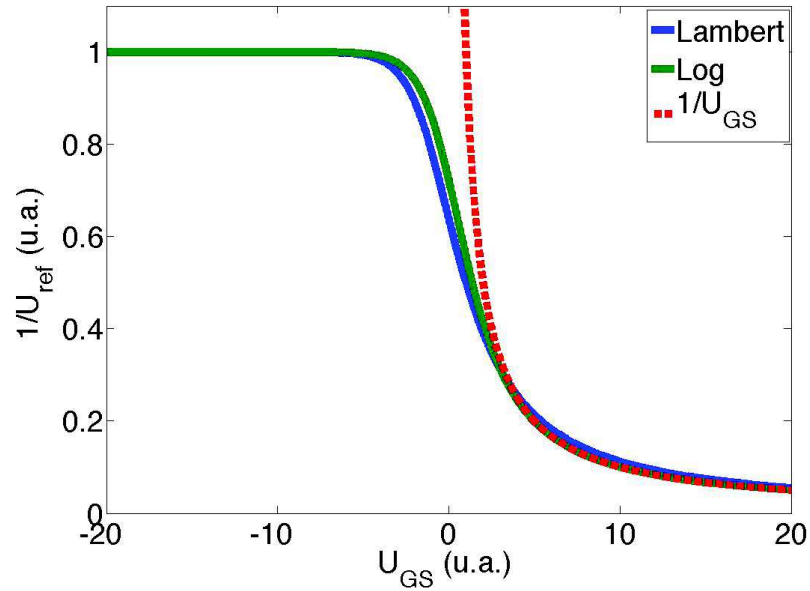


FIGURE 2.2 – Inverse de la tension de référence U_{ref} en fonction du potentiel U_{GS} appliqué. La courbe bleue correspond au modèle (2.8) qui utilise la fonction de Lambert. La courbe verte correspond au second modèle (2.10). La courbe rouge correspond au modèle au dessus du seuil pour référence.

Deuxième partie

Simulation numérique du modèle Dyakonov-Shur

Chapitre 3

Instabilité de Dyakonov-Shur, saturation et ressaut hydraulique

Nous avons montré comment le modèle Dyakonov-Shur permet de détecter un rayonnement térahertz incident en le convertissant en un signal de rectification basse fréquence aisément détectable avec un dispositif de mesure électronique conventionnel. Cependant, le modèle Dyakonov-Shur fut initialement formulé dans le but de résoudre le problème de l'émission de rayonnement électromagnétique térahertz. On a ainsi montré comment la formulation initiale du système d'équation hydrodynamique admet des solutions stationnaires instables. Pour cela, une méthode perturbative, considérant la contribution oscillante comme faible devant la contribution stationnaire, au premier ordre a été développée. Cependant, cette méthode prédit la croissance infinie des ondes plasmas ce qui n'est évidemment pas physique. Cependant, les contributions d'ordres supérieurs ayant été négligées, on peut s'attendre à ce que ces termes contribuent à la saturation de l'amplification des ondes plasmas. Dans ce chapitre, nous allons étudier les régimes d'oscillations à forte amplitude, où les effets non-linéaires du modèle sont non-négligeables. Puis nous allons étudier l'effet de la présence de résistivité au courant dans le canal sur la stabilité des modes propres d'oscillation plasma. Enfin, nous étudierons quelques modèles de viscosité, essayant de rendre compte des collisions électron-électron dans le canal.

3.1 Régime d'oscillation aux temps longs

3.1.1 Limites de la méthode perturbative

Lorsque toutes les formes de dissipations envisageables dans le canal d'un MOSFET sont négligeables, le modèle Dyakonov-Shur initial prédit l'instabilité de certaines valeurs de courants stationnaires. Cette instabilité se traduit comme on l'a montré partie 1.3 en la croissance initialement exponentielle de petites fluctuations de densité et de courant, menant à la génération spontanée d'ondes plasma. Au bout d'un certain temps, l'amplitude de ces fluctuations devient comparable à la charge appliquée au transistor

et la méthode perturbative utilisée pour prédire la croissance des fluctuations n'est alors plus valide. La non-linéarité du système devient en effet non-négligeable et des phénomènes de dissipation dus à ces interactions non-linéaires apparaissent alors. Notamment, lorsque l'amplitude d'oscillation est suffisante, un phénomène similaire au ressaut hydraulique peut se produire. Ce phénomène, observé dans le cadre des vagues en eau peu profonde, est attendu et vérifié numériquement dans le modèle Dyakonov-Shur. En effet, les deux modèles sont mathématiquement équivalents et des phénomènes similaires doivent donc être observés. Lors de l'apparition du ressaut, la densité d'électrons et le courant au sein du canal deviennent discontinus. La formation de cette discontinuité s'accompagne d'une forte dissipation énergétique. Ainsi, dans le régime Dyakonov-Shur instable, l'amplification des ondes plasmas finit par être compensée et donne lieu à un régime final oscillant à amplitude constante.

L'apparition du ressaut vient imposer une contrainte sur le choix de la méthode de discrétisation utilisée. En effet, beaucoup d'algorithmes deviennent instables à l'apparition de discontinuités, ce qui, ajouté à l'instabilité intrinsèque du modèle Dyakonov-Shur, rend sa modélisation ardue. Il convient alors d'utiliser une méthode appropriée à la modélisation de discontinuités et ne prévenant pas leurs apparitions. En effet, certains algorithmes à l'origine instables peuvent être stabilisés par différentes méthodes souvent au prix d'une légère diffusion de la discontinuité. Celle-ci se retrouve alors étalée sur quelques points de discrétisation spatiale. Parfois, cette stabilisation plus ou moins arbitraire peut malheureusement engendrer des différences de comportement notables entre la solution numérique calculée et la solution réelle du système. Il est donc d'autant plus important dans ce cas d'établir des critères de vérification de la solution numérique.

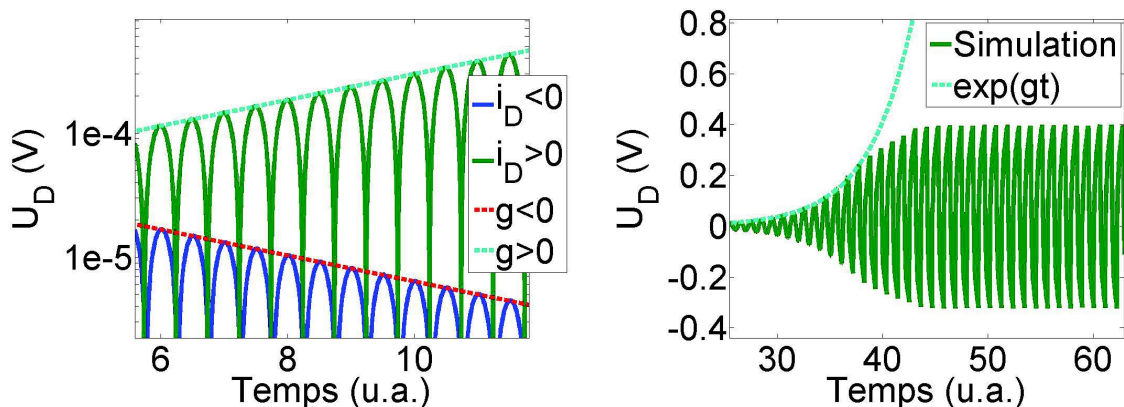


FIGURE 3.1 – Comparaison entre l'évolution numérique temporelle du potentiel de drain (traits pleins) et de la croissance exponentielle prédite par une théorie perturbative (traits pointillés). A courant de drain négatif correspond un gain négatif et la perturbation disparaît. Pour un courant positif, l'amplitude de la perturbation croît exponentiellement comme attendu avec un facteur de gain g en accord avec la valeur théorique (une droite en échelle lin-log à gauche). L'onde plasma finit par saturer à une amplitude d'environ 0.5V, de l'ordre de la différence de potentiel statique (ici 1V) imposée entre la grille et la source.

Sur la figure 3.1, on peut voir le résultat d'une simulation numérique du modèle Dyakonov-Shur pour un courant de drain positif et négatif. L'état initial choisi est la somme de la solution stationnaire du système et d'une légère perturbation. Dans le cas d'un courant de drain négatif, on a vu partie 1.3 que l'on s'attend à voir une dissipation exponentielle de la contribution perturbative initiale avec le temps. L'amplitude théorique de la perturbation (en traits pointillés rouge) à l'instant t est alors donnée par $\exp(gt) = \exp(-|g|t)$ et correspond à l'amplitude calculée par intégration numérique. Au contraire, pour un courant de drain positif, on s'attend à voir un accroissement exponentiel de la contribution perturbative. L'amplitude calculée par intégration numérique de la perturbation n'est, comme attendu, en accord avec la prédiction théorique (en pointillés bleu clair) que pour des temps relativement courts. On observe en effet que l'amplitude d'oscillation calculée numériquement finit par saturer et ne croît donc pas exponentiellement *ad infinitum*. En effet, au bout d'un certain temps, l'amplitude de la perturbation devient de l'ordre de la contribution stationnaire. Dans ce cas, on s'attend à ce que la méthode perturbative ne soit plus correcte car celle-ci néglige des termes d'ordres supérieurs donc la contribution devient équivalente aux termes considérés. La méthode numérique ne reposant pas sur une méthode perturbative, celle-ci ne néglige pas les interactions non-linéaires du modèle et devrait prédire le bon résultat.

3.1.2 Etude de la discontinuité

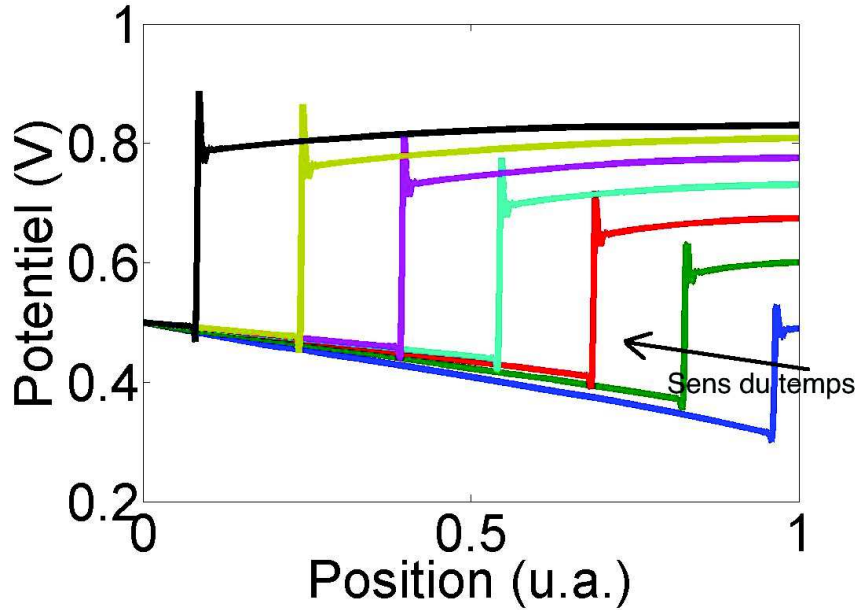


FIGURE 3.2 – Propagation du ressaut hydraulique dans le canal, du drain à droite vers la source à gauche avec le temps, une fois le régime final atteint. Le phénomène de Gibbs (oscillations autour de la discontinuité) est visible du fait de la méthode de discrétisation spatiale utilisée.

Lorsque l'amplitude de la contribution perturbative devient importante, un ressaut hydraulique peut se former dans le canal (voir figure 3.2). Pour cela, il est nécessaire de remplir l'une des conditions suivantes [11] :

- la dérivée de la densité électronique est positive et l'onde de densité se propage vers la source (à gauche).
- la dérivée de la densité électronique est négative et l'onde de densité se propage vers le drain (à droite).

Dans les autres cas, le ressaut se dissipe et la solution du système Dyakonov-Shur redevient continue (voir figure 3.3).

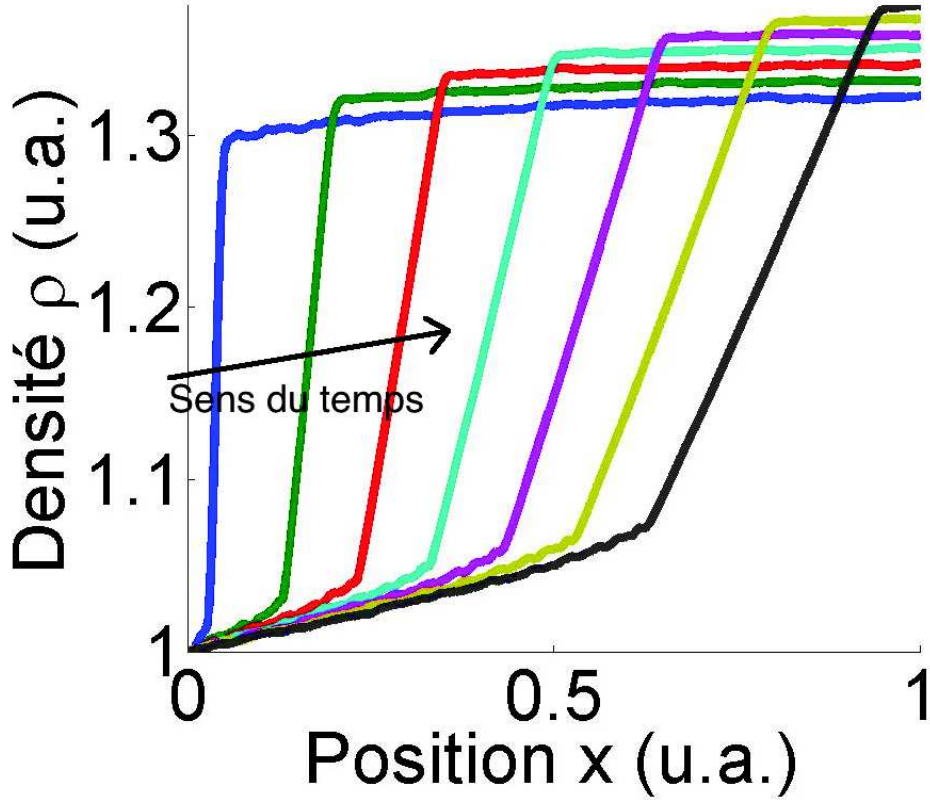


FIGURE 3.3 – Dissipation du ressaut hydraulique après que celui-ci soit réfléchi à la source. Comme la réflexion change le sens de propagation de la discontinuité, celle-ci n'est plus stable et se dissipe. Le ressaut se reforme et se dissipe alors périodiquement avec une fréquence correspondant à une fréquence de résonance Dyakonov-Shur du transistor. Après la réflexion du ressaut, le phénomène de Gibbs disparaît presque instantanément car la solution redevient continue.

Du fait de la méthode de discrétisation spatiale utilisée, on peut voir sur la figure 3.2 le phénomène de Gibbs aux abords de la discontinuité. C'est la présence de ce phénomène qui rend certains schémas de discrétisation instables car les oscillations de Gibbs sont amplifiées par ceux-ci et finissent par diverger.

La vitesse de propagation v_R théorique du ressaut peut être facilement calculée dans le modèle Dyakonov-Shur et est donnée par (voir calcul en partie 5.2.1) :

$$v_R = \bar{v} + \text{sign}(\delta U \delta v) s \sqrt{\frac{\bar{U}}{U_{GS}}}$$

où $\bar{U}$ et $\bar{v}$ sont respectivement les tension et vitesse moyennes aux abords de la discontinuité et δU et δv les différences de tension et de vitesse autour de la discontinuité.

Cette vitesse peut être comparée à la vitesse de propagation numérique mesurée lors de la simulation et celles-ci sont sensiblement les mêmes avec un écart relatif n'excédant pas 1%, ce qui conforte la validité du schéma (voir table 3.1) et donc de la modélisation effectuée. En effet, suivant le schéma utilisé, la vitesse de propagation d'une discontinuité peut être facilement sous ou sur évaluée. Cet écart dans la vitesse de propagation de la discontinuité est alors d'autant plus important que la discontinuité est importante et peut facilement dépasser les 20%. Dans ce cas, le modèle numérique ne converge pas vers la solution réelle du système.

Vitesse numérique (u.a.)	-2,602	-2,636	-2,675	-2,719	-2,761	-2,801
Vitesse théorique (u.a.)	-2,628	-2,656	-2,694	-2,734	-2,776	-2,815
Ecart relatif	1,0%	0,76%	0,69%	0,58%	0,54%	0,49%

TABLE 3.1 – Table de comparaison de la vitesse théorique de propagation d'un ressaut avec la vitesse numérique observée pour la figure 3.2

Un tel écart se produit dans le cas des modèles numériques dit upwind/downwind (voir partie 5.2.1). Dans ces schémas, le sens de propagation de l'information est pris en compte et introduit une asymétrie entre les différentes directions spatiales de propagation. Une conséquence directe de ce biais est que lorsque la vitesse de propagation de l'information est différente des deux côtés d'une discontinuité, comme dans le cas du modèle Dyakonov-Shur, la vitesse de propagation d'une discontinuité dépend directement du modèle utilisé. Dans un cas, celle-ci sera surévaluée et la discontinuité se propagera bien plus vite qu'elle ne le devrait, et dans l'autre sous évaluée et la discontinuité se propagera trop lentement. Pour obtenir la bonne vitesse de propagation, il est nécessaire d'utiliser un schéma dit centré, qui est symétrique dans toutes les directions de propagation.

3.2 Amortissement par résistivité et diminution du gain

Comme on l'a montré en partie 1.3, le critère de stabilité de la solution stationnaire du système Dyakonov-Shur est donné par le signe du gain perturbatif. Lorsqu'il est positif, celui-ci entraîne la génération d'ondes plasmas. Après un temps t , l'amplitude de l'onde plasma est multipliée par un facteur $\exp(gt) > 1$ pour $g > 0$. Dans le cas d'une résistivité nulle, on a montré que le gain g est donné par :

$$g = \frac{s}{2L} \left(1 - \left(\frac{v_s}{s} \right)^2 \right) \ln \left(\frac{s + v_s}{s - v_s} \right)$$

où on rappelle que v_s est la vitesse de dérive stationnaire induite par le courant de drain et s l'unité naturelle de vitesse des ondes plasma dans le canal. Lorsque la résistivité

du transistor est non-nulle, le calcul du gain n'est pas aussi aisé. On peut toutefois déterminer une valeur approchée du gain dans le cas $|v_s| \ll s$ [12] :

$$g_\tau = g - \frac{1}{2\tau} \quad (3.1)$$

où on rappelle que τ est le temps de libre parcours moyen des électrons avant une collision avec le réseau cristallin. Pour calculer le gain perturbatif du transistor de manière générale, on utilisera alors une approche numérique. Pour cela, deux méthodes ont été envisagées. La première consiste à résoudre le système Dyakonov-Shur perturbatif dans l'espace de Fourier pour la variable temporelle et direct pour la variable spatiale. On extrait ainsi les modes perturbatifs de la même manière que dans le cas sans résistivité et on lit directement le gain dans la partie imaginaire de la pulsation complexe obtenue. La seconde méthode consiste à intégrer dans l'espace direct à la fois temporel et spatial le système Dyakonov-Shur. En partant d'une solution initiale arbitrairement proche de la solution stationnaire du système, on peut mesurer le taux d'amplification de l'écart entre la solution obtenue et la solution stationnaire. Dans la partie suivante, nous déterminerons la condition de génération d'onde THz en utilisant ces deux méthodes, en commençant par la méthode de Fourier.

3.2.1 Modes perturbatifs et pulsation complexe dans l'espace de Fourier

Pour calculer les modes perturbatifs d'oscillations ainsi que leurs pulsations complexes, on développe comme précédemment les champs de densité et de vitesse en une somme d'une contribution stationnaire et d'une petite contribution perturbative. On linéarise les termes perturbatifs du système obtenu en négligeant les termes quadratiques et supérieurs ρ_p^2 , $\rho_p v_p$ et v_p^2 . On obtient alors le système d'équation suivant :

$$\begin{pmatrix} \frac{\partial}{\partial t} & 0 \\ 0 & \frac{\partial}{\partial t} + \frac{1}{\tau} \end{pmatrix} \begin{pmatrix} \rho_p \\ v_p \end{pmatrix} + \frac{\partial}{\partial x} \left[\begin{pmatrix} v_s & \rho_s \\ 1 & v_s \end{pmatrix} \begin{pmatrix} \rho_p \\ v_p \end{pmatrix} \right] = 0 \quad (3.2)$$

En prenant la transformée de Fourier temporelle de ce système, on obtient :

$$\begin{pmatrix} -i\omega & 0 \\ 0 & -i\omega + \frac{1}{\tau} \end{pmatrix} \begin{pmatrix} \tilde{\rho} \\ \tilde{v} \end{pmatrix} + \frac{\partial}{\partial x} \left[\begin{pmatrix} v_s & \rho_s \\ 1 & v_s \end{pmatrix} \begin{pmatrix} \tilde{\rho} \\ \tilde{v} \end{pmatrix} \right] = 0 \quad (3.3)$$

On a vu que contrairement au cas sans résistivité, les champs stationnaires ρ_s et v_s dépendent ici de la position x et l'équation (3.3) n'est donc pas trivialement intégrable. On peut cependant réécrire l'équation (3.3) de la manière suivante :

$$\frac{\partial}{\partial x} \left[\begin{pmatrix} v_s & \rho_s \\ 1 & v_s \end{pmatrix} \begin{pmatrix} \tilde{\rho} \\ \tilde{v} \end{pmatrix} \right] + \begin{pmatrix} 0 & 0 \\ 0 & \frac{1}{\tau} \end{pmatrix} \begin{pmatrix} \tilde{\rho} \\ \tilde{v} \end{pmatrix} = i\omega \begin{pmatrix} \tilde{\rho} \\ \tilde{v} \end{pmatrix} \quad (3.4)$$

afin de reconnaître là un problème aux valeurs propres linéaire. En choisissant une discrétisation des différents opérateurs qui agissent sur le vecteur discrétisé $\begin{pmatrix} \tilde{\rho} \\ \tilde{v} \end{pmatrix}$, on obtient directement les vecteurs propres de l'équation (3.4) et leurs valeurs propres associées par diagonalisation. Différentes méthodes ont été utilisées pour cela et sont décrites et comparées en partie 5.2.3.

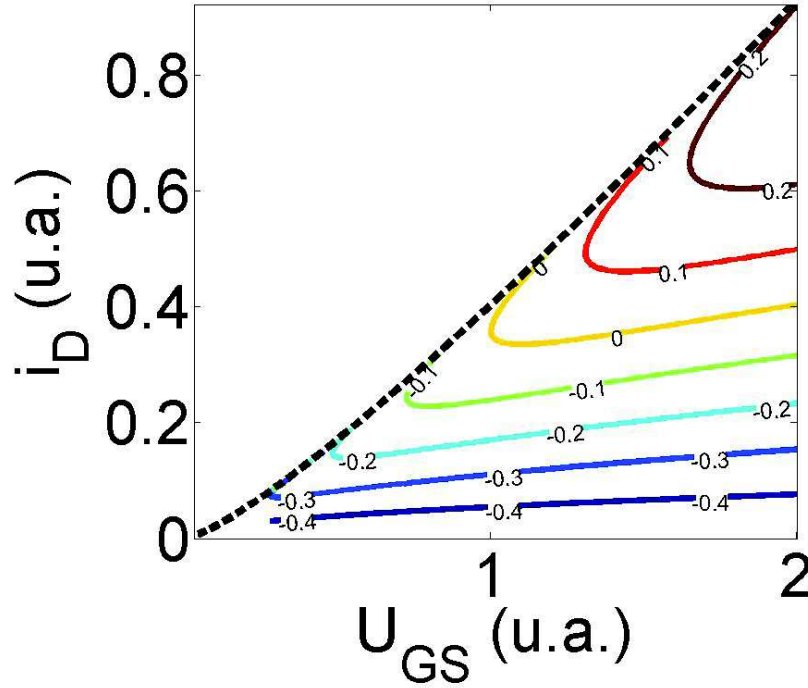


FIGURE 3.4 – Carte de niveau de gain g en fonction du potentiel grille-source en abscisse et du courant de drain en ordonnée. Sous la ligne jaune, le gain est négatif. La solution stationnaire y est donc stable. Au-dessus de la ligne jaune, le gain est positif. La solution stationnaire y est instable.

Les courbes de niveau Fig.3.4 représentent les lignes de niveau telles que le produit $g\tau$ (grandeur adimensionnée) soit constant en fonction du potentiel grille-source (en unité arbitraire) et du courant de drain (unité arbitraire également) imposés. Sur cette carte, on peut distinguer 3 zones :

- la première zone en bas (en dessous de la ligne pointillée noire et de la ligne jaune) représente la zone où le transistor est stable. Dans cette zone le gain est négatif et les fluctuations autour de la solution stationnaire sont amorties.
- la deuxième zone en haut à droite (sous la ligne noire et au dessus de la jaune) représente la zone où le transistor est instable. Dans cette zone le gain est positif et les fluctuations sont amplifiées avec le temps. C'est le régime que l'on souhaite

atteindre pour émettre un rayonnement THz car le transistor se met à osciller spontanément.

- la troisième zone en haut (au dessus de la ligne noire) représente une zone inaccessible au modèle. En effet, cette zone correspond à un courant de drain supérieur au courant critique de pincement. Cette zone est inaccessible car on a montré partie 1.4 que l'imposition d'un courant de drain supérieur au courant critique est incompatible avec le modèle Dyakonov-Shur.

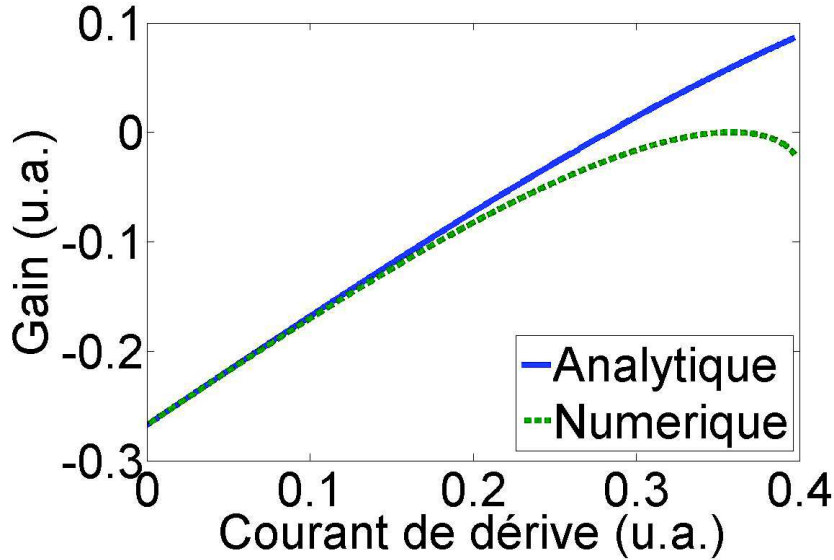


FIGURE 3.5 – Comparaison entre la formule analytique (3.1) et le calcul numérique du gain du transistor en fonction du courant de dérivation. La prédiction analytique est supérieure au calcul numérique car les approximations permettant d'obtenir la formule (3.1) ne sont plus valables (faible courant de dérivation).

Une coupe de cette carte à U_{GS} fixé permet d'obtenir la figure 3.5 qui compare le gain prédit théoriquement à faible courant avec le gain numérique calculé. On constate que lorsque les approximations analytiques sont valables, pour de faibles courants tels que $|v_s(0)| \ll s$, les deux méthodes prédisent la même valeur du gain comme le montre la superposition des deux courbes, ce qui conforte la validité des deux méthodes. En revanche, lorsque le courant augmente, le gain calculé numériquement croît moins vite que la prédiction analytique. L'effet du courant est donc plus important que ce que l'approximation de faible courant prend en compte. Ainsi, on constate qu'il semble plus difficile d'obtenir un gain positif qu'attendu et donc d'atteindre le régime d'instabilité de Dyakonov-Shur.

En dessous d'une certaine valeur de tension grille-source, le gain du transistor est négatif pour toute valeur de courant de drain Fig.3.4. Cette tension grille-source minimum à atteindre est due à la compétition entre l'énergie électrostatique fournie au

transistor et l'énergie dissipée par la résistivité du transistor. En effet, pour pouvoir espérer atteindre le régime de Dyakonov-Shur, il est nécessaire de fournir au transistor une énergie supérieure à l'énergie dissipée par résistivité. Dans ce cas, l'énergie peut être accumulée dans le transistor et se manifeste dans l'apparition d'ondes plasmas au sein du canal. L'énergie électrostatique apportée à un électron est de l'ordre de $E_{\text{elec.}} = qU_{GS}$. L'énergie dissipée par résistivité est de l'ordre de $E_{\tau} = \frac{m^* L^2}{\tau^2}$. Ainsi, pour atteindre le régime instable, il est au moins nécessaire d'avoir $E_{\text{elec.}} \gg E_{\tau}$. Grâce à la méthode numérique développée, on peut obtenir la valeur minimale du potentiel U_{GS} et donc de l'énergie $E_{\text{elec.}}$ telle qu'il existe une solution stationnaire instable. On obtient :

$$\frac{E_{\text{elec.}}}{E_{\tau}} > R^2 \approx 3,488 \quad (3.5)$$

Si on appelle $Q = \frac{s\tau}{L}$ le facteur de qualité du transistor, l'inégalité (3.5) peut s'écrire :

$$Q > R \approx 1,868$$

Lorsque cette inégalité n'est pas vérifiée, la solution stationnaire des équations de Dyakonov-Shur est stable pour tout courant de drain. Ainsi, en-dessous d'une certaine valeur de tension U_C , il ne peut pas y avoir d'émission d'onde électromagnétique par le transistor. Au-dessus de la tension critique U_C , il existe un intervalle de courant dans lequel la solution stationnaire est instable. En dehors de cet intervalle, la solution stationnaire reste néanmoins stable. On peut extraire la tension minimale critique d'instabilité de la relation (3.5) pour obtenir $U_C = \frac{R^2 E_{\tau}}{q}$. A cette tension, les ondes plasma dans le transistor sont amorties sauf pour une valeur spécifique du courant de drain où ces ondes sont conservées et oscillent à amplitude fixe en l'absence d'excitation extérieure. Grâce à la courbe caractéristique courant/tension (voir figure 1.9), on déduit la tension drain-source correspondant à ce courant :

$$U_{DS} \approx 0,2826 U_C$$

Pour un canal de 200 nm dans de l'arséniure de gallium (GaAs) à température ambiante (300K), on obtient U_C de l'ordre de 0,6 volt.

3.2.2 Résolution directe FDTD du gain

Par la méthode de Fourier, nous avons pu calculer les pulsations complexes et modes de densité $\tilde{n}$ et de vitesse $\tilde{v}$ perturbatifs. Nous en avons déduit un critère d'instabilité menant à la génération spontanée d'ondes plasmas au sein du canal et donc à terme à l'émission d'ondes électromagnétiques par un transistor. Cependant, cette méthode étant perturbative, elle n'est valide que tant que les ondes plasmas ont une amplitude relativement faible. Dans le régime stable, cette condition est toujours vérifiée puisque

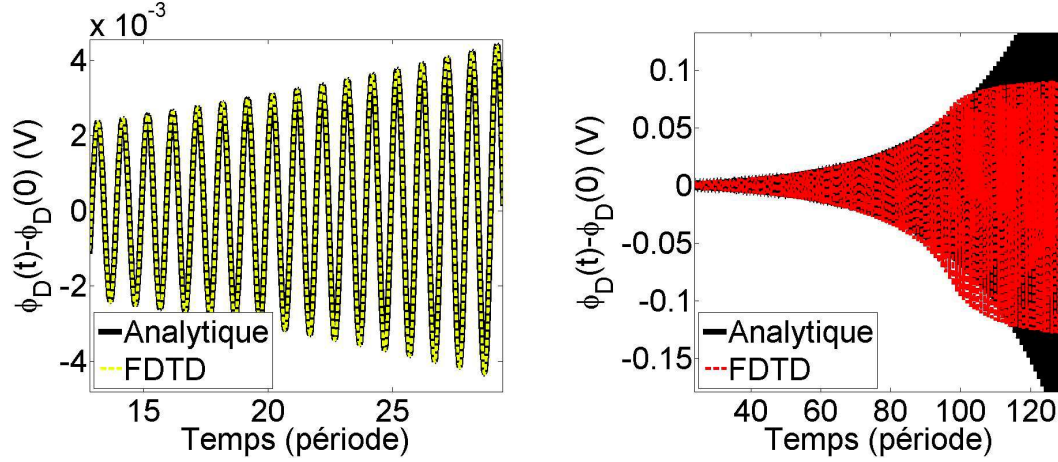


FIGURE 3.6 – Comparaison du potentiel de drain calculé par la méthode FDTD (en traits pointillés) et par la méthode analytique (en trait noir pleins). Lorsque l’amplitude d’oscillation du potentielle est faible, les deux méthodes sont en accord. Au bout d’un certain temps, l’amplitude d’oscillation a suffisamment crû pour que les effets non-linéaires entre en jeu. L’amplification de l’onde plasma sature et l’amplitude d’oscillation se stabilise.

les ondes plasmas sont amorties. Ainsi, si on n’excite que faiblement le canal du transistor, les ondes générées garderont une amplitude inférieure à l’excitation initiale et finissent par disparaître. Au contraire, dans le régime instable, la moindre fluctuation engendre l’excitation d’ondes plasmas dans le canal. Une fois générées, l’amplitude de ces ondes croît exponentiellement avec le temps. Au bout d’un certain temps, l’onde plasma atteint alors une amplitude considérable et la méthode perturbative n’est plus valable. Pour rendre compte de cette saturation et évaluer les champs ρ et v à temps long, nous utilisons une méthode d’intégration numérique FDTD décrite en section 5.2.1.

Sur la figure 3.6, on peut voir une comparaison entre la solution FDTD et la solution perturbative obtenue par transformation de Fourier dans le cas d’un régime instable. On peut constater que comme attendu, la solution perturbative ne coïncide avec la solution FDTD que pendant un certain temps. Lorsque l’amplitude de l’onde plasma dans le canal devient trop importante, la solution FDTD finit par saturer comme attendu. La cohérence entre les deux méthodes, lorsque l’approche perturbative est valide, justifie la validité de celles-ci.

3.3 Introduction des collisions électron-électron : canal visqueux

3.3.1 Origines microscopique de la viscosité et contribution mathématiques

Le modèle Dyakonov-Shur considère l'ensemble des électrons du canal d'un MOS-FET comme un gaz d'électrons libres en mouvement ballistique. Cela signifie que les électrons se déplacent individuellement dans le canal comme si ils étaient dans le vide. La valeur des champs de densité et de vitesse du mouvement collectif sont donnés par la moyenne locale du nombre et vitesse des électrons respectivement, dans un petit volume de longueur caractéristique faible devant la longueur du canal. Le modèle initial néglige toute forme d'interaction directe entre électrons. Notamment, les collisions électron-électron sont inexistantes. En effet, d'après le modèle Dyakonov-Shur, les électrons n'interagissent entre eux que par le potentiel électrostatique généré par la densité d'électron ρ . Cette densité étant une propriété macroscopique moyenne, elle ne contient aucune information quant à la position microscopique réelle des électrons. L'approche hydrodynamique efface donc les propriétés microscopiques du gaz d'électron en rendant la matière continue à toute échelle et donc à l'échelle des électrons. Ces derniers se déplacent donc individuellement dans un champ moyen insensible aux propriétés discrètes du gaz d'électron à petite échelle. Il ne peut donc se produire aucune collision entre deux électrons proches car la notion même d'électrons individuels disparaît. Dans la réalité, ces collisions existent et peuvent avoir un impact important sur le comportement global du fluide.

Les collisions électron-électron entraînent une corrélation du mouvement des électrons voisins. À l'échelle macroscopique cela se traduit par l'apparition de contraintes au sein du fluide. Ces contraintes sont interprétées comme une forme de friction entre deux particules fluides voisines, volumes de fluide infinitésimaux, du fait du gradient du champ de vitesse du fluide. Elles donnent alors lieu à une diffusion de la vitesse du fluide. Le mouvement d'une partie du fluide entraîne les parties voisines et ainsi de suite, traduisant sa viscosité. Cette friction explique la formation des écoulements de Couette entre deux plaques planes parallèles et idéalement infinies et de Poiseuille dans un tuyau dont les extrémités sont à différentes pressions (voir figure 3.7). En l'absence de viscosité, des écoulements discontinus (en bloc) sont en effet envisageables (voir figure 3.8).

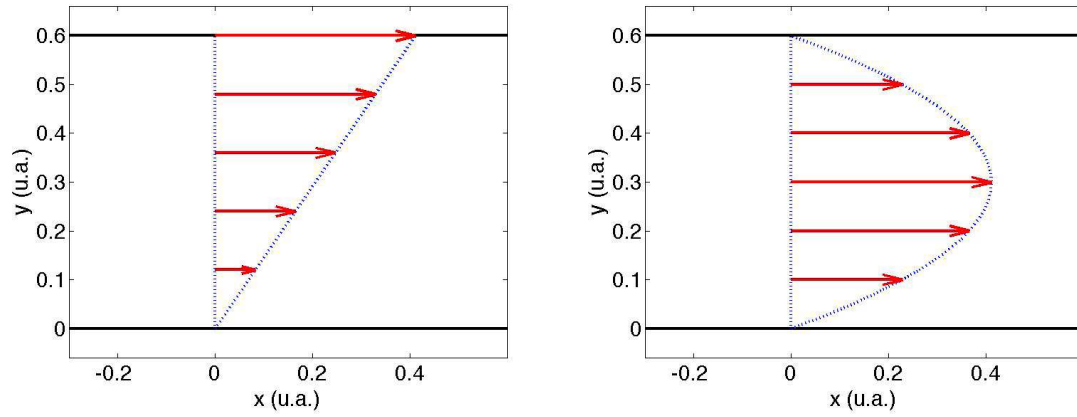


FIGURE 3.7 – Champ de vitesse des écoulements de Couette et Poiseuille respectivement à gauche et à droite. Le champ de vitesse varie continument du fait de la diffusion de l'impulsion des électrons dans le fluide. Pour l'écoulement de Couette, le profil du champ de vitesse est linéaire et est obtenu du fait des vitesses différentes des plaques englobant le fluide. Pour l'écoulement de Poiseuille, le profil est parabolique, les bords sont fixes et le fluide est mis en mouvement par la différence de pression d'un bout à l'autre du "tube".

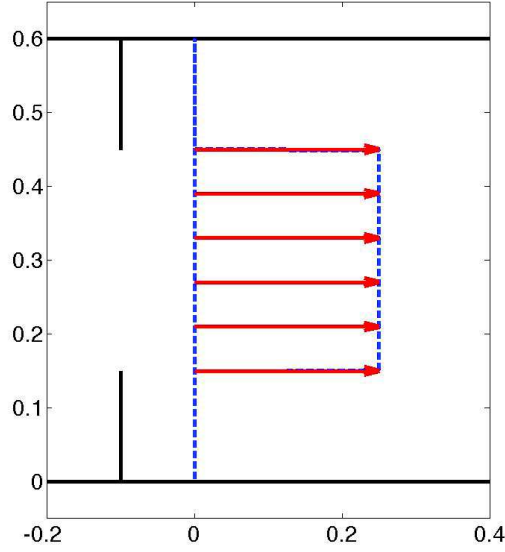


FIGURE 3.8 – Écoulement non visqueux à travers un trou. Le fluide devant le trou est mis en mouvement à la vitesse v_x . Du fait de l'absence de viscosité, l'écoulement dans l'axe du trou n'entraîne pas le fluide en dehors de cet axe. Il se forme une interface effective entre le fluide en mouvement et le fluide stationnaire.

La prise en compte des effets des collisions électron-électron est donnée par l'ajout de la divergence du tenseur de contrainte déviateur dans les équations de Navier-Stokes. Ce tenseur rend compte de la relation entre le taux de déformation du fluide (gradient de vitesse) et les contraintes générées dans le fluide. La relation exacte dépend du fluide et de ses lois de constitutions. Pour de faibles déformations (faible gradient du champ de vitesse dans une certaine mesure), on utilise généralement une relation de constitution linéaire (3.6) où le tenseur des contraintes $\vec{\sigma}$ est une combinaison linéaire du gradient du champ de vitesse (tenseur de rang 2 symétrique) et de sa divergence (scalaire).

$$\vec{\sigma} = \lambda \left(\vec{\nabla} \cdot \vec{v} \right) \vec{1} + \mu \left(\vec{\nabla} \otimes \vec{v} + \vec{v} \otimes \vec{\nabla} \right) \quad (3.6)$$

Les deux constantes de proportionnalité sont données par deux viscosités : la viscosité dynamique μ et la viscosité de volume λ . Dans le cadre d'un modèle de contrainte visqueuse linéaire et à coefficient constant, on obtient alors la force visqueuse dans le fluide suivante :

$$\vec{\nabla} \cdot \vec{\sigma} = \mu \Delta \vec{v} + (\mu + \lambda) \vec{\nabla} \left(\vec{\nabla} \cdot \vec{v} \right) \quad (3.7)$$

En 1D, les deux contributions de l'équation (3.7) sont équivalentes et un seul coefficient de viscosité effectif est nécessaire. On obtient alors le nouveau système Dyakonov-Shur

visqueux :

$$\begin{cases} \frac{\partial}{\partial t}\rho + \frac{\partial}{\partial x}(\rho v) = 0 \\ \frac{\partial}{\partial t}v + v\frac{\partial}{\partial x}v = -\frac{\partial}{\partial x}\rho + \mu'\frac{\partial^2}{\partial x^2}v - \frac{v}{\tau} \end{cases} \quad (3.8)$$

La force de viscosité fait donc intervenir une dérivée spatiale d'ordre 2 du champ de vitesse. La résolution unique du système (3.8) n'est alors possible que si on ajoute une condition au bord supplémentaire aux conditions (1.5). Un choix généralement considéré est l'imposition d'une dérivée nulle du champ de vitesse à la source [13]. Dans la partie à venir, nous allons étudier la stabilité de la solution stationnaire en présence de viscosité.

3.3.2 Application de la méthode perturbative : contribution stationnaire

Nous allons d'abord nous intéresser à la stabilité d'un courant stationnaire dans le canal. En premier lieu, nous cherchons à extraire une solution stationnaire du système. Une première remarque à faire sur le nouveau système d'équation (3.8) est que de manière générale, le coefficient de viscosité n'est pas constant pour un fluide donné mais peut éventuellement dépendre de sa densité, sa température ou autre paramètre interne du fluide. La spécification de cette dépendance de la viscosité sur certains paramètres du fluide est généralement donnée par l'équation d'état du fluide. Une conséquence d'une viscosité variable est que la force de viscosité s'écrit en réalité :

$$\vec{\nabla} \cdot \vec{\sigma} = \frac{\partial}{\partial x} \left(\mu' \frac{\partial}{\partial x} v \right) \vec{e}_x$$

On considèrera néanmoins que l'effet de la viscosité est décrit par le système (3.8). Cela revient donc à négliger la contribution $\frac{\partial \mu'}{\partial x} \frac{\partial v}{\partial x}$. Nous nous intéresserons à deux modèles de viscosité pour leur simplicité :

- une viscosité μ constante
- une viscosité μ inversement proportionnelle à la densité.

Contrairement au cas non-visqueux, on ne connaît pas de solution stationnaire explicite du système (3.8) en considérant simultanément la contribution de la viscosité du canal et de sa résistivité. En revanche, lorsque l'on suppose une résistivité nulle, il est facile de montrer qu'avec les conditions aux bords spécifiées, la solution stationnaire est la solution constante. Pour pouvoir effectuer l'approche perturbative, on intègre la solution stationnaire numériquement dans le cas général. La simplicité du second modèle de viscosité vient alors du fait que dans ce cas, on peut néanmoins intégrer le système (3.8) une première fois explicitement afin de n'avoir qu'une équation différentielle du premier ordre à intégrer numériquement (voir plus loin (3.10)).

En intégrant le système d'équation (3.8) selon x en prenant $\frac{\partial}{\partial t} = 0$ pour une faible viscosité, on obtient la figure 3.9. Le courant de drain est choisi supérieur au courant

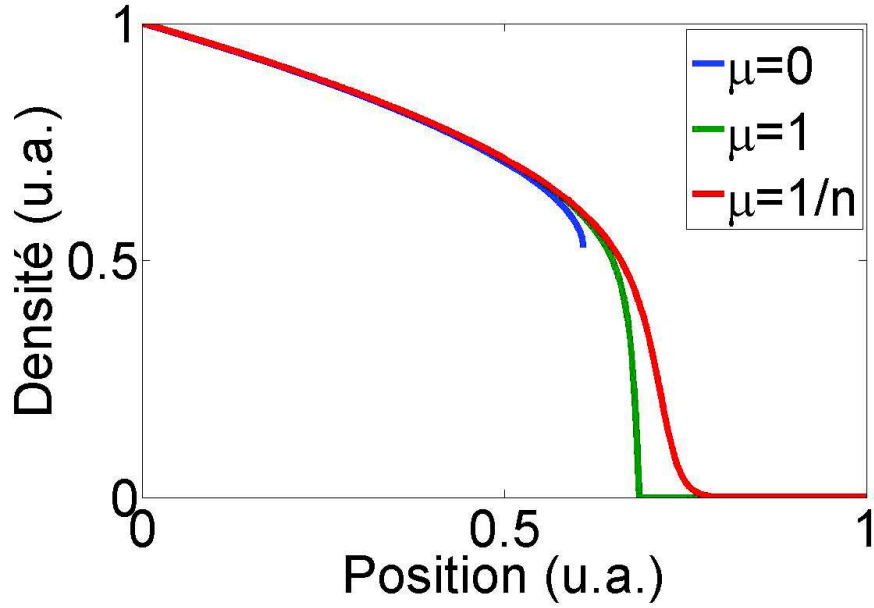


FIGURE 3.9 – Solution stationnaire du système (3.8) pour une faible viscosité. En bleu, on a la solution pour une viscosité nulle, en vert constante et en rouge inversement proportionnelle à la densité.

critique de sorte à ce que le canal soit pincé (en bleu) sans viscosité. Lorsque l'on introduit une viscosité non-nulle dans le canal, on constate que la singularité au point de pincement disparaît. En effet, les solutions stationnaires visqueuses semblent définies pour toute longueur de canal. Dans le cas du modèle de viscosité en $\frac{1}{\rho}$, on montrera plus loin qu'une solution stationnaire bien définie existe effectivement pour toute longueur de canal. Au contraire, bien qu'une solution mathématique existe pour toute longueur de canal dans le cas du modèle de viscosité constante, on montrera que celle-ci cesse de représenter une grandeur physique réelle.

Avant le pincement de la solution sans viscosité, les trois solutions de la figure 3.9 semblent très similaires. Cela est en accord avec la faible valeur de viscosité choisie. En effet, dans la limite d'une viscosité nulle, on s'attend à retrouver la solution initiale sans viscosité. Dans ce cas, son effet sera notamment important uniquement autour des points de forte courbure de la solution stationnaire, compensant le faible coefficient de viscosité, comme par exemple lorsqu'on se rapproche du point de pincement ou au contact de source (voir plus loin).

Analyse de la solution $\mu = \frac{\mu_0}{\rho}$

Dans le cas où la viscosité est inversement proportionnelle à la densité, il est facile de montrer qu'une solution stationnaire existe pour tout courant de drain et toute longueur de canal. En effet, dans ce cas, le système stationnaire s'écrit :

$$\begin{cases} \frac{d}{dx}(\rho v) = 0 \\ \frac{d}{dx}\rho + v\frac{d}{dx}v + \frac{v}{\tau} - \frac{\mu_0}{\rho}\frac{d^2}{dx^2}v = 0 \end{cases} \quad (3.9)$$

La première équation du système (3.9) s'intègre trivialement et on obtient $\rho v = j_L$ constant dans le canal, égal à la valeur imposée en L . En injectant dans la seconde équation, on obtient alors :

$$\frac{d}{dx} \left(-\mu_0 \frac{d}{dx}v + \frac{j_L x}{\tau} + j_L v + \frac{1}{2}\rho^2 \right) = 0$$

d'où on déduit :

$$-\mu_0 \frac{d}{dx}v + \frac{j_L x}{\tau} + j_L v + \frac{j_L^2}{2v^2} = \frac{j_L^2}{\rho_0} + \frac{1}{2}\rho_0^2$$

où ρ_0 est la valeur de la densité imposée en 0. En multipliant par ρ , on peut réécrire cette équation en fonction de ρ et on aboutit à :

$$\mu_0 j_L \frac{d}{dx}\rho + \frac{1}{2}\rho^4 + \left(\frac{j_L x}{\tau} - \frac{j_L^2}{\rho_0} - \frac{1}{2}\rho_0^2 \right) \rho^2 + j_L^2 \rho = 0 \quad (3.10)$$

Lorsque la viscosité μ_0 est nulle, on reconnaît le polynôme donnant la solution stationnaire non-visqueuse ρ_0 . On en déduit que ρ_0 est une borne inférieure de la solution visqueuse ρ_μ . En effet, si $\rho_\mu = \rho_0$ en un point x du canal, la dérivée $\frac{d}{dx}\rho_\mu$ est nulle en ce point. La dérivée de ρ_0 étant négative, la solution avec viscosité ρ_μ ne peut pas traverser la solution sans viscosité ρ_0 depuis une valeur supérieure à celle-ci. Une fois le point de pincement atteint, la borne inférieure de ρ_μ est alors donnée par l'axe des abscisses par le même raisonnement. Par ailleurs, l'étude du signe du polynôme (1.9) garanti que la dérivée de ρ_μ est négative pour toute valeur de ρ_μ supérieure à sa borne inférieure. On en déduit que $\rho_\mu(x)$ est une fonction strictement décroissante pour tout x dans le canal. Les valeurs admissibles de ρ_μ sont donc bornées à la fois par au-dessus et par en-dessous. L'équation (3.10) étant non-singulière pour toute valeur de ρ finie, une solution existe pour toute longueur de canal.

On peut également montrer que le comportement asymptotique de la solution ρ_μ au-delà du point de pincement est donné par :

$$\rho_\mu(x) \propto \exp\left(-\frac{j_L x}{\mu_0}\right) \quad (3.11)$$

En effet, la fonction ρ_μ étant strictement décroissante et bornée par 0, celle-ci est tend vers une limite l comprise entre ρ_μ et 0. La dérivée de ρ_μ tend alors vers 0 lorsque ρ_μ tend vers l . Si $l > 0$ on aboutit à une contradiction car on a

$$0 = \left. \frac{d}{dx} \rho_\mu \right|_{\rho_\mu \rightarrow l} < 0 \quad \text{par l'équation 3.10}$$

On a donc $l = 0$. Lorsque ρ_μ tend vers 0, on peut alors négliger les puissances de ρ supérieures à la plus basse puissance de ρ dans l'équation (3.10). Celle-ci se réduit alors à :

$$\mu_0 j_L \frac{d}{dx} \rho + j_L^2 \rho = 0$$

dont la solution est donnée par l'équation (3.11).

Cas de la viscosité constante

Dans le cas du modèle de viscosité constante, le système stationnaire s'écrit :

$$\begin{cases} \frac{d}{dx} (\rho v) = 0 \\ \frac{d}{dx} \rho + v \frac{d}{dx} v + \frac{v}{\tau} - \mu_0 \frac{d^2}{dx^2} v = 0 \end{cases} \quad (3.12)$$

Dans ce cas, on s'attend à ce que la vitesse des électrons croisse plus vite que dans le cas précédent. En effet, lorsque cette vitesse augmente, la densité d'électron diminue. Dans le modèle précédent, cela entraîne un accroissement de la contribution visqueuse. Une plus faible courbure du champ de vitesse est alors nécessaire afin de vérifier l'équation (3.9) que l'équation (3.12). Ainsi, on s'attend à ce que la densité d'électron dans le canal décroisse plus rapidement dans ce cas que dans le précédent, ce qui est bien ce qui est observé figure 3.9. Cette fois-ci, il n'existe alors pas de solution physique au système (3.12). En effet, on peut montrer que lorsque la densité d'électron tend vers 0, celle-ci coupe l'axe des abscisses de façon linéaire. Pour cela, on réécrit la deuxième équation du système (3.12) en fonction du champ de densité ρ sachant le produit $\rho v = j_L$ constant. On obtient alors :

$$\mu j_L \frac{d^2}{dx^2} \rho + \left(\rho^2 - \frac{j_L^2}{\rho} - \frac{2\mu j_L}{\rho} \frac{d}{dx} \rho \right) \frac{d}{dx} \rho + \frac{j_L}{\tau} \rho = 0 \quad (3.13)$$

Lorsque ρ tend vers 0, la contribution dominante dans la parenthèse de l'équation (3.13) est clairement le terme

$$\frac{j_L^2}{\rho} + \frac{2\mu j_L}{\rho} \frac{d}{dx} \rho \quad (3.14)$$

Supposons que ce terme soit la contribution dominante de toute l'équation. On obtient alors la solution approximative suivante :

$$\rho \approx -\frac{j_L}{2\mu} (x - x_0) \quad (3.15)$$

où x_0 est le point où la densité s'annule à cet ordre. Dans ce cas, ρ est de l'ordre de $(x - x_0)$ linéaire autour de x_0 , sa dérivée $\frac{d}{dx}\rho$ de l'ordre d'une constante et sa dérivée seconde identiquement nulle. On obtient bien en injectant la solution (3.15) dans (3.13) que le terme (3.14) est bien le terme dominant de l'équation lorsque ρ tend vers 0. Cette solution est donc cohérente avec les hypothèses de départ.

La solution bien que définie mathématiquement n'est malheureusement plus physique car la densité d'électron est une grandeur strictement positive. Le problème de la modélisation vient alors du fait, comme on l'a remarqué plus tôt, que la relation (1.3) n'est plus valable lorsque la densité d'électron dans le canal tend vers 0. Il n'est donc pas surprenant que le modèle (3.12) ne soit plus physique lorsque ρ tend vers 0. Dans ce cas, on peut considérer que le canal se pince lorsque la densité ρ s'annule.

Apparition d'une couche limite à la source

Au contact de source, on impose que la dérivée de la vitesse soit nulle pour une viscosité non-nulle. Si la solution est stationnaire, la conservation du courant impose que cette contrainte soit équivalente à une dérivée de densité nulle à la source également. Or, dans la limite d'une viscosité nulle, la solution stationnaire est donnée par le système (1.9). Ainsi, dans cette limite, les dérivées des densité et vitesse ne sont pas nulles à la source, sauf pour un courant de drain nul. Lorsque la viscosité tend vers zéro, il se forme en réalité une couche limite de plus en plus fine à la source. Dans la limite μ tend vers 0, la viscosité n'a d'effet que dans cette couche limite. Dans cette couche limite, la dérivée de la vitesse passe d'une valeur nulle au contact de source, à la valeur limite en l'absence de viscosité à la fin de la couche limite. L'épaisseur de cette couche limite tendant vers 0 avec la viscosité, on retrouve bien dans la limite la solution attendue. Sur la figure 3.9, l'épaisseur de la couche limite est déjà trop faible pour être observée. Cependant, la couche limite est bien présente comme le montre la figure 3.10. Sur cette dernière figure, on peut voir la diminution de l'épaisseur de la couche limite avec la diminution de la viscosité. La dérivée de la densité reste cependant bien nulle en $x = 0$ pour toute valeur de viscosité.

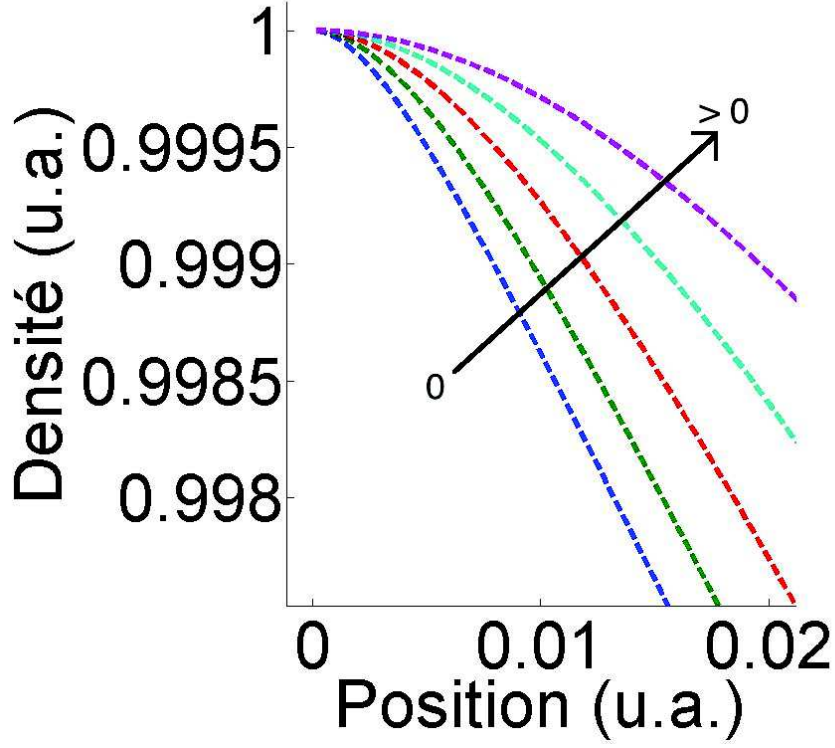


FIGURE 3.10 – Densité électronique stationnaire près de la source en fonction de la position (abscisse) et de différentes viscosité (couleurs). La flèche indique le sens des μ croissants. Pour toute valeur de μ , la condition de dérivée nulle est respectée à la source. Lorsque μ tend vers 0, la solution tend vers la solution non-visqueuse et la couche limite passant d’une dérivée nulle en 0 à non-nulle devient de plus en plus fine et son épaisseur tend également vers 0.

3.3.3 Application de la méthode perturbative : disparition de la discontinuité

Par l’analyse de Fourier, on sait que la dissipation visqueuse est d’autant plus importante que l’échelle de variation spatiale du champ est faible. Concrètement, les variations les plus localisées spatialement, donc associées à de faibles longueurs d’onde, sont plus rapidement atténuées que les variations spatialement lentes, ou de grandes longueurs d’onde. Lorsque le coefficient de viscosité est faible, son effet sera alors principalement de lisser les champs de densité et de vitesse. En conséquence, dès que l’on considère une viscosité non-nulle dans le canal, les ressauts hydrauliques disparaissent et sont remplacés par une variation éventuellement très rapide spatialement mais continue des champs [13]. La figure 3.11 présente le résultat d’une simulation numérique prenant

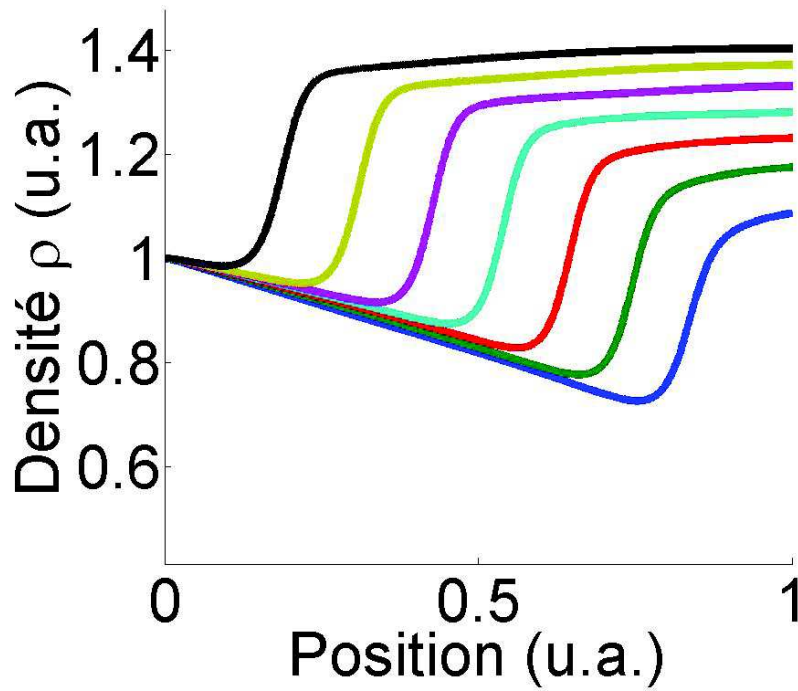


FIGURE 3.11 – Propagation du ressaut dans le canal en présence de viscosité. La viscosité prévient la formation de la discontinuité. Le champ de densité reste continu.

une viscosité non-nulle du canal en compte. On peut voir que le ressaut hydraulique est adouci, conformément au résultat attendu. En effet, celui-ci est stabilisé avant de former une discontinuité. Le champ de densité reste alors continu, même lorsque l'amplitude d'oscillation de la densité devient de l'ordre de l'amplitude de la solution stationnaire.

Chapitre 4

Détection térahertz non perturbative

Dans le chapitre 1.5, nous avons montré comment, à l'aide de la méthode des perturbations, un transistor à effet de champ peut convertir une oscillation térahertz en un signal direct. Cette démonstration repose sur l'hypothèse que l'amplitude d'oscillation du potentiel de l'antenne U_A est faible devant une tension de référence U_0 . Lorsque ce n'est pas le cas, on peut s'attendre à ce que la formule de rectification (1.13) obtenue ne soit plus vérifiée. Notamment, nous allons montrer que le facteur de résonance f change de forme à forte puissance, montrant l'apparition de nouveaux phénomènes de rectification.

4.1 Résonances sub-harmoniques

4.1.1 Résolution FDTD du système Dyakonov-Shur à forte puissance

Le système Dyakonov-Shur est un système d'équation différentiel non-linéaire. En conséquence, sa représentation dans l'espace de Fourier réciproque montre un couplage explicite des différentes fréquences de vibrations du système, contrairement aux systèmes linéaires pour lesquels ces fréquences sont découplées. C'est ce couplage qui est à l'origine de la rectification DC du potentiel drain-source et constitue le point fort de la détection térahertz par un transistor à effet de champ. Cependant, c'est ce même couplage qui implique que la fréquence fondamentale calculée n'est en réalité pas la fréquence la plus basse pouvant résonner dans le canal d'un MOSFET. En effet, nous allons montrer qu'à plus forte intensité d'excitation, des pics de résonances apparaissent dans le facteur de résonance f de la rectification à des pulsations plus basses que ω_0 . Ces pics sont appelés résonances sous-harmoniques.

La figure 4.1 montre la rectification du potentiel drain-source ΔU_{DS} en fonction de la fréquence d'excitation du canal en abscisse et de l'amplitude d'excitation en différentes couleurs. L'unité de fréquence choisie est telle que la fréquence fondamentale de résonance Dyakonov-Shur vaille 1. A faible puissance d'excitation (courbe bleu foncé

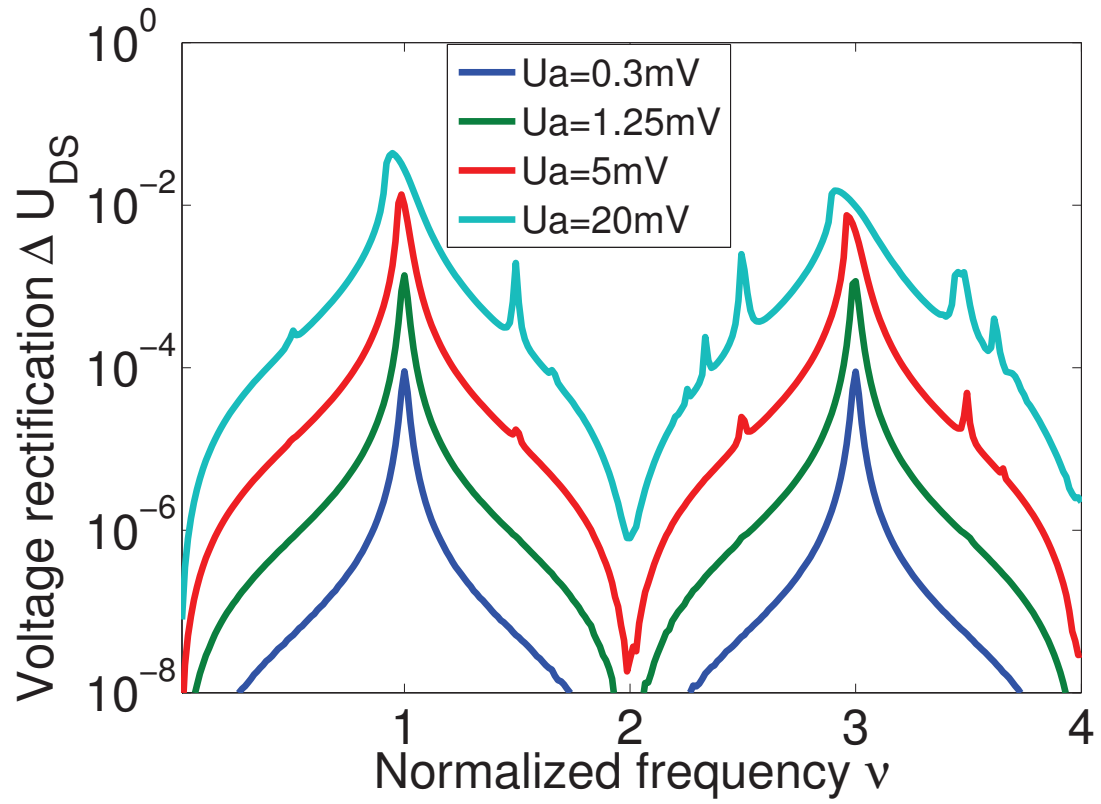


FIGURE 4.1 – Simulation numérique FDTD de la rectification du potentiel drain-source en fonction de la fréquence de l’onde incidente (en abscisse) et de la puissance incidente (différentes courbes). La fréquence de résonance fondamentale Dyakonov-Shur est ici normalisée à 1.

en bas) on n'observe bien que les fréquences de résonances attendues 1 et 3 (puis 5, 7, ... à plus haute fréquence). Lorsque la puissance d'excitation augmente, on voit alors le spectre de résonance s'enrichir. En effet, de plus en plus de pics de résonances semblent apparaître avec l'augmentation de la puissance. On voit en effet apparaître de nouvelles résonances vers les fréquences $\frac{3}{2}$, $\frac{5}{2}$ et $\frac{7}{2}$ sur la courbe rouge à puissance relativement élevée puis encore d'autres résonance à plus forte puissance encore.

On montrera que ces nouvelles fréquences de résonances apparaissent en réalité à toutes les pulsations telles que $\omega = \frac{2p+1}{q}\omega_0$ où ω_0 est la pulsation de résonance fondamentale. L'entier p décrit les fréquences de résonances initiales lorsque $q = 1$. L'entier q décrit l'ordre sous-harmonique qui nécessite une puissance de plus en plus grande afin d'être observé lorsque cette ordre augmente. En effet, cet ordre sous-harmonique provient des effets non-linéaires du système qui induisent une génération d'harmonique supérieures dans le canal à l'image du doublage de fréquence dans les cristaux non-linéaire qui permettent notamment de créer un faisceau laser bleu à partir d'un faisceau initial dans l'infra-rouge proche. Le mélange d'ordre q permet alors de convertir la pulsation initiale ω en une nouvelle pulsation d'oscillation $q\omega$ dans le canal. C'est lorsque cette nouvelle pulsation entre en résonance dans le canal qu'un pic apparaît.

4.1.2 Développement perturbatif général

Lorsque l'amplitude d'oscillation dans le canal devient importante, l'approche linéaire ou quasi-linéaire ne suffit plus afin de correctement modéliser le comportement du transistor. Si on veut continuer à utiliser une approche perturbative, il va falloir aller à un ordre d'approximation plus grand. La méthode FDTD permet d'observer à forte puissance l'apparition de nouveaux pics de résonances 4.1 et on va montrer que la méthode perturbative permet d'expliquer l'origine de ces pics, leur position et leur ordre d'apparition. Pour cela, on va supposer que les champs de densité et de vitesse peuvent être développés en une série de la forme suivante :

$$\left\{ \begin{array}{l} \rho = \sum_{p=0}^{+\infty} \epsilon^p \rho_p \\ v = \sum_{p=0}^{+\infty} \epsilon^p v_p \end{array} \right. \quad (4.1)$$

où ϵ est le paramètre perturbatif, que l'on identifiera plus tard, et les suites (ρ_p, v_p) sont les champs de densité et de vitesse à l'ordre p . En injectant un tel développement dans

le système (1.8), on obtient le système suivant :

$$\begin{cases} \sum_{p=0}^{+\infty} \epsilon^p \left(\frac{\partial}{\partial t} \rho_p + \frac{\partial}{\partial x} \left(\sum_{q=0}^p \rho_{p-q} v_q \right) \right) = 0 \\ \sum_{p=0}^{+\infty} \epsilon^p \left(\frac{\partial}{\partial t} v_p + \frac{\partial}{\partial x} \left(\rho_p + \frac{1}{2} \sum_{q=0}^p v_{p-q} v_q \right) + \frac{v_p}{\tau} \right) = 0 \end{cases} \quad (4.2)$$

Ce nouveau système d'équation doit être vérifié pour toute valeur de ϵ ce qui n'est possible que si chaque terme d'ordre p s'annule indépendamment. On a alors :

$$\begin{cases} \frac{\partial}{\partial t} \rho_p + \frac{\partial}{\partial x} \left(\sum_{q=0}^p \rho_{p-q} v_q \right) = 0 \\ \frac{\partial}{\partial t} v_p + \frac{\partial}{\partial x} \left(\rho_p + \frac{1}{2} \sum_{q=0}^p v_{p-q} v_q \right) + \frac{v_p}{\tau} = 0 \end{cases}$$

pour tout $p \in \mathbb{N}$. On voit alors une hiérarchie d'équations s'établir. En effet, à l'ordre 0, on reconnaît le système Dyakonov-Shur dont les inconnues sont ρ_0 et v_0 . On peut alors résoudre ce système en premier indépendamment en choisissant par exemple la solution stationnaire (1.9). Cette solution est celle attendue dans le cas d'un transistor stable et qui n'est soumis à aucune onde incidente. Aux ordres p strictement positifs, le système s'écrit :

$$\begin{pmatrix} \frac{\partial}{\partial t} & 0 \\ 0 & \frac{\partial}{\partial t} + \frac{1}{\tau} \end{pmatrix} \begin{pmatrix} \rho_p \\ v_p \end{pmatrix} + \frac{\partial}{\partial x} \left(\begin{pmatrix} v_0 & \rho_0 \\ 1 & v_0 \end{pmatrix} \begin{pmatrix} \rho_p \\ v_p \end{pmatrix} \right) = -\frac{\partial}{\partial x} \sum_{q=1}^{p-1} \begin{pmatrix} \rho_{p-q} v_q \\ \frac{1}{2} v_{p-q} v_q \end{pmatrix} \quad (4.3)$$

où l'on peut identifier à gauche de l'égalité un opérateur linéaire L eq.(4.4) agissant sur les termes d'ordre p donnant lieu à une équation d'onde et à droite un opérateur non-linéaire agissant sur les termes d'ordre q strictement inférieurs à p . On peut donc résoudre le système complet ordre par ordre de manière unique en considérant à l'ordre p les termes d'ordres inférieurs comme des termes sources et en imposant les bonnes conditions aux bords. Ce sont ces conditions aux bords qui permettent d'identifier le paramètre de perturbation ϵ en fonction des paramètres physiques du système.

En mode détection, l'absorption d'une onde incidente par une antenne branchée en série avec le générateur grille-source fournit un potentiel oscillant à la fréquence de l'onde absorbée. On a alors :

$$U_{GS} = U_0 + \Re(U_A \exp(-i\omega t))$$

qui du fait de l'approximation du canal graduel se traduit par :

$$\rho(t, 0) = \rho_0 + \Re(\rho_A \exp(-i\omega t))$$

où l'on peut identifier $\epsilon = \left| \frac{\rho_A}{\rho_0} \right|$. Ainsi, à puissance nulle, on a bien $\epsilon = 0$ et la solution est la solution stationnaire.

4.1.3 Mélange de fréquence comme source de signal sous-harmonique

On a établi que le développement perturbatif en puissance de ϵ conduit à une hiérarchie d'équations d'onde dont les termes sources à l'ordre p sont donnés par des produit de termes d'ordre inférieur. On sait par l'analyse spectrale que le produit de deux fonctions dans l'espace direct se traduit par un produit de convolution ou mélange de fréquence dans l'espace réciproque. Ce mélange de fréquence est à l'origine du signal de rectification aux fréquences Dyakonov-Shur à faible amplitude, mais également aux fréquences sub-harmoniques à forte amplitude. En effet, supposons qu'une onde incidente monochromatique de pulsation ω soit absorbée dans le canal du transistor. Au premier ordre, l'équation (4.3) ne possède pas de terme non-linéaire source. Cependant, une excitation est tout de même présente à la pulsation ω du fait des conditions aux bords. On a en effet :

$$\rho_1(t, 0) = \Re(\exp(-i(\omega t - \varphi)))$$

au premier ordre. Le spectre de Fourier des champs perturbatifs au premier ordre comprend alors deux raies (théoriquement infiniment fine tendant vers deux distributions de Dirac) aux pulsations ω et $-\omega$. L'opérateur

$$L[\bullet] = \begin{pmatrix} \frac{\partial}{\partial t} & 0 \\ 0 & \frac{\partial}{\partial t} + \frac{1}{\tau} \end{pmatrix} \bullet + \frac{\partial}{\partial x} \left(\begin{pmatrix} v_0 & \rho_0 \\ 1 & v_0 \end{pmatrix} \bullet \right) \quad (4.4)$$

possédant les résonances aux fréquences Dyakonov-Shur, l'amplitude de ces champs à l'ordre 1 sera exaltée si ω correspond à une de ces fréquences de résonance. Du fait du mixage de fréquence, on aura au second ordre trois raies dans le spectre de Fourier aux pulsations -2ω , 0ω et 2ω . Les termes sources au second ordre étant proportionnels au carré des amplitudes du premier ordre, ce second ordre sera fortement excité même si l'opérateur L n'est pas résonant à ces pulsations à condition que ω soit une pulsation de résonance. Le terme à 0ω correspondant à une fréquence nulle, c'est celui-ci qui définit le signal de rectification et celui-ci présente alors comme attendu une résonance aux fréquences Dyakonov-Shur.

Supposons maintenant que ω soit exactement à la moitié d'une fréquence de résonance. Au premier ordre, l'opérateur L n'est pas résonant et l'amplitude d'oscillation sera faible. Ainsi, au second ordre, l'excitation à 0ω sera faible et le signal ne semble pas faire apparaître de pic de résonance. Cependant, les excitations à -2ω et 2ω bien que très faibles seront exaltées par les résonances de l'opérateur L et donneront lieu à d'importantes oscillations. Au troisième ordre, la forme du couplage non-linéaire dans l'équation (4.3) induit un mélange de fréquence entre les ordres 1 et 2. Ce mélange conduit à des termes sources aux fréquences -3ω , $-\omega$, ω et 3ω . Du fait des fortes oscillations au second ordre, ces termes sources seront importants et la réponse au troisième ordre le sera également et ainsi de suite aux ordres suivants. On voit alors une structure pyramidale tableau 4.1 apparaître dans le système. Le pic de résonance dans le signal

ordre p	composantes du spectre de Fourier	ordres sources
0	0ω	aucun
1	$-1\omega \quad 1\omega$	condition aux bords
2	$-2\omega \quad 0\omega \quad 2\omega$	ordre : $1 \oplus 1$
3	$-3\omega \quad -1\omega \quad 1\omega \quad 3\omega$	ordre : $1 \oplus 2$
4	$-4\omega \quad -2\omega \quad 0\omega \quad 2\omega \quad 4\omega$	ordres : $1 \oplus 3$ et $2 \oplus 2$
5	$-5\omega \quad -3\omega \quad -1\omega \quad 1\omega \quad 3\omega \quad 5\omega$	ordres : $1 \oplus 4$ et $2 \oplus 3$
$\vdots$	$\vdots$	$\vdots$

TABLE 4.1 – Table de décomposition du spectre de Fourier ordre par ordre. Une résonance à $n\omega$ fournira un pic de résonance dans la rectification à l'ordre $2n$.

sera alors associé à l'ordre 4 et donc proportionnel non pas au carré de l'amplitude d'oscillation du potentiel généré par l'antenne mais à sa puissance quatrième. Il est alors nécessaire d'avoir une puissance d'excitation suffisante afin de voir apparaître ce pic de résonance.

Le raisonnement se généralise pour toute pulsation ω telle que $q\omega$ soit une fréquence de Dyakonov-Shur $\omega_{DS} = (2p + 1)\omega_0$. En effet, une résonance apparaîtra à l'ordre q et se répercutera sur la rectification à partir de l'ordre $2q$. Un pic de résonance apparaît alors pour toute pulsation

$$\omega = \frac{2p+1}{q}\omega_0 \quad \forall (p, q) \in \mathbb{N} \otimes \mathbb{N}^*$$

avec une amplitude proportionnelle à U_A^{2q} ou de manière équivalente $\mathcal{P}^q$. Ainsi, plus l'ordre sous-harmonique q est grand, plus une puissance $\mathcal{P}$ élevée est nécessaire afin d'observer la résonance.

4.1.4 Comparaison directe des deux approches

La méthode perturbative attribue l'origine des résonances sub-harmoniques à la génération de seconde harmonique et harmoniques supérieures dans le canal. L'efficacité de cette génération d'harmonique supérieure dépend de l'excitation générée par l'antenne et de l'ordre de l'harmonique. Plus cet ordre est élevé, plus la puissance incidente doit l'être. L'approche perturbative ne possédant pas de paramètre libre ajustable, ses prédictions sont uniques et doivent être comparées aux résultats obtenus indépendamment par FDTD. Sur la figure 4.2, on peut voir les contributions à la rectification du potentiel à l'ordre 2 et 4 de la méthode perturbative. Comme attendu, la première contribution d'ordre 2 ne comprend qu'une résonance à la fréquence fondamentale Dyakonov-Shur (et ses multiples impairs) mais pas à ses sous-multiples. En revanche, la contribution d'ordre 4 comprend bien à la fois les résonances précédentes ainsi que de nouvelles résonances exclusivement à la moitié des fréquences Dyakonov-Shur. A partir de l'ordre

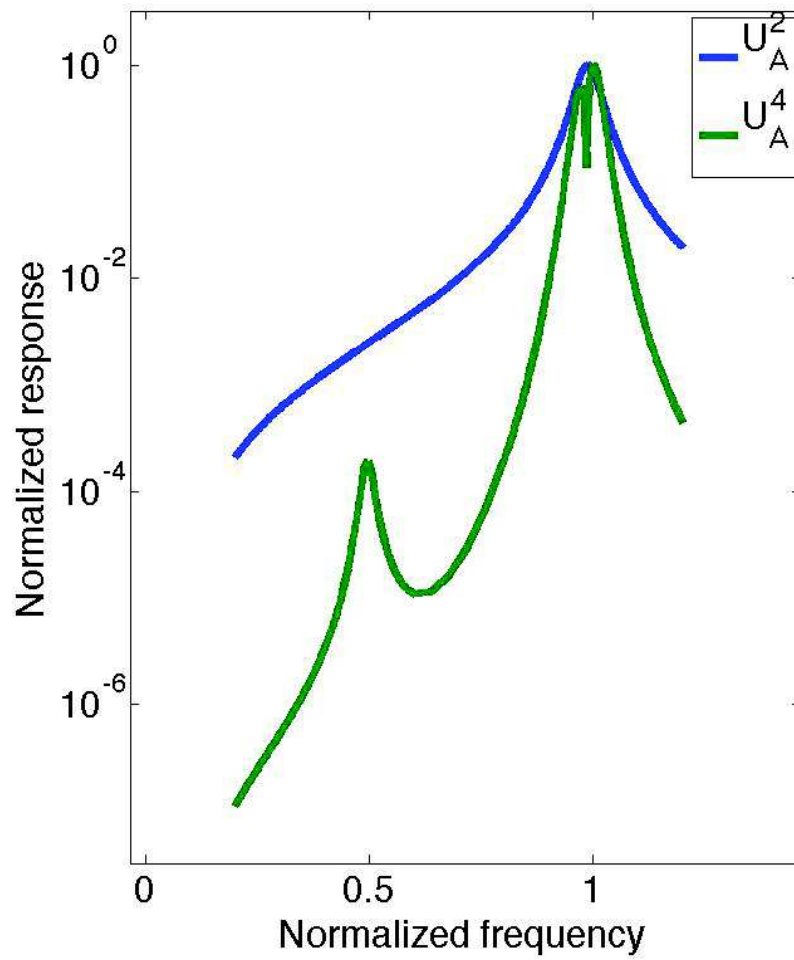


FIGURE 4.2 – Contribution perturbative des ordres 2 et 4 à la rectification du potentiel drain-source. A l'ordre 2, seule la résonance fondamentale (fréquence 1) apparaît ici. A l'ordre 4, un nouveau pic à demi-résonance apparaît.

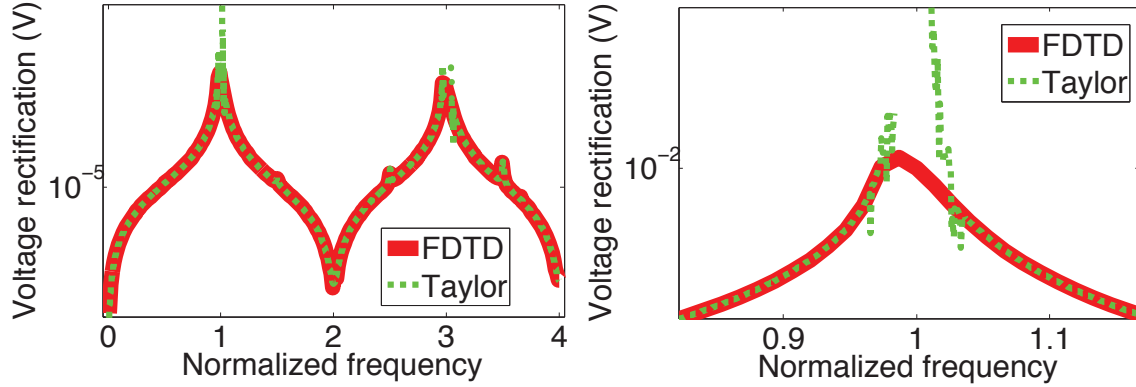


FIGURE 4.3 – Comparaison de la simulation FDTD avec l’approche perturbative (développement de Taylor en fonction de la puissance d’excitation). A droite, on voit la divergence de la méthode perturbative autour de la résonance fondamentale pour cette forte amplitude d’excitation. La méthode perturbative a été ici calculée jusqu’à l’ordre 32.

6, on voit apparaître les résonances au tiers des fréquences Dyakonov-Shur et ainsi de suite.

Si on somme toutes les contributions aux différents ordres, on obtient le spectre de réponse total figure 4.3 où il est comparé au spectre obtenu en FDTD. On remarque alors que la méthode perturbative semble posséder un rayon de convergence fini et dépendant de la fréquence d’excitation. En effet, la somme semble diverger sur certains intervalles de fréquences pour l’amplitude d’excitation choisie alors qu’elle converge pour des amplitudes plus faibles. Lorsque l’on reste dans le rayon de convergence de la méthode perturbative, on obtient néanmoins un accord qui semble parfait entre approche perturbative et simulation FDTD. En effet, l’amplitude du pic de résonance à demi-harmonique est la même avec les deux méthodes 4.4. L’interprétation du mélange de fréquence semble donc bien être à l’origine des pics de résonances sous-harmoniques.

4.2 Décalage du spectre vers le rouge

L’apparence de nouvelles fréquences de résonances sur la figure 4.1 n’est pas le seul nouveau phénomène observé à forte puissance. En effet, il semble que la fréquence de résonance fondamentale descende à plus basse fréquence. Ce décalage est davantage mis en évidence sur la figure 4.5. Malheureusement, ce décalage n’est explicitement observable autour de la fréquence de résonance fondamentale qu’au delà du rayon de convergence de la méthode perturbative. Ainsi, ce phénomène n’est pas quantitativement vérifié par cette méthode. Cependant, un argument qualitatif permet néanmoins de justifier l’apparition de ce phénomène par la méthode perturbative.

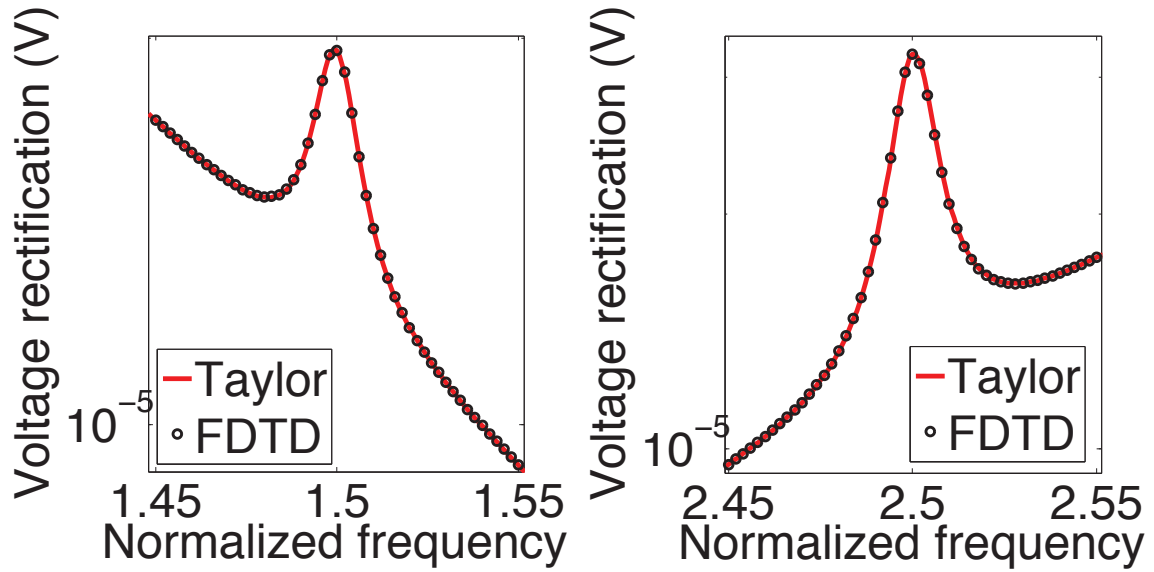


FIGURE 4.4 – Amplitude des pics sous-harmoniques aux fréquences demi-résonantes. La méthode perturbative (développement de Taylor sur la puissance incidente) concorde parfaitement avec la simulation FDTD directe.

En effet, comme on peut le constater figure 4.6, la contribution de l'ordre 4 à la fréquence fondamentale est asymétrique. Ainsi, tant que la puissance d'excitation est dans le rayon de convergence de la série, cette asymétrie tend à baisser la fréquence fondamentale. En effet, la contribution d'ordre 4 est positive, donc augmente le signal, juste en dessous de la fréquence fondamentale et négative, donc supprime le signal, juste au dessus de cette dernière. L'effet global est donc de décaler le maximum du pic de rectification vers les basses fréquences.

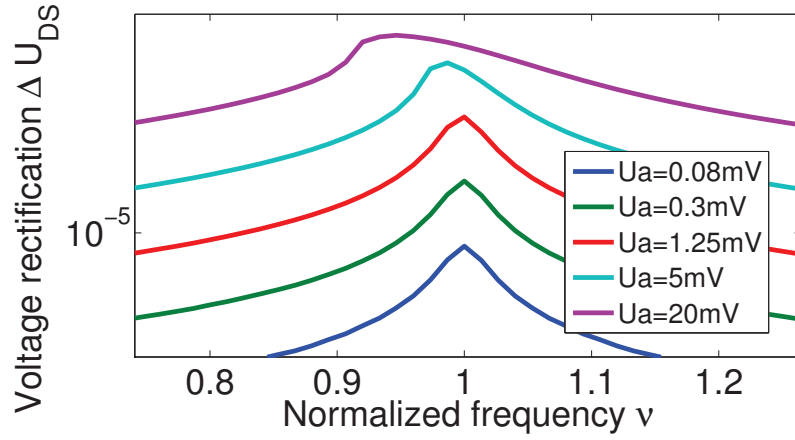


FIGURE 4.5 – Décalage de la fréquence du pic de résonance vers les basses fréquences avec l’augmentation de la puissance d’excitation.

L’apparente convergence de la série perturbative pour de faibles puissances semble suggérer que la rectification du potentiel ΔU_{DS} est une fonction analytique de l’amplitude d’excitation U_a . Son rayon de convergence fini s’explique alors par l’apparition d’un ou des deux phénomènes suivant :

- l’existence de pôles dans le plan complexe
- l’existence de points de branchement

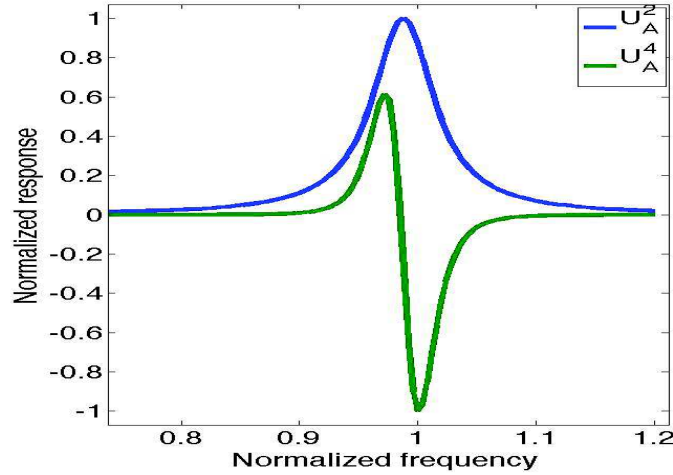


FIGURE 4.6 – Contribution des corrections perturbatives à la rectification d’ordre 2 (contribution fondamentale) et 4 (première correction) en fonction de la fréquence d’excitation. L’asymétrie de la correction à l’ordre 4 suggère le décalage du pic de rectification à plus basse fréquence.

La simulation FDTD 4.7 semble favoriser la deuxième option. Deux phénomènes poussent à cette conclusion. Le premier est que le spectre de réponse simulé en FDTD devient discontinu en fonction de la fréquence à très forte puissance. Si on admet que la rectification est également une fonction analytique de la fréquence, celle-ci doit varier continument avec cette dernière. Le prolongement analytique de la rectification en fonction de la fréquence d'un coté de la discontinuité à l'autre montre alors l'existence de plusieurs valeurs de rectification possible pour certaines puissances et fréquences d'excitations données. Le passage d'une valeur à une autre est alors possible continument en variant les paramètres dans le plan complexe autour d'un point de branchement. Le deuxième phénomène est la non-convergence d'un développement de Padé de la rectification.

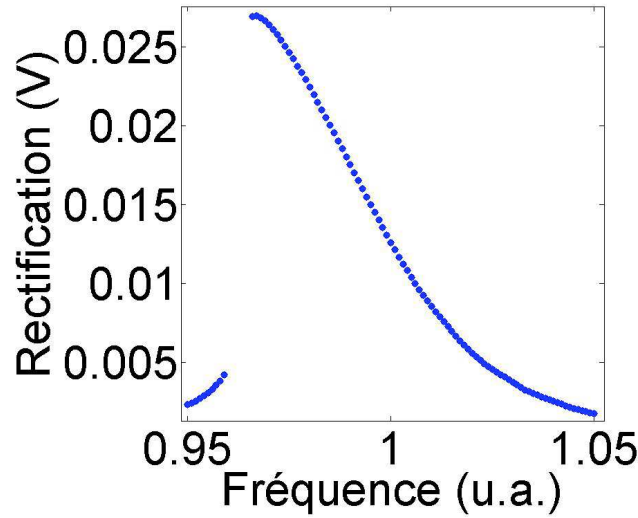


FIGURE 4.7 – Rectification du potentiel drain source à très forte amplitude d'excitation. Le facteur de résonance f semble devenir discontinu comme le montre sa transition brutale à une fréquence d'environ 0,96 fois la fréquence de résonance fondamentale.

Le développement de Padé est une sorte de généralisation du développement de Taylor. Il remplace le polynôme de Taylor par une fonction rationnelle, fraction de deux polynômes. Celui-ci possède alors deux ordres, les deux exposants maximums (p, q) des polynômes respectivement au numérateur et au dénominateur. Le développement de Padé d'ordre $(p, 0)$ correspond alors au développement de Taylor d'ordre p . Un choix approprié de l'ordre (p, q) du développement de Padé peut avoir un rayon de convergence supérieure au développement de Taylor. En effet, le fait que le développement de Padé soit rationnel permet de factoriser la singularité source de la divergence du développement de Taylor. Si le rayon de convergence du développement perturbatif utilisé était limité par la présence d'une singularité, un développement de Padé aurait du permettre d'obtenir une meilleure convergence du développement perturbatif. Au contraire, si le

rayon de convergence est limité par un point de branchement, le contour de celui-ci par un coté ou l'autre dans le plan complexe donne deux valeurs différentes. Le développement de Padé introduit naturellement une coupure grâce à un alignement de certains de ces pôles. Le long de cette coupure, le développement de Padé ne converge toujours pas. Si cette coupure se trouve sur l'axe des réels, alors l'utilisation du développement de Padé ne permet pas l'amélioration de la convergence de la méthode perturbative ce qui est bien constaté.

4.3 Extension de la plage de linéarité

La théorie quasi-linéaire de la rectification du potentiel drain-source permet d'aboutir à la relation (1.13), qui montre une dépendance de la rectification en U_A , l'amplitude d'oscillation du potentiel de l'antenne, au carré. Cette loi quadratique avec l'amplitude U_A se traduit par une loi linéaire avec la puissance incidente. Ainsi, la rectification ΔU_{DS} est ici proportionnelle à la puissance de l'onde incidente sur le transistor. Le calcul permettant d'aboutir à cette conclusion étant perturbatif, on s'attend à voir apparaître des termes correctifs à l'expression (1.13), proportionnels à des puissances $k > 1$ de la puissance incidente $\mathcal{P}$. Ces termes seront donc notamment important lorsque la puissance d'excitation augmente. Grâce à l'approche numérique, nous pouvons directement évaluer l'influence de ces termes d'ordres supérieurs et mesurer leur importance. Notamment, sous certaines conditions, on a pu observer que ces termes sont étonnamment faibles pour des rapports longueur du canal sur longueur d'absorption de l'onde plasma spécifiques. Ainsi, on peut voir sur la figure 4.8, l'apparition de lignes le long desquelles l'écart entre la formule de rectification analytique ne prenant que l'ordre 2 en compte et le résultat numérique est très faible. Le long de ces lignes, les corrections perturbatives d'ordres supérieurs sont faibles et la formule analytique d'ordre 2 est suffisante. La plage de linéarité est donc étendue pour des longueurs de canal spécifiques. La position exacte de ces lignes diffèrent selon les paramètres de la simulation.

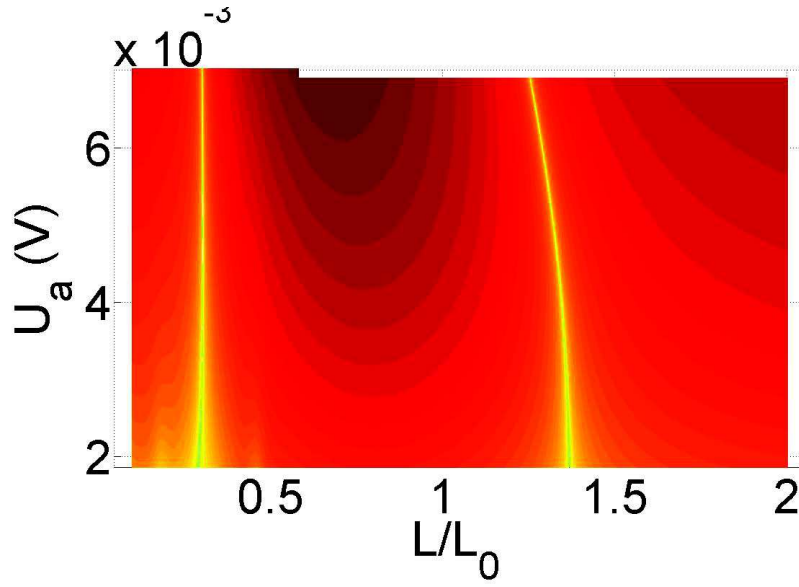


FIGURE 4.8 – Ecart relatif entre la correction d'ordre 2 analytique et la correction totale calculée numériquement en fonction du rapport de la longueur du canal sur la longueur de propagation de l'onde plasma en abscisse et de l'amplitude d'excitation de l'antenne en ordonnée. L'apparition de lignes de faible écart montre que les corrections d'ordres supérieurs peuvent être négligées pour des longueurs de canal spécifiques. Dans ces cas, la formule analytique donnant la rectification du potentiel est valable sur la plus grande plage de fréquence.

Troisième partie

Analyse numérique

Chapitre 5

Algorithmes numériques

Le système Dyakonov-Shur et ses variantes constituent des équations différentielles aux dérivées partielles. Afin de les résoudre, plusieurs méthodes numériques sont envisageables. Parmi ces méthodes, nous allons décrire celles qui sont les plus répandues et que nous avons implémenté. Pour la discrétisation temporelle, ces méthodes sont :

- les méthodes d'Euler,
- la méthode de Crank-Nicholson,
- les méthodes de Runge-Kutta.

Pour la discrétisation spatiale, ces méthodes sont :

- la méthode des différences finies,
- la méthode des éléments finis,
- la méthode spectrale.

5.1 Méthode de discrétisation temporelle

Le système Dyakonov-Shur est un problème aux valeurs initiales. Cela signifie que l'on peut résoudre les équations de Dyakonov-Shur entre un instant initial, arbitrairement imposable à 0, et un instant final T , en divisant l'intervalle de temps $[0; T]$ en N sous intervalles $I_1 = [t_0; t_1]$, $I_2 = [t_1; t_2]$, $\dots$, $I_N = [t_{N-1}; t_N]$ tels que $0 = t_0 < t_1 < \dots < t_N = T$. Lorsque les intervalles $\{I_n\}$ sont suffisamment courts, différents algorithmes permettent d'évaluer la solution du problème à l'instant t_n à partir du temps t_{n-1} avec une précision acceptable. La méthode la plus simple est la méthode d'Euler avant qui est basée sur la définition de l'intégrale de Riemann.

Soit

$$\frac{d}{dt}f = F(f(t), t) \quad (5.1)$$

une équation différentielle générale. On a le théorème fondamental de l'analyse :

$$f(t_n) - f(t_{n-1}) = \int_{t_{n-1}}^{t_n} dt F(f(t), t)$$

Les méthodes d'Euler consiste à considérer l'intervalle I_n comme infinitésimal et d'approximer l'intégrale précédente par :

$$\int_{t_{n-1}}^{t_n} dt F(f(t), t) = (t_n - t_{n-1}) F(f(t_E), t_E) + \mathcal{O}((t_n - t_{n-1})^2)$$

où t_E est un temps compris dans l'intervalle I_n . La méthode d'Euler dite explicite consiste à prendre $t_E = t_{n-1}$. De cette manière, on obtient explicitement $f(t_n)$ connaissant $f(t_{n-1})$. On peut alors calculer progressivement l'évolution de f du temps initial au temps final.

Certains critères de stabilité montrent que les schémas explicites, notamment le schéma d'Euler explicite, ne sont généralement stables que pour un pas de temps Δt suffisamment faible. Pour cela, il est souvent préférable d'utiliser un schéma implicite. Le schéma d'Euler retardé en est un exemple. Celui-ci consiste à prendre $t_E = t_n$ et non t_{n-1} . Ainsi, on obtient une contrainte sur la valeur $f(t_n)$ en fonction des $f(t_{n-1})$ de la forme :

$$G(f(t_n), t_n) = H(f(t_{n-1}), t_{n-1})$$

où G et H sont deux fonctions dépendant de la forme de l'équation différentielle à intégrer. Le schéma est implicite car pour obtenir la valeur $f(t_n)$, il est nécessaire de calculer la réciproque de la fonction G , qui n'est pas toujours triviale. Pour les équations linéaires, cette fonction est une fonction affine de $f(t_n)$ à coefficients potentiellement dépendants de t_n . Dans tous les cas, la réciproque de G est triviale à obtenir et on peut aisément intégrer l'équation (5.1) pas à pas. Si l'équation différentielle est non-linéaire, on ne connaît pas de forme générale à la réciproque de G . On peut néanmoins utiliser un algorithme d'extraction de racines tel que l'algorithme de Newton pour trouver $f(t_n)$ sans inverser G explicitement. L'avantage des schémas implicites est qu'ils sont généralement beaucoup plus stables que les schémas explicites. Ces schémas divergent donc moins souvent que les schémas explicites. Cependant, dans le cadre des équations Dyakonov-Shur, et notamment en présence du ressaut hydraulique, le phénomène de Gibbs prévient la convergence de l'algorithme d'inversion implicite et rend les schémas implicites testés divergents. Pour cette raison, nous avons décidé d'utiliser un schéma Runge-Kutta d'ordre 4 explicite.

Les schémas de Runge-Kutta font partis de la classe des algorithmes de types prédicteurs-correcteurs. Ces schémas évaluent la fonction $f(t_n)$ de façon itérative, en faisant une première prédiction de cette valeur et en la corrigeant progressivement afin d'améliorer la précision de l'estimation obtenue.

5.2 Méthodes de discrétisation spatiale

5.2.1 La méthode des différences finies

Concept de la méthode pour la discrétisation temporelle

La méthode des différences finies est une des méthodes de résolution d'équation différentielles les plus répandues. En effet, il s'agit d'une des méthodes les plus simples à comprendre et implémenter. Celle-ci consiste à remplacer les dérivées d'ordre quelconque d'une fonction en un point par une combinaison linéaire des valeurs de la fonction à dériver sur un ensemble de points discret et espacés d'une distance finie. Concrètement, prenons la définition formelle de la dérivée de $f(x)$:

$$\left. \frac{d}{dx} f \right|_x = \lim_{h \rightarrow 0} \frac{f(x+h) - f(x)}{h}$$

Un schéma de type différence finie consiste à remplacer la limite du pas h tendant vers 0 par une limite finie. On obtient donc :

$$\left. \frac{d}{dx} f \right|_{x_n} \approx \frac{f_{n+1} - f_n}{x_{n+1} - x_n}$$

où f_n est une abréviation pour $f(x_n)$ et x_n un ensemble de points discrets de l'intervalle sur lequel on cherche à résoudre l'équation différentielle donnée. L'ensemble des points x_n est ordonné, ce qui implique que $x_m > x_n$ pour tout indice $m > n$. Ce schéma est le schéma le plus simple de différence finie. Il est d'ordre 1, ce qui signifie que l'erreur de l'approximation est bornée par une certaine constante multipliée par le pas $(x_{n+1} - x_n)$. Ce schéma est qualifié de différence finie avant car pour évaluer la dérivée de f en x_n , on utilise la valeur de f au point suivant x_{n+1} . Le seul autre schéma d'ordre 1 est la différence finie arrière où l'on considère cette fois-ci la valeur de f au point précédent :

$$\left. \frac{d}{dx} f \right|_{x_n} \approx \frac{f_n - f_{n-1}}{x_n - x_{n-1}} \quad (5.2)$$

De manière générale, on peut construire un schéma d'ordre quelconque k en considérant une combinaison linéaire des valeurs de f sur un ensemble E_n^k de $k+1$ points adjacents contenant le point d'évaluation x_n . Dans ce cas, pour un ensemble de points équidistants, l'erreur d'approximation est bornée sous certaines conditions par une constante fois le pas Δx à la puissance k . Ce pas étant idéalement faible, l'élever à une puissance k suffisante diminue l'erreur bien plus vite que de réduire le pas spatial en augmentant le nombre de points total. Il y a néanmoins deux problèmes qui peuvent se poser pour un schéma de différence fini.

Le premier problème qui peut se poser est que le schéma peut ne pas converger si la fonction f n'est pas au moins k fois différentiables sur le plus petit fermé I_n^k simplement connexe contenant les points E_n^k . En effet, le calcul des coefficients de la combinaison

linéaire aboutissant à l'approximation d'une dérivée d'ordre donné provient du calcul de l'interpolation polynomiale de plus petit degré passant exactement par les points de l'ensemble E_n^k . A l'ordre 1, l'interpolation est donnée par la droite passant par les deux points extrêmes (x_n, f_n) et (x_{n+1}, f_{n+1}) pour la différence finie avant. La dérivée en x_n est alors approximée par la pente de cette droite, d'où l'expression (5.2). Si la fonction n'est pas k fois différentiable, alors l'interpolation polynomiale de degré k présentera des oscillations autour des points singulier où la fonction n'est pas différentiable. C'est le phénomène de Gibbs. Dans ce cas, la convergence de l'interpolation sera simple au mieux et l'évaluation des différentes dérivées ne convergera pas.

Le second problème qui peut se poser est que même si la fonction recherchée est k fois différentiable, des oscillations peuvent apparaître dans l'interpolation polynomiale selon la distribution de points utilisée. C'est le phénomène de Runge. Ces oscillations apparaissent du fait de la contribution inéquivalente de différents points. Notamment, dans le cas de points équidistants, les points centraux ont un poids asymptotiquement exponentiellement plus grand que les points les plus externes de E_n^k . Ainsi, de petites variations des valeurs de la fonction en ces points peuvent induire de grandes variations de l'interpolation obtenue et donc de la valeur des différentes dérivées associées. Pour minimiser ce phénomène, il est nécessaire d'utiliser une distribution de points plus dense sur les bords de I_n^k pour calculer l'interpolation, tels que les points de Chebyshev.

Formulation du modèle discret

Les équations de Dyakonov-Shur sont un système d'équations différentielles aux dérivées partielles non-linéaires dans un espace à $1 + 1$ dimensions. On rappelle les équations dans le cadre initial sans viscosité ni résistivité :

$$\begin{cases} \frac{\partial}{\partial t} \rho = -\frac{\partial}{\partial x} (\rho v) \\ \frac{\partial}{\partial t} v = -\frac{\partial}{\partial x} \left(\rho + \frac{1}{2} v^2 \right) \end{cases} \quad (5.3)$$

Pour résoudre ce système d'équation à l'aide des différences finies, on discrétise chaque champ $\rho(t, x)$ et $v(t, x)$ sur un ensemble $E = \{x_1 \dots x_N\}$ de N points. Chaque champ peut être représenté comme un vecteur de dimension N :

$$\begin{cases} \rho(t, x) \rightarrow \begin{pmatrix} \rho(t, x_1) \\ \vdots \\ \rho(t, x_N) \end{pmatrix} = \begin{pmatrix} \rho_1(t) \\ \vdots \\ \rho_N(t) \end{pmatrix} \\ v(t, x) \rightarrow \begin{pmatrix} v(t, x_1) \\ \vdots \\ v(t, x_N) \end{pmatrix} = \begin{pmatrix} v_1(t) \\ \vdots \\ v_N(t) \end{pmatrix} \end{cases}$$

En injectant cette discrétisation dans le système (5.3), on obtient le système approché :

$$\left\{ \begin{array}{l} \frac{\partial}{\partial t} \begin{pmatrix} \rho_1(t) \\ \vdots \\ \rho_N(t) \end{pmatrix} o = -D_x \begin{pmatrix} \rho_1(t) v_1(t) \\ \vdots \\ \rho_N(t) v_N(t) \end{pmatrix} \\ \frac{\partial}{\partial t} \begin{pmatrix} \rho_1(t) \\ \vdots \\ \rho_N(t) \end{pmatrix} = -D_x \begin{pmatrix} \rho_1(t) + \frac{1}{2} v_1^2(t) \\ \vdots \\ \rho_N(t) + \frac{1}{2} v_N^2(t) \end{pmatrix} \end{array} \right. \quad (5.4)$$

où D_x est la matrice carré de dimension N de dérivation. La i -ième ligne de D_x représente les coefficients de la combinaison linéaire approximant la dérivée d'une fonction au point x_i . Cette matrice est creuse car pour un schéma d'ordre k , seules les k diagonales les plus proches de la diagonale principale ont une valeur potentiellement non-nulle. Par exemple, D_x s'écrit avec un schéma avant d'ordre 1 :

$$D_x = \begin{pmatrix} -\frac{1}{x_2 - x_1} & \frac{1}{x_2 - x_1} & & & & \\ & -\frac{1}{x_3 - x_2} & \frac{1}{x_3 - x_2} & & & \\ & & \ddots & \ddots & & \\ & & & \ddots & 1 & -\frac{1}{x_N - x_{N-1}} \\ & & & & \frac{1}{x_N - x_{N-1}} & \end{pmatrix}$$

? ? Condition à spécifier ? ?

qui se simplifie pour une distribution de points équidistants de Δx :

$$D_x = \frac{1}{\Delta x} \begin{pmatrix} -1 & 1 & & & & \\ & -1 & & & & \\ & & 1 & & & \\ & & & \ddots & & \\ & & & & \ddots & \\ & & & & -1 & 1 \\ & ? & ? & \text{Condition à spécifier} & ? & ? \end{pmatrix}$$

où toutes les entrées non spécifiées ou symbolisées sont nulles. Du fait du schéma, seules les $N-1$ premières lignes de la matrice D_x sont imposées par celui-ci. Une condition doit être imposée afin d'obtenir la dernière ligne. Pour un schéma arrière d'ordre 1, c'est la première ligne qui nécessiterait une condition supplémentaire, les $N-1$ dernières lignes étant données par les $N-1$ premières lignes du schéma avant. Une autre solution pour spécifier la ligne manquante est de changer de schéma sur les points nécessaire. Par exemple, dans le cas d'un schéma avant, on peut imposer un schéma arrière sur le dernier point x_N ce qui a pour effet de dupliquer l'avant dernière ligne.

Description du schéma

Pour résoudre le système Dyakonov-Shur, nous avons choisi un schéma spatial d'ordre 4 centré. Il y a plusieurs raisons à ce choix. Comme nous l'avons dit, un schéma d'ordre plus élevé peut potentiellement converger plus rapidement, à nombre de points total fini, qu'un schéma d'ordre plus faible. Ainsi, pour économiser du temps de calcul et de la mémoire, il convient d'utiliser un ordre aussi grand que possible. Cependant, le système Dyakonov-Shur possède parmi ces solutions mathématiques des solutions discontinues. L'apparition de discontinuités dans un schéma d'ordre élevé entraîne l'apparition du phénomène de Gibbs, soit l'apparition d'oscillations fictives des champs modélisés autour de ces discontinuités. Ces oscillations peuvent alors rendre le schéma instable pour différentes raisons, notamment :

- la dérivée potentiellement élevée des oscillations peut favoriser leur amplification numérique par les erreurs d'arrondi
- le caractère intrinsèquement instable du modèle pour certaines valeurs de courant peut coupler ces oscillations avec les modes propres Dyakonov-Shur entraînant leur amplification.
- le caractère numériquement instable des modes propres Dyakonov-Shur peut amplifier ces oscillations même si le régime sondé est théoriquement stable.

L'origine de la troisième forme d'instabilité vient du fait que la discrétisation finie utilisée entraîne une évaluation approximative des pulsations complexes des modes propres des ondes plasmas. En particulier, les modes de courtes longueurs d'ondes seront associés aux pulsations complexes ω les plus éloignées de leurs valeurs asymptotiques, tandis que les modes de grande longueur d'onde auront bien mieux convergé. En effet, les oscillations de courtes longueurs d'onde sont les plus mal discrétisées car peu de points de discrétisation couvrent une longueur d'onde. Hors, les oscillations de Gibbs sont justement des oscillations courtes. Celles-ci seront donc potentiellement amplifiées par le gain numérique effectif calculé pour un nombre de points de discrétisation donné. Dans ce cas, augmenter la résolution de la discrétisation ne résout pas le problème car la longueur d'onde de l'oscillation de Gibbs décroît linéairement avec le raffinement de la discrétisation. Ainsi, cette oscillation se couplera toujours à un mode dont l'amplitude complexe effective n'a pas convergé vers sa limite. Si le gain associé est positif, les oscillations de Gibbs seront instables.

Pour minimiser le couplage des oscillations de Gibbs avec ces instabilités, il convient d'utiliser un schéma d'ordre faible. Pour cela, l'ordre 4 semble le meilleur compromis entre temps de calcul et stabilité. Le caractère centré du schéma vient du fait que la matrice de différentiation D_x évalue la dérivée au point x_n en utilisant une combinaison linéaire antisymétrique des points x_{n-2} à x_{n+2} , sauf pour les 2 premiers et derniers points pour lesquels certaines de ces valeurs ne sont pas disponibles. En ces points, le schéma est décalé d'autant de points que nécessaire vers l'intérieur du domaine. Par exemple, on utilise les points x_1 à x_5 pour l'évaluation des dérivées en x_1 et x_2 , ce qui revient à utiliser le même polynôme interpolant sur ces points. Les conditions aux bords correctes sont ensuite imposées manuellement à chaque pas de temps. L'importance du

caractère centré du schéma vient de la possibilité d'apparition d'un ressaut hydraulique, ce que l'on va expliquer dans la partie suivante.

Gestion du ressaut hydraulique

Sous certaines conditions, une onde plasma discontinue, le ressaut hydraulique, peut se former dans le système Dyakonov-Shur. Cette discontinuité est une des grandes difficultés qui se présentent lors de l'intégration numérique du système d'équations. Même si un schéma est stable par rapport à l'apparition du ressaut, celui-ci peut ne pas correctement le modéliser. En effet, utiliser un schéma asymétrique peut permettre de réduire le phénomène de Gibbs. Pour cela, on peut par exemple de modifier le schéma de discrétisation en utilisant des points d'interpolation E_n se trouvant du même côté de la discontinuité. Le problème est que dans ce cas, le domaine spatial est divisé en deux régions séparée par la discontinuité, où les points d'une région n'ont accès qu'à l'information se trouvant dans leurs régions respective. Du point de vue des points autour de la discontinuité, celle-ci disparaît et les valeurs des champs de densité et de vitesse des points de l'autre côté de la discontinuité semblent extrapolées par continuité depuis les valeurs des champs de leurs régions. Une corrélation des deux domaines doit être ajoutée. Pour cela on peut s'inspirer des schémas dits upwind/downwind.

Les schémas upwind/downwind sont des schémas de discrétisation asymétriques dont l'orientation dépend du sens de propagation physique de l'information. Les schémas dits upwind prennent en compte dans l'évaluation d'une dérivée plus de points du côté d'où vient l'onde que vers où elle se déplace. Au contraire, les schémas downwind considèrent plus de points du côté vers lequel l'onde se dirige que du côté d'où elle vient. Ainsi, si on considère une onde se déplaçant dans le sens des x positifs, le schéma d'Euler arrière est upwind tandis que le schéma avant est downwind. Pour connecter les domaines de part et d'autre d'une discontinuité, on peut intégrer le premier point à droite de la discontinuité dans le domaine à gauche de celle-ci, faisant du schéma un schéma upwind si l'onde se propage vers la droite et downwind sinon. De façon équivalente, on peut à la place intégrer le dernier point à gauche de la discontinuité au domaine à droite de celle-ci. Un critère qui semble physiquement justifié pour déterminer quel point intégrer dans l'autre domaine et de rendre le schéma upwind. En effet, si une onde se déplace dans une direction, la valeur que prend un champ en un point à un instant correspond généralement à la valeur que ce champ a pris en un point d'où vient l'onde à un instant précédent. Ainsi, le schéma upwind reflète cette observation physique mais rien n'empêche en principe d'obtenir la même information à l'aide d'un schéma downwind en prolongeant en extrapolant la fonction dans la direction d'où l'onde vient.

Dans le cadre du système Dyakonov-Shur, ce genre de schéma ne modélise cependant pas correctement la propagation du ressaut hydraulique. En effet, on peut calculer la vitesse du ressaut en supposant une solution de la forme :

$$\begin{cases} \rho(t, x) = \rho_- + (\rho_+ - \rho_-) H(ut - x) \\ v(t, x) = v_- + (v_+ - v_-) H(ut - x) \end{cases} \quad (5.5)$$

où on a :

- $H(x)$ est la fonction de Heaviside telle que $H(-|x|) = 0$ et $H(|x|) = 1$
- u est la vitesse de propagation du ressaut
- et $f_{\pm}$ est la valeur de la fonction f après (+) et avant (−) le ressaut.

En injectant les équations (5.5) dans le système (1.8) et en le supposant vrai au sens des distributions, on obtient les contraintes suivantes :

$$\begin{cases} (u - \bar{v})^2 = \bar{\rho} \\ \delta \rho^2 = \bar{\rho} \delta v^2 \end{cases}$$

où $\delta f = f_+ - f_-$ et $\bar{f} = \frac{f_+ + f_-}{2}$ pour toute grandeur f . Ainsi, la vitesse u de propagation du ressaut dépend de façon symétrique des valeurs des champs de densité et vitesse de part et d'autre de celui-ci, suggérant le besoin d'utiliser un schéma centré. En effet, si on utilisait un schéma asymétrique, on obtient la contrainte suivante : $(u_{\pm} - v_{\pm})^2 = \rho_{\pm}$. Pour une onde se déplaçant vers la droite, u_+ correspond à la vitesse downwind et u_- à la vitesse upwind et inversement pour une onde se déplaçant vers la gauche. La différence potentiellement non bornée de la valeur des champs de part et d'autre du ressaut fait que la différence des valeurs de vitesse $u_{\pm}$ et u peut être arbitrairement grande. Ainsi, seul un schéma centré peut correctement propager le ressaut hydraulique.

Régularisation du ressaut hydraulique

Dans la plupart des cas, le schéma de discrétisation implémenté permet l'intégration numérique du système Dyakonov-Shur en un temps raisonnable. Cependant, dans certains cas, notamment lorsque l'on considère un modèle unifié de la charge, le phénomène de Gibbs rend la simulation numérique instable et l'algorithme diverge. Pour éviter de réduire d'avantage le degré d'approximation choisi, nous avons décidé d'implémenter un filtre spatial destiné à supprimer uniquement les oscillations de Gibbs. A chaque pas de temps, avant d'intégrer le système, on applique un filtre spatial destiné à supprimer uniquement les variations de faibles longueur d'onde tel que les oscillations de Gibbs. Ce filtre est obtenu en réalisant une moyenne locale des champs ρ et v . Pour cela, on commence par remplacer la valeur f_n du champ f en un point de discrétisation x_n par la moyenne suivante :

$$f'_n = \frac{f_{n+1} + 2f_n + f_{n-1}}{4}$$

Cette moyenne a la propriété de réduire fortement les oscillations en dents-de-scie des champs ρ et v . Puis une deuxième étape de correction de courbure est effectuée :

$$f''_n = \frac{6f'_n - (f'_{n+1} + f'_{n-1})}{4}$$

Ces deux étapes peuvent être combinées en une seule pour obtenir :

$$f''_n = \frac{-f_{n-2} + 4f_{n-1} + 10f_n + 4f_{n+1} - f_{n+2}}{16}$$

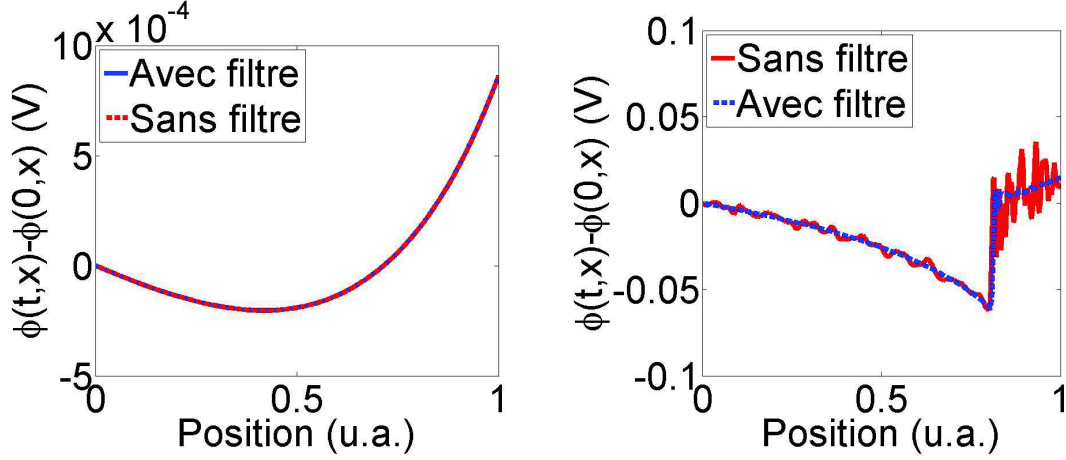


FIGURE 5.1 – Comparaison des simulations numérique FDTD avec l'utilisation d'un filtre spatial et sans. On constate qu'en l'absence de discontinuité, les deux simulations sont équivalentes. En revanche, lorsque le ressaut apparaît, les oscillations de Gibbs provoquent l'apparition d'un signal en dents-de-scie. Le filtrage spatial appliqué semble à même de supprimer ce signal fictif sans déformer le signal attendu.

La deuxième étape sert à compenser l'atténuation apportée par la première. En effet, bien que la première étape atténue majoritairement les oscillations courtes longueurs d'onde fictives, elle atténue également dans une moindre mesure les oscillations grandes longueurs d'onde réelles. En effet, si on applique uniquement la première étape à une fonction de la forme $f(x) = \sin(kx)$, on obtient $f'(x) \approx \left(1 - \frac{k^2 \Delta x^2}{4}\right) f(x)$ pour $k \ll \frac{1}{\Delta x}$. Cette régularisation étant effectuée à chaque pas de temps Δt , on obtient un amortissement total de $\left(1 - \frac{k^2 \Delta x^2}{4}\right)^{\frac{t}{\Delta t}} \approx \exp\left(-\frac{k^2 \Delta x^2}{4 \Delta t} t\right)$ après un temps t . La première étape introduit donc un amortissement numérique diminuant le gain g des modes d'oscillation de $\frac{k^2 \Delta x^2}{4 \Delta t}$. La seconde étape sert à compenser ce facteur afin de l'amplitude d'oscillation des modes plasma propres ne soit plus amortie par le filtre spatial.

En l'absence de discontinuités, on peut voir sur la figure 5.1, que l'effet du filtre spatial n'est pas perceptible. En revanche, lorsqu'une discontinuité apparaît, le phénomène de Gibbs provoque une oscillation en dents-de-scie parfaitement observable sur la simulation originale non filtrée. Les oscillations de Gibbs dans la simulation avec le filtre spatial en place sont en revanche correctement atténuées. La solution filtrée prend des valeurs tout à fait réaliste et semble passer exactement là où on s'y attendrait visuellement. Le filtre semble donc bien remplir son rôle de stabilisation en atténuant uniquement les oscillations de Gibbs.

5.2.2 La méthode des éléments finis

La méthode des éléments finis est une méthode très populaire, notamment pour ce qui est de la résolution d'équations aux dérivées partielles elliptiques dans un domaine à géométrie complexe. Celle-ci est basée sur un formalisme très différent de celui de la méthode des différences finies. Cette méthode est basée sur la discrétisation d'une formulation faible du problème. Partant d'une équation différentielle quelconque, la formulation forte, on obtient la formulation faible en multipliant l'équation de départ par une fonction quelconque $\varphi(x)$ et en intégrant le tout sur le domaine de définition de l'équation différentielle. La formulation faible du système (1.8) est donnée par :

$$\begin{cases} \int_0^L dx \varphi \left[\frac{\partial}{\partial t} \rho + \frac{\partial}{\partial x} (\rho v) \right] = 0 \\ \int_0^L dx \varphi \left[\frac{\partial}{\partial t} v + \frac{\partial}{\partial x} \left(\rho + \frac{1}{2} v^2 \right) + \frac{v}{\tau} \right] = 0 \end{cases} \quad (5.6)$$

Pour résoudre la formulation faible du problème, on décompose les champs ρ et v sur un ensemble de fonctions de bases $e_k(x)$ à support compacte. Suivant le type de noeuds choisis, la forme des fonctions de base diffère. Le cas le plus simple correspond à une division du domaine de définition de l'équation différentiel en sous-intervalles dont seuls les bornes constituent les noeuds de calcul. La fonction de base $e_k(x)$ associée au point x_k est alors la fonction linéaire par morceau définie par :

$$e_k(x) = \begin{cases} \frac{x - x_{n-1}}{x_n - x_{n-1}} & \text{si } x \in [x_{n-1}; x_n] \\ \frac{x_{n+1} - x}{x_{n+1} - x_n} & \text{si } x \in [x_n; x_{n+1}] \\ 0 & \text{sinon} \end{cases}$$

En évaluant la formulation faible sur les fonctions de base ($\varphi = e_k$), on obtient une discrétisation des opérateurs dérivée similaire à la formulation en différence finie. La matrice obtenue n'est cependant pas exactement identique à celle obtenue précédemment mais le formalisme d'intégration temporelle utilisé dans le schéma de différence finie se transpose trivialement ici. Lorsqu'on considère un modèle unifié de la charge, seule cette méthode permet la simulation FDTD.

5.2.3 Les méthodes spectrales

Les méthodes spectrales sont des méthodes d'approximation ne reposant plus sur une discrétisation spatiale du domaine d'intégration mais sur la troncature d'une somme pondérée de fonctions de bases dont la limite restitue la fonction à calculer. Un exemple de méthode spectrale est simplement la transformée de Fourier qui remplace les fonctions de carré sommable sur un intervalle par une série de coefficients f_n . La somme

infinie d'exponentielles complexes, qui sont des fonctions sinusoïdales complexes, pondérées par les coefficients f_n converge dans une certaine mesure vers la fonction de départ en tous points du domaine. Lorsque les coefficients de Fourier décroissent suffisamment vite avec l'ordre n , la sommation jusqu'à un certain ordre N permet de restituer la fonction originale avec une précision suffisante. D'autres types de fonctions de base peuvent également être considérées comme les polynômes de Legendre ou autres familles de polynômes orthogonaux selon une fonction de poids w . Le degré de précision de l'approximation numérique est alors donné par le degré maximal du polynôme de base considéré. L'intérêt de ce genre de schéma est que les opérations de différentiation/intégration peuvent être calculées exactement sur les fonctions de bases choisies. La limite de résolution, plus petite échelle de variation spatiale mesurable, est donnée non plus par la distance entre deux points de discrétisation mais par la plus petite échelle de variation spatiale des fonctions de bases. Pour la transformée de Fourier, celle-ci est par exemple donnée par la période spatiale ou longueur d'onde λ du plus grand vecteur d'onde k considéré.

Bien que calculer les coefficients spectraux associés à une base donnée puisse paraître fastidieux, de nombreux algorithmes et méthodes permettent de calculer ces derniers de manière très efficace ainsi que la valeur d'une fonction connaissant ses coefficients sur une base. On dispose par exemple de la transformée de Fourier rapide, ou FFT de l'anglais Fast Fourier Transform, pour la base de Fourier.

Les fonctions de bases "globales" généralement considérées (base de Fourier, Chebyshev, ...) sont des fonctions dites lisses. Elles font en effet partis des fonctions holomorphes dans le plan complexe et donc indéfiniment différentiables. L'inconvénient d'utiliser ce genre de fonction pour notre problème est que ces fonctions sont très mal adaptées pour représenter des fonctions singulières, comme les fonctions discontinues ou non différentiables. En effet, le phénomène de Gibbs est un phénomène parfaitement mathématique et non numérique. Son apparition dans les simulations ne reflète que son existence mathématique sous-jacente. En conséquence, les méthodes spectrales appliquées à la résolution FDTD n'ont pas pu être implémentée car souffrant de divergence à la moindre apparition du ressaut hydraulique. Cependant, ces méthodes ont été les plus efficaces pour résoudre le problème aux valeurs propres des oscillations plasmas perturbatives en présence de résistivité. On rappelle le système (3.4) à résoudre afin d'évaluer la stabilité d'un courant stationnaire à travers le transistor :

$$\frac{\partial}{\partial x} \left[\begin{pmatrix} v_s & \rho_s \\ 1 & v_s \end{pmatrix} \begin{pmatrix} \tilde{\rho} \\ \tilde{v} \end{pmatrix} \right] + \begin{pmatrix} 0 & 0 \\ 0 & \frac{1}{\tau} \end{pmatrix} \begin{pmatrix} \tilde{\rho} \\ \tilde{v} \end{pmatrix} = i\omega \begin{pmatrix} \tilde{\rho} \\ \tilde{v} \end{pmatrix}$$

muni des conditions aux bords suivantes :

$$\begin{cases} \tilde{\rho}(\omega, 0) = 0 \\ v_s(L) \tilde{\rho}(\omega, L) + \rho_s(L) \tilde{v}(\omega, L) = 0 \end{cases}$$

Ainsi, le problème se résume à calculer les valeurs propres $i\omega$ et fonctions propres $\begin{pmatrix} \tilde{\rho} \\ \tilde{v} \end{pmatrix}$ de l'opérateur $M = D_x * \begin{pmatrix} v_s & \rho_s \\ 1 & v_s \end{pmatrix} + \begin{pmatrix} 0 & 0 \\ 0 & \frac{1}{\tau} \end{pmatrix}$ L'opérateur D_x correspond à la

représentation de la dérivée sur une base donnée. Dans l'espace réciproque associé à la position x , cet opérateur est diagonal et vaut ik sur la ligne associée à la fonction de base $e_k = \exp(ikx)$. L'avantage de cette formulation, est que si on se donne une base e_n , il est généralement facile d'obtenir la représentation matricielle de l'opérateur dérivée dans cette base. Afin, de calculer les valeurs propres et vecteurs propres associés, il suffit de diagonaliser la matrice M , ce qui peut se faire relativement rapidement à l'aide la fonction `eig` intégrée dans `matlab`.

La base que nous avons choisi est la base de Chebyshev pour ces propriétés de convergence très rapides. En effet, la complexité algorithmique de la diagonalisation d'une matrice de rang N varie en $\mathcal{O}(N^3)$. Ainsi, le temps de calcul nécessaire à la diagonalisation de M vaut environ 2.10^{-2} secondes pour $N = 100$ mais 2,4 secondes pour $N = 1000$. Afin de générer la carte fig.3.4, il faut effectuer cette opération près d'un million de fois pour avoir une résolution de 1000×1000 pixels. Il est donc fondamental d'utiliser une base permettant la convergence rapide de la valeur propre du mode fondamental et de ses fonctions propres associées afin de réduire au mieux la taille de la matrice M . Pour cela, la base de Chebyshev est la meilleure que nous ayons implémenté. En effet, dans cette base, la valeur propre du mode fondamentale converge à la précision numérique près avec moins de 20 points de discrétisations seulement pour un temps de calcul de moins d'un millièème de seconde, permettant le calcul de la carte fig.3.4 en quelques minutes seulement pour une bonne résolution. Pour arriver à une convergence relative de seulement 10^{-6} , qui n'est même pas la précision simple d'un nombre flottant, avec un algorithme de différence finie, il faut plus de 1000 points de discrétisation par pixel. La supériorité de la méthode spectrale est donc sans appel, même si une précision numérique relative de seulement 10^{-4} est plus que suffisante.

Conclusion

Dans ce manuscrit, nous avons présenté un mécanisme prédisant l'émission d'une part et la détection d'autre part d'ondes électromagnétiques. Nous avons montré comment ce mécanisme permet la conversion d'un signal térahertz incident, bien trop rapide pour un appareil de mesure conventionnel, est converti en un signal basse fréquence aisément mesurable par couplage non-linéaire. Nous avons développé la théorie initiale plus en profondeur et formalisé les corrections perturbatives à tout ordre. Ces corrections ont une importance cruciale dans les régimes de détection résonante, où les effets diffusifs sont moindres. Le calcul analytique de ces termes étant fastidieux un algorithme numérique a été implémenté. Afin de vérifier l'exactitude des corrections perturbatives obtenues, une comparaison de cette méthode avec un algorithme de résolution direct non-perturbatif a été effectué. Le système étudié admettant des solutions discontinues, un algorithme stable à l'apparition de discontinuités a dû être développé. La cohérence des différents résultats obtenus suggère l'exactitude des méthodes développées. Ainsi, nous avons été en mesure de montrer l'apparition de nouveaux phénomènes de rectification non-linéaire pour la détection d'ondes électromagnétiques. Notamment, le résultat le plus surprenant est certainement l'apparition de nouvelles fréquences de résonance sous-harmoniques à forte puissance incidente. Ces résonances ont été montrées comme résultant d'une interaction non-linéaire dont l'efficacité croît avec la puissance incidente, expliquant ainsi l'absence de ces résonances à faible intensité. Nous avons également pu établir un critère nécessaire pour atteindre le régime d'instabilité plasma. Dans ce régime, les ondes plasmas spontanément générées dans le canal du transistor émettent un rayonnement électromagnétique.

Afin de vérifier ces prédictions expérimentalement, il est nécessaire de disposer de transistors dits de forte mobilité. Cette forte mobilité traduit la faible résistivité au courant du transistor. Ces transistors auraient un facteur de qualité grand devant 1, permettant la résonance et donc l'interférence des ondes plasmas au sein du canal du transistor. L'apparition et le développement de nouvelles technologies basées sur de nouveaux matériaux, tels que les transistors au graphène, semble suggérer la possibilité de vérifier expérimentalement ces prédictions dans un futur relativement proche. De plus, l'utilisation de transistors résonants et l'accordabilité de ses fréquences de résonance permettra une très grande discrimination fréquentiel et versatilité du dispositif.

Annexe A

A.1 Inversion de la relation ρ/U_{GS} unifiée

Selon le modèle [9], le potentiel électrostatique dans le canal est donné en fonction de la densité par :

$$U_{GS} = V_T + \frac{q}{C} (\rho - \rho_T) + \eta V_{th} \ln \left(\frac{\rho}{\rho_T} \right)$$

On obtient la densité en fonction du potentiel comme suit :

$$\begin{aligned} \ln \left(\frac{\rho}{\rho_T} \right) + \frac{q}{C\eta V_{th}} \rho &= \frac{U_{GS} - V_T + \frac{q\rho_T}{C}}{\eta V_{th}} \\ \frac{\rho}{\rho_T} \exp \left(\frac{q\rho}{C\eta V_{th}} \right) &= \exp \left(\frac{U_{GS} - V_T + \frac{q\rho_T}{C}}{\eta V_{th}} \right) \\ \frac{q\rho}{C\eta V_{th}} \exp \left(\frac{q\rho}{C\eta V_{th}} \right) &= \frac{q\rho_T}{C\eta V_{th}} \exp \left(\frac{U_{GS} - V_T + \frac{q\rho_T}{C}}{\eta V_{th}} \right) \end{aligned}$$

En posant $\frac{q\rho}{C\eta V_{th}} = w$ et $\frac{q\rho_T}{C\eta V_{th}} \exp \left(\frac{U_{GS} - V_T + \frac{q\rho_T}{C}}{\eta V_{th}} \right) = z$, on obtient la définition de la fonction W de Lambert :

$$w \exp(w) = z$$

dont la solution est donnée par :

$$w = W_n(z) \quad \forall n \in \mathbb{Z}$$

où W_n est la n -ième branche de la fonction W de Lambert. La seule solution purement réelle est donnée par la branche principale de la fonction W de Lambert, pour $n = 0$. On obtient alors :

$$\frac{q\rho}{C\eta V_{th}} = W_0 \left(\frac{q\rho_T}{C\eta V_{th}} \exp \left(\frac{U_{GS} - V_T + \frac{q\rho_T}{C}}{\eta V_{th}} \right) \right)$$

d'où le résultat annoncé :

$$\rho = \frac{C\eta V_{th}}{q} W_0 \left(\frac{q\rho_T}{C\eta V_{th}} \exp \left(\frac{U_{GS} - V_T + \frac{q\rho_T}{C}}{\eta V_{th}} \right) \right)$$

Bibliographie

- [1] M.Dyakonov and M.Shur, "Shallow water analogy for a ballistic field effect transistor : new mechanism of plasma wave generation by dc current," *PHYSICAL REVIEW LETTERS*, vol. 71, october 1993.
- [2] P. Yu and M. Cardona, *Fundamentals of semiconductors*. Springer, fourth ed., march 2010.
- [3] D. Lide, *Handbook of chemistry and physics*. CRC press, 2003-2004.
- [4] U. Mishra and J. Singh, *Semiconductor device physics and design*. Springer, 2008.
- [5] W. Knap and J. Lusakowski, "Terahertz emission by plasma waves in 60 nm gate high electron mobility transistors," *APPLIED PHYSICS LETTERS*, vol. 84, march 2004.
- [6] M.Dyakonov and M.Shur, "Detection, mixing, and frequency multiplication of terahertz radiation by two-dimensional electronic fluid," *IEEE TRANSACTIONS ON ELECTRON DEVICES*, vol. 43, march 1996.
- [7] Y. Wu, K. Jenkins, A. Valdes-Garcia, D. Farmer, Y. Zhu, A. Bol, C. Dimitrakopoulos, W. Zhu, F. Xia, P. Avouris, and Y. Lin, "State-of-the-art graphene high-frequency electronics," *Nano letters*, 2012.
- [8] W.Knap, V.Kachorovskii, Y.Deng, S.Rumyantsev, J.-Q.Lü, and R.Gaska, "Nonresonant detection of terahertz radiation in field effect transistors," *JOURNAL OF APPLIED PHYSICS*, vol. 91, june 2002.
- [9] Y. H. Byun, K. Lee, and M. Shur, "Unified charge control model and subthreshold current in heterostructure field effect transistors," *IEEE ELECTRON DEVICE LETTERS*, vol. 11, january 1990.
- [10] T. Fjeldly, M. Shur, and T. Ytterdal, *Introduction to device modeling and circuit simulation*. New York, NY, USA : John Wiley & Sons, Inc., 1st ed., 1997.
- [11] a. E. M. L. L. D. Landau, *Fluid mechanics*, vol. 6. Pergamamon press, 1959.
- [12] A.Dmitriev, A.Furman, and V.Yu.Kachorovskii, "Nonlinear theory of the current instability in a ballistic field-effect transistor," *PHYSICAL REVIEW B*, vol. 54, november 1996.
- [13] A.Dmitriev, A.Furman, V.Yu.Kachorovskii, G.Samsonidze, and Ge.Samsonidze, "Numerical study of the current instability in a two-dimension electron fluid," *PHYSICAL REVIEW B*, vol. 55, april 1997.