



Simulation des plasmas de tokamak avec XTOR : régimes des dents de scie et évolution vers une modélisation cinétique des ions.

Leblond David

► To cite this version:

Leblond David. Simulation des plasmas de tokamak avec XTOR : régimes des dents de scie et évolution vers une modélisation cinétique des ions.. Physique des plasmas [physics.plasm-ph]. Ecole Polytechnique X, 2011. Français. NNT : . pastel-00618453

HAL Id: pastel-00618453

<https://pastel.hal.science/pastel-00618453>

Submitted on 1 Sep 2011

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.

SIMULATION DES PLASMAS DE TOKAMAK AVEC XTOR : RÉGIMES DES DENTS DE SCIE ET ÉVOLUTION VERS UNE MODÉLISATION CINÉTIQUE DES IONS

THÈSE

présentée pour obtenir le grade de

DOCTEUR DE L'ÉCOLE POLYTECHNIQUE

spécialité

PHYSIQUE DES PLASMAS

par

DAVID LEBLOND

soutenue le 06/07/2011 devant le jury composé de :

Mme Laurence Rezeau,	<i>Présidente</i>
M. Patrick Mora,	<i>Examineur</i>
M. Peter Beyer,	<i>Examineur</i>
M. Xavier Garbet,	<i>Rapporteur</i>
M. Olivier Sauter,	<i>Rapporteur</i>
M. Hinrich Lütjens,	<i>Directeur de thèse</i>
M. Jean-François Luciani,	<i>Directeur de thèse</i>

Table des matières

1	Introduction	1
1.1	Fusion nucléaire contrôlée et critère de Lawson	1
1.2	Confinement magnétique et tokamaks	2
1.3	Modélisation multi-modèles et codes hybrides	4
1.4	Présentation du travail de thèse	6
I	Régimes d'établissement des dents de scie en plasmas ohmiques	9
2	La MHD appliquée aux tokamaks : cas des dents de scie	10
2.1	Généralités sur la MHD	10
2.1.1	Dérivation du modèle : de Vlasov au mono-fluide	11
2.1.2	Hypothèses et domaine de validité	14
2.1.3	Principe d'Énergie : conditions nécessaires de stabilité	17
2.2	Les tokamaks vus par la MHD	18
2.2.1	Géométrie et configuration magnétique	20
2.2.2	Instabilités MHD dans les tokamaks	23
2.3	Les dents de scie	26
2.3.1	Découverte et description du phénomène	27
2.3.2	Théorie de Kadomtsev : reconnection du kink interne	28
2.3.3	Confrontation aux expériences, effets bi-fluides et idéaux	30
2.3.4	MHD étendue et effets non-linéaires	32
2.3.5	Approche semi-empirique et questions ouvertes	33
2.4	Conclusion	34
3	Dynamique du kink interne avec XTOR	36
3.1	XTOR : un code MHD non-linéaire modulaire pour les tokamaks	36
3.1.1	Solveur et schéma numérique implicite	37
3.1.2	Systèmes d'équations	39
3.1.3	Opérateurs MHD et pré-conditionneur	46
3.1.4	Géométrie et conditions aux limites	48
3.2	Étude numérique des dents de scie avec XTOR : cadre et objectif	51
3.2.1	Modèle physique $\eta\chi$ MHD et paramètres clés	51
3.2.2	Conditions initiales	52
3.2.3	Modes représentés	54
3.3	Régimes du kink en $\eta\chi$ MHD	54
3.3.1	Régimes dans l'espace des paramètres $(\chi_{\perp}, \beta_{p1})$	54
3.3.2	Régime oscillant	56
3.3.3	Régime stationnaire hélicoïdal	58

3.3.4	Régime de faibles oscillations	58
3.3.5	Transition entre les régimes $\eta\chi$ MHD	60
3.4	Ajout des effets diamagnétiques	61
3.4.1	Évolution du régime oscillant en dents de scie	61
3.4.2	Évolution des seuils de transition et premières comparaisons avec JET	63
3.5	Conclusion	65
II Contribution au développement de XTOR-K		67
4	XTOR-K : un code hybride MHD/PIC « full-f, full orbit »	68
4.1	Intégration d'une population cinétique dans un code MHD	68
4.1.1	Physique des ions rapides dans un tokamaks	69
4.1.2	Modèles cinétiques et hypothèse δf	70
4.2	Algorithme de couplage pour XTOR-K	71
4.2.1	Schéma Newton-Krylov/Picard et asymétrie temporelle	71
4.2.2	Condition suffisante de stabilité linéaire	74
4.2.3	Comportement pour les modes MHD fondamentaux	76
4.3	Conclusion	81
5	Intégration numérique des trajectoires	82
5.1	Schéma leap-frog et méthode de Boris	82
5.1.1	Algorithme de Boris	82
5.1.2	Adaptation en géométrie torique	85
5.1.3	Stabilité et précision	88
5.2	Implémentation pour XTOR	88
5.2.1	Unités et normalisation	88
5.2.2	Paramètres, entrées et sorties.	90
5.2.3	Mapping, changement de base, et interpolations des champs	92
5.2.4	Détail des opérations et distribution du temps de calcul	94
5.3	Trajectoires et conservation des invariants	98
5.3.1	Trajectoires des charges dans un tokamak	98
5.3.2	Conservation des invariants du mouvement	100
5.4	Conclusion	110
6	Stratégie de réduction du bruit sur le tenseur de pression	111
6.1	Définition et calcul du tenseur de pression	112
6.1.1	Expression du tenseur de pression numérique	112
6.1.2	Dépôt du tenseur numérique	112
6.2	Régularisation de la distribution initiale dans l'espace des phases	113
6.2.1	Anisotropie de l'initialisation dans l'espace des phases	114

6.2.2	Low Discrepancy Sequences et séquence de Sobol	116
6.2.3	Variation en N_{part}^{-1} du bruit après initialisation	121
6.2.4	Fonction de distribution discrète à poids égaux	122
6.3	Aspect temporel et filtrage du tenseur de pression	124
6.3.1	Structure temporelle du signal Π_n	124
6.3.2	Filtre de Butterworth	126
6.4	Résultats préliminaires : kink interne avec particules chaudes	132
7	Conclusion Générale	137
	Références	139

Remerciements

Je tiens à remercier Olivier Sauter et Xavier Garbet d'avoir accepté d'être les rapporteurs de cette thèse ; leurs remarques et corrections ont grandement contribué à la qualité du manuscrit. Je remercie également Laurence Rezeau, Patrick Mora, et Peter Beyer d'avoir bien voulu participer au jury.

Je suis reconnaissant au personnel du Centre de Physique Théorique de l'École Polytechnique de leur accueil chaleureux, du directeur, Patrick Mora, aux informaticiens, Florence Hamet et Stéphane Aircadi, sans oublier les secrétaires, Florence Auger, Malika Lang, Jeannine Thomas et Fadila Debbou.

Je ne puis manquer de nommer ici Jeanne, Claudia, Mounir, Jeff et Clémentine, sans qui mes déjeuners auraient été plus mornes et mes week-ends plus tristes. Leur soutien et leur amitié m'ont été et me sont toujours infiniment précieux. Leur contribution à ce travail ne se compte pas en nombre de pages : elle est inestimable.

Je n'oublie pas Abdou, Valentina, Benjamin, et Liu, qui ont su ramener un peu de fraîcheur et de soleil dans un bureau de thésards vieillissants, perdant peu à peu ses occupants dans une lente et fatale hémorragie.

À ma grande sœur Aude, pour avoir soutenu sa thèse quelques mois avant moi, (cf. [0]) me fournissant comme il se doit un modèle sororal salubre, j'adresse toute ma gratitude.

Je remercie ma grand-mère, Jacqueline Delmas, qui nous a quittés il y a quelques mois, pour l'exemple qu'elle m'a montré ces dernières années : celui d'une force de volonté inébranlable et d'une soif de savoir inextinguible, que seule la maladie a finalement vaincues.

Au sein de l'équipe Plasmas Magnétisés du CPhT, j'ai bénéficié de conditions de travail excellentes. En quatre ans, j'ai pu découvrir ce qu'est la recherche scientifique sur un sujet de pointe, en travaillant avec des chercheurs dont les qualités scientifiques et humaines sont évidentes. Ce bref passage dans le monde académique a été une étape parfois difficile mais importante pour moi. Elle a été grandement facilitée par le soutien et la compréhension de mes deux directeurs de thèse, Hinrich Lütjens et Jean-François Luciani, que je remercie

pour leur grande gentillesse,
leur très grande compétence,
et leur très, très grande patience.

Résumé

Récemment, plusieurs codes MHD utilisés pour la simulation des modes macroscopiques dans les tokamaks ont opéré une transition vers des codes hybrides MHD-cinétiques. Ces codes ont pour but de coupler les effets des particules rapides, cinétique, sur la stabilité du plasma vu comme un fluide. L'effet des α de fusion à 3.5 MeV, qui assureront la majorité du chauffage dans ITER, sur les dents de scie, instabilité universelle des tokamaks, est un exemple d'application intéressant : la dynamique du cycle rampe/crash des dents de scie suscite encore des questions, et l'influence expérimentale des α est largement documentée.

Dans une première partie, nous présentons une étude numérique des régimes d'établissement des dents de scie dans un plasma ohmique avec le code XTOR-2F. Cette étude est à notre connaissance la première à explorer la dynamique au long terme ($\sim 10^6$ temps d'Alfvén) du kink interne (responsable du crash) en géométrie toroïdale complète. Utilisant la MHD avec transport anisotrope et résistivité ($\eta\chi$ MHD), nous trouvons deux régimes distincts. Le premier, à bas β_p ou bas $\chi_\perp$, est un régime d'oscillations stables. Le second est un régime stationnaire, avec un kink saturé et une déformation du cœur du plasma avec hélicité $m = n = 1$. Les dérives diamagnétiques, qui stabilisent le kink, décalent le régime oscillant vers des valeurs de pression plus hautes ; les oscillations deviennent comparables à des dents de scie expérimentales. Ainsi, avec des paramètres typiques de JET, la $\eta\chi$ MHD ne prévoit pas de dents de scie, mais les effets ω^* permettent de les retrouver.

Dans la seconde partie, nous détaillons les contributions faites au développement du code hybride XTOR-K. On a choisi un modèle cinétique « full-f, full-orbit » : on intègre exactement le mouvement de toutes les particules – par opposition aux modèles δf et gyrocinétiques utilisés en général. À chaque pas de temps, l'algorithme de couplage Newton-Krylov/Picard (MHD/cinétique) concentre les itérations sur la partie Newton-Krylov. La partie Picard, qui utilise un sous-pas de temps particulaire, est facilement parallélisée, et fait 3 itérations au plus. Le schéma reste stable vis-à-vis des modes MHD fondamentaux pour des pas de temps de l'ordre de τ_A et des pressions particulières importantes. L'avancée des particules est faite par l'algorithme de Boris, adapté en géométrie torique. Les invariants du mouvement ne dérivent pas numériquement, et sont bien conservés à partir de 4 sous-pas de temps par période cyclotronique. La méthode full-f entraîne un bruit important sur le tenseur de pression, qu'on essaye de réduire en utilisant une séquence de Sobol pour maximiser la régularité de l'espace des phases à l'initialisation. On utilise également un filtre passe-bas temporel (Butterworth) pour éliminer les fréquences de bruit élevées.

On montre en conclusion une première simulation de kink sur quelques milliers de τ_A avec le code hybride et 800000 particules. Après filtrage, le temps de calcul et le niveau de bruit sont assez faibles pour envisager des simulations plus précises avec jusqu'à 10^9 particules sur un grand nombre de processeurs dans un avenir proche.

Abstract

Some MHD codes used for simulating macroscopic modes in tokamaks have recently undergone a transition into hybrid MHD/kinetic codes. Their objective is to consider simultaneously the kinetic effects of fast particles and the fluid-modeled plasma stability. In that line of thought, using an hybrid code to explore the effects of the 3.5 MeV fusion α , which will be the main source of heating in ITER, on the sawtooth instability, which is a pervasive phenomenon in tokamak plasmas, would be a good test for such a code. The ramp/crash cycle dynamics of sawteeth is not fully understood, and the influence of fast particles on the internal kink is well known.

In the first part of this work, we present a numerical study of sawteeth regimes in an ohmic plasma with the XTOR-2F extended MHD code. To the extent of our knowledge, this study is the first to explore the long-term ($\sim 10^6 \tau_A$) dynamics of the internal kink (which is responsible for the crash) in full toroidal geometry. Using resistive MHD with anisotropic thermal transport ($\eta\chi$ MHD) yields two regimes. The first one, at low pressure or low perpendicular transport is a regime of stable oscillations. The second one is a stationary regime, with a saturated kink, and a permanent tridimensionnal kinking of the plasma core with $m = n = 1$ helicity. Diamagnetic drifts, which have a stabilizing effect on the kink, shift the oscillating regime towards higher pressure; the oscillations time scales get closer to those of experimental sawteeth. Thus, with parameters similar to those of an ohmic discharge on JET, $\eta\chi$ MHD predicts a saturated regime, but with ω^* effects, sawteeth are recovered.

In the second part, we explain our contributions to developing the hybrid code XTOR-K. We chose a "full-f, full-orbit" kinetic model, *i.e.* we compute the movement of all the particles exactly, contrary to the δf and gyrokinetic models that are generally used. During each time step, the Newton-Krylov/Picard (MHD/kinetic) coupling algorithm concentrates most iterations on the Newton-Krylov solver. The Picard part uses a particular time substep; it is easily parallelized and iterates at most 3 times. The scheme remains stable up to times steps of the order of τ_A and important particular pressure. Particles orbits integration is done via the Boris algorithm, adapted into toroidal geometry. There is no numerical drift on the constants of the movement, and these show good conservation down to 4 substeps per gyrotronic rotation. The full-f method creates an important noise on the pressure tensor, that we can try to limit by using a Sobol sequence to maximize the regularity of the initial distribution in phase space. We also use a low-pass numerical Butterworth filter to eliminate the high noise frequencies.

Finally, we show a first hybrid simulation of a kink on a few thousand τ_A with 800000 particles. After filtering, computing time and noise level are low enough to consider running more precise simulations with up to 10^9 particles on a larger number of processors in the close future.

1 Introduction

1.1 Fusion nucléaire contrôlée et critère de Lawson

La maîtrise de la fusion nucléaire comme source d'énergie est un enjeu énergétique majeur. Seule source d'énergie concentrée pouvant remplacer la fission et les hydrocarbures en termes de quantité d'énergie disponible, elle présente encore les avantages d'être « propre » (pas d'émission de gaz à effet de serre, peu de produits radioactifs¹) et quasi-inépuisable.

Ces atouts indéniables sont contre-balancés par un obstacle de taille : les sections efficaces optimales des réactions de fusion maximaux sont situées à très haute température, vers 100 keV (1 eV = 11600 K). La réaction ayant le taux de réaction de plus grand est la réaction deutérium - tritium :



La température visée dans un réacteur est de 10 keV (seules les particules les plus rapides fusionnent). Dans ces conditions, les atomes sont tous complètement ionisés, et les électrons sont libres. Les réactifs forment donc un plasma.

Pour obtenir un état stationnaire, où la fusion est entretenue dans le réacteur, la densité et la température du plasma doivent rester élevées. Le problème qui se pose alors est celui du confinement. Un confinement classique, par exemple dans une enceinte matérielle, entraînerait des pertes de particules et de chaleur rédhibitoires. Les conditions nécessaires à la fusion seraient perdues aussitôt². Il faut donc réussir à confiner le plasma dans un espace restreint, en minimisant tout contact avec les parois du réacteur pour limiter les pertes.

Le problème peut être contourné en envisageant un réacteur qui ne fonctionnerait pas en régime stationnaire, mais dans lequel l'énergie proviendrait de réactions successives régulières. C'est la voie du confinement inertiel. Nous nous concentrons sur la démarche opposée, le confinement magnétique, qui consiste à établir les conditions d'une réaction de fusion permanente, idéalement auto-entretenu. L'énergie des particules issues de la fusion doit compenser les pertes et maintenir la température du plasma. Une fois la réaction lancée, le plasma n'a alors plus besoin d'apport d'énergie. Pour quantifier cette condition, Lawson, en 1957, a introduit la notion de

¹Remarquons tout de même que si la réaction de fusion elle-même ne produit que des éléments stables, les particules énergétiques ($\sim$ MeV) qui en sont issues activeront une partie des matériaux entourant la réaction. La quantité de matériaux radioactifs produits sera toutefois beaucoup plus faible que pour la fission.

²La difficulté à maintenir les conditions de fusion est cependant un avantage de taille pour la sécurité d'un réacteur. En effet, la moindre perturbation du système entraînant la perte des conditions de fusion et l'arrêt de la réaction, il n'y a aucun risque d'emballement.

temps de confinement de l'énergie τ_E . Ce paramètre empirique mesure la vitesse de refroidissement du système :

$$\tau_E = \frac{W}{P_{loss}}$$

Où W est le contenu énergétique du système, et P_{loss} est la puissance perdue. Le paramètre τ_E quantifie la qualité du confinement. En supposant que le plasma est un mélange 50-50 de deutérium-tritium (sans cendres ni impuretés) à la température T , l'énergie W s'écrit $W = 3n_e k_b T$. Le critère de Lawson énonce que la réaction est rentable si et seulement si le chauffage dû aux noyaux d'hélium (α) issus de la fusion compense au moins les pertes. Les neutrons quittent le milieu et c'est leur énergie qui est récupérée en sortie du réacteur. Le chauffage par les α s'écrit : $P_{ch} = n_D n_T \langle \sigma v \rangle E_\alpha = 1/4 n_e \langle \sigma v \rangle E_\alpha$. n_x est la densité de l'espèce x , E_α l'énergie des α de fusion, k_b la constante de Boltzmann, et $\langle \sigma v \rangle$ le taux de réaction. Finalement, le critère de Lawson s'écrit :

$$n_e \tau_E \geq \frac{12}{E_\alpha} \frac{k_b T}{\langle \sigma v \rangle}$$

En prenant en compte la dépendance de $\langle \sigma v \rangle$ en fonction de T , on exprime souvent ce critère en fonction du triple produit $n_e T \tau_E$. En injectant les valeurs numériques correspondant à la réaction choisie, on obtient :

$$n_e T \tau_E \geq 10^{21} \text{ keV s.m}^{-3}$$

Il s'agit d'un critère empirique et opératoire, qui fournit une limite inférieure de fonctionnement pour un réacteur de fusion.

1.2 Confinement magnétique et tokamaks

Un champ magnétique est un bon candidat pour assurer le confinement matériel du plasma. En effet, la trajectoire hélicoïdale d'une particule chargée autour d'une ligne de champ magnétique la maintient à une distance ρ_L de cette ligne de champ (ρ_L est le rayon de Larmor). Il suffit donc que les lignes de champs occupent une région bornée de l'espace pour que les particules du plasma soient effectivement confinées dans cette région. On peut montrer que la frontière d'une telle configuration doit être topologiquement équivalente à un tore. Cependant, le champ créé dans une telle configuration a nécessairement un gradient et une courbure, qui font dériver les particules et détruisent le confinement. Cet effet est compensé si les particules transitent non seulement autour de l'axe de révolution du tore (dans la direction

toroïdale) mais aussi autour de l'axe magnétique³ (dans la direction *poloïdale*). Dans cette configuration, les lignes de champ s'enroulent à la façon d'hélices autour de surfaces toriques.

Ce genre de configuration a été utilisée pour la première fois à la fin des années 50. Ce sont deux chercheurs soviétiques, Tamm et Sakharov, qui ont proposé le principe de la machine, appelée tokamak⁴. Dans un tokamak, le champ magnétique a une composante toroïdale, créée par des bobines qui forment un solénoïde refermé sur lui-même, et une composante poloïdale, due au courant toroïdal induit dans le plasma. En 1968, les soviétiques annoncent qu'ils ont atteint des températures électroniques de 1 keV dans le tokamak T3. La nouvelle est accueillie avec scepticisme par les chercheurs occidentaux, qui sont loin de tels résultats. Une équipe de chercheurs britannique, envoyée en URSS, confirme le résultat quelques années plus tard. À partir de moment, les tokamaks deviennent les machines de référence pour la fusion contrôlée par confinement magnétique⁵. La découverte d'une corrélation entre taille du tore et temps de confinement a mené à des projets de plus en plus grands et de plus en plus nombreux.

La génération de tokamaks des années 80 a permis d'approcher le critère de Lawson de façon significative. Dans ces machines, on a mis en œuvre des techniques de chauffage auxiliaires, par opposition au simple chauffage ohmique dû au courant toroïdal : injection de neutres rapides, résonances ondes-particules à fréquence cyclotronique ionique, électronique ou hybride. JET, en Grande-Bretagne, a atteint des rendements⁶ proches de l'équilibre énergétique ($Q = 0.7$ en 1997 ; l'équilibre, ou break-even, correspond à $Q = 1$, et on estime que $Q = 30$ serait nécessaire pour un réacteur). JT-60, au Japon, détient le record du rapport $nT\tau_E$ le plus élevé. TORE SUPRA, en France, a permis de confiner un plasma moins énergétique pendant des durées de temps longues (6 minutes), grâce à des bobines magnétiques supraconductrices. On peut également citer TFTR aux USA, et ASDEX Upgrade en Allemagne. L'expertise accumulée grâce à ces machines et à plusieurs dizaines d'années de recherche théorique a conduit au projet international de réacteur expérimental ITER, actuellement en construction à Cadarache (Var). ITER devrait démontrer la faisabilité industrielle de la fusion en atteignant des rendements élevés ($Q = 10$) pendant environ 400 secondes.

De nombreuses difficultés subsistent néanmoins. La stabilité globale du plasma impose des limitations et un contrôle permanent des paramètres macroscopiques (n, T) pour empêcher des instabilités de se développer et de détruire le confinement.

³*i.e.*, la particule passe alternativement côté « intérieur » et « extérieur » du tore.

⁴« tokamak » est une contraction du nom russe de la machine : *тороидальная камера с магнитными катушками*, littéralement « chambre toroïdale à bobines magnétiques ».

⁵D'autres machines, comme les stellarators ou les sphéromaks, existent également, mais la majorité des recherches s'oriente dans la direction des tokamaks.

⁶Le rendement Q est défini comme le rapport de la puissance de fusion sur la puissance consommée par la machine.

De l'autre côté du spectre, le transport de chaleur, qui contrôle le paramètre τ_E , est dû aux micro-instabilités et aux collisions. Sa modélisation implique d'étudier la turbulence du plasma. L'interaction plasma-paroi constitue un troisième problème : la conception de la couverture (*i.e.*, les composants faisant face au plasma) est un défi de taille, à cause du flux permanent de neutrons très énergétiques dirigés vers elle. Il est crucial d'éviter que le plasma soit pollué par des impuretés venant de la couverture. Ces composants doivent également régénérer le tritium nécessaire à la réaction. Une autre difficulté réside dans la nécessité de générer des courants importants (de l'ordre du MA) en régime continu. Enfin, l'influence des particules chaudes issues de la fusion est une problématique de plus en plus présente. Des particules chaudes créées par chauffage (résonance cyclotronique ionique ou injection de neutres) sont déjà présentes dans la plupart des machines, et utilisées pour le contrôle des profils de courant et de chauffage. Les particules α issues de la fusion sont destinées à être la principale source de chauffage du plasma dans ITER, et leur confinement est un enjeu crucial.

1.3 Modélisation multi-modèles et codes hybrides

Les plasmas de tokamaks constituent des systèmes physiques très riches, qui mélangent électromagnétisme et physique des fluides. Les échelles caractéristiques recouvrent un large domaine, et exigent de combiner plusieurs théories physiques pour arriver à une modélisation intégrée.

Ainsi, la question de la stabilité macroscopique appelle logiquement une théorie fluide. La magnétohydrodynamique (MHD) est une théorie mono-fluide qui a été beaucoup développée pour la physique des tokamaks, et qui a prouvé son efficacité⁷. Elle décrit en principe les phénomènes collectifs rapides, sur des temps de l'ordre de τ_A , le temps d'Alfvén⁸, et des longueurs de l'ordre de a , le petit rayon du tore⁹. La MHD idéale fournit des relations de dispersion pour un plasma homogène parfaitement conducteur, et permet donc de caractériser les ondes qui peuvent s'y propager. En précisant la topologie magnétique de la configuration, on peut calculer la variation d'énergie potentielle due à une petite perturbation. Un principe d'énergie prévoit alors quelles perturbations vont déboucher sur des instabilités. La dynamique et le développement non-linéaire de ces instabilités exigent des calculs plus lourds, mais comme le développement d'une instabilité idéale¹⁰ entraîne généralement la perte du confinement, le seuil de stabilité est déjà une donnée cruciale.

La MHD fournit donc des critères suffisants d'instabilité, *i.e.* des critères nécessaires

⁷Ce en dépit du fait que les plasmas de tokamaks sont *a priori* hors du domaine de validité de la MHD, comme on le verra.

⁸typiquement, $\tau_A \sim 10^{-6}$ s dans nos cas.

⁹typiquement, 1 m dans nos cas.

¹⁰*i.e.*, pour un fluide parfaitement conducteur

de stabilité. Ces critères fixent les limites autorisées pour les paramètres macroscopiques du système (par exemple, le n du critère de Lawson). L'instabilité la plus violente prévue par la MHD est le kink externe, qui se manifeste par un déplacement hélicoïdal du plasma jusqu'à sa frontière. Les expériences sont systématiquement conçues dans les limites dictées par la MHD pour éviter cette instabilité. On trouve également des instabilités comme les kinks internes, qui n'affectent pas la frontière du plasma, ou les instabilités d'interchange ou de ballooning, plus localisées.

La robustesse des critères MHD a été éprouvée au cours des cinquante dernières années. Cependant, les phénomènes non MHD sont également cruciaux pour le confinement, en particulier parce que la physique implique que les expériences restent dans des domaines proches des frontières de stabilité MHD. Les instabilités sont alors influencées par des effets exclus de la MHD idéale. Certains peuvent être inclus dans le modèle pour passer à une MHD étendue : effets bi-fluides, résistifs, de Hall. D'autres, comme les effets dits cinétiques, proviennent d'un comportement non fluide des particules du plasma, et dépassent ce cadre.

Ainsi, les particules chaudes (créées par la fusion, le chauffage par injection de neutres ou résonance cyclotronique ionique) ne sont pas modélisables par une théorie fluide. À haute température, l'excursion radiale des particules devient comparable aux échelles de variation du champ magnétique, ce qui implique l'abandon de l'hypothèse d'une pression scalaire, et l'échec des modèles de type MHD.

Les problématiques physiques liées aux particules chaudes sont très importantes. Les particules chaudes doivent être confinées pour assurer le chauffage du plasma et l'auto-entretien de la fusion. Elles servent également à contrôler la stabilité du plasma par l'intermédiaire des profils de courant et de pression. Elles soulèvent des problèmes comme l'extraction des cendres. Leur interaction avec le « bulk », le cœur thermique du plasma, crée de nouveaux modes (fishbones, TAEs), modifie la stabilité des modes pré-existants, et peut entraîner des pertes fatales au confinement ou dangereuses pour la couverture.

Le traitement des effets cinétiques exige de s'intéresser directement à la dynamique de la fonction de distribution des particules, décrite par l'équation de Vlasov. Les grandeurs de type fluide, plus parlantes, s'expriment comme moments de la fonction de distribution. Les théories cinétiques donnent accès aux effets microscopiques (collisions, turbulence) qui régissent le transport. Elles sont plus riches mais aussi nettement plus complexes que les théories fluides, et plus exigeantes en ressources numériques, ce qui explique leur faible utilisation jusqu'au dix dernières années. Le mouvement peut être décrit avec différents degrés de précision : par des vitesses de dérives plus ou moins complexes (théorie drift-cinétique), en résolvant l'équation du mouvement des particules moyennée sur leur période cyclotronique (théorie gyrocinétique). Dans ces cas-là, on ne s'intéresse qu'au mouvement du centre-guide des particules. On peut aussi intégrer l'équation du mouvement complète, ce qui inclut automatiquement les effets de grand rayon de Larmor (pour des particules chaudes,

ce rayon atteint plusieurs % du petit rayon).

Plutôt que d'essayer de modéliser tout le plasma par une théorie unique, qui sera « trop » complète pour une partie importante du plasma, il semble judicieux de combiner plusieurs modèles, chacun adapté à une fraction du plasma. Durant la dernière quinzaine d'années, certains codes MHD ont évolué pour intégrer des théories cinétiques : on parle alors de codes hybrides. Un code hybride est caractérisé par le modèle cinétique choisi, et par le couplage entre partie fluide et partie cinétique, qui doit assurer des résultats self-consistants. Les informations fournies par ces codes sont indispensables pour confirmer la viabilité des scénarios de fonctionnement d'ITER.

Dans ce travail, on s'intéresse à la stabilité macroscopique d'un plasma de tokamak, et plus particulièrement à l'instabilité dite de dents de scie, au moyen du code MHD étendue XTOR. On décrit également notre contribution au développement d'une extension cinétique ajoutée à ce code, dans le but d'explorer l'influence des particules rapides sur la stabilité globale du plasma en général, et sur les dents de scie en particulier.

1.4 Présentation du travail de thèse

L'équipe du CPhT dispose depuis longtemps d'un code MHD étendue dynamique pour les tokamaks : XTOR. XTOR est un code fortement modulaire : l'utilisation d'un solveur implicite permet d'élargir facilement le modèle MHD de base (qui inclut transport thermique et résistivité) pour inclure des effets bi-fluides variés. Pour atteindre une modélisation encore plus complète, le code est également en train d'évoluer vers un code hybride MHD-particules. Le but est de pouvoir inclure ou exclure les effets choisis, y compris les particules chaudes, de façon à comprendre leurs influences sur la dynamique du plasma. Lors de cette thèse, on a mené deux travaux distincts : d'une part, on a initié un travail de fond visant à caractériser les régimes susceptibles de donner lieu à des dents de scie dans le cadre de la MHD étendue ; d'autre part, on a contribué à la mise en place d'outils nécessaires à l'évolution d'XTOR vers le code hybride XTOR-K (la transformation du code et sa validation sont des processus de longue haleine, qui s'étendent sur plusieurs années).

Le chapitre 2 est consacré à la présentation du cadre théorique de la MHD. On commence par rappeler les principales caractéristiques de la MHD : dérivation des équations, hypothèses et domaine de validité, principe d'énergie. On présente ensuite la configuration des tokamaks, avec leur ordering et les paramètres clés pour la MHD. On détaille les principaux modes prédits par la MHD, et on questionne la validité de la MHD dans le cas des tokamaks. Enfin, on s'intéresse ensuite plus en détail à la modélisation des dents de scie en MHD, en vue de leur simulation avec XTOR. On présente le phénomène, ainsi que les principales théories avancées pour son explication. Leur confrontation aux expériences a conduit à chercher des modèles plus complets ; à l'heure actuelle, de nombreuses questions restent en suspens.

Dans le chapitre 3, on utilise **XTOR** pour établir un cadre de référence pour la dynamique au long terme des kinks internes, qui jouent un rôle majeur dans les dents de scie. À l’heure actuelle, le modèle de dents de scie le plus utilisé procède plutôt d’une approche semi-empirique visant à établir un scaling de leur période et de leur amplitude. Il n’existait pas auparavant à notre connaissance d’étude MHD en géométrie complète sur des temps de l’ordre du temps résistif (plusieurs dizaines de cycles). On commence par présenter en détail le code **XTOR**. On précise ensuite le cadre de l’étude : domaine de validité et objectif. On utilise d’abord par le modèle le plus simple, celui de la MHD résistive avec transport anisotrope (noté $\eta\chi$ MHD). Ce modèle correspond à des plasmas ohmique de petits tokamaks, comme ceux dans lesquels on a observé les premières dents de scie. Les paramètres principaux sont η , χ , et β_{p1} . Le résultat le plus important de cette partie est d’exhiber une transition entre deux régimes : un régime oscillant, et un régime d’instabilité héliçoïdale saturée. Postérieurement à cette thèse, le cadre $\eta\chi$ MHD présenté ici a pu être étendu aux effets diamagnétiques (cf. [49]), et on a constaté que la transition perdurait. Le régime oscillant évolue pour présenter des temps caractéristiques proches de ceux des dents de scie.

Le chapitre 4 présente la structure du code hybride **XTOR-K**. On détaille la physique que recouvre la partie cinétique du code, et on rappelle le principe de fonctionnement d’un code hybride. On explique le modèle choisi, avec intégration de l’équation du mouvement complète particule par particule. Ce choix est original, par rapport aux méthodes de type δf généralement utilisées. Dans une méthode δf , on fait l’hypothèse que seule une faible fraction (δf) de la distribution des particules varie. Cela présente l’avantage de réduire le bruit sur la pression (puisqu’il est proportionnel à δf) ; en revanche, l’hypothèse peut être mise à mal si le système s’éloigne beaucoup de l’équilibre de départ. Notre choix (« full-f ») entraîne l’apparition d’un sous-pas de temps particulaire, qui accroît le temps de calcul de la partie cinétique. Cette partie est cependant facilement parallélisée. En contrepartie, le tenseur de pression est exact et simple à évaluer, bien que fortement bruité. On étudie ensuite l’algorithme de couplage, reposant sur une méthode Newton-Krylov/Picard. La stabilité linéaire du schéma en fonction de la fraction particulaire et du pas de temps est étudiée pour les trois modes MHD fondamentaux, par un calcul semi-qualitatif, et illustrée par des tests numériques. Ce bon comportement permet de faire reposer la majorité des itérations sur le solveur implicite de la partie MHD, qui est très efficace ; ainsi, la pénalisation due à l’apparition du sous-pas de temps particulaire est minimisée.

Le chapitre 5 concerne le pousseur de particules, c’est-à-dire l’algorithme qui intègre les trajectoires des particules dans le champ électromagnétique. Tout d’abord, on expose le schéma de type leap-frog choisi : il s’agit de la méthode de Boris, adaptée en géométrie torique. Cette méthode possède de bonnes qualité de stabilité et précision, évitant notamment l’accumulation d’erreur numérique sur les invariants du

mouvement (*i.e.* le moment magnétique, l'énergie, et le moment cinétique toroïdal). On se concentre par la suite sur les détails de l'implémentation dans XTOR-K : normalisation, paramètres, interpolations et changements de maillages. Enfin, on explore les caractéristiques de fonctionnement du poussoir, en étudiant la conservation des invariants du mouvement.

Le chapitre 6 expose les stratégies envisagées pour minimiser le bruit numérique sur le tenseur de pression des particules. À cause de l'absence d'hypothèse de type δf , notre tenseur a en effet un bruit inhérent important qui doit impérativement être minimisé. Si le chargement des particules est aléatoire, le bruit doit logiquement varier en $N_{part}^{-1/2}$, N_{part} étant le nombre de particules numériques. Pour accélérer cette convergence et diminuer le bruit, on peut optimiser ce chargement ; on peut également appliquer des filtres numériques sur le tenseur de pression pour éliminer certaines composantes du bruit. Après avoir rappelé la définition du tenseur, on détaille ensuite une stratégie d'initialisation anisotrope visant à régulariser au maximum la distribution initiale des particules dans l'espace des phases, en utilisant des séquences à faible discrédance (LDS). On met en place un filtre temporel passe-bas (filtre numérique de Butterworth) dans le but de compenser la perte d'information sur la distribution initiale au cours du mouvement, et d'éliminer les fréquences non pertinentes. Les spécifications du filtre sont décidées en fonction de la structure temporelle du signal de pression.

Dans la continuité de la première partie de la thèse, la simulation du kink interne en présence de particules chaudes constituerait un test de choix pour le code hybride. Les effets des particules chaudes sur les dents de scie sont à la fois importants et bien documentés. Une telle simulation joindrait les deux parties de notre travail. Nous fournissons à la fin du dernier chapitre des résultats préliminaires, qui montrent que l'état actuel du code permet de l'envisager dans les mois qui viennent.

Première partie

Régimes d'établissement des dents de scie en plasmas ohmiques

2 La MHD appliquée aux tokamaks : cas des dents de scie

Dans ce chapitre, nous précisons le contexte théorique et numérique de l'étude des dents de scie que nous avons menée avec **XTOR**.

Pour ce faire, nous commençons par les généralités nécessaires sur la MHD : dérivation du système (modèle fluide et électromagnétisme), domaine de validité (grande collisionnalité et échelles grandes devant celles du mouvement cyclotronique), et principe d'énergie pour obtenir les conditions nécessaires de stabilité d'un plasma. On expose ensuite les principales caractéristiques des tokamaks vus sous l'angle de la MHD : ordres de grandeurs, paramètres importants, instabilités MHD et critères nécessaires de stabilité associés. On remarque et on discute le fait que la MHD fournit des critères qui se révèlent efficaces malgré la faible collisionnalité des tokamaks. On présente ensuite le phénomène des dents de scie en détail, et on rappelle quels outils théoriques la MHD propose pour aborder ce problème. Au-delà de la MHD résistive, un consensus existe sur l'importance d'effets bi-fluides et non-linéaire, pour répondre aux questions encore ouvertes, comme le déclenchement rapide ou la dynamique de la rampe. On rappelle également la diversité des résultats expérimentaux, et on discute des perspectives pour une étude numérique des dents de scie en conclusion.

2.1 Généralités sur la MHD

La magnétohydrodynamique (MHD) est une théorie mono-fluide décrivant un fluide chargé. Elle a été développée à partir de la première moitié du XX^{ème} siècle. Ses fondements - et son nom - sont dus à l'astrophysicien suédois Hannes Alfvén. Il a notamment étudié les ondes électromagnétiques se propageant dans un plasma, et énoncé le théorème du gel en MHD idéale. À partir de la fin des années 40, les recherches se multiplient dans le domaine de la fusion contrôlée ; la MHD, initialement pensée pour l'étude d'objets astrophysiques (comme les champs magnétiques terrestre et solaire), est rapidement appliquée aux plasmas de fusion, qui lui fournissent un cadre expérimental. Elle prouve rapidement son efficacité et sa richesse, ce qui explique qu'elle soit encore aujourd'hui une théorie incontournable dans l'étude des plasmas de tokamak.

En tant que théorie fluide, la MHD présente l'avantage d'une relative simplicité, comme nous le constaterons en dérivant les équations du système à partir des équations fluides qui régissent un fluide chargé. Cette simplification est bien sûr due à des hypothèses physiques que nous détaillerons. Ces hypothèses définissent le domaine de validité de la MHD.

2.1.1 Dérivation du modèle : de Vlasov au mono-fluide

Le système de base pour modéliser un plasma est constitué de deux équations de Vlasov, qui décrivent la dynamique des fonctions de distribution ionique et électronique. Dans cette partie, les calculs sont en SI et les notations standards : $\mathbf{E}$ et $\mathbf{B}$ sont les champs électrique et magnétique, $\mathbf{J}$ la densité de courant, f la fonction de distribution, $\mathbf{r}$, $\mathbf{u}$, et t les variables d'espace, de vitesse et de temps. q est la charge, m la masse ; l'indice s dénote l'espèce, ions ou électrons. La démarche qui suit est présentée plus en détail dans de nombreux ouvrages, par exemple par Freidberg, dans [1], chap. 2, ou Hazeltine et Meiss, dans [2], chap. 6 :

$$\frac{\partial f_s}{\partial t} + \mathbf{u} \cdot \nabla f_s + \frac{q_s}{m_s} (\mathbf{E} + \mathbf{u} \times \mathbf{B}) \cdot \nabla_{\mathbf{u}} f_s = \left(\frac{\partial f_s}{\partial t} \right)_c \quad (1)$$

L'opérateur $\left(\frac{\partial f_s}{\partial t} \right)_c$ est un terme collisionnel. On ajoute les équations de Maxwell¹¹ :

$$\nabla \cdot \mathbf{E} = \frac{\sigma}{\epsilon_0} \quad (2)$$

$$\nabla \times \mathbf{E} = -\frac{\partial \mathbf{B}}{\partial t} \quad (3)$$

$$\nabla \cdot \mathbf{B} = 0 \quad (4)$$

$$\nabla \times \mathbf{B} = \epsilon_0 \mu_0 \frac{\partial \mathbf{E}}{\partial t} + \mu_0 \mathbf{J} \quad (5)$$

Pour exprimer les équations de Maxwell en fonction de nos variables, on a posé :

$$\begin{aligned} \sigma &= n_i - n_e \\ \mathbf{J} &= e(n_i \mathbf{v}_i - n_e \mathbf{v}_e) \end{aligned}$$

Les densités n_s sont définies juste en-dessous. Pour simplifier, on considère dans la suite de ce chapitre que $q_i = -q_e = e$. Nous avons donc un système où les inconnues sont les fonctions f_e et f_i , qui dépendent de sept variables $(\mathbf{u}, \mathbf{r}, t)$.

Passage à un système fluide

Nous allons prendre les moments des équations cinétiques pour arriver à un modèle bi-fluide (électrons et ions). Une explication très complète de ce calcul a été fournie par Braginskii en 1965 (cf. [3]).

Le moment d'ordre $k > 0$ d'une fonction de distribution est défini par :

¹¹Notons qu'on utilise les équations classiques ; on se situe donc hors du domaine relativiste.

$$\mathbf{M}_k \equiv \frac{1}{n_s} \int \mathbf{u}^k f_s d\mathbf{u}$$

Le moment d'ordre 0 est la densité n_s (l'intégrale de f_s sur la variable $\mathbf{u}$), celui d'ordre 1 la vitesse fluide $\mathbf{v}_s$, celui d'ordre 2 le tenseur de pression $\mathbf{P}_s$.

Le moment d'ordre 0 des équations de Vlasov donne l'équation de la conservation de la quantité de matière pour chaque espèce :

$$\frac{\partial n_s}{\partial t} + \nabla \cdot n_s \mathbf{v}_s = 0 \quad (6)$$

Le moment d'ordre 1 correspond à la conservation de la quantité de mouvement :

$$\rho_s \frac{\partial \mathbf{v}_s}{\partial t} + \rho_s (\mathbf{v}_s \cdot \nabla) \mathbf{v}_s = q_s n_s (\mathbf{E} + \mathbf{v}_s \times \mathbf{B}) - \nabla \cdot \mathbf{P}_s + \mathbf{F}_s \quad (7)$$

Où on a introduit $\rho_s = m_s n_s$ et $\mathbf{F}_s$, la quantité de mouvement reçue par collisions.

Le moment d'ordre 2 correspond à la conservation de l'énergie ; nous l'écrivons sous la forme d'une équation sur p_s , la pression scalaire (partie isotrope du tenseur de pression $\mathbf{P}_s$), avec l'intention de négliger les termes anisotropes rapidement) :

$$\frac{3}{2} \left(\frac{dp_s}{dt} \right)_s + \frac{5}{2} p_s \nabla \cdot \mathbf{v}_s = -(\mathbf{\Pi}_s : \nabla \mathbf{v}_s) - \nabla \cdot \mathbf{Q}_s + H_s \quad (8)$$

Dans cette équation, $\mathbf{\Pi}_s$ est la partie anisotrope du tenseur de pression, $\mathbf{Q}_s$ est le flux de chaleur dû au mouvement thermique des particules, et H_s est la chaleur créée par les collisions avec les autres particules. On a introduit l'opérateur de dérivation $(d/dt)_s \equiv \partial/\partial t + \mathbf{v}_s \cdot \nabla$. Les équations 2 - 8 forment alors un système bi-fluide pour décrire le plasma¹². Elles peuvent être retrouvées par exemple dans [2], chap. 6.

Fermeture mono-fluide et hypothèse adiabatique

Notre système n'est pas fermé : il contient plus d'inconnues que d'équations. En effet, l'équation de la dynamique du moment d'ordre k fait toujours intervenir le moment d'ordre $k + 1$. Pour clore le système, il faut faire une hypothèse sur un moment ; dans notre cas, le flux de chaleur (on fera également des simplifications sur la pression). Dans la version hybride, c'est le le tenseur de pression calculé à partir de la dynamique des particules par la partie cinétique du code qui clôt le système. Dans un modèle purement fluide, on utilise souvent l'hypothèse adiabatique, la plus simple. Cette hypothèse nous permet d'exprimer le système sous sa forme mono-fluide classique, en définissant :

¹²On rappelle que les équations de Maxwell ne sont pas indépendantes ; elles ne sont pas toutes nécessaires à la résolution du système

$$\begin{aligned}
n &= n_i = n_e \\
\rho &= m_i n_i + m_e n_e \approx m_i n_i \\
\mathbf{v} &= \frac{m_i n_i \mathbf{v}_i + m_e n_e \mathbf{v}_e}{m_i n_i + m_e n_e} \approx \mathbf{v}_i \\
\mathbf{J} &= e(n_i \mathbf{v}_i - n_e \mathbf{v}_e) \\
p &= p_e + p_i = nT \\
T &= T_e + T_i
\end{aligned}$$

En plus des définitions, on a fait trois hypothèses : la quasi-neutralité, $n_i = n_e$, l'égalité des pressions scalaires, $p_e = p_i$, et la nullité de la masse des électrons (en fait, $m_e \ll m_i$). Le système se réduit alors à un système mono-fluide, équation par équation.

La conservation de la densité ionique, multipliée par m_i , donne :

$$\frac{\partial \rho}{\partial t} + \nabla \cdot (\rho \mathbf{v}) = 0 \quad (9)$$

En multipliant les deux équations 6 par les charges et en soustrayant, on obtient :

$$\nabla \cdot \mathbf{J} = 0 \quad (10)$$

En additionnant les équations 7 (les $\mathbf{F}_s$ s'annulent et on néglige les $\mathbf{\Pi}_s$), on obtient :

$$\rho \frac{d\mathbf{v}}{dt} = -\nabla p + \mathbf{J} \times \mathbf{B} \quad (11)$$

En ré-écrivant l'équation 7 pour les électrons, avec $\mathbf{F}_e = m_e n \nu (\mathbf{v}_e - \mathbf{v}_i)$, et ν la fréquence de collisions électrons-ions, on obtient :

$$\mathbf{E} = -\mathbf{v}_e \times \mathbf{B} + \frac{\nabla p_e}{en} + \frac{\nabla \cdot \mathbf{\Pi}_e}{en} + \frac{\nu m_e}{e} (\mathbf{v}_i - \mathbf{v}_e) + \frac{m_e}{e} \frac{d\mathbf{v}_e}{dt}$$

En se souvenant que $\mathbf{v}_e = \mathbf{v} - \mathbf{J}/(en)$ et en substituant, on obtient la loi d'Ohm généralisée :

$$\mathbf{E} = -\mathbf{v} \times \mathbf{B} + \frac{\mathbf{J} \times \mathbf{B}}{en} + \frac{\nabla p_e}{en} + \frac{\nabla \cdot \mathbf{\Pi}_e}{en} + \underbrace{\frac{\nu m_e}{ne^2}}_{\equiv \eta} \mathbf{J} + \frac{m_e}{e} \frac{d\mathbf{v}_e}{dt} \quad (12)$$

Dans le cadre de la MHD idéale, on néglige les termes jusqu'à obtenir :

$$\mathbf{E} = -\mathbf{v} \times \mathbf{B} \quad (13)$$

Dans le cas de la MHD résistive, on garde en plus le terme en η , la résistivité :

$$\mathbf{E} = -\mathbf{v} \times \mathbf{B} + \eta \mathbf{J} \quad (14)$$

Enfin, en additionnant les équations 8, en négligeant les termes anisotropes des tenseurs, et en prenant les termes de source et de flux de chaleur (H_s et $\nabla \cdot \mathbf{Q}_s$) nuls, on obtient pour la pression :

$$\frac{dp}{dt} = -\Gamma p \nabla \cdot \mathbf{v} \quad (15)$$

Avec l'équation 9, cette dernière équation est équivalente à l'hypothèse adiabatique classique :

$$\frac{d}{dt}(p\rho^{-\Gamma}) = 0 \quad (16)$$

Où Γ est le coefficient adiabatique. On a éliminé le moment d'ordre supérieur ($\mathbf{Q}_s$), et donc fermé le système. Cette fermeture permet d'aboutir à un système mono-fluide simple. Dans le code MHD XTOR, l'option de fermeture de base est similaire, mais elle inclut en plus des termes de transport anisotrope sur la pression (qui modélisent $\nabla \cdot \mathbf{Q}_s$). Nous verrons dans le chapitre suivant qu'il existe d'autres fermetures, qui découlent d'hypothèses plus fines sur le flux de chaleur.

Pour arriver au système MHD classique, il reste à ajouter les équations de Maxwell, réduites à :

$$\nabla \times \mathbf{E} = -\frac{\partial \mathbf{B}}{\partial t} \quad (17)$$

$$\nabla \cdot \mathbf{B} = 0 \quad (18)$$

$$\nabla \times \mathbf{B} = \mu_0 \mathbf{J} \quad (19)$$

Les équations 9, 11, 13 (resp. 14), 15, 17 et 19 forment le modèle de la MHD idéale (resp. résistive). Nous allons maintenant récapituler les différentes hypothèses faites pour arriver à ce système.

2.1.2 Hypothèses et domaine de validité

Les hypothèses que nous avons effectuées se traduisent par des contraintes sur les temps et les longueurs caractéristiques des phénomènes que la MHD peut décrire. Grossièrement, les phénomènes qu'on veut pouvoir décrire ont le temps et la longueur caractéristiques :

$$\begin{aligned}\tau &\sim \frac{a}{v_{thi}} \\ \lambda &\sim a\end{aligned}$$

Avec v_{thi} la vitesse thermique des ions, et a une dimension caractéristique du système (pour un tokamak, typiquement le petit rayon).

1. Depuis le début nous utilisons une théorie fluide. Nous avons donc implicitement fait l'hypothèse que les fonctions de distribution étaient maxwelliennes, *i.e.* que les dynamiques des ions et des électrons étaient dominées par les collisions. Cela explique pourquoi le tenseur de pression $\mathbf{P}$ se réduit à une pression scalaire p : les collisions le rendent rapidement isotrope. Cette condition implique que nos temps caractéristiques soient grands devant les temps de collisions. On note τ_{sp} le temps de collisions des espèces s et p , et on rappelle que $\tau_{ee} \sim \tau_{ie}$ et que $\tau_{ee} \ll \tau_{ii}$ à températures égales. La contrainte de temps devient donc : $\tau \gg \tau_{ii}$, donc $v_{thi}\tau_{ii}/a \ll 1$. Nos longueurs caractéristiques doivent aussi être grandes devant les libres parcours moyens (définis par $\lambda_s \equiv v_{ths}\tau_{ss}$). La contrainte la plus restrictive est : $v_{Te}\tau_{ee}/a \ll 1$.
2. Dans l'équation de Maxwell-Ampère, on a négligé le courant de déplacement $\epsilon_0\mu_0\partial_t\mathbf{E}$. Cela revient à prendre des vitesses de phase et des vitesses thermiques petites devant c : $\omega/k \ll c$, $v_{the} \ll c$, $v_{thi} \ll c$.
3. On a posé $\epsilon_0\nabla \cdot \mathbf{E}/(en) \ll 1$, qui implique la quasi-neutralité¹³ $n_i = n_e$, grâce à l'équation de Maxwell-Gauss. Cette hypothèse revient à supposer que les électrons compensent instantanément les séparations de charges macroscopiques pour maintenir la neutralité du plasma. Autrement dit, nos longueurs caractéristiques doivent être grandes devant la longueur de Debye électronique : $\lambda \gg [(\epsilon_0 k_B T_e)/(ne^2)]^{1/2} \equiv \lambda_{de}$, et nos temps caractéristiques doivent être grands devant l'inverse de la pulsation plasma électronique : $\tau \gg [(m_e e^2)/(n\epsilon_0)]^{1/2} \equiv \omega_{pe}^{-1}$.
4. Dans la loi d'Ohm, on a négligé le terme d'inertie électronique $d\mathbf{v}_e/dt$. On a négligé ce terme à cause du facteur m_e qui le précédait : à nouveau, on considère que la faible masse des électrons leur permet de réagir instantanément. Il faut donc que nos temps caractéristiques soient grands devant l'inverse de ω_{pe} , mais aussi devant l'inverse de la fréquence de Larmor électronique $\omega_{ce}^{-1} \equiv m_e/(eB)$ (puisque'on a conservé le terme en $\mathbf{v} \times \mathbf{B}$). De plus, les longueurs doivent être grandes devant la longueur de Debye mais aussi devant le rayon de Larmor électronique : $\lambda \gg \rho_{Le} \equiv v_{Te}/\omega_{ce}$.

¹³Et non $\mathbf{E} = 0$ ou $\nabla \cdot \mathbf{E} = 0$.

5. Dans la loi d'Ohm, les termes de Hall ($\mathbf{J} \times \mathbf{B}/(en)$) et de compressibilité électronique ($\nabla p_e/(en)$) ont été négligés. Ces termes sont d'ordre 1 en ρ_{Li}/a . Ils appartiennent donc à la catégorie des effets de rayon de Larmor fini : ils peuvent être négligés si $\rho_{Li}/a \ll 1$. Notons que cette condition est équivalente à : $\omega/\omega_{ci} \ll 1$. Ces termes peuvent être importants pour des particules rapides, et ne sont pas pris en compte par la MHD.
6. Dans le cas de la MHD idéale, le terme résistif $\eta \mathbf{J}$ de la loi d'Ohm est négligé aussi. Cela implique que le plasma ne soit pas trop collisionnel ; nous ne quantifions pas cette condition, puisque nous introduirons toujours la résistivité dans nos simulations.

En résumé, les hypothèses que nous avons faites reviennent à exiger : (i) une forte collisionnalité du plasma, (ii) des dimensions grandes devant les rayons de Larmor et les longueurs de Debye (iii) des temps grands devant les inverses de fréquences plasmas et cyclotroniques et (iv) dans le cas de la MHD idéale, une collisionnalité suffisamment faible pour négliger le terme en η de la loi d'Ohm.

Cas des tokamaks : pertinence de la MHD à faible collisionnalité

Du point de vue de la modélisation des tokamaks, l'avantage de la MHD est qu'en tant que théorie fluide, elle est plus simple à implémenter et nécessite moins de ressources numériques qu'une théorie cinétique. On l'a dit, l'efficacité des critères MHD pour la stabilité des plasmas de tokamaks a été mise à l'épreuve pendant plus de 40 ans ; elle est donc bien établie¹⁴. Cela est d'autant plus remarquable que les plasmas de tokamaks sont clairement en-dehors du domaine de validité de ce modèle. L'hypothèse la plus fondamentale, celle de forte collisionnalité, qui autorise un modèle fluide et l'utilisation d'une pression scalaire, est en effet fautive dans les tokamaks.

Pour expliquer ce phénomène, Freidberg donne l'argument suivant (cf. [1], chap. 2) : dans les directions perpendiculaires au champ, le mouvement de gyration des particules joue le rôle des collisions dans le sens où leur mouvement est effectivement isotrope. On comprendrait alors pourquoi la MHD modélise efficacement les modes perpendiculaires au champ, comme les shear Alfvén. La dynamique parallèle au champ reste problématique : dans cette direction, le mouvement des particules est libre. Les mouvements dans la direction parallèle sont mal modélisés par la MHD ; ils ne seront correctement pris en compte que par un code cinétique. Un tel code permettrait également de ne pas faire d'hypothèse sur le caractère scalaire de la pression, ni sur son transport par un terme de type $\Gamma \nabla \cdot \mathbf{v}$. Une modélisation correcte du tenseur de pression et de sa dynamique, au moins pour une population de particules rapides,

¹⁴Par exemple, il est intéressant de remarquer que seules les décharges dites avancées ont dépassé la limite de Troyon, malgré les effets non idéaux importants dans les machines modernes (chauffages auxiliaires, etc).

appelle donc un code hybride MHD/cinétique.

Interprétation des hypothèses dans le cas des dents de scie

En reprenant une à une les hypothèses mentionnées, on peut les mettre en relation avec la physique des dents de scie, que nous étudierons plus loin :

1. La première hypothèse est la restriction fondamentale : comme on vient de le dire, un plasma de tokamak est non collisionnel, et la dynamique parallèle, absente en MHD idéale, reste très mal décrite en MHD étendue. En particulier, les particules piégées donnent une contribution non isotrope importante au tenseur de pression, et elles ont un effet important sur les dents de scie, notamment les particules chaudes. Or ces particules disparaissent en MHD, car le piégeage est supprimé par les collisions. On voit ici l'intérêt d'utiliser une modélisation plus complète des particules chaudes dans ce contexte.
2. La deuxième hypothèse n'est pas restrictive pour les dents de scie.
3. Idem.
4. L'inertie électronique est importante dans la dynamique de la reconnection, lorsque la résistivité est faible.
5. Les effets de rotation bi-fluides (rotation diamagnétique) et les effets de Larmor finis peuvent être importants. Le code XTOR-2F reprend certains de ces effets.

2.1.3 Principe d'Énergie : conditions nécessaires de stabilité

Le système d'équations de la MHD idéale peut être linéarisé facilement et ramené au premier ordre par rapport à une perturbation des champs autour d'un équilibre. On obtient des équations algébriques, et après substitution, des relations de dispersion qui caractérisent les ondes pouvant se propager dans un plasma : la « shear Alfvén », l'onde magnétoacoustique rapide et l'onde magnétoacoustique lente.

En notant ξ le déplacement dû à la perturbation, on peut exprimer l'équation du mouvement (11) linéarisée au premier ordre sous la forme :

$$\rho \frac{\partial^2 \xi}{\partial t^2} = \mathbf{F}(\xi)$$

Où $\mathbf{F}$ est l'opérateur MHD idéale linéarisé autour des valeurs d'équilibre des champs. Il est important de remarquer que $\mathbf{F}$ est auto-adjoint. Quand on introduit une perturbation en $e^{i\omega t}$, l'équation ci-dessus devient une équation aux valeurs propres et vecteurs propres. Le fait que $\mathbf{F}$ soit auto-adjoint implique que ses valeurs propres ω^2 sont réelles. Cela simplifie grandement le calcul des seuils d'instabilités.

Le spectre de $\mathbf{F}$ contient à la fois des valeurs propres discrètes et des continus (seulement quand $\omega^2 > 0$, donc en cas de stabilité).

Le principe d'énergie est directement issu de la conservation de l'énergie. Considérons un déplacement ξ associé à une instabilité ; la variation d'énergie potentielle associée à ξ s'écrit :

$$\delta W(\xi, \xi^*) = -\frac{1}{2} \int \xi^* \cdot \mathbf{F}(\xi) d\mathbf{r}$$

Où ξ^* est l'adjoint (le transposé du conjugué) de ξ . En connaissant la forme de $\mathbf{F}$, qui dépend de la configuration magnétique, on peut calculer ce δW analytiquement.

Le principe d'énergie affirme alors que le système est stable (*i.e.*, pas de croissance de la perturbation) si et seulement si :

$$\delta W(\xi, \xi^*) \geq 0 \quad \forall \xi$$

En effet, s'il existe un déplacement associé à une variation d'énergie potentielle négative, alors la conservation de l'énergie implique que le déplacement engendre une variation d'énergie cinétique positive, et la perturbation croît.

Insistons sur le fait que le principe d'énergie donne une *condition nécessaire de stabilité*, ce qui équivaut à une *condition suffisante d'instabilité*. Par son application aux tokamaks, la MHD fournit des contraintes sur les paramètres macroscopiques des machines.

Des instabilités peuvent apparaître dans des zones stables du point de vue du principe d'énergie : elles seront dues à des termes non inclus dans la MHD idéale. Comme nous nous intéresserons dans la suite à la MHD résistive, nous insistons en particulier sur le fait que le terme de résistivité de la loi d'Ohm n'est pas stabilisant, comme on pourrait s'y attendre. En tant que terme de diffusion, on pourrait penser qu'il concourt à dissiper l'énergie disponible pour une instabilité. En fait, son effet majeur est de rendre faux le théorème du gel, qui spécifie, en MHD idéale, que le plasma est gelé dans le champ. Par conséquent, les lignes de champ peuvent se briser et se reconnecter en MHD résistive. La topologie magnétique peut être radicalement modifiée, ce qui permet de nouvelles instabilités.

2.2 Les tokamaks vus par la MHD

Les tokamaks sont les machines les plus répandues et étudiées dans le domaine de la fusion contrôlée par confinement magnétique. Ce sont des machines à confinement toroïdal, *i.e.*, le confinement magnétique est assuré par des lignes de champ magnétique qui s'enroulent à la manière d'hélices sur des surfaces toriques emboîtées (appelées dans la suite surfaces de champ ou surfaces de flux). Une ligne de champ

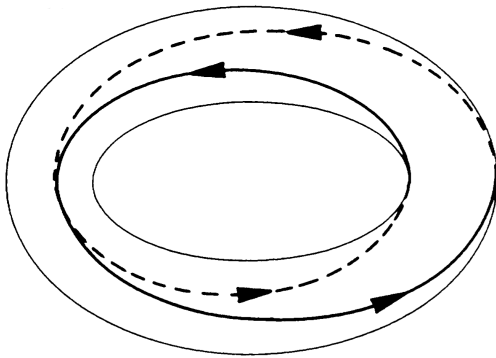


FIG. 1 – Ligne de champ s'enroulant autour d'une surface rationnelle ($q = 2$) (figure tirée de [15]).

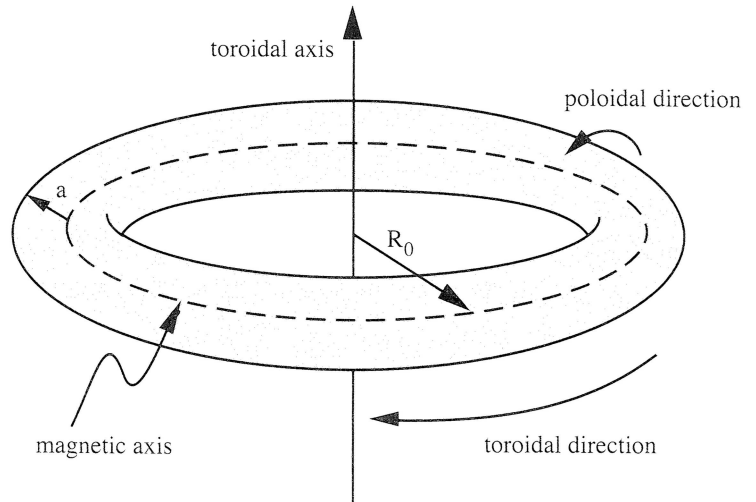


FIG. 2 – Axes toroïdal et poloïdal (figure tirée de [2]).

peut recouvrir entièrement la surface, ou être rationnelle : dans ce cas, après un nombre fini de tours autour de l'axe toroïdal (et poloïdal¹⁵), la ligne de champ se referme sur elle-même. Le nombre de tours toroïdaux par tour poloïdal est mesuré par le facteur de sécurité q . La figure 2.2 fournit un exemple. Les surfaces rationnelles ont une grande importance en MHD car elles sont le lieu de résonance d'instabilités. L'équation 11 (sans les termes anisotropes) implique qu'à l'équilibre, les lignes de champ et de courant sont tangentes aux surfaces de champ en tout point, et que la pression est constante sur une surface de champ.

La composante toroïdale du champ est créée par des bobines magnétiques formant un solénoïde fermé sur lui-même¹⁶. La composante poloïdale est créée par un courant toroïdal induit dans le plasma. La composante poloïdale assure le confinement radial du plasma grâce à la pression magnétique, et la composante toroïdale assure de bonnes conditions de stabilité. Un champ magnétique vertical créé par des bobines horizontales est utilisé pour contrôler la forme et la position du plasma. La figure 3 fournit un schéma de principe¹⁷.

Après avoir présenté la configuration magnétique des tokamaks et les paramètres pertinents pour appliquer la MHD à ces machines, nous pourrions situer le rôle et

¹⁵L'« axe » poloïdal, aussi appelé axe magnétique, est en fait le cercle horizontal de centre O et de rayon le grand rayon du tore R_0 , cf. la figure 2.2.

¹⁶Dans une machine, le nombre fini de bobines brise l'invariance par rotation de la configuration autour de l'axe de symétrie de révolution du tore ; notre code ne modélise pas les bobines, mais n'utilise pas non plus d'hypothèse d'invariance en ϕ .

¹⁷Figure trouvée sur le site : <http://www1.cfi.lu.lv/teor/main.html>.

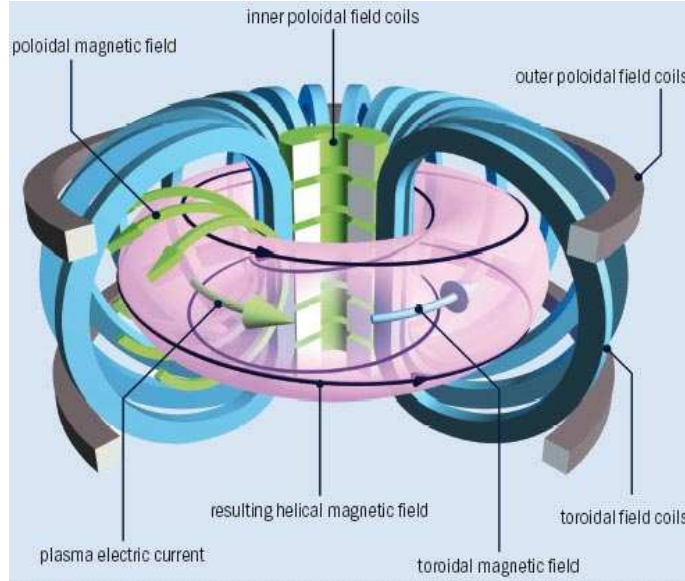


FIG. 3 – Schéma de principe d'un tokamak.

l'apport de la MHD dans l'étude des tokamaks. Remarquons que nous laissons de côté les techniques de chauffage auxiliaires pour nous concentrer sur les machines ohmiques (*i.e.*, dans lesquelles le chauffage est assuré par l'effet Ohm dû au courant toroïdal).

2.2.1 Géométrie et configuration magnétique

Dans cette section, on définit les notations d'usage dont on se sert dans le reste du texte.

Paramètres géométriques

On introduit le système de coordonnées toroïdales, voir figure 4. Les angles ϕ et θ dénotent les directions toroïdale et poloïdale, et r la direction radiale (*i.e.*, celle du petit rayon du tore).

Le petit rayon du tore est noté a , le grand rayon R_0 . Le rapport $A \equiv a/R_0$ est appelé rapport d'aspect. Dans la pratique on raisonne souvent avec l'inverse de A , noté ϵ . On voit que faire $\epsilon \rightarrow 0$ revient à considérer la limite cylindrique du tore (un tore à grand rayon infini tend vers un cylindre). Les configurations cylindriques sont analytiquement beaucoup plus simples que les toroïdales pour l'étude de la stabilité et du confinement. Le confinement radial dans un tokamak est analogue à celui dans un « Z-pinch » : il est assuré par un fort courant toroïdal, c'est-à-dire par un champ magnétique poloïdal. On utilise très fréquemment l'approximation « à grand rapport

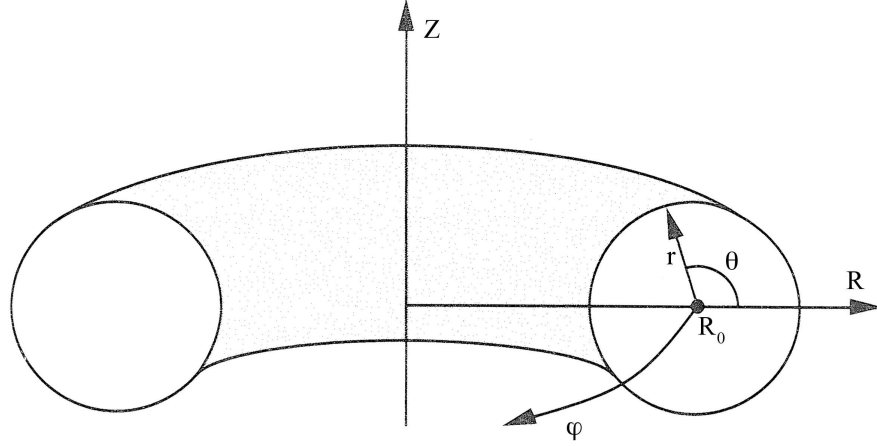


FIG. 4 – Coordonnées toroïdales de base (figure tirée de [2]).

d'aspect », *i.e.* $\epsilon \ll 1$ lorsqu'on veut pousser les calculs analytiquement, ou se faire une idée rapide d'un problème. Cela permet également de conserver les effets de géométrie torique, qui sont très importants. Le Z-pinch a de mauvaises propriétés MHD, c'est pourquoi on ajoute le champ magnétique toroïdal, qui stabilise le plasma (comme dans un θ -pinch).

Remarquons que dans les tokamaks récents, cette approximation est nettement violée ($\epsilon \sim 0.3$). Les tokamaks récents¹⁸ ont également des sections poloïdales non circulaires. Les paramètres κ , triangularité, et δ , élongation, décrivent leur forme. Ces paramètres ont une influence sur la stabilité du plasma. Dans ces cas, la coordonnée r sera remplacée par une nouvelle coordonnée, généralement notée s , reliée au flux poloïdal normalisé qu'on définit dans la section suivante.

Champ et flux magnétiques

On note les composantes du champ magnétique¹⁹ B_ϕ (composante toroïdale) et B_θ (composante poloïdale) :

$$\mathbf{B} = B_\phi \mathbf{e}_\phi + B_\theta \mathbf{e}_\theta$$

On utilise fréquemment les flux magnétiques pour exprimer le champ :

$$\mathbf{B} = \nabla \Psi_\theta \times \nabla \phi + G \nabla \phi$$

¹⁸Plus exactement, les surfaces de champs des tokamaks récents

¹⁹avec les notations analogues pour tous les autres vecteurs.

Avec

$$\Psi_\theta \equiv (2\pi)^{-1} \int_{S_\theta} \mathbf{B} \cdot d\mathbf{S}$$

Où S_θ est un disque centré à l'origine du repère de la figure 4 et contenu dans le plan horizontal. On peut voir Ψ_θ comme un label des surfaces magnétiques (à chacune correspond un seul disque S_θ , et une seule valeur de Ψ_θ , qui est monotone), c'est pourquoi on les appelle aussi surfaces de flux. On exprime souvent les grandeurs physiques constantes sur les surfaces de flux en fonction de Ψ_θ ; cela simplifie grandement les notations lorsque les sections poloïdales ne sont pas circulaires et les surfaces de flux ne sont pas définies par $r = cte$. Par exemple, Ψ_ϕ , le flux magnétique toroïdal (défini par analogie avec Ψ_θ), les courants toroïdaux et poloïdaux I_θ et I_ϕ ($I_\theta = \int_{S_\theta} \mathbf{J} \cdot d\mathbf{S} = -2\pi G$ etc), et la pression p sont des fonctions de Ψ_θ .

Paramètres physiques

Un paramètre crucial est le rapport entre énergie cinétique et énergie magnétique moyenne (avec $\langle x \rangle = 1/V \int_V x(\mathbf{r}) d\tau$) :

$$\beta \equiv \frac{2\mu_0 \langle p \rangle}{\langle B^2 \rangle}$$

À cause des rôles différents des composantes du champ, il est utile de distinguer le bêta poloïdal :

$$\beta_p \equiv \frac{2\mu_0 \langle p \rangle}{\langle B_p^2 \rangle}$$

Avec $B_p \equiv \mu_0 I_{\phi 0} / (2\pi a)$, et $I_{\phi 0}$ le courant toroïdal total, et le bêta toroïdal :

$$\beta_t \equiv \frac{2\mu_0 \langle p \rangle}{\langle B_{\phi 0}^2 \rangle}$$

Où $B_{\phi 0}$ est le champ magnétique toroïdal sur l'axe magnétique. Le critère de Lawson favorise les β les plus élevés possibles; cependant, la MHD impose des fortes contraintes pour que le plasma soit stable. Dans les expériences pertinentes pour nos simulations (tokamaks à basse pression), β est de l'ordre de quelques %.

Nous introduisons le facteur de sécurité q , qui mesure le pas de l'hélice dessinée par les lignes de champ :

$$q \equiv \frac{\Delta\phi}{2\pi}$$

Où $\Delta\phi$ est la variation d'angle toroïdal effectuée par une ligne de champ quand elle effectue un tour dans la direction poloïdale. q est constant sur une surface de flux. Il mesure le nombre de tours toroïdaux que la ligne de champ dessine en un tour

poloïdal. C'est un paramètre très important pour la stabilité du plasma : comme on l'a dit, les surfaces rationnelles (pour lesquelles $q = m/n$) sont le lieu de résonances d'instabilités. Ces résonances sont dues au fait que sur une surface rationnelle où $q = m/n$, une instabilité de nombres d'ondes toroïdaux et poloïdaux n et m peut se propager sans dépenser d'énergie pour tordre les lignes de champ ; on veut donc limiter leur nombre au maximum, surtout pour des m, n petits.

On définit de plus la dérivée logarithmique du facteur de sécurité, le shear magnétique :

$$\hat{s} \equiv \frac{\Psi}{q} \frac{dq}{d\Psi}$$

Ordering et variations des champs dans un tokamak

Dans chaque famille de machines à confinement toroïdal, la configuration magnétique est définie par un ordering des champs et par les variations des composantes des champs en fonction de r . Dans un tokamak à basse pression, l'ordering est le suivant :

$$\begin{aligned} \frac{B_\theta}{B_\phi} &\sim \epsilon \\ \beta_\phi &\sim \epsilon^2 \\ \beta_\theta &\sim 1 \\ q &\sim \frac{rB_\phi}{RB_\theta} \sim 1 \end{aligned}$$

Dans ce travail, nous nous intéressons aux tokamas à faible β (comme JET, ITER, ou Tore Supra). La composante toroïdale B_ϕ varie logiquement comme $1/R$, d'où la distinction entre côté « fort champ » ($\pi/2 < \theta < 3\pi/2$) et « faible champ » ($0 < \theta < \pi/2$ et $3\pi/2 < \theta < 2\pi$). Contrairement à B_ϕ , B_θ s'annule à l'intérieur du plasma, sur l'axe poloïdal, qui est pour cette raison appelé axe magnétique. Les variations de B_θ en fonction de r dépendent du profil de la densité de courant toroïdale j_ϕ . La figure 5 donne un exemple typique de la variation des champs sur l'axe $\theta = 0$ (la configuration est symétrique en ϕ).

2.2.2 Instabilités MHD dans les tokamaks

Principe d'énergie dans un tokamak

Dans un tokamak, on peut exprimer δW comme la somme de trois termes :

$$\delta W = \delta W_P + \delta W_S + \delta W_V$$

Ces termes représentent les contributions énergétiques du plasma, de l'interface plasma-vide, et du champ magnétique du vide, respectivement. Les deux derniers

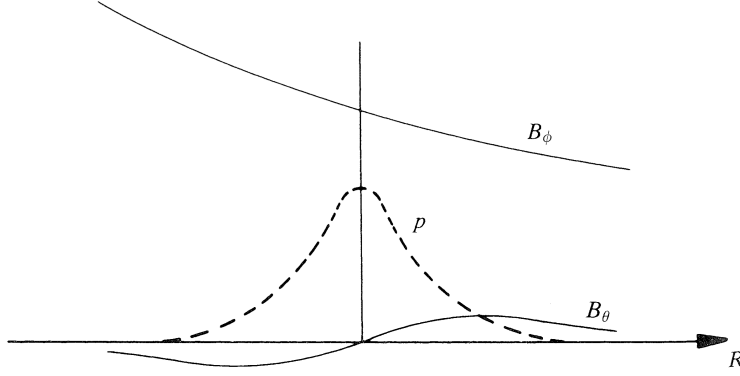


FIG. 5 – Variation des champs magnétique et de la pression dans un tokamak (figure tirée de [15]). L'axe vertical marque l'abscisse de l'axe magnétique.

termes sont non nuls uniquement dans le cas d'une instabilité externe (*i.e.*, qui se propage jusqu'à l'interface plasma-vide).

On développe :

$$\begin{aligned}\delta W_P &= \frac{1}{2} \int_P \left[\frac{\|\mathbf{Q}\|^2}{\mu_0} - \xi_\perp^* \cdot (\mathbf{J} \times \mathbf{Q}) + \Gamma p \|\nabla \cdot \xi\|^2 + (\xi_\perp \cdot \nabla p) \nabla \cdot \xi_\perp^* \right] d\mathbf{r} \\ \delta W_S &= \frac{1}{2} \int_S \|\mathbf{n} \cdot \xi_\perp\|^2 \mathbf{n} \cdot \left[\nabla \left(p + \frac{B^2}{2\mu_0} \right) \right] ds \\ \delta W_V &= \frac{1}{2} \int_V \frac{\|\mathbf{B}_{1v}\|^2}{\mu_0} d\mathbf{r}\end{aligned}$$

On a : $\mathbf{Q} = \nabla \times (\xi \times \mathbf{B})$ la perturbation magnétique, $\mathbf{n}$ la normale orientée à l'interface plasma-vide, et $\mathbf{B}_{1v}$ le champ magnétique dans le vide. Les doubles crochets dénotent une intégration de part et d'autre de l'interface plasma-vide.

La contribution du plasma peut se réécrire :

$$\begin{aligned}\delta W_p &= \frac{1}{2} \int_p \left[\frac{\|\mathbf{Q}_\perp\|^2}{\mu_0} + \frac{B^2}{\mu_0} \|\nabla \cdot \xi_\perp + 2\xi_\perp \cdot \kappa\|^2 + \Gamma p \|\nabla \cdot \xi\|^2 \right. \\ &\quad \left. - 2(\xi_\perp \cdot \nabla p)(\kappa \cdot \xi_\perp^*) - J_\parallel \left(\xi_\perp^* \times \frac{\mathbf{B}}{B} \right) \cdot \mathbf{Q}_\perp \right] d\mathbf{r}\end{aligned}$$

Avec $\kappa = \mathbf{B} \cdot \nabla \mathbf{B} / B^2$. Cette expression permet de comprendre quels effets sont stabilisants et déstabilisants en MHD idéale.

Les trois premiers termes sont positifs, donc stabilisants. Le premier représente l'énergie nécessaire pour tordre les lignes de champ ; le deuxième celle nécessaire pour compresser le champ magnétique, et le troisième celle nécessaire pour compresser le plasma.

Les deux derniers termes peuvent être négatifs, et donc déclencher des instabilités. Le premier est proportionnel au gradient de pression, et le second à la densité de courant toroidale. On parlera de modes de pression ou de modes de courant, suivant la source d'énergie qui déstabilise ces modes.

Instabilités et critères nécessaires de stabilité

L'instabilité d'interchange est une instabilité locale de pression. Elle peut se développer lorsque la courbure du champ magnétique est concave vers le gradient de densité du plasma ; dans un tokamak, c'est le cas côté faible champ (on parle de courbure favorable ou défavorable). Dans ce cas, deux tubes de flux peuvent s'échanger et atteindre tous les deux un état énergétique plus bas, l'un grâce au gradient de pression, l'autre grâce à la courbure de ses lignes de champs qui diminue. Les interchanges les plus instables ont $k_{\parallel}/k_{\perp} \ll 1$ et $k_{\perp}a \gg 1$. Les critères de Suydam (1D) et de Mercier (2D) sont les critères de stabilité associés ; ils portent sur des grandeurs locales, et non intégrées du plasma.

Les modes de ballonnement sont des modes de pression macroscopiques. Ce sont des modes évoluant lentement le long des lignes de champ, avec des perturbations concentrées dans les zones de courbure défavorable. Comme pour les interchanges, les plus instables ont $k_{\parallel}/k_{\perp} \ll 1$ et $k_{\perp}a \gg 1$. Le critère de stabilité mène à une limite en β ; il est donc très important.

Les modes de kink sont des modes de courant, appelés ainsi à cause du déplacement hélicoïdal qu'ils entraînent. Les kinks externes sont les modes les plus dangereux pour le confinement ; les kink internes (le plus souvent avec un nombre d'onde azimuthal $m = 1$, avec $k_{\theta} = m/r$), sont responsables des dents de scie²⁰. Le caractère externe ou interne d'un kink dépend de la localisation de sa surface de résonance, définie par $q = m/n$, $n = Rk_{\phi}$ étant le nombre d'onde toroïdal. Ces modes ont en général $k_{\parallel}/k_{\perp} \ll 1$ et $k_{\perp}a \sim m \sim 1$. La stabilité des kinks impose des limites sur les variations de la densité de courant toroïdal (ou de façon équivalente sur q au bord et au centre), sur le courant toroïdal total, ainsi que sur β , pour empêcher un mode de kink-ballonnement de se développer (critère de Troyon). Ainsi, les kinks externes imposent que le facteur de sécurité soit élevé au bord ($q \sim 3 - 5$ en pratique), c'est-à-dire que le courant décroisse vite au bord.

²⁰Il s'agit en fait de la version résistive du kink interne.

2.3 Les dents de scie

Les dents de scie sont un phénomène universel de la physique des tokamaks, où plusieurs effets se combinent probablement. Apparues tout d'abord dans les petits tokamaks ohmiques, que la MHD résistive devrait assez bien modéliser, elles ont ensuite été observées sous des formes très diverses et dans des conditions opératoires variées dès que les moyens de les détecter ont été disponibles. Elles sont présentes dans toutes les machines, et dépassent largement le cadre théorique de la MHD résistive. Leur comportement en présence de chauffages auxiliaires et de particules rapides dans ces machines doit être maîtrisé pour garantir qu'elles ne menaceront pas les paramètres physiques nécessaires pour un réacteur. Cela exige des modèles incluant au minimum des effets bi-fluides, et éventuellement des effets cinétiques.

Expérimentalement, elles se traduisent par des réorganisations périodiques de la température et de la densité électronique au cœur du plasma. Elles n'entraînent généralement pas la perte du confinement (sauf dans des cas extrêmes : dents de scie monstres ou couplage avec des autres instabilités comme les ELMs²¹ ou les NTMs²²). En cela, elles constituent un phénomène original : la majorité des instabilités macroscopiques sont fatales pour le confinement. Il y a là une similitude avec les ELMs, qui est une motivation supplémentaire pour l'étude des dents de scie.

Après avoir présenté le phénomène expérimental plus en détail, nous montrerons comment les interprétations théoriques et les simulations numériques des dents de scie ont évolué au fur et à mesure que des résultats expérimentaux plus complets étaient obtenus, pour situer dans son contexte l'étude que nous présenterons dans le chapitre 3. Hastie adopte une démarche historique similaire et très éclairante dans [5]. Partant des années 70, nous exposerons d'abord la théorie de reconnection complète de Kadomtsev, dans le cadre de la MHD résistive, où le kink interne est la seule instabilité en cause. Nous rappellerons les écarts entre les prédictions de Kadomtsev et les mesures obtenues dans les années 80, et les modifications et alternatives qui furent proposées en conséquence : instabilité secondaire, instabilité idéale. Nous expliquerons ensuite comment les modèles MHD ont continué à évoluer dans les années 90, par l'introduction de termes supplémentaires (bi-fluides et non-linéaires). Nous terminerons en donnant une idée de la position actuelle de la communauté scientifique sur les dents de scie, qui est plus axée sur une stratégie de contrôle et de prédiction semi-empirique, laissant de côté un certain nombre de questions, que nous évoquerons.

Un code hybride devrait à terme permettre d'obtenir une modélisation intégrée du phénomène, incluant MHD résistive, effets bi-fluides et particules rapides, pour passer à une démarche plus prédictive. C'est pourquoi les dents de scie sont un objet d'étude privilégié pour un code comme XTOR, qui évolue vers un algorithme hybride.

²¹Edge Localized Modes.

²²Neoclassical Tearing Modes.

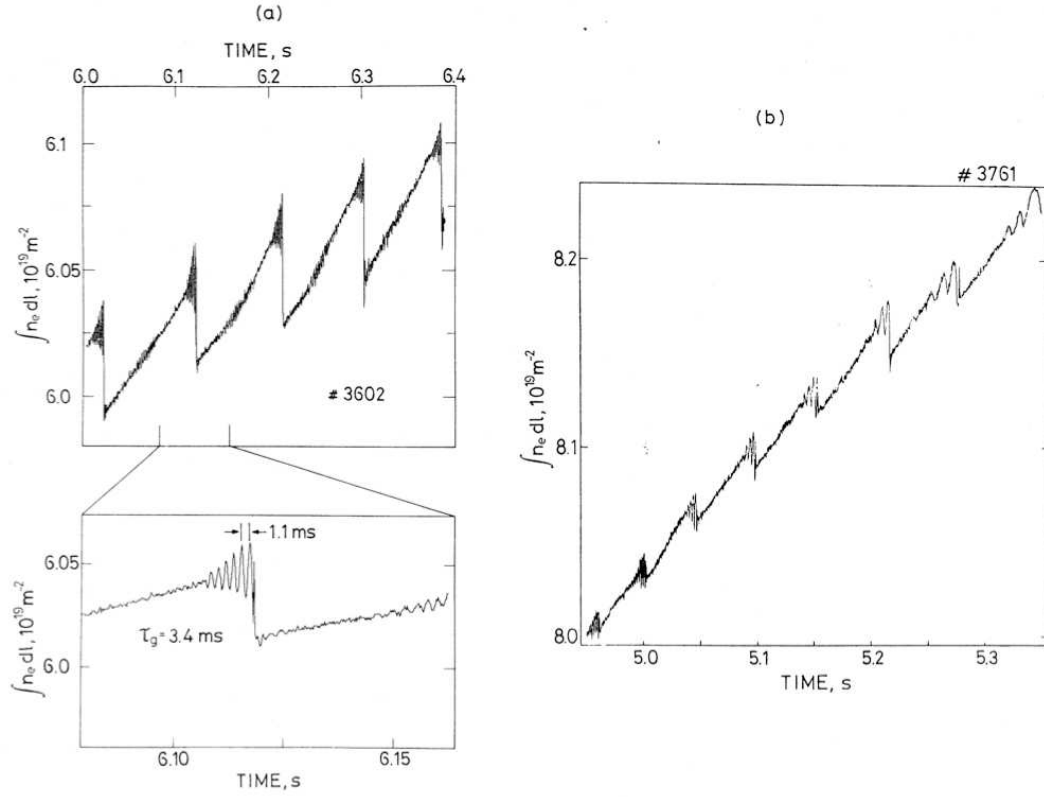


FIG. 6 – Exemple de dents de scie sur JET (densité électronique vs. temps ; figure tirée de [5]).

2.3.1 Découverte et description du phénomène

Les dents de scie ont été observées pour la première fois en 1974 dans le tokamak ST, à Princeton (cf. [4]). Ce tokamak était l'un des premiers à disposer d'un système de mesure de l'émission des X mous (« soft X rays »), qui permet de mesurer l'évolution temporelle du profil de température électronique. Les dents de scie se traduisent par des oscillations asymétriques. Chaque cycle est constitué d'une montée lente de T_e et n_e au cœur du plasma, suivie par une chute brutale. Par « cœur », on entend « l'intérieur de la surface $q = 1$ » ; les dents de scie sont liées à l'existence de cette surface, et apparaissent en général quand q descend en-dessous de 1 au centre du plasma. L'asymétrie donne aux courbes temporelles de température et de densité une forme caractéristique en dents de scie, d'où le nom du phénomène. On montre un exemple de diagnostic typique sur la figure 6.

Durant la phase lente (ramp-up, ou rampe), les profils de densité et de température électronique se piquent au centre du plasma sous l'effet des sources de chauffage. La

chaleur et la densité accumulées sont brutalement évacuées vers l'extérieur pendant la phase de chute, ou crash (par conséquent, si on mesure T_e à l'extérieur de la surface $q = 1$, la rampe devient une phase de décroissance et le crash une phase de croissance brutale). Après le crash, les profils de n_e et T_e sont plats au cœur du plasma. L'amplitude des oscillations de densité et température est de l'ordre de 10% pour des machines ohmiques. Le facteur de sécurité oscille inversement (décroissance pendant la rampe, croissance pendant le crash), avec une amplitude moins importante, de 1% à 5%. La période totale sur JET pour une décharge ohmique est de l'ordre de 100 ms, avec un temps de crash de l'ordre de $\tau_{crash} \sim 100 \mu s$. La période croît avec la taille de la machine. Dans ITER, on prévoit une période d'environ 20 - 40 s (cf. [6] ; notons que dans ITER, le chauffage auxiliaire allonge la durée de la rampe).

Il existe une grande variété de dents de scie : inversées, saturées, monstres... Sur la plupart d'entre elles, mais pas sur toutes, le crash est précédé d'oscillations croissantes dites précurseurs, et parfois suivies d'oscillations amorties, postcurseurs. Les chauffages auxiliaires (injection de neutres, résonance cyclotronique ionique et électronique) et les α de fusion ont un effet stabilisant : la période totale est allongée. La quantité de matière et de chaleur emmagasinée au centre du plasma peut alors devenir très grande, et il arrive que le crash détruise le confinement à la fin d'une dent de scie très longue (dent de scie monstre). Les particules rapides peuvent aussi faire évoluer les dents de scie vers des fishbones (l'influence des particules rapides sur les dents de scie et les fishbones a fait l'objet d'un article de revue par Porcelli, cf. [7]). Les fishbones sont des déplacements hélicoïdaux ($m = n = 1$) du cœur du plasma, qui croissent puis décroissent sans causer de crash, contrairement aux dents de scie. Dans la suite, sauf contre-indication, les ordres de grandeurs donnés correspondent à des décharges ohmiques, qui sont celles qui correspondent au domaine de validité de notre étude.

2.3.2 Théorie de Kadomtsev : reconnection du kink interne

Le kink interne est rapidement apparu comme un acteur majeur du cycle, même si son rôle n'est pas parfaitement clair. Ce mode croît juste avant le crash, ce qui se traduit parfois par les oscillations précurseurs/postcurseurs. La première théorie des dents de scie, proposée par Kadomtsev dès les années 70, repose sur la reconnection complète de ce kink.

Le kink interne

Le kink interne est une instabilité avec $m = n = 1$, due au gradient de pression et au courant, qui engendre un déplacement hélicoïdal constant en r à l'intérieur de sa surface de résonance, la surface $q = 1$. Il est stable si $q > 1$ dans tout le plasma. Le kink idéal a d'abord été étudié dans le cas cylindrique par Shafranov (cf. [8]), puis par Bussac *et al.*, (cf. [9]), qui ont montré que dans le cas toroïdal, avec un grand

rapport d'aspect, le mode a un seuil en β_{p1} (contrairement au cas cylindrique où le mode est instable dès que $q_0 < 1$). β_{p1} est le bêta poloïdal à l'intérieur de la surface $q = 1$:

$$\beta_{p1} = \frac{8\pi}{B_\theta(r_1)} [\langle p \rangle_{r_1} - p(r_1)]$$

Où r_1 est le rayon de la surface $q = 1$, et les crochets $\langle \rangle_{r_1}$ désignent la moyenne sur le volume contenu par cette surface. Dans les plasmas ohmiques, le seuil se trouve en général autour de $\beta_{p,cr} \sim 0.2 - 0.3$ (dans les grandes machines, où l'approximation à grand rapport d'aspect est fautive, le seuil est plutôt vers 0.15 - 0.20). Le taux de croissance idéal est de l'ordre de $\gamma \sim 10^{-3} \text{ s}^{-1}$ pour une machine ohmique. Selon le principe d'énergie, une instabilité idéale se développe si l'énergie potentielle associée est négative; Bussac a calculé que l'énergie potentielle est proportionnelle à :

$$\delta W_{MHD} \sim \left(\frac{r_1}{R} \right)^2 \hat{s}_1^{-1} [(\beta_{p,cr})^2 - \beta_{p1}^2]$$

Rosenbluth *et al.*, en étudiant la saturation non-linéaire du kink idéal, ont conclut (cf. [10]) que l'instabilité entraîne l'apparition d'une nappe de courant sur la surface $q = 1$. La dynamique des électrons dans cette nappe est influencée par la reconnection des lignes de champs permise par la résistivité, l'inertie électronique, le terme de Hall, et la viscosité électronique. Ces effets déterminent la dynamique du crash.

Le seuil en β_{p1} disparaît lorsqu'on prend en compte la résistivité dans le kink. Le kink résistif a été étudié en premier par Furth *et al.*, cf. [12] dans le cas des plasmas astrophysiques. Le taux de croissance du kink résistif est plus faible que celui du kink idéal; il est proportionnel à $\eta^{1/3}$ et, dans un tokamak, à $q'(r_1)^{2/3}$. Le kink idéal est dominant quand δW_{MHD} est négatif. Lorsque δW_{MHD} est proche de 0, le kink résistif devient dominant. Dans le cas des tokamaks, on est en dessous du seuil de stabilité idéale, c'est donc le kink résistif qui nous intéresse. Pour δW_{MHD} positif, c'est le mode de tearing qui devient dominant, avec un taux de croissance encore plus faible, proportionnel à $\eta^{3/5}$.

Les dents de scie de Kadomtsev : reconnection complète du flux hélicoïdal

Il est tentant et élégant de considérer que le cycle entier peut être expliqué par le kink interne. Le premier modèle de dents de scie a été développé sur cette idée par Kadomtsev en 1975 (cf. [14]). Ce modèle, très classique, considère que le crash est dû à la reconnection complète du flux magnétique intérieur à la surface $q = 1$ avec le flux extérieur. Il s'agit ici du flux à travers la surface hélicoïdale bornée d'un côté par l'axe magnétique et de l'autre par une ligne de champ de la surface $q = 1$. Le flux reconnecté forme un îlot qui grandit jusqu'à occuper tout le cœur du plasma, après

avoir poussé le cœur dense et chaud vers l'extérieur. À la fin du cycle, un équilibre axisymétrique est rétabli, avec $q_0 > 1$. Le chauffage ohmique recommence alors à piquer les profils de densité et de courant, jusqu'à ce que q_0 redescende sous l'unité et que le kink croisse à nouveau, etc.

La reconnection a lieu dans une couche de faible épaisseur, typiquement (cf. [15]) :

$$\delta = r_1 \sqrt{\frac{\tau_A}{\tau_\eta}}$$

Où τ_A est le temps d'Alfvén, $\tau_\eta = \mu_0 r_1^2 / \eta$ est le temps de diffusion résistif, et r_1 est le rayon de la surface $q = 1$. On déduit que le temps caractéristique de crash devrait être :

$$\tau_{crash} = \sqrt{\tau_\eta \tau_A}$$

Rapidement, des simulations numériques MHD en géométrie cylindrique ont montré un bon accord avec le modèle de Kadomtsev (cf. Sykes et Wesson dans [16], et Denton *et al.* dans [17]), prévoyant une reconnection complète du flux et un temps de crash similaire. Ces études ont été étendues à une géométrie toroïdale par Aydemir *et al.* dans [18] et Park *et al.* dans [19], en incluant des effets de résistivité néo-classique et d'hyper-résistivité.

2.3.3 Confrontation aux expériences, effets bi-fluides et idéaux

Les diagnostics sur les dents de scie se sont multipliés rapidement au début des années 80, et les prédictions du modèle de Kadomtsev ont été mises à mal par les mesures expérimentales. Cela a contribué à l'émergence de nouvelles théories, alternatives ou complémentaires à ce modèle.

Temps de crash : instabilité secondaire

Le désaccord le plus flagrant concerne le temps de crash prévu par Kadomtsev. S'il est plus faible que le temps résistif, il reste en général nettement au-dessus des valeurs expérimentales. Dans les petites machines ohmiques comme ST, qui doivent être correctement modélisées par la MHD résistive, l'ordering est acceptable (typiquement $\tau_A \approx 1\mu s$, $\tau_\eta \approx 10ms$, et $\tau_{crash} \approx 100\mu s$). Pour des machines comme JET, où r_1 est plus grand et η plus faible, le temps obtenu est au moins un ordre de grandeur trop grand. Au début des années 80, Dubois *et al.* ont observé que le crash intervenait avant que la reconnection du flux soit complète, ce qui explique que le temps de crash de Kadomtsev soit trop grand (cf. [20]). Cela a conduit à faire l'hypothèse d'une instabilité secondaire, qui interviendrait alors que l'îlot $m = n = 1$ est encore petit et qui entraînerait le crash très rapide de la température. Dans ce cas, il y

aurait une reconnection partielle et non complète du flux. La nature cette instabilité secondaire a été beaucoup discutée.

Déclenchement rapide : instabilité de quasi-interchange

Dans les années 80 également, on a commencé à observer des dents de scie sans oscillations précurseurs. Le problème du déclenchement rapide (« fast trigger ») se pose alors : le mode causant le crash atteint des temps de croissance importants ($\gamma^{-1} \sim 100 \mu\text{s}$ sur JET, par exemple) en un temps très court (environ $100 \mu\text{s}$, toujours sur JET). Ce genre de comportement évoque une instabilité idéale ; les instabilités résistives croissent sur des temps plus longs. Wesson, en 1986 (cf. [25]), a théorisé qu'un profil de q très proche de 1 dans toute la surface $q = 1$ permettait à une instabilité idéale d'hélicité $m = n = 1$ de se développer. Le déplacement ayant la même hélicité que les lignes de champ, la torsion des lignes de champ, qui stabilise normalement l'instabilité idéale, disparaît. Dans ce cas, l'ordering de l'étude de Bussac (faible bêta, grand rapport d'aspect) n'est plus valable, et Aydemir a montré qu'un mode de pression idéal pouvait se développer et déclencher le crash avec un temps caractéristiques proche des expériences, pour des valeurs de bêta bien en dessous du seuil idéal (cf. [26]). Wesson décrit un cycle caractérisé par l'apparition d'une « bulle froide » qui se déplacerait vers le centre, le cœur chaud se déformant pour ressembler à un croissant. Une telle dynamique a notamment été observée dans JET. Ce type d'instabilité a été nommé quasi-interchange.

Mesures du profil de q et stabilité de la rampe

Un avantage de l'hypothèse de la reconnection partielle est que le profil de q après le crash est lié à la quantité de reconnection magnétique pendant le crash. Dans le cas d'une reconnection complète, on devrait avoir $q_0 > 1$ à la fin d'un cycle. Or, à partir des années 80, et contrairement aux prédictions de Kadomtsev, on a mesuré des valeurs de q_0 loin de l'unité (0.7 - 0.8) dans plusieurs machines (*e.g.* dans TEXTOR, TFTR et JET, cf. respectivement [27], [28], et [29]). Dans ces cas, q_0 n'atteint jamais la valeur 1 pendant le cycle. Ces mesures semblent exclure le modèle de Kadomtsev et celui de Wesson. Cependant, dans d'autres cas on a mesuré des profils de q_0 proches de 1 après le crash (*e.g.* dans ASDEX, TEXT, CTA et JET, cf. respectivement [30], [31], [32], et [33]). Ces divergences illustrent la richesse physique du phénomène, et expliquent la difficulté qu'il y a à trouver une théorie unificatrice.

L'existence d'une phase de stabilité pendant la rampe pose alors problème. En effet, si le facteur de sécurité est inférieur à 1, le kink interne résistif devrait être instable. On a pourtant observé des rampes beaucoup plus longues que le temps de croissance du kink résistif. On a envisagé une stabilisation due à un shear nul sur $q = 1$, mais le profil de courant nécessaire ne semble pas pouvoir perdurer durant toute la rampe. En revanche, il est plausible que des effets bi-fluides, notamment diamagnétiques, stabilisent efficacement le kink pendant la rampe.

La très grande diversité des diagnostics disponibles à partir des années 80 a donc conduit très rapidement à des remises en cause du modèle de Kadomtsev. Plusieurs nouvelles théories ont été envisagées, mais déjà, la grande variété des phénomènes observés semblait compromettre l'existence d'une explication simple et universelle des dents de scie, et aucune théorie ne s'est imposée comme telle. Dans la décade suivante, d'autres effets furent encore proposés, avant que la communauté scientifique ne se tourne vers une démarche plus empirique, laissant de côté certains problèmes, jugés non redhibitoires.

2.3.4 MHD étendue et effets non-linéaires

Dans les années 90, de nombreux nouveaux effets ont été ajoutés aux modèles MHD pour essayer de mieux modéliser les dents de scie. Ainsi, on a étudié comment la phase reconnection du kink pouvait être accélérée : par les termes d'inertie électronique (cf. [21]), l'effet Hall (cf. [22], [23]), les effets de viscosité électronique (cf. [24]). Ces effets interviennent sur la dynamique des électrons dans la couche de reconnection. Dans la MHD, ils se traduisent par des termes supplémentaires dans la loi d'Ohm. Les deux derniers introduisent les effets de dérive diamagnétique, qui dépendent de la nature des particules ; on passe donc de la MHD pure à un modèle bi-fluide. La loi d'Ohm généralisée se ramène alors à l'équation de la dynamique électronique. L'influence des particules piégées a également été explorée ; ces problématiques sont trop vastes pour être récapitulées ici.

Dans cette même période, l'absence de crash dans les fishbones a conduit à penser que le changement de topologie dû à la reconnection (absent dans les fishbones) pouvait être un composant essentiel de l'accélération du crash. Le crash serait alors en deux étapes, une première instabilité linéaire se développant jusqu'à l'apparition d'un îlot de taille suffisante, qui déclencherait une seconde instabilité, non-linéaire et plus rapide. Une telle hypothèse repose sur l'existence d'un îlot magnétique plus large que la couche de reconnection ; le transport parallèle dans l'îlot pourrait alors devenir un transport perpendiculaire effectif, à cause du changement de topologie magnétique. Cet effet n'est pas pris en compte en théorie linéaire. Cependant, la faible épaisseur de la couche de reconnection (quelques millimètres) rend des tels îlots difficiles à détecter.

Dans cette hypothèse, la théorie linéaire déciderait du seuil d'instabilité (puisque le crash est toujours initié par la reconnection partielle du kink résistif), mais pas du temps de crash, qui relèverait d'une théorie non-linéaire. Notons cependant que jusqu'ici, le calcul d'un seuil théorique d'instabilité linéaire pour le déclenchement des dents de scie n'a pas abouti à des résultats convaincants, malgré l'ajout de nombreux effets, notamment bi-fluides.

L'hypothèse d'une instabilité non-linéaire évoquée ci-dessus a aussi été avancée pour résoudre le problème du déclenchement rapide. Durant la rampe, le plasma

évolue lentement sous l'effet du chauffage, et il passe par le seuil de stabilité marginale de l'instabilité. Si celle-ci est linéaire, le taux de croissance au seuil est par définition nul, et on devrait bien voir le mode croître, et non apparaître avec un taux déjà élevé. L'hypothèse d'une instabilité non-linéaire est donc logique. Il faut alors que le mode résistif soit stabilisé si le facteur de sécurité est en-dessous de l'unité ; des effets diamagnétiques peuvent avoir ce rôle.

2.3.5 Approche semi-empirique et questions ouvertes

Le modèle de Porcelli et Boucher

Numériquement, tous les effets évoqués ci-dessus ont fait l'objet de simulations en MHD étendue, mais sur seulement quelques cycles, pour des raisons de temps de calcul. Dans une optique plus empirique, Porcelli, Boucher et Rosenbluth ont établi en 1996 un modèle pour le scaling de la période et l'amplitude des dents de scie (cf. [39]). Dans ce modèle, la rampe est simulée par un code de transport ; le modèle fournit un critère de déclenchement (fondé sur un seuil de stabilité linéaire), et l'état final est calculé sur une base heuristique de reconnection incomplète (avec une instabilité secondaire qui se déclenche quand l'îlot atteint une taille critique). Le critère de déclenchement est triple : il y a crash si l'une des trois inégalités suivantes est vraie

$$\begin{aligned} -\delta W_{core} &> c_h \omega_{Dh} \tau_A \\ -\delta W &> 0.5 \omega_i^* \tau_A \\ \hat{s}_1 &> \hat{s}_{crit} \end{aligned}$$

Les notations sont les suivantes : $\delta W_{core} = \delta W_{MHD} + \delta W_{KO}$, $\delta W = \delta W_{core} + \delta W_{fast}$, où δW_{KO} est le terme de Kruskal-Oberman, contenant les effets dûs aux particules piégées (dominant quand $q \sim 1$), et δW_{fast} le terme dû aux particules rapides. ω_i^* est la fréquence diamagnétique ionique, et ω_{Dh} la fréquence de rebond des particules rapides piégées. La première inégalité signifie que les particules rapides piégées ne décrivent pas un grand nombre d'orbites dans le temps caractéristique de la perturbation. La seconde signifie que le kink n'est pas stabilisé par les effets diamagnétiques ioniques. La troisième est présentée ici sous une forme simplifiée (elle se décompose en fait en deux inégalités séparées) ; elle signifie que le mode n'est pas stabilisé par des effets cinétiques dans la couche de reconnection. Comme on le voit, l'argument est empirique : le critère de déclenchement est une juxtaposition de conditions portant sur des effets physiques différents.

Ce modèle a été très utilisé et discuté. Certaines expériences ont semblé en accord avec les critères de déclenchement, notamment dans JET, cf. [34], et dans TCV, cf. [35] et [36], montrant un accord qualitatif assez bon (erreur de l'ordre de 20% sur la période prévue). D'autres ont trouvé des désaccords, notamment sur ASDEX

Upgrade, cf. [38]. Ces cas ont prouvé que le calcul des termes individuels des critères est délicat, notamment dans le choix des constantes c_h , qui sont de l'ordre de l'unité et qui dépendent des paramètres du tokamak. Le modèle est globalement considéré comme fonctionnel.

Questions ouvertes

Il subsiste de nombreux problèmes ouverts sur les dents de scie. Citons en particulier :

- le scaling de la période, du temps de crash, et de l'amplitude avec les paramètres du tokamak ; d'une façon plus générale, le dimensionnement et l'existence même du régime des dents de scie
- l'accélération du crash
- la difficulté à établir un seuil de stabilité linéaire pour le déclenchement des dents de scie
- le déclenchement rapide
- la quantité de flux reconnecté pendant le crash
- le rôle de la stochasticité magnétique dans le crash

Les nombreux effets à prendre en compte et la diversité des phénomènes observés ont conduit la communauté scientifique à adopter une approche axée sur le contrôle semi-empirique. Le but n'est pas de comprendre parfaitement le phénomène, mais de pouvoir définir une stratégie crédible pour opérer les machines futures sans que les dents de scie soient fatales au confinement. On peut soit allonger la période des dents de scie pour qu'elle soit plus longue que la durée totale de la décharge, soit au contraire faire diminuer cette période pour garder une amplitude d'oscillation faible.

2.4 Conclusion

Après avoir cadré la théorie de la MHD et précisé son application aux tokamaks, on s'est intéressé en détail au cas des dents de scie, dans la perspective d'une étude numérique avec un modèle MHD étendue.

L'omniprésence des dents de scie dans les tokamaks a conduit à mettre en place des méthodes expérimentales pour contrôler leur stabilité (injection de neutres, etc). De par leur similitude avec les ELMs et leurs interactions avec les particules chaudes qu'on observera dans ITER (transition en fishbones ou stabilisation de la période pour donner des dents de scie monstres), l'étude des dents de scie est un enjeu important pour la fusion contrôlée. À l'heure actuelle, la compréhension théorique du phénomène est encore incomplète, et l'approche privilégiée procède plutôt d'un contrôle semi-empirique. Il existe un consensus sur l'influence de nombreux effets, mais pas sur un tableau d'ensemble, que la très grande variété des phénomènes

expérimentaux rend difficile à imaginer. Il est clair cependant que des effets hors MHD sont nécessaires pour expliquer, entre autres, la stabilisation du kink pendant la rampe et l'accélération brutale qui cause le crash : dérives diamagnétiques, effets cinétiques, effets non-linéaires.

Cette constatation appelle logiquement l'utilisation d'un code modulaire – idéalement, hybride MHD/cinétique – pour leur simulation, qui permettrait d'identifier les influences des différents effets physiques. Ainsi, l'objectif de l'étude menée avec le code XTOR, qui est en train d'opérer la transition vers un code hybride, est d'initier une modélisation intégrée du phénomène pour dépasser l'approche actuelle et se diriger vers une démarche plus prédictive.

3 Dynamique du kink interne avec XTOR

Dans ce chapitre, on présente une étude numérique de la dynamique à long terme du kink interne dans les tokamaks, menée avec le code MHD étendue XTOR. Nous cherchons à reproduire une dynamique oscillatoire comparable aux dents de scie, avec une longue phase de stabilité du kink pendant la rampe. On utilise la MHD résistive avec transport thermique, puis on inclut des effets diamagnétiques.

On commence par présenter le code utilisé : schéma numérique, solveur, pré-conditionneur, normalisation. Ensuite, on expose le cadre de l'étude effectuée, et son domaine de validité, qui correspond expérimentalement à des petits tokamaks ohmiques. On présente les différents régimes (oscillatoires ou non) de la dynamique du kink interne obtenus dans le cadre de la $\eta\chi$ MHD, puis on évoque l'influence des effets diamagnétiques sur ces régimes.

3.1 XTOR : un code MHD non-linéaire modulaire pour les tokamaks

XTOR est un code de MHD étendue dynamique 3D non-linéaire pour la physique des tokamaks²³. Initialement, XTOR utilisait un schéma semi-implicite (toujours utilisable, cf. [40]). Ce code a été utilisé pour modéliser des instabilités génériques (cf. [41],[42],[43]), et également pour simuler des décharges expérimentales sur TORE SUPRA (cf. [44], [45], [46]).

Le développement est maintenant centré autour d'un schéma implicite. Un schéma semi-implicite est optimal pour les problèmes se traduisant par un opérateur self-adjoint, comme la MHD idéale. L'ajout de nouveaux effets physiques, brisant cette symétrie, a conduit à mettre au point la version implicite, plus robuste et modulaire. Un schéma implicite assure une bonne stabilité, et des pas de temps importants. Cependant, le choix d'une méthode de résolution itérative conduit à optimiser soit le nombre d'itérations (comme c'est le cas dans XTOR), soit le pas de temps. La mise au point du schéma implicite a permis de faire évoluer le code MHD en ajoutant des effets bi-fluides (XTOR-2F, cf. [47]). Plus précisément, on a ajouté les termes de vitesses diamagnétiques, habituellement notés $\mathbf{v}_{i/e}^*$, qui correspondent à la prise en compte des dérivées de gradient de pression des particules²⁴. Ces effets viennent s'ajouter au modèle de base d'XTOR, contenant les équations de la MHD résistive (avec soit pression, soit densité et température) avec transport thermique anisotrope. Le code hybride XTOR-K, incluant une population de particules rapides modélisées à part, est en cours de tests.

²³Remarquons à ce sujet qu'XTOR est le plus beau programme du monde, il fait tout même le café.

²⁴On peut bien sûr toujours utiliser le code sans ces effets ; par conséquent, dans la suite on dira XTOR-2F pour désigner le code implicite, même lorsqu'on l'utilise en $\eta\chi$ MHD.

On commence par présenter le solveur utilisé et le nouveau schéma implicite numérique choisi. On présente le système d'équations et sa normalisation, puis le pré-conditionneur, crucial pour l'efficacité du code. On termine par la géométrie choisie, qui est la plus générale possible pour pouvoir modéliser des plasmas à sections poloïdales non circulaires, et les conditions aux limites.

3.1.1 Solveur et schéma numérique implicite

Les équations de la MHD dans les tokamaks constituent un problème très raide. En effet, les fréquences des ondes d'Alfvén couvrent une large gamme de valeurs, et les phénomènes étudiés sont lents devant le temps caractéristique du modèle (le temps d'Alfvén). On doit donc mener des simulations sur des temps longs devant le pas de temps numérique. Pour pouvoir augmenter ce pas de temps au maximum, le mieux est de choisir une méthode de résolution implicite. La résolution du problème implicite en géométrie toroïdale est cependant exigeante en termes de ressources, puisqu'elle implique d'inverser une grande matrice à chaque pas de temps²⁵. Le solveur itératif utilisé dans XTOR permet d'éviter cela en ne travaillant que sur le produit matrice×vecteur. Cependant, le choix d'un bon pré-conditionneur est alors crucial pour que le solveur soit efficace.

Discrétisation du système

D'une façon très générale, on cherche à résoudre un problème non-linéaire du type :

$$\dot{x} = F(x)$$

x est un vecteur qui contient les variables du modèle (dans notre cas, $\mathbf{v}$, $\mathbf{B}$, p , et éventuellement ρ), et F est l'opérateur MHD, auquel s'ajoutent divers effets physiques. On discrétise le problème ainsi :

$$x_{n+1} - x_n = \Delta t F \left(\frac{x_{n+1} + x_n}{2} + c(x_{n+1} - 2x_n + x_{n-1}) \right) \quad (20)$$

Toutes les variables sont traitées de la même façon, ce qui rend très simple l'ajout de nouveaux effets physiques. Le schéma est donc fortement modulaire.

Le paramètre c permet de rendre le schéma plus ou moins implicite : si $c = 0$, on obtient un schéma réversible. En général, on prendra $c > 0$ (typiquement $c = 1$), ce qui permet de conserver une précision d'ordre 2 en temps, et de d'amortir les hautes fréquences non résolues.

²⁵Pour cette raison, dans les années 90, de nombreuses méthodes semi-implicites furent créées, notamment pour XTOR.

Stabilité linéaire

Pour déterminer les propriétés de stabilité linéaire du schéma, on écrit que le vecteur x est vecteur propre de la matrice de transfert linéaire, avec la valeur propre $i\omega$. En posant $q = \omega\Delta t$, on a :

$$x_{n+1} - x_n = iq \left(\frac{x_{n+1} + x_n}{2} + c(x_{n+1} - 2x_n + x_{n-1}) \right)$$

En posant le facteur d'amplification $\lambda = x_{n+1}/x_n$, on obtient l'équation caractéristique :

$$\lambda(\lambda - 1) = iq \left(\frac{\lambda(\lambda + 1)}{2} + c(\lambda - 1)^2 \right)$$

On constate que le module des solutions de l'équation est inférieur à 1 pour tout q si $c > 0$. Autrement dit, quand $c > 0$, le schéma est linéairement inconditionnellement stable. Le cas $c = 0$ correspond à la stabilité marginale ($\|\lambda\| = 1$).

Solveur Newton-Krylov

La façon la plus directe de résoudre le problème serait de discrétiser le système à chaque pas de temps, de construire la matrice correspondante, et de l'inverser avec une méthode numérique. Cette méthode, bien que robuste, est sans doute intrac-table avec les ordinateurs actuels. XTOR utilise une méthode itérative pour résoudre le système à chaque pas de temps : l'algorithme Newton-Krylov. Cette méthode travaille directement sur le produit matrice×vecteur, ce qui évite de devoir construire la matrice.

À chaque pas de temps, il faut inverser le système :

$$x_{n+1} - x_n - \Delta t F \left(\frac{x_{n+1} + x_n}{2} + c(x_{n+1} - 2x_n + x_{n-1}) \right) = 0$$

pour trouver la valeur de x_{n+1} . On introduit $\Delta_n = x_{n+1} - x_n$ et $\bar{x} = \left(\frac{1}{2} - 2c\right) x_n + c x_{n-1}$; le système s'écrit de façon plus compacte :

$$G(\Delta_n, \bar{x}) \equiv \Delta_n - \Delta t F \left[\left(\frac{1}{2} + c \right) \Delta_n + \bar{x} \right] = 0$$

Formellement, il faut trouver le zéro de la fonction G de la variable Δ_n à chaque pas de temps t_n . Comme indiqué plus haut, on utilise une méthode de Newton-Krylov avec pré-conditionnement, *i.e.*, on résout numériquement le système itératif :

$$M^{-1}G'(\bar{x})\Delta_n^{k+1} - M^{-1}G(\Delta_n^k, \bar{x}) = 0$$

Où M est le pré-conditionneur, k est l'index de l'itération de la méthode de Newton, et G' est la dérivée de G par rapport à Δ_n . À chaque itération de Newton, l'équation ci-dessus est résolue itérativement par une méthode de Krylov, et Δ_n^{k+1} est calculé en fonction de l'ancien Δ_n^k . On itère jusqu'à ce qu'une condition de convergence soit atteinte (soit $\|\Delta_n^{k+1} - \Delta_n^k\| < \epsilon$, soit $\|G(\Delta_n^k, \bar{x})\| < \eta$).

On utilise le package NITSOL (Newton Iterative solveur, cf. [48]) qui implémente cette méthode. Ce solveur offre le choix entre différentes méthodes pour l'étape Krylov : General Minimum Residual (GMRES), Bi-Conjugate Gradient Stabilized (BiCGSTAB), et Transpose-Free Quasi-Minimum Residual (TFQMR). Dans notre cas, on utilise la méthode GMRES, qui s'est avérée la plus rapide.

La librairie NITSOL permet aussi d'éviter l'évaluation analytique de la matrice jacobienne $G'(\bar{x})$ à chaque pas de temps implicite : cette matrice est interpolée numériquement, par différentiation (méthode « matrix-free »). Cela évite de devoir linéariser le système (mais dans notre cas, le calcul du pré-conditionneur l'exigera partiellement).

Le choix d'un pré-conditionneur adéquat s'avère crucial pour l'efficacité de la méthode ; dans notre cas, le préconditionneur choisi est lié au problème physique, et sera présenté dans la section suivante.

Typiquement, le solveur effectue une dizaine d'itérations de Krylov, et une ou deux itérations de Newton par pas de temps.

3.1.2 Systèmes d'équations

L'option de base de XTOR, basé sur le modèle de la $\eta\chi$ MHD, résoud le système suivant :

$$\rho \partial_t \mathbf{v} = -\rho(\mathbf{v} \cdot \nabla) \mathbf{v} + \mathbf{J} \times \mathbf{B} - \nabla p + \nu \nabla^2 \mathbf{v} \quad (21)$$

$$\partial_t \mathbf{B} = \nabla \times (\mathbf{v} \times \mathbf{B}) - \nabla \times \eta(\mathbf{J} - \mathbf{J}_{\text{boot}}) \quad (22)$$

$$\partial_t p = -\Gamma p \nabla \cdot \mathbf{v} - \mathbf{v} \cdot \nabla p + \nabla \cdot \chi_{\perp} \nabla p + \nabla \cdot \left[\mathbf{B} \left(\chi_{\parallel} \frac{(\mathbf{B} \cdot \nabla) p}{B^2} \right) \right] + \Theta \quad (23)$$

Ce système est directement issu des équations de la MHD résistive données au premier chapitre (9 - 19), avec toutefois les distinctions suivantes :

- Le terme de viscosité $\nu \nabla^2 \mathbf{v}$ a été inclus
- Le courant de bootstrap J_{boot} apparaît explicitement
- Le flux de chaleur est modélisé par un transport diffusif anisotrope de la pression, commandé par les coefficients $\chi_{\perp}$ et $\chi_{\parallel}$.

- On a fait l'hypothèse que $\rho = cte$.

La version bi-fluide du code, **XTOR-2F**, inclut les effets des dérives diamagnétiques en plus de la $\eta\chi$ MHD. Les équations résolues par **XTOR-2F** forment le système :

$$\rho \partial_t \mathbf{v} = -\rho(\mathbf{v} \cdot \nabla) \mathbf{v} - \rho(\mathbf{v}_i^* \cdot \nabla) \mathbf{v}_\perp + \mathbf{J} \times \mathbf{B} - \nabla p + \nu \nabla^2 \mathbf{v} \quad (24)$$

$$\partial_t \mathbf{B} = \nabla \times (\mathbf{v} \times \mathbf{B}) + \alpha \nabla \times \frac{\nabla_\parallel p_e}{\rho} - \nabla \times \eta(\mathbf{J} - \mathbf{J}_{\text{boot}}) \quad (25)$$

$$\begin{aligned} \partial_t p = & -\Gamma p \nabla \cdot \mathbf{v} - \mathbf{v} \cdot \nabla p - \alpha \Gamma \frac{p_i}{\rho} \nabla p \cdot \nabla \times \frac{\mathbf{B}}{B^2} + \nabla \cdot \chi_\perp \nabla p + \\ & \nabla \cdot \left[\mathbf{B} \left(\chi_\parallel \frac{(\mathbf{B} \cdot \nabla) p}{B^2} \right) \right] + \Theta \end{aligned} \quad (26)$$

$$\partial_t \rho = -\rho \nabla \cdot \mathbf{v} - \mathbf{v} \cdot \nabla \rho - \alpha \nabla p_i \cdot \nabla \times \left(\frac{\mathbf{B}}{B^2} \right) + \nabla \cdot D_\perp \nabla \rho + \Sigma \quad (27)$$

Notons que dans ces systèmes, $\mathbf{v}$ est en réalité un développement au premier ordre de la vitesse fluide en $\delta = \rho_{Li}/a$, ρ_{Li} étant le rayon cyclotronique ionique ; nous l'explicitons dans le paragraphe suivant. Le courant de bootstrap est donné par une formule simplifiée : $\mathbf{J}_{\text{boot}}(\mathbf{t}) = f_{bs} \|\mathbf{J}_{\text{boot},0}\| (\nabla p(t))_r / p'_0 \mathbf{B} / \|\mathbf{B}\|$. Dans le second système, D est un terme de diffusion de matière. Les indices $\parallel$ et $\perp$ font référence à la direction de $\mathbf{B}$, les indices i et e aux ions et aux électrons. Θ et Σ sont des termes de source de chaleur et de densité. Ils peuvent prendre la forme $\Theta = -\nabla \cdot \chi_\perp \nabla p_0$, $\Sigma = -\nabla \cdot D_\perp \nabla \rho_0$; dans ce cas ils tendent à restaurer les profils initiaux de pression et de densité (avec les temps caractéristiques de diffusion $a^2/\chi_\perp$ et $a^2/D_\perp$) ; Θ peut aussi s'écrire $\Theta = \eta J^2$, pour modéliser une source ohmique.

Le profil initial de η est obtenu en imposant une condition aux limites sur le champ électrique toroïdal : $E_{\phi,0} = cte = \eta_0 (\langle \mathbf{J}_{\phi,0} - \mathbf{J}_{\text{boot},\phi,0} \rangle)$, où $\langle \dots \rangle$ désigne la moyenne le long d'une surface de flux d'équilibre. On effectue ensuite un mapping entre résistivité et température de façon à vérifier la loi de Spitzer $\eta \propto T^{-3/2}$ (ou $p^{-3/2}$ si T est prise constante). Au cours du run, on peut soit garder le profil de résistivité constant, soit refaire le mapping à chaque nouveau pas de temps. En général, on prend $\chi_\perp$ proportionnel à η , tandis que le profil de $\chi_\parallel$ est constant au cours du temps.

Normalisations

La normalisation utilisée dans **XTOR** est la suivante :

- l'unité de *longueur* est le petit rayon, a . Par conséquent, dans nos unités, $A = R_0/a = R_0$.

- l'unité de *champ magnétique* B_u est choisie de façon à ce que $B(R_0) = B(A) = A = B_0/B_u$. Par conséquent, $B_u = B_0/A$.
- l'unité de *densité* est ρ_0 , la densité sur l'axe à $t = 0$. ρ_0 , comme B_0 , est un paramètre libre permettant de rescaler la densité.
- l'unité de *vitesse* est $v_A = B_u/(\mu_0\rho_0)$.
- l'unité de *temps* est le temps d'Alfvén, $\tau_A = a/v_A = \sqrt{\mu_0\rho_0}A/B_0$.
- l'unité de *courant* est $J_u = B_u/(\mu_0a)$, pour éviter les constantes μ_0 dans les calculs.
- l'unité de *champ électrique* est v_AB_0/A^2 .
- l'unité de *pression* est $\rho_0v_A^2$
- l'unité de *résistivité* est μ_0av_A

Remarquons que dans ces unités, le nombre de Lundquist S vaut $S = \eta^{-1}$ (ces deux nombres sont sans dimensions). Enfin, il reste deux paramètres libres, ρ_0 et B_0 . Ils représentent respectivement la densité et le champ magnétique initiaux sur l'axe magnétique.

Le coefficient α est un paramètre sans dimension qui apparaît devant tous les termes contenant les effets diamagnétiques du modèle lorsqu'on passe en normalisation XTOR. Il mesure l'importance de ces effets. Si on prend $\alpha = 0$, on retrouve le modèle de la $\eta\chi$ MHD. Notons qu'on peut choisir d'introduire exclusivement les termes diamagnétiques ioniques ou électroniques. α s'exprime comme le rapport de deux temps caractéristiques du système :

$$\alpha = (\omega_{ci}\tau_A)^{-1}$$

Où ω_{ci} est la fréquence cyclotronique ionique, dans les unités du système : $\omega_{ci} = eB_0/(Am_p)$. Insistons sur le fait que la pulsation cyclotronique est mesurée ici dans l'unité de champ magnétique du système, et non dans le champ toroïdal. Par conséquent, $\alpha = m_p/e(\mu_0\rho_0)^{-1/2}$. Fixer α revient donc à choisir ρ_0 .

Dérivation du système bi-fluide

On détaille ici le calcul du système comprenant les effets diamagnétiques. Nous insistons sur le fait que même si on a gardé les mêmes variables qu'en $\eta\chi$ MHD, ce modèle est bien bi-fluide, et non mono-fluide. Dans notre représentation, les vitesses ioniques et électroniques sont développées de la façon suivante :

$$\begin{aligned} \mathbf{v}_i &= \mathbf{v} + \mathbf{v}_i^* \\ \mathbf{v}_e &= \mathbf{v} + \mathbf{v}_e^* - \frac{\mathbf{J}_{\parallel}}{en_e} \end{aligned}$$

Avec :

$$\begin{aligned}\mathbf{v} &= \frac{\mathbf{E} \times \mathbf{B}}{B^2} + \mathbf{v}_{\parallel \mathbf{i}} \\ \mathbf{v}_s^* &= \frac{\mathbf{B} \times \nabla p_s}{q_s n_s B^2} = \frac{\boldsymbol{\beta} \times \nabla p_s}{q_s n_s}\end{aligned}$$

On a introduit la notation $\boldsymbol{\beta} \equiv \mathbf{B}/B^2$, à ne pas confondre avec le paramètre macroscopique β . Remarquons bien qu'avec ces définitions, on n'a pas posé $\mathbf{v} = (\rho_i \mathbf{v}_i + \rho_e \mathbf{v}_e)/(\rho_i + \rho_e)$, *i.e.*, $\mathbf{v}$ n'est pas définie comme étant la vitesse fluide²⁶. Remarquons de plus que la partie perpendiculaire du courant est entièrement représentée par le courant diamagnétique : $\mathbf{J}_\perp = \mathbf{J}^* = Zen(\mathbf{v}_i^* - \mathbf{v}_e^*)$, qui est dû aux dérives de gradient et de courbure (Z est la charge des ions). Les vitesses $\mathbf{v}^*$ sont les vitesses de dérives diamagnétiques. De cette façon, on exclut naturellement tous les effets d'ordre plus élevés que les dérives diamagnétiques.

Pour arriver aux équations 24 - 27, on part des équations fluides ioniques et électroniques (conservation de la quantité de mouvement, 7, de la quantité de matière, 6, et de l'énergie, 8). On fait le changement de variables pour passer à la représentation $(\mathbf{v}, \mathbf{v}_e^*, \mathbf{v}_i^*)$. Ce calcul est également donné dans [49].

On introduit la notation $d/dt \equiv \partial/\partial t + \mathbf{v} \cdot \nabla$.

L'équation de la conservation de la masse pour les ions (6) s'écrit avec nos notations :

$$\begin{aligned}\frac{dn_i}{dt} + n_i \nabla \cdot \mathbf{v} + n_i \nabla \cdot \mathbf{v}_i^* &= 0 \\ \Leftrightarrow \frac{dn_i}{dt} + n_i \nabla \cdot \mathbf{v} + \nabla \cdot \frac{\boldsymbol{\beta} \nabla p}{q_i} &= 0 \\ \Leftrightarrow \frac{dn_i}{dt} + n_i \nabla \cdot \mathbf{v} + \boldsymbol{\kappa}_i \cdot \nabla p &= 0\end{aligned}$$

Où on a introduit $\boldsymbol{\kappa}_s \equiv 1/q_s \nabla \times \boldsymbol{\beta}$. En considérant que $m_e \ll m_i$, nous négligeons l'équation similaire pour les électrons. Avec nos normalisations, un facteur α apparaît devant le troisième terme ; en ajoutant les termes de source et de transport, on obtient bien l'équation 27.

De même, l'équation de la conservation de la quantité de mouvement s'écrit (7) :

$$\rho_i \frac{\partial \mathbf{v}}{\partial t} + \rho_i ((\mathbf{v} + \mathbf{v}_i^*) \cdot \nabla) \mathbf{v} = q_i n_i (\mathbf{E} + \mathbf{v} \times \mathbf{B}) - \nabla p_i$$

²⁶En fait, au lieu de chercher l'expression de la vitesse fluide à partir des équations, on introduit une variable $\mathbf{v}$ correspondant au premier ordre de cette vitesse en ρ_{Li}/a , et on établit le système autour de cette variable.

Où on a utilisé la « diamagnetic cancellation », cf. par exemple [2], chap. 6 : $d\mathbf{v}_i^*/dt + \nabla \cdot \Pi_i \approx \nabla p_i$, et $(\mathbf{v}_i^* \cdot \nabla)\mathbf{v}_{\parallel i}$ s'annule avec certaines composantes de $\nabla \cdot P_i$. Pour les électrons,

$$\rho_e \frac{\partial \mathbf{v}}{\partial t} + \rho_e \left((\mathbf{v} + \mathbf{v}_e^* - \frac{\mathbf{J}_{\parallel}}{en_e}) \cdot \nabla \right) \mathbf{v} = -en_e (\mathbf{E} + \mathbf{v}_e \times \mathbf{B}) - \nabla p_e$$

Et l'addition des deux fournit bien l'équation 24, en ajoutant le terme de viscosité. On doit utiliser la quasi-neutralité, et négliger le terme en $\mathbf{J}_{\parallel}$; cette hypothèse, qui revient à négliger la dérive de polarisation, reviendra également dans l'équation de pression.

Pour arriver à l'équation de Faraday, nous partons de la conservation de la quantité de mouvement électronique 12 (sans les termes anisotropes) :

$$\mathbf{E} = -\mathbf{v}_e \times \mathbf{B} + \eta \mathbf{J} - \frac{\nabla p_e}{n_e e} - \frac{m_e}{e} \frac{d\mathbf{v}_e}{dt}$$

Soit avec nos notations :

$$\mathbf{E} = -\mathbf{v} \times \mathbf{B} - \mathbf{v}_e^* \times \mathbf{B} + \underbrace{\frac{\mathbf{J}_{\parallel} \times \mathbf{B}}{n_e e}}_0 + \eta \mathbf{J} - \frac{\nabla p_e}{n_e e} - \frac{m_e}{e} \frac{d\mathbf{v}_e}{dt}$$

Or :

$$\begin{aligned} \mathbf{v}_e^* \times \mathbf{B} &= -\frac{\mathbf{B} \times \nabla p_e}{n_e e B^2} \times \mathbf{B} = -\frac{\nabla_{\perp} p_e}{n_e e}, \text{ donc :} \\ -\mathbf{v}_e^* \times \mathbf{B} - \frac{\nabla p_e}{n_e e} &= -\frac{\nabla_{\parallel} p_e}{n_e e} \end{aligned}$$

Et par conséquent :

$$\mathbf{E} = -\mathbf{v}_{\perp} \times \mathbf{B} - \frac{\nabla_{\parallel} p_e}{n_e e} + \eta \mathbf{J} - \frac{m_e}{e} \frac{d\mathbf{v}_e}{dt}$$

Il s'agit de notre loi d'Ohm généralisée. À droite du signe égal se trouvent plusieurs termes. Le degré de précision du modèle est défini par le choix de ceux qu'on néglige et de ceux qu'on conserve. Il faut que le degré de précision soit cohérent avec l'expression choisie des vitesses, *i.e.*, il faut garder jusqu'aux termes dus aux dérivées diamagnétiques.

– Le premier terme correspond à la MHD idéale; il faut bien sûr le conserver.

- Le deuxième terme correspond à la compressibilité électronique. Il peut aussi s'écrire $\mathbf{J}_\perp \times \mathbf{B}/(en_e)$, avec $\mathbf{J}_\perp = -\nabla p \times \mathbf{B}/B^2$. $\mathbf{J}_\perp$ est le courant dû aux dérives diamagnétiques. Il convient donc de conserver ce terme.
- Le troisième terme est le terme résistif. Il convient de le garder également, puisque c'est lui qui introduit la résistivité dans le système.
- Le quatrième terme est un terme d'inertie électronique. On le néglige pour le moment devant le terme résistif. Ce terme peut avoir une influence importante, notamment pendant le processus de reconnection. Le négliger revient à s'intéresser à des longueurs inférieures à la profondeur de peau électronique et au rayon de Larmor électronique, ainsi qu'à des fréquences petites devant la fréquence plasma électronique ($\omega_e = ne^2/(m_e\epsilon_0)$) et devant la fréquence cyclotronique électronique.

On obtient une expression de $\mathbf{E}$ qu'il suffit d'introduire dans Maxwell-Faraday pour obtenir l'équation 25²⁷. Enfin, reste à établir l'équation de la pression. Cette équation est la plus délicate, puisqu'elle vient des équations fluides 8, avec cette fois des flux de chaleurs différents pour les électrons et les ions. On relie ces flux à la variation de l'entropie :

$$n_s \frac{\partial S_s}{\partial t} + n_s \mathbf{v}_s \cdot \nabla S + \frac{\nabla \cdot \mathbf{Q}_s}{T_s} = 0$$

Avec :

$$S = \ln \frac{T_s^{3/2}}{n_s}, \text{ et } \mathbf{Q}_s = \frac{5}{2} \frac{p_s}{q_s} \boldsymbol{\beta} \times \nabla T_s$$

Dans notre représentation de vitesse, on obtient, pour les ions :

$$n_i \frac{dS_i}{dt} + n_i \mathbf{v}_i^* \cdot \nabla S_i + \nabla \cdot \frac{\mathbf{Q}_s}{T_s} = 0$$

Après simplifications, en notant que $p_i \equiv n_i T_i$ et grâce à $\mathbf{v}_i^* \cdot \nabla p_i = 0$, on obtient que :

²⁷En général, on dérive la loi d'Ohm généralisée sans faire d'hypothèse sur la forme de la vitesse, comme dans par exemple [2], p. 246. On voit alors apparaître un terme dit de Hall, sous la forme $\mathbf{J} \times \mathbf{B}/(ne)$. Notre représentation conduit à poser $\mathbf{J}_\perp = Zen(\mathbf{v}_i^* - \mathbf{v}_e^*) = -\nabla p \times \mathbf{B}/B^2$. En injectant ce terme dans le terme de Hall, on retrouve bien le terme diamagnétique $\nabla_\parallel p_e/(ne)$ qu'on a dans notre loi d'Ohm. Autrement dit, notre représentation, qui consiste à considérer jusqu'au dérives diamagnétiques, revient à réduire le terme de Hall à un terme de compressibilité électronique. Cette représentation de vitesse exclut les dérives de polarisation.

$$n_i \frac{dS_i}{dt} + \frac{5}{2} \frac{p_i}{T_i} \boldsymbol{\kappa}_i \cdot \nabla T_i = 0 \quad (28)$$

D'une façon analogue, on obtient pour les électrons :

$$n_e \frac{dS_e}{dt} - \frac{5}{2} \frac{p_e}{T_e} \boldsymbol{\kappa}_e \cdot \nabla T_e + \frac{5}{2} \frac{\mathbf{J}_{\parallel} \cdot \nabla_{\parallel} T_e}{q_e T_e} = \frac{\mathbf{J}_{\parallel} \cdot \nabla_{\parallel} p_e}{q_e n_e} = \frac{\mathbf{J}_{\parallel} \cdot \mathbf{E}_{\parallel}}{T_e} \quad (29)$$

On s'est servi de l'équation 25 pour $\nabla_{\parallel} p_e$; les termes en $\mathbf{J}_{\parallel}$ viennent de l'expression de $\mathbf{v}_e$.

Nous écrivons ensuite l'équation de la température (fournie par la définition de l'entropie) :

$$\frac{3}{2} \frac{n_s}{T_s} \frac{dT_s}{dt} = n_s \frac{dS_s}{dt} + \frac{dn_s}{dt} \quad (30)$$

Et nous injectons 28/29 et 30 dans la relation (fournie par la définition de la pression) :

$$\frac{dp_s}{dt} = n_s \frac{dT_s}{dt} + T_s \frac{dn_s}{dt} = \Gamma T_s \frac{dn_s}{dt} + \frac{2}{3} n_s T_s \frac{dS_s}{dt}$$

Qui nous fournit finalement les équations de pression bi-fluides :

$$\begin{aligned} \frac{dp_i}{dt} + \Gamma p_i \nabla \cdot \mathbf{v} + \Gamma T_i \boldsymbol{\kappa}_i \cdot \nabla p_i + \Gamma p_i \boldsymbol{\kappa}_i \cdot \nabla T_i &= 0 \\ \frac{dp_e}{dt} + \Gamma p_e \nabla \cdot \mathbf{v} - \Gamma T_e \boldsymbol{\kappa}_e \cdot \nabla p_e + \Gamma p_e \boldsymbol{\kappa}_e \cdot \nabla T_e - \frac{\Gamma}{q_e} \mathbf{J}_{\parallel} \cdot \nabla_{\parallel} T_e &= \mathbf{E}_{\parallel} \cdot \mathbf{J}_{\parallel} \end{aligned}$$

On additionne ces équations avec les hypothèses suivantes : $T_e = T_i$, $p_e = p_i = p/2$, $\nabla_{\parallel} T_e = 0$ et $\mathbf{J}_{\parallel} = 0$ (c'est cette dernière qui exclut les dérives de polarisation) pour obtenir :

$$\partial_t p = -\Gamma p \nabla \cdot \mathbf{v} - \mathbf{v} \cdot \nabla p - \frac{\Gamma p}{2 \rho} \nabla p \cdot \nabla \times \left(\frac{\mathbf{B}}{B^2} \right)$$

Comme précédemment, la normalisation introduit le facteur α devant le dernier terme. L'ajout des sources et du transport redonne l'équation 26.

3.1.3 Opérateurs MHD et pré-conditionneur

Le choix du pré-conditionneur affecte les performances du solveur d'une façon extrêmement importante, en réduisant le nombre d'itérations du solveur à chaque pas de temps. Résoudre le système itérativement sans le pré-conditionner est en général aussi coûteux (en temps CPU) que d'utiliser une méthode explicite. Dans le cas d'XTOR, on fait le choix d'un pré-conditionneur physique, c'est-à-dire que la matrice d'opérateurs M^{-1} définie plus haut est liée au système d'équations qu'on veut résoudre. Plus précisément, cette matrice est basée sur une linéarisation du système d'équations. Trois pré-conditionneurs sont disponibles dans XTOR-2F :

- Le pré-conditionneur le simple, utilisé initialement dans XTOR,
- Un pré-conditionneur plus complexe, utilisé pour les cas verticalement symétriques, qui contient des termes de résistivité et de transport,
- Le pré-conditionneur le plus général, qui inclut les termes bi-fluides et des cas asymétriques.

Les deux dernières options ont été implémentées récemment, et leurs performances sont en cours de test. Pour cette raison, on présente plus en détail le pré-conditionneur standard (première option).

Dans le cas le plus simple, le pré-conditionneur s'écrit sous la forme d'une matrice d'opérateurs 8×8 , et d'un produit :

$$M^{-1} = M_{\eta}^{-1} M_D^{-1} M_{\chi}^{-1} M_{MHD}^{-1}$$

Où :

$$\begin{aligned} M_{MHD} &= (L_{MHD}, 1) \\ M_{\eta} &= (1, L_{\eta}, 1, 1) \\ M_{\chi} &= (1, 1, L_{\chi}, 1) \\ M_D &= (1, 1, 1, L_D) \end{aligned}$$

Les opérateurs $L_{MHD, \eta, \chi, D}$ sont les linéarisations des opérateurs MHD idéale, résistive, de transport thermique, et de transport de matière. Ils s'appliquent respectivement aux variables $(\mathbf{v}, \mathbf{B}, p)$, $\mathbf{B}$, p , et ρ , d'où les expressions des opérateurs $M_{MHD, \eta, \chi, D}$, qui s'appliquent au vecteur contenant toutes les variables. Notons bien que malgré ces expressions, tous les M sont bien des matrices de dimension 8×8 . En effet, les deux premières colonnes (dans par exemple $M_{\chi} = (1, 1, L_{\chi}, 1)$) représentent les variables $\mathbf{v}$ et $\mathbf{B}$, qui sont de dimension 3. Remarquons également que pour traiter le problème le plus général possible, nous utilisons une représentation co- et contravariante des champs, avec des termes de métrique non triviaux. En prenant $\tilde{c} = c + 1/2$,

et en indiquant les quantités d'équilibre (autour desquelles on linéarise) par l'indice 0, on écrit :

$$\begin{aligned} L_\chi &= 1 - \tilde{c}\Delta t \left(\nabla \cdot \chi_\perp^0 \nabla + \nabla \cdot \left[\mathbf{B}_0 \left(\frac{\chi_\parallel^0}{B_0^2} (\mathbf{B}_0 \cdot \nabla) \right) \right] \right) \\ L_D &= 1 - \tilde{c}\Delta t \nabla \cdot D_\perp^0 \nabla \\ L_\eta &= 1 - \tilde{c}\Delta t \nabla \times \eta^0 \nabla \times \end{aligned}$$

Et :

$$L_{MHD} = \begin{pmatrix} g_{ij}^0 - \tilde{c}\Delta t (\nabla \nu^0 \nabla) & -\tilde{c}\Delta t \mathbf{L}_v^B & -\tilde{c}\Delta t \mathbf{L}_v^p \\ -\tilde{c}\Delta t \mathbf{L}_B & 1 & 0 \\ \tilde{c}\Delta t L_p & 0 & 1 \end{pmatrix}$$

Où les opérateurs $\mathbf{L}_v^{B/p}$, $\mathbf{L}_B$ et L_p sont les linéarisations des opérateurs $\mathbf{F}_v$, $\mathbf{F}_B$, F_p , qui forment les termes de droite des équations 24-26. Par conséquent :

$$\begin{aligned} L_v^B \delta \mathbf{B} &= (\nabla \times \mathbf{B}_0) \times \delta \mathbf{B} + (\nabla \times \delta \mathbf{B}) \times \mathbf{B}_0 \\ L_v^p \delta p &= -\nabla(\delta p) \\ L_B \delta \mathbf{v} &= \nabla \times (\delta \mathbf{v} \times \mathbf{B}_0) \\ L_p \delta \mathbf{v} &= -\delta \mathbf{v} \cdot \nabla p_0 - \Gamma p_0 \nabla \cdot \delta \mathbf{v} \end{aligned}$$

g_{ij}^0 est le tenseur métrique associé à l'équilibre. En inversant la relation qui définit L_{MHD} :

$$L_{MHD} \begin{pmatrix} \delta \mathbf{v} \\ \delta \mathbf{B} \\ \delta p \end{pmatrix} = \begin{pmatrix} \mathbf{F}_v \\ \mathbf{F}_B \\ F_p \end{pmatrix}$$

on obtient :

$$\begin{aligned} \delta \mathbf{v} &= L_0^{-1} \left(\mathbf{F}_v + \tilde{c}\Delta t [L_v^B(\mathbf{F}_B) + L_v^p(F_p)] \right) \\ \delta \mathbf{B} &= \mathbf{F}_B + \tilde{c}\Delta t L_B(\delta \mathbf{v}) \\ \delta p &= F_p + \tilde{c}\Delta t L_p(\delta \mathbf{v}) \end{aligned}$$

Où on a défini l'opérateur L_0 :

$$L_0 = g_{ij}^0 - \tilde{c}\Delta t(\nabla\nu^0\nabla) - (\tilde{c}\Delta t)^2(L_v^B L_B + L_v^p L_p)$$

Un opérateur très proche ($L_0 = g_{ij}^0 - c\nabla^2 - \Delta t_0^2(L_v^B L_B + L_v^p L_p)$) était utilisé dans la version semi-implicite d'XTOR pour stabiliser le schéma (c étant une constante numérique choisie par l'utilisateur). Lors du passage au schéma implicite, cet opérateur a été inclus dans le pré-conditionneur. On inverse les matrices L_{MHD} , L_η , L_χ et L_D par une décomposition LU par blocs. La plus grande matrice, L_{MHD} , est bloc penta-diagonale.

Contrairement au code semi-implicite, dans lequel les parties transport (thermique et de matière) et résistivité des équations étaient avancées séparément du reste du système, le code implicite inclut toutes les contributions dans le solveur Newton-Krylov.

3.1.4 Géométrie et conditions aux limites

Maillages et métriques issus de CHEASE

Les choix géométriques sont les mêmes que pour la version XTOR semi-implicite ; ils sont détaillés dans [40]. Géométriquement, XTOR utilise une représentation semi-spectrale. Dans la direction radiale, la coordonnée r est liée aux surfaces du flux magnétique poloïdal à l'équilibre, Ψ . Ce flux est issu de l'équilibre 2D calculé par CHEASE (voir [51] pour une présentation détaillée) qui est utilisé comme condition initiale par XTOR-2F. CHEASE résout l'équation de Grad-Shafranov :

$$\nabla \cdot \frac{1}{R^2} \nabla \Psi = \frac{j_\phi}{R} = -p'(\Psi) - \frac{1}{R^2} F F'(\Psi)$$

Où j_ϕ est la densité de courant toroïdale, R est le grand rayon, et F est le flux magnétique toroïdal. La coordonnée radiale utilisée par XTOR est alors définie par²⁸ : $r = \sqrt{\Psi/\Psi_{edge}}$. Ψ_{edge} est le flux au bord du plasma. Les dérivées dans cette direction sont calculées par différences finies. En réalité, XTOR utilise deux grilles imbriquées en r , notées (r_l) (grille entière) et $(r_{l+1/2})$ (grille demi-entière). Certains champs sont évalués sur la grille entière ($v_r, B_r, J_\theta, J_\phi, \eta$), d'autres sur la grille demi-entière ($v_\theta, v_\phi, B_\theta, B_\phi, J_r, p, T, \rho$). La dérivée d'une grandeur exprimée sur la grille entière est exprimée sur la grille demi-entière, et réciproquement. Le nombre de points est fixé par le paramètre $lmax$.

L'angle θ est l'angle poloïdal, et l'angle ϕ est l'angle toroïdal. Pour ces deux coordonnées, les grilles sont équi-angulaires. Les dimensions des grilles poloïdale et toroïdale sont fixées par les paramètres $mmax$ et $nmax$, respectivement. Pour les dimensions angulaires, les calculs, notamment les dérivées, sont effectués dans l'espace

²⁸CHEASE transmet également à XTOR les champs p, p', F, FF' , et r'

de Fourier. L'avantage d'utiliser l'angle polaire pour θ , par rapport à un système de coordonnées où les lignes de champs sont des droites, est que les termes de la métrique ne dépendent pas de la dérivée seconde de Ψ , ce qui augmente leur précision. Notons que le tenseur métrique est non trivial dans le plan poloïdal (*i.e.*, non diagonal). Il est nécessaire de prendre en compte le caractère co- ou contra-variant des composantes des champs. Le tenseur métrique intervient dans le calcul des opérateurs utilisés par le solveur. Il est fourni en sortie de **CHEASE** avec le reste de l'équilibre. Plus précisément, **XTOR** nécessite les grandeurs $g^{rr}, g^{r\theta}, g^{\theta\theta}, g^{\phi\phi}, D^{29}$, sur les grilles entières et demi-entières. D est le réciproque du Jacobien, $D = [\nabla r, \nabla\theta, \nabla\phi]$ et $D^2 = g_{ij}g^{ij}$.

Structure des modes

Il est important de signaler qu'on applique régulièrement un projecteur après les opérations algébriques sur les champs, qui a pour effet de projeter le champ sélectionné sur l'ensemble de modes poloïdaux et toroïdaux avancés par **XTOR-2F**. En général, cet ensemble contient quelques dizaines de modes toroïdaux au maximum, et quelques dizaines de modes poloïdaux par mode toroïdal, au maximum. Ces modes sont choisis pour fournir une résolution suffisante pour le problème étudié. Grâce à ces projections, les matrices de préconditionnement restent relativement petites – sauf dans le cas du préconditionneur total, avec asymétrie haut-bas –, par rapport à celles qui seraient nécessaires pour résoudre le même problème avec des différences finies dans le plan poloïdal. La décomposition LU mentionnée plus haut constitue donc un choix raisonnable en terme de temps de calcul.

Le projecteur sert aussi à empêcher l'aliasing dans les directions poloïdale et toroïdale. Pour cette raison, on prévoit toujours des tableaux 3/2 fois plus grands que les modes effectivement modélisés. Le nombre de modes toroïdaux modélisé est $nsmax$; pour chaque mode toroïdal, on calcule tous les modes poloïdaux m tels que : $0 < n - minf < m < n + msup$, sauf pour $n = 0$ pour lequel on calcule les modes $m < mn0$. Pour ne pas avoir de problèmes d'aliasing, on cherche à vérifier :

$$\frac{nsmax}{2} - 1 + msup \lesssim \frac{2}{3}mmax \quad (31)$$

$$\frac{nsmax}{2} - 1 \lesssim \frac{2}{3}nmax \quad (32)$$

$$(33)$$

Et il faut bien sûr vérifier impérativement : $nsmax/2 - 1 + msup < mmax/2 - 1$, et $mn0 < mmax/2 - 1$. La figure 7 résume la structure des modes utilisés pendant un run d'**XTOR**.

²⁹Les composantes covariantes sont calculées à partir des contravariantes.

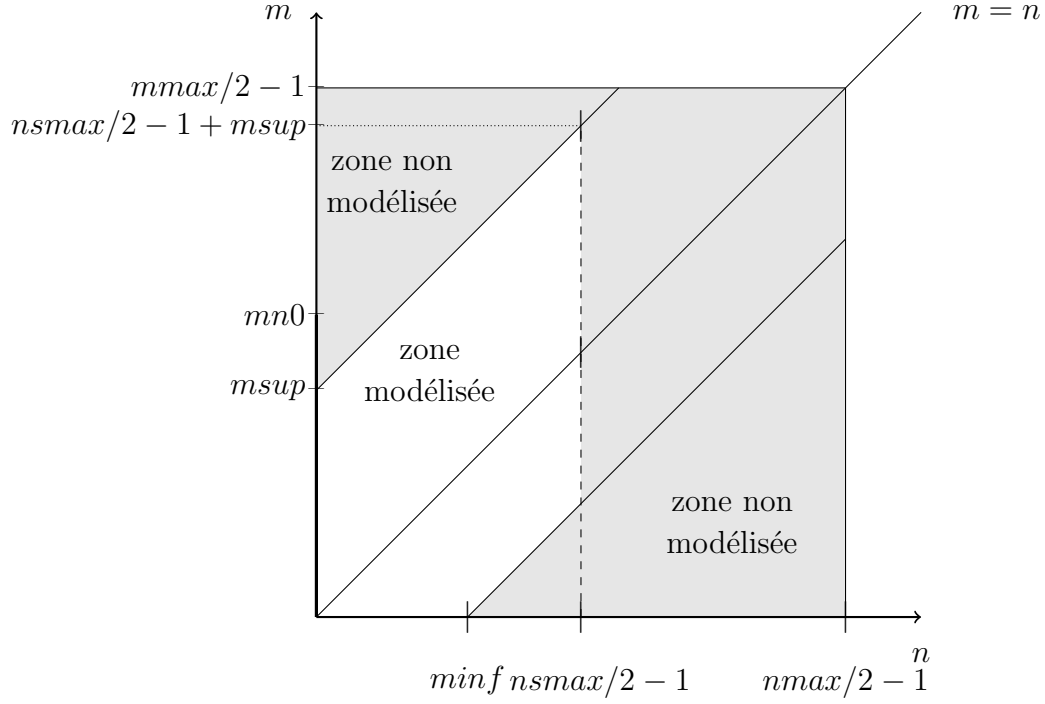


FIG. 7 – Structure du spectre modélisé par XTOR

Conditions aux limites

La géométrie du plasma dans XTOR est toroïdale, avec une section poloïdale non nécessairement circulaire. En MHD idéale, le code impose une frontière fixe sous la forme d'une coque infiniment conductrice à la surface du plasma. La surface du plasma est donc une surface de flux. En MHD résistive, on utilise le champ électrique toroïdal E_ϕ comme condition aux limites pour l'imposer : $\langle \eta J_\phi \rangle = \langle E_\phi \rangle = cte$. Il n'est pas possible de simuler une séparatrice dans le plasma.

Les conditions aux limites sont choisies avec soin pour maintenir la cohérence des calculs. Remarquons que les composantes θ contra-variantes des champs vectoriels sont singuliers sur l'axe magnétique (ils se comportent comme $1/r$). Pour éviter ce problème, on utilise les variables v^θ/D , B^θ/D , et J^θ/D ($D \propto 1/r$ sur l'axe magnétique). Par symétrie, on utilise les variables $v_\theta D$, $B_\theta D$, et $J_\theta D$ pour les composantes covariantes.

Les conditions de continuité impliquent des points de maillage surnuméraires au-delà des bords et en-deça de l'axe magnétique ; les champs sont interpolés linéairement sur ces points pour assurer les conditions voulues. On recommande au lecteur de se référer à [40] pour plus de détails.

3.2 Étude numérique des dents de scie avec XTOR : cadre et objectif

À l'aide du code décrit en première partie, nous menons une étude numérique des dents de scie en partant d'un modèle de base, la $\eta\chi$ MHD, puis en incorporant les effets de dérives diamagnétiques. Cette étude a fait l'objet de deux publications : [52] pour la partie $\eta\chi$ MHD et [49] pour l'ajout des effets diamagnétiques. Pour cela, nous étudions les régimes de fonctionnement à long terme (*i.e.*, sur un temps comparable au temps résistif, $\tau_\eta \sim \eta^{-1} = S$) du kink interne. Notre première étape est d'établir un cadre de référence du comportement du kink en utilisant la $\eta\chi$ MHD. On pourra ensuite choisir un ensemble de cas représentatifs et ajouter les effets diamagnétiques.

Physiquement, la $\eta\chi$ MHD devrait contenir assez de physique pour modéliser correctement les petits tokamaks à chauffage ohmique, comme ST à Princeton ou TFR à Fontenay-aux-roses, où on a observé des dents de scie pour la première fois. Dans ces machines, le nombre de Lundquist est de l'ordre de $S = 10^6$. Pour des machines plus grandes, à résistivité plus faible, les termes négligés dans la loi d'Ohm (effets bi-fluides, inertie électronique, compressibilité électronique) deviennent importants dans le processus de reconnection, et le modèle est insuffisant. Notre but n'est pas de simuler correctement le crash, dans lequel la reconnection joue un rôle majeur, et qui exige un modèle plus raffiné. Nous nous concentrons sur la dynamique de la rampe, pendant laquelle le système reste stable.

Notre but étant d'étudier les régimes de fonctionnement stationnaires du kink interne, nous devons attendre que la période et l'amplitude des éventuelles oscillations soient stabilisées. Cela intervient en général après un temps de simulation comparable au temps de diffusion résistive. Un faible nombre de Lundquist permet de faire tourner des cas d'une telle longueur en un temps raisonnable. Sur un processeur de type de type Intel[®] Core à 2.33 GHz, un cas met environ quatre jours à tourner sur $5 \times 10^6 \tau_A$ (Δt variant entre 0.1 et $10 \tau_A$).

3.2.1 Modèle physique $\eta\chi$ MHD et paramètres clés

Au cours de nos simulations, nous maintenons le profil de η constant (le profil de $\chi_\perp$ lui est proportionnel), nous négligeons l'évolution de la densité, et nous choisissons comme terme source $\Theta = -\nabla \cdot \chi_\perp \nabla p_0$.

Avec ces hypothèses, les paramètres clés du système sont :

- $\eta = S^{-1}$, qui contrôle l'évolution du profil de densité de courant, et donc la relaxation du profil de q après le crash. Notons que η étant constant en t , il n'y a pas de couplage dynamique entre résistivité et pression par la loi de Spitzer.
- $\chi_\perp$, qui contrôle la dynamique perpendiculaire de la pression par le terme de

source, en ramenant le profil de pression vers le profil initial.

- β_{p1} , le bêta poloïdal contenu dans la surface $q = 1$, qui influe sur le taux de croissance du kink interne.

Les termes de sources fournissent deux temps caractéristiques de l'évolution du plasma. D'une part, le transport perpendiculaire tend à ramener le profil de pression au profil initial, et donc à piquer le profil de pression au centre. Le temps caractéristique de cet effet est $\tau_{\chi\perp} = a^2/\chi\perp = \chi\perp^{-1}$ (puisque $a = 1$ dans nos unités). D'autre part, le temps résistif $\tau_\eta \sim S$ joue sur le taux de croissance du kink ($\gamma \propto \eta^{1/3}$, cf. [15], p. 332), et sur la relaxation du profil de q après chaque crash. Le kink se manifeste par un déplacement hélicoïdal du cœur du plasma vers l'extérieur, et tend donc à aplanir le profil de pression au centre. C'est de cette opposition entre kink et transport que des oscillations peuvent venir. Le kink est dû à la fois au gradient de pression et au courant.

Le taux de croissance du kink est aussi lié à β_{p1} : à faible β_{p1} , on est dans le domaine résistif du kink, le taux de croissance est faible ; quand β_{p1} grandit, le kink idéal prend le dessus, avec un taux de croissance plus important (typiquement quelques $10^{-3} \tau_A$).

Remarquons que le terme de transport n'aplanit pas le profil de pression par diffusion, mais il reconstruit le profil d'équilibre. La pression dans l'îlot associé au kink est aplanie par le transport parallèle (en général, $\chi_{\parallel} = 100$ dans nos unités) ; quand l'îlot envahit le cœur du plasma, le profil de pression au centre devient plat.

Les questions que nous nous posons sont les suivantes :

1. Avec le modèle de $\eta\chi$ MHD, quels sont les régimes de fonctionnement du kink interne ? En particulier, existe-t-il un domaine dans notre espace de paramètres où on observe l'hystérésis caractéristique des dents de scie ?
2. Quelles différences observe-t-on en ajoutant les effets de dérives diamagnétiques ?

3.2.2 Conditions initiales

Nous utilisons deux configurations initiales différentes, fournies par CHEASE, qui présentent les caractéristiques suivantes :

- Section poloïdale circulaire.
- Rapport d'aspect $A = 2.7$.
- Profil de pression parabolique ; on peut le rescaler, ce qui permet de faire varier β_{p1} .
- Profil de β_p quasi-plat.
- Profil de q :

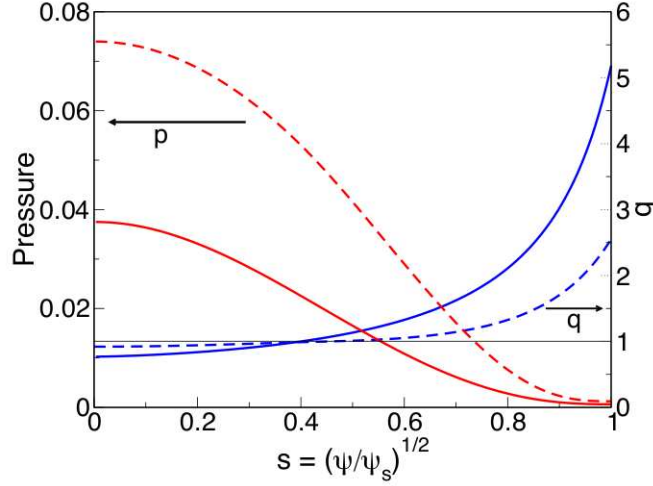


FIG. 8 – Profils de q (bleu) et p (rouge) initiaux dans les cas high shear (trait plein) et low shear (pointillés).

1. cas « low shear » : q vaut 0.9 au centre et croît jusqu'à 2.5 au bord. Le profil est plat sur la surface $q = 1$ ($\hat{s}_{q=1} = 0.1$).
 2. cas « high shear » : q vaut 0.7 au centre, et croît selon un profil parabolique jusqu'à 5 au bord. Le shear sur $q = 1$ vaut $\hat{s}_{q=1} = 0.4$.
- Seuil idéal du kink interne (déterminé numériquement³⁰) :
1. cas « low shear » : $\beta_{p1} = 0.21$
 2. cas « high shear » : $\beta_{p1} = 0.33$

Nous illustrons ces caractéristiques sur la figure 8.

Le cas high shear est plus proche de la majorité des expériences. Le cas low shear mérite cependant d'être étudié. Le fait d'avoir un profil de q plat baisse le taux de croissance du kink résistif (proportionnel à $q'^{2/3}$ sur la surface $q = 1$, cf. [15], p. 333) ; cette hypothèse a été envisagée pour expliquer l'absence de développement de ce mode pendant la rampe. Le profil très plat de q serait alors dû au précédent crash. Rappelons que plusieurs expériences ont rapporté des facteurs de sécurité proches de 1 au cœur du plasma pour des décharges ohmiques avec dents de scie, ce qui est consistant avec cette configuration.

Notons que ces deux configurations sont naturellement instables vis-à-vis du kink interne résistif (même si à faible β_{p1} , le taux de croissance est très faible) puisque

³⁰On a fait tourner les cas en MHD idéale en faisant varier β_{p1} , et on a observé la croissance ou non des modes $n = 0$ et $n = 1$.

$q_0 < 1$. On aurait obtenu les mêmes résultats en partant d'une configuration stable et en faisant décroître q en dessous de 1 au centre. Partir d'une configuration instable permet simplement d'économiser du temps de simulation.

3.2.3 Modes représentés

Nous prenons les paramètres géométriques suivants :

$$\begin{aligned} lmax &= 151 & ; & \quad nsmx = 8 \\ nmax &= 12 & ; & \quad mmax = 32 \\ msup &= 7 & ; & \quad minf = 4 \\ mn0 &= 7 \end{aligned}$$

On utilise donc 4 modes toroïdaux, et de 7 à 11 modes poloïdaux par mode toroïdal.

3.3 Régimes du kink en $\eta\chi$ MHD

3.3.1 Régimes dans l'espace des paramètres $(\chi_\perp, \beta_{p1})$

Le cadre de référence de la dynamique du kink est établi par une étude paramétrique en β_{p1} et $\chi_\perp$, pour $S = 10^6$, dans le cas low shear. Les plages de paramètres sont les suivantes :

- $\beta_{p1} = 0.07, 0.14, 0.21, 0.23$
- $\chi_\perp = 1.0 \times 10^{-6}, 3.0 \times 10^{-6}, 1.0 \times 10^{-5}, 3.0 \times 10^{-5}, 1.0 \times 10^{-4}, 3.0 \times 10^{-4}$

Les valeurs de β_{p1} sont choisies pour explorer du domaine résistif du kink ($\beta_{p1} = 0.07$) jusqu'au domaine idéal ($\beta_{p1} = 0.24$). Les valeurs de $\chi_\perp$ sont centrées autour de la valeur $\chi_\perp S = 30$, qui est pertinente pour notre domaine de validité. On prend $\nu = 5.0 \times 10^{-6}$ et $\chi_\parallel = 100$.

L'étude paramétrique révèle trois régimes de fonctionnement du kink interne, répartis comme le montre la figure 9.

Influence du profil de départ

Pour étudier l'influence du profil de départ, on fait un scan en β_{p1} dans le cas high shear, avec $\chi_\perp S = 30$. Les valeurs sont $\beta_{p1} = 0.11, 0.22, 0.33, 0.44$. Le comportement du système est qualitativement identique dans les cas high shear et low shear. En effet, le système ne se « souvient » pas de l'état initial après le premier crash : le facteur de sécurité n'a jamais le temps de relaxer jusqu'au profil de départ. Le profil

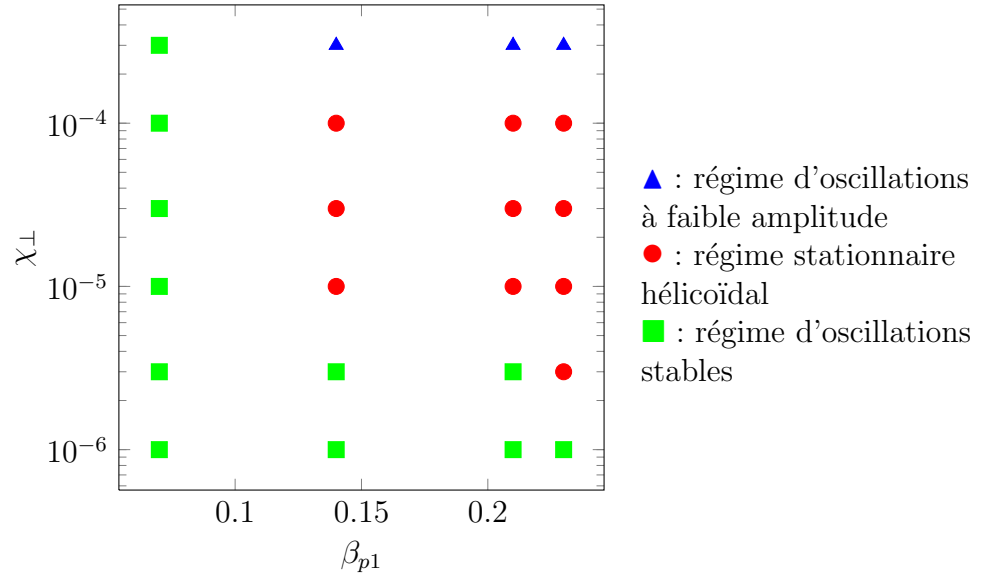


FIG. 9 – Distribution des régimes du kink dans le plan $(\chi_{\perp}, \beta_{p1})$.

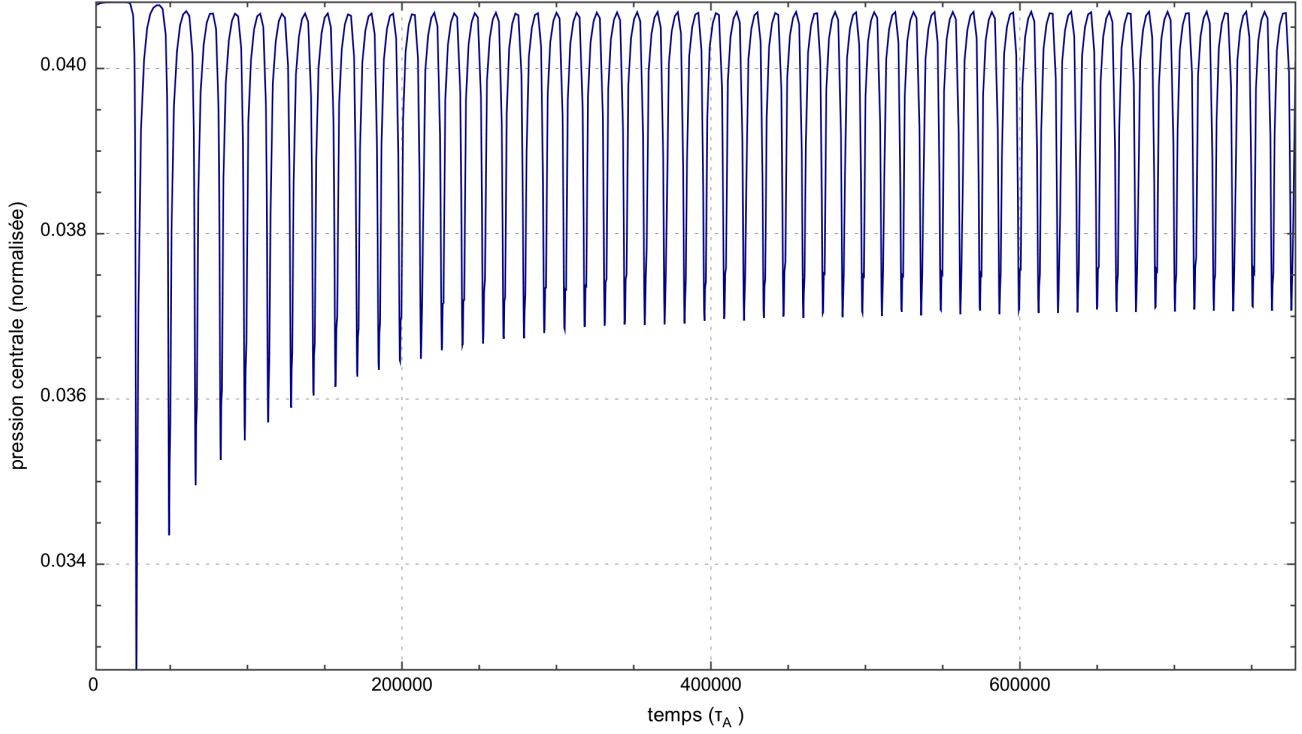


FIG. 10 – Évolution de la pression centrale pour le cas $\beta_{p1} = 0.14, \chi_{\perp} = 3.0 \times 10^{-6}$ (régime oscillant).

du facteur de sécurité reste toujours presque plat à l'intérieur de la surface $q = 1$; ainsi, même dans le cas high shear, on a affaire à une instabilité « low shear » (comme dans [53]), dans la mesure où après le premier crash, le shear reste toujours faible à l'intérieur de $q = 1$. L'état final du système, que nous voulons caractériser, est entièrement gouverné par les sources de courant et de pression.

3.3.2 Régime oscillant

Le régime représenté par les carrés verts sur la figure 9 est caractérisé par des oscillations régulières de la pression centrale, comme le montre la figure 10.

La dynamique d'un cycle est la suivante : pendant la phase de rampe, la pression augmente au centre à cause de la source de transport perpendiculaire Θ . Lorsque le gradient de pression est suffisamment important, un kink se développe, et pousse le cœur du plasma contre la surface $q = 1$. Une nappe de courant et une couche de reconnection apparaissent autour de cette surface. Le processus de reconnection

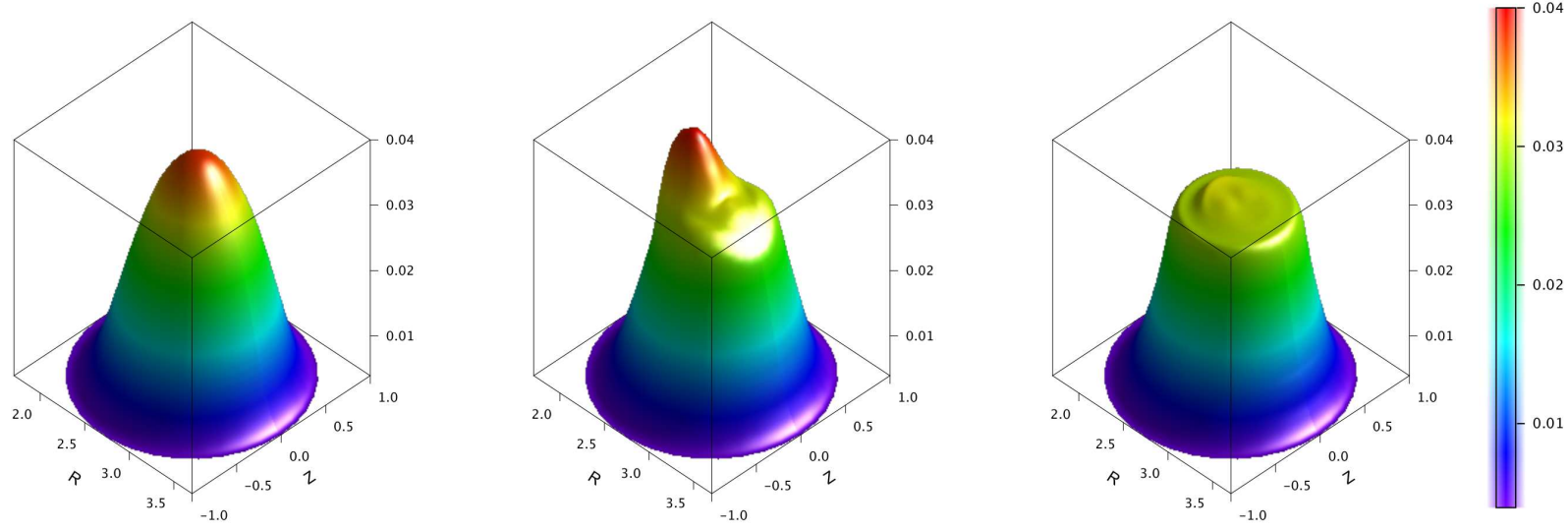
Time = 39808.2 τ_a Time = 45313.7 τ_a Time = 46351.3 τ_a 

FIG. 11 – Le graphe de gauche correspond à la pression juste avant un crash (profil piqué au maximum). Le graphe du milieu montre le kink qui déforme le cœur du plasma et le pousse contre la surface $q = 1$. Le graphe de droite montre la pression après le crash, avec le profil plat à l'intérieur de $q = 1$.

modifie la topologie des lignes de champ en créant un îlot magnétique $m = n = 1$. Le transport parallèle se met alors à contribuer au déplacement perpendiculaire de la pression. À cause de la largeur du spectre utilisé dans les simulations, on observe des couplages du kink avec des harmoniques (m, n) plus élevées. Cela se traduit par la formation de chaînes d'îlots, qui créent une zone stochastique autour de la surface $q = 1$. Ce phénomène contribue également à augmenter le transport perpendiculaire de la pression. Celle-ci est évacuée vers l'extérieur, à travers la surface $q = 1$, ce qui se manifeste par un crash, qui termine le cycle. Après le crash, le profil de pression au centre est complètement plat. La pression recommence alors à augmenter sous l'effet du terme de source, et ainsi de suite. On montre l'évolution de la pression³¹ en 3D sur la figure 11.

La dynamique de la pression est caractéristique du modèle de Kadomtsev, et incompatible avec l'hypothèse de l'instabilité de quasi-interchange. On ne voit pas de bulle froide caractéristique, mais bien le kink déplacer le plasma en croissant. Une telle dynamique a entre autres été étudiée en détail dans [54]. On observe le profil de q très plat et très proche de l'unité prévu par le modèle. L'amplitude des oscillations de pression est de l'ordre de 15%. Les simulations présentent le même défaut que le

³¹Dans le cas $\beta_{p1} = 0.07, \chi_{\perp} = 3.0 \times 10^{-6}$.

modèle de Kadomtsev : le temps de crash est beaucoup plus grand que les valeurs expérimentales. Ainsi, sur la figure 10, le temps de crash et le temps de rampe sont comparables : $\tau_{rampe} \approx 12000\tau_A$ et $\tau_{crash} \approx 3000\tau_A$, soit $\tau_{crash}/\tau_{rampe} \approx 16\%$. Le temps de crash est donc de l'ordre de $(\tau_\eta\tau_A)^{1/2}$, conformément au modèle de Kadomtsev.

On doit donc conclure que si l'on observe effectivement des cycles réguliers du kink, on ne peut pas pour autant les appeler des dents de scie. Leur dynamique est fondamentalement différente. D'une part, le temps de crash est trop long, ce qui est attendu, étant donné que le modèle ne contient pas les effets physiques nécessaires pour modéliser correctement la phase de reconnection (dérives diamagnétiques, inertie électronique, effet Hall). D'autre part, le modèle ne reproduit pas d'hystérésis, et le système ne passe pas par l'état de rampe avec un kink stable.

3.3.3 Régime stationnaire hélicoïdal

Dans ce régime (représenté par les ronds rouge sur la figure 9), la pression centrale évolue vers un état stationnaire tridimensionnel d'hélicité $m = n = 1$ (voir figure 13). La configuration initiale étant instable vis-à-vis du kink interne, l'équilibre stationnaire est précédé de quelques oscillations du kink, rapidement amorties. Nos simulations représentent un kink saturé. Rappelons que dans le cas d'une configuration instable pour le kink idéal, un tel état est toujours obtenu. On représente sur les figures 12 et 13 l'évolution de la pression centrale et de l'énergie cinétique des modes toroïdaux ($W_{kinetic}^n = \int_{V_{plasma}} \sum_{m \in E_n} v^2$) pour un cas standard ($\beta_{p1} = 0.14$, $\chi_\perp = 3.0 \times 10^{-5}$, $\chi_\parallel = 100$). On montre la somme des énergies sur tous les m , mais les modes $m/n = 1$ sont bien sûr dominants.

L'état stationnaire est associé à l'apparition d'une couche de reconnection et d'une cellule de convection tridimensionnelle permanentes autour de la surface $q = 1$. Le courant et la pression introduits au cœur du plasma par les termes de source sont constamment évacués vers l'extérieur par ces structures. Ainsi, il n'y a pas d'accumulation de pression ni de courant à l'intérieur de la surface $q = 1$.

3.3.4 Régime de faibles oscillations

Dans ce dernier régime, (les triangles bleus sur la figure 9) on observe des oscillations irrégulières, d'amplitude faible (sur la pression centrale, l'amplitude est en moyenne de 3%). Ces oscillations peuvent éventuellement disparaître, puis recroître. La dynamique du cycle est la même que pour le régime oscillant, mais les variations de pression deviennent négligeables.

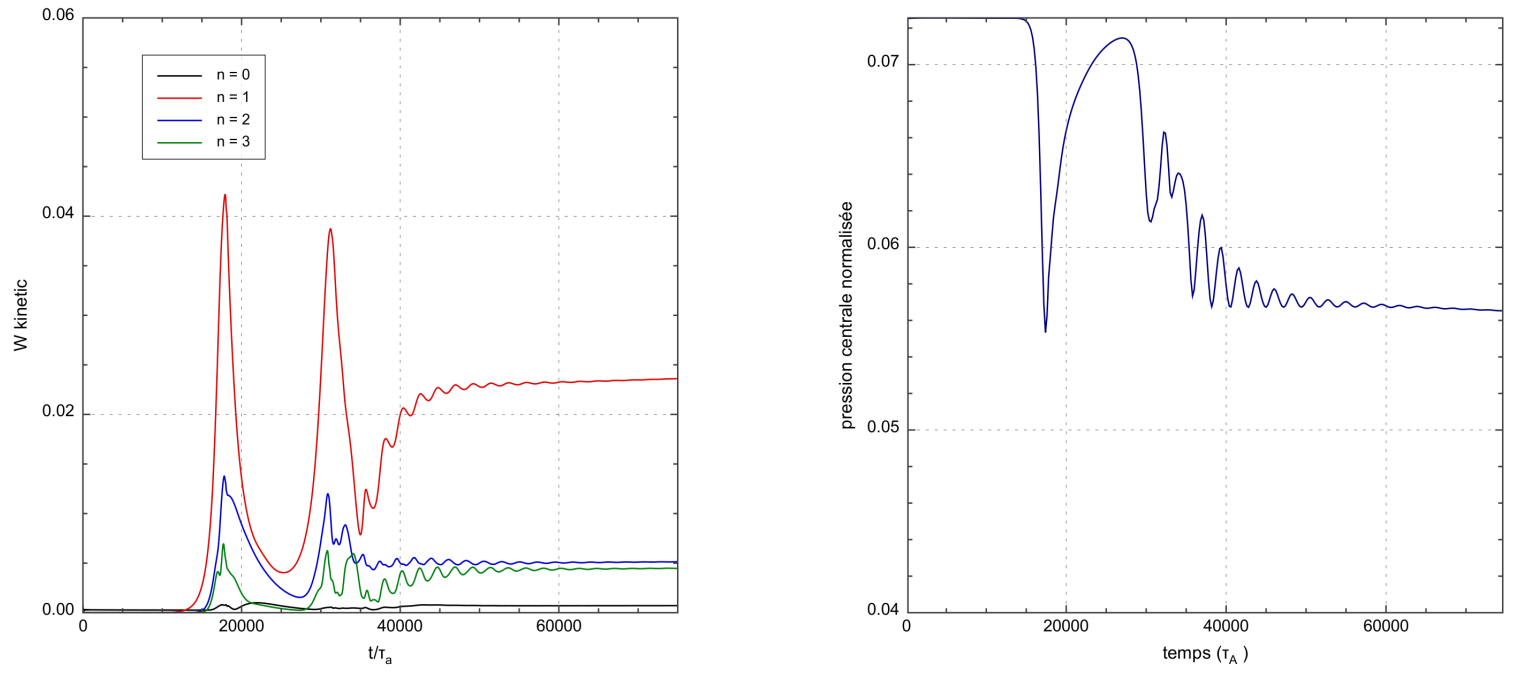


FIG. 12 – Évolution de l'énergie cinétique des modes toroïdaux (gauche) et de la pression centrale (droite) dans un cas saturé (cas $\beta_{v1} = 0.14$, $\chi_{\perp} = 3.0 \times 10^{-5}$, $\chi_{\parallel} = 100$).

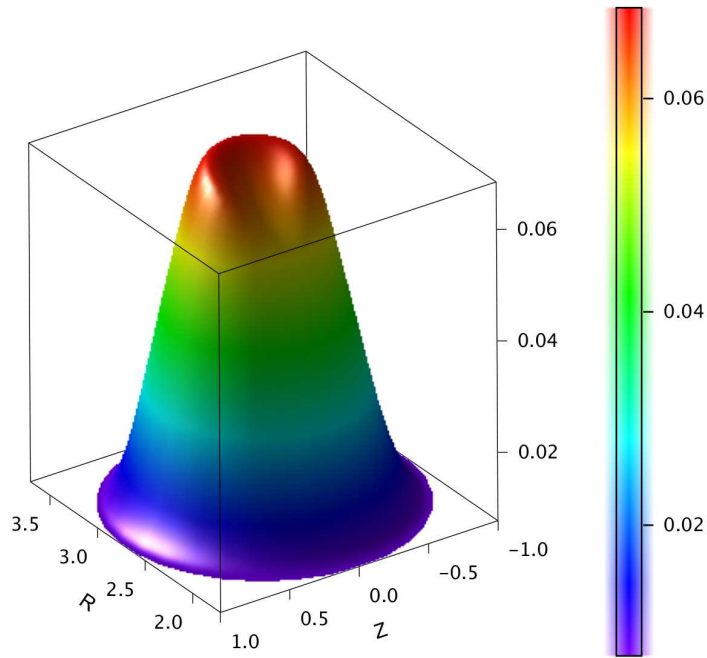


FIG. 13 – État stationnaire d'hélicité $m = n = 1$ de la pression centrale (cas $\beta_{p1} = 0.14$, $\chi_{\perp} = 3.0 \times 10^{-5}$, $\chi_{\parallel} = 100$).

3.3.5 Transition entre les régimes $\eta\chi$ MHD

Dans cette section, on met en évidence les conditions physiques nécessaires pour passer d'un régime à un autre.

Nous commençons par analyser le diagramme de répartition des régimes dans l'espace $(\chi_{\perp}, \beta_{p1})$. Dans nos simulations, le système est gouverné par deux effets concurrents, la relaxation du profil de courant (et donc de q), et le développement du kink low shear dû au gradient de pression, qui crée une cellule de convection. Le fait de varier soit β_{p1} soit $\chi_{\perp}$ sans modifier S revient à changer $\tau_{\chi_{\perp}}$ sans changer τ_{η} , *i.e.*, à jouer sur la source de pression en maintenant la dynamique du courant inchangée.

Sur le diagramme, on constate que le régime oscillant correspond à des valeurs faibles pour β_{p1} et $\chi_{\perp}$, *i.e.*, un temps de transport $\tau_{\chi_{\perp}}$ grand. Si on augmente β_{p1} ou $\chi_{\perp}$ à partir de ce régime, on réduit $\tau_{\chi_{\perp}}$ et on passe au régime stationnaire avec kink saturé. On interprète cela en considérant que le rôle de $\chi_{\perp}$ est de restaurer le profil de β_{p1} d'équilibre, et donc de renforcer l'instabilité; quant à β_{p1} , il est lié au taux de croissance du kink. Ainsi, le diagramme traduit une bifurcation du système d'un régime oscillant vers un régime saturé du kink lorsque l'instabilité de pression est assez forte. Le système se bloque dans l'état saturé lorsque la cellule de convection est assez efficace. Sinon, le profil de q a le temps de relaxer un peu, et le système oscille à cause du développement périodique du kink low shear.

Notons que le diagramme n'est pas symétrique par rapport à la diagonale, c'est-à-dire que β_{p1} et $\chi_{\perp}$ n'ont pas des rôles exactement similaires. En effet, même lorsque β_{p1} est faible mais que $\chi_{\perp}$ est élevé, $\tau_{\chi_{\perp}}$ est assez court pour empêcher la dynamique du courant de modifier le profil de pression. Dans ce cas, on observe le régime de faibles oscillations.

Si la confirmation de l'existence d'un régime oscillant du kink interne sur des temps longs dans un certain espace de paramètres est intéressante, il n'en reste pas moins que la dynamique de ces oscillations diffère fondamentalement de celle des dents de scie. Cela n'est pas surprenant attendu qu'un simple modèle MHD n'inclut pas suffisamment de termes pour modéliser la reconnection précisément, d'où nos temps de crash trop élevés.

Ainsi, par exemple, le coefficient $\chi_{\parallel}$ ne semble pas jouer un rôle majeur. Il devrait participer au transport perpendiculaire pendant la reconnection, à cause du changement de topologie des lignes de champ. Faisant varier $\chi_{\parallel}$ entre 0.1 et 100, (avec $\chi_{\perp} = 3.0 \times 10^{-6}$ et $S = 10^6$), un nouveau scan en β_{p1} ne montre qu'une seule différence, pour $\beta_{p1} = 0.14$, pour $\chi_{\parallel} = 0.1$ et $\chi_{\parallel} = 1$. Le système passe du régime saturé au régime oscillant. La faible amplitude des oscillations indique qu'il reste proche de la frontière. La diminution de $\chi_{\parallel}$ rend la cellule de convection moins efficace, et permet au système de sortir de l'état stationnaire pour retomber dans la dynamique périodique accumulation/crash. Ce résultat indique que les valeurs de $\chi_{\parallel}$

utiisées dans nos simulations ne changent pas la dynamique de façon sensible dans la zone de stochasticité.

3.4 Ajout des effets diamagnétiques

Cette étude numérique $\eta\chi$ MHD a permis de dégager une fenêtre de paramètres qui a ensuite servi à mener, en collaboration avec Federico Halpern lors de son séjour au sein de l'équipe Plasmas Magnétisés du CPhT, une étude des dents de scie en régime diamagnétique. La dynamique cyclique reproduite par la $\eta\chi$ MHD est fondamentalement différente de celle des dents de scie. Le modèle utilisé ne contient pas assez de physique pour reproduire un cycle d'hystérésis, avec un état métastable pendant la rampe, comme dans les expériences. Des termes stabilisants, comme les dérives diamagnétiques, pourraient permettre d'atteindre un tel état. Les simulations incluant les effets diamagnétiques sont faites avec la valeur centrale $\chi_{\perp} S = 30$, $S = 10^6$. On utilise la configuration high shear. La densité évolue suivant l'équation 27.

3.4.1 Évolution du régime oscillant en dents de scie

Dans les cas oscillants, l'ajout d'effets diamagnétiques se traduit par l'apparition d'oscillations pré- et postcurseurs sur la pression centrale, juste avant et après le crash. De telles oscillations sont souvent observées expérimentalement. Dans nos simulations, elles sont dues à un mode rémanent $m = n = 1$, qui tourne autour de l'axe magnétique. Malheureusement, la période du cycle est trop courte pour pouvoir observer une phase de rampe bien définie : le mode rémanent n'a pas le temps d'être amorti.

Pour pouvoir étudier la dynamique de la phase de rampe plus en détail, on a effectué une simulation avec un nombre de Lundquist de 10^7 . Cela a pour effet de rallonger la longueur du cycle. À cause du temps nécessaire pour une simulation à S grand, on s'est concentré sur un cas : $\beta_{p1} = 0.22$. Rappelons que dans le cas high shear, le seuil idéal est à $\beta_{p1} = 0.33$. On a rescalé le coefficient de transport pour garder $\chi_{\perp} S = 30$.

Moyennant ces modifications, on obtient des cycles de kink dont la forme s'approche des dents de scie observées dans les tokamaks. On distingue bien une longue phase de rampe, durant laquelle le kink ne se développe pas, suivie d'un crash rapide entouré d'oscillations pré- et postcurseurs. On montre l'évolution de la pression sur la figure 14.

La dynamique du cycle est plus complexe que dans le cas de la $\eta\chi$ MHD. Pendant la rampe, les effets diamagnétiques stabilisent le kink. À la fin de la rampe, une instabilité se développe, qui se traduit par les oscillations précurseurs sur la

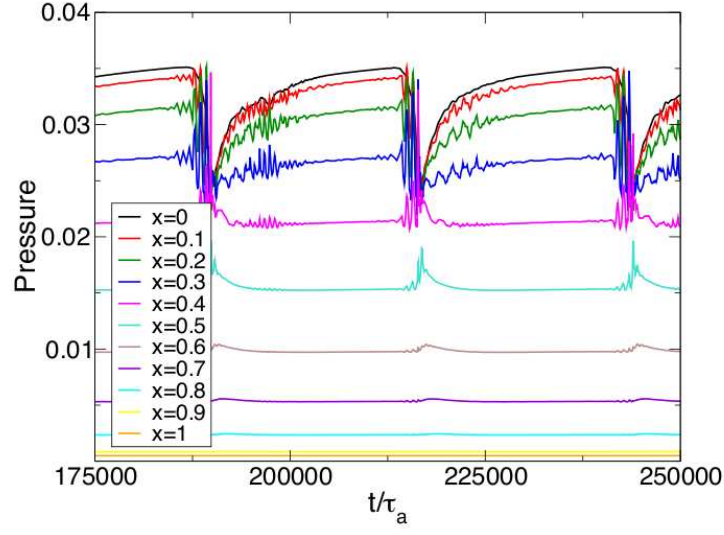


FIG. 14 – Évolution de la pression en différents r (x dans la légende) dans le cas $S = 10^7$, avec effets diamagnétiques ioniques et électroniques.

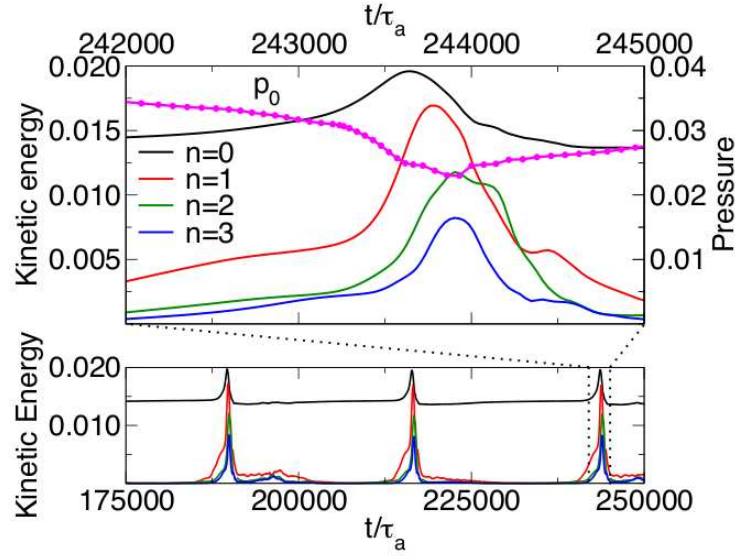


FIG. 15 – Évolution de l'énergie cinétique des modes toroïdaux pendant un crash.

pression. Ces oscillations voient leur amplitude augmenter. Le cœur du plasma est poussé contre la surface $q = 1$ par le kink : il se déplace vers l'extérieur en tournant à la vitesse diamagnétique ionique, et les oscillations représentent ce déplacement observés sur une corde fixe. En d'autres termes, la pression ne diminue pas pendant cette phase, au contraire. Juste avant le crash, il y a un changement drastique dans la dynamique du kink. Le taux de croissance augmente très brutalement. La figure 15 détaille l'évolution de l'énergie cinétique des modes toroïdaux pendant un crash. On constate que l'énergie des modes est multipliée par 3 environ en $500 \tau_A$, au début du crash. On rappelle qu'une telle accélération est généralement acceptée comme étant à l'origine du crash. Différentes études ont proposé des interprétations pour cette accélération, fondées sur les effets diamagnétiques électroniques, comme [55], [56], et [57]. Alors que le temps de rampe est de $\tau_{rampe} \sim 25000 \tau_A$, le temps de crash est de $\tau_{crash} \sim 2500 \tau_A$. Notons que même si la dérive diamagnétique a un effet sur le kink, elle ne suffit pas à modéliser exactement la reconnection. Notre but est toujours la dynamique de la rampe, et non du crash ; la prise en compte de l'inertie électronique, ou d'un éventuel régime semi-collisionnel serait nécessaire pour atteindre cet objectif.

3.4.2 Évolution des seuils de transition et premières comparaisons avec JET

Les effets des dérives diamagnétiques sur les transitions entre régimes sont montrés sur la figure 16, qui expose les résultats d'un scan en β_{p1} et α . Rappelons que α apparaît comme coefficient devant les termes diamagnétiques : il mesure donc la force des ces effets (cf. les équations 25-27). Le cas $\alpha = 0$ correspond à la $\eta\chi$ MHD.

On constate que les dérives diamagnétiques décalent le seuil de transition entre régime oscillant et régime stationnaire vers des plus hautes valeurs de β_{p1} . D'une façon plus générale, tout effet stabilisant pour le kink tend à décaler le régime hélicoïdal saturé vers un régime oscillant. On avait bien la même tendance en $\eta\chi$ MHD.

On peut comparer les nouveaux régimes obtenus avec un régime ohmique typique de JET. On renormalise les résultats avec des grandeurs typiques :

- $\alpha = 0.08$ (cf. [58]), avec une densité typique des plasmas ohmiques de $n = 2.0 \times 10^{19} \text{m}^{-3}$.
- $\beta_p \sim \beta_{p,crit}/2$
- $R_0 = 3.1 \text{m}$
- $B_0 = 2.7 \text{T}$
- $T_e = 2 \text{keV}$
- $S > 10^7$.

Notre résultat principal est de constater que pour des paramètres proches des valeurs expérimentales de JET, l'ajout des effets diamagnétiques qui stabilisent le

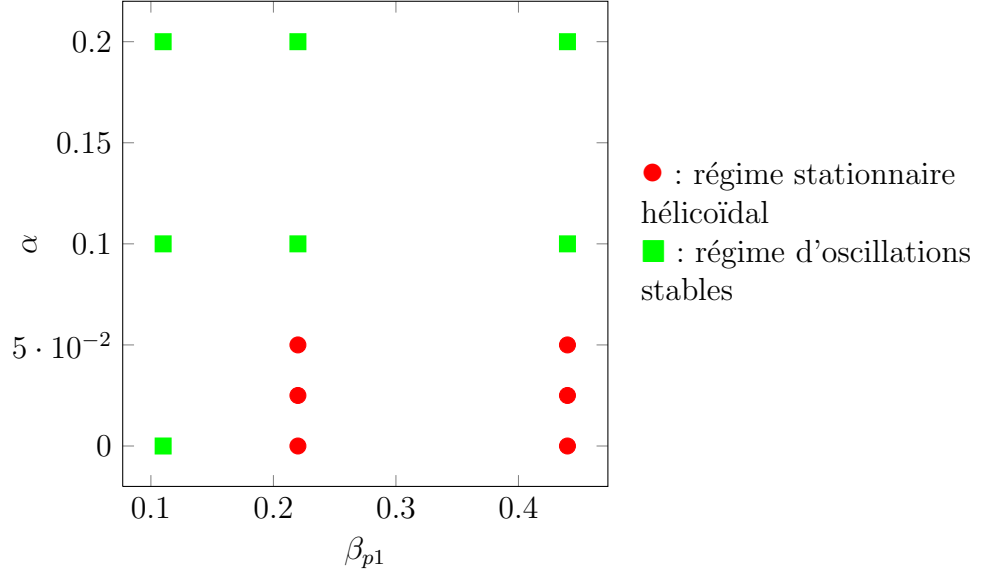


FIG. 16 – Distribution des régimes du kink dans le plan (β_{p1}, α) , avec effets diamagnétiques ioniques et électroniques.

kink décale le seuil en β_{p1} , et on retrouve un régime d'oscillations comparables à des dents de scie, alors que la $\eta\chi$ MHD ne reproduit des oscillations stables qu'à très faible β_{p1} . Remarquons de plus qu'en réalité le nombre de Lundquist de JET est supérieur à celui que nous avons utilisé (pris plus faible pour des raisons de temps de calcul) ; par comparaison, l'effet réel des termes diamagnétiques est plus grand que celui que nous avons simulé. Autrement dit, le graphe effectif se trouve décalé vers des valeurs de α plus grandes, et on s'éloigne de la frontière avec le régime saturé.

Échelles de temps

Le rapport entre la période des dents de scie et le temps résistif en régime ohmique est de $\tau_\eta/\tau_{dds} \approx 300 - 500$. Dans nos simulations, la résistivité est artificiellement augmentée pour des raisons de temps de calcul. Il faut donc rescaler notre temps résistif pour qu'il corresponde à sa valeur expérimentale, $S \approx 2 \times 10^8$. On obtient alors une période de dents de scie de 0.2 - 0.3 s. Il est intéressant d'étudier le temps de crash, dans le cas où la reconnection est accélérée par les effets diamagnétiques électroniques. On avait un τ_{crash} de $600\tau_A$ environ, ce qui donne $\tau_{crash} \approx 200\mu s$; l'accord avec les valeurs de JET est raisonnable. Rappelons que d'autres effets, tels que l'inertie électronique et la viscosité électronique, ne sont pas pris en compte. On s'attend à ce que ces effets accélèrent encore le processus de reconnection.

Facteur de sécurité

Dans toutes nos simulations, le facteur de sécurité est resté proche de 1 pendant tout le cycle. Cela est cohérent puisque le temps de relaxation est le temps résistif, et que la période de nos dents de scie est beaucoup plus faible. Rappelons que des dents de scie avec q proche de 1 après le crash ont été observés expérimentalement. Cependant, les exemples de reconnection partielle pendant le crash, avec des valeurs de q de 0.6-0.8, sont plus nombreux, notamment dans JET, mais aussi dans des décharges ohmiques à section poloidale circulaire, comme dans TEXTOR et TFTR. La modélisation exacte du crash, pendant lequel la reconnection a un rôle majeur, nécessite sans doute de nouveaux effets, tels que des orbites ioniques finies, une dynamique électronique semi-collisionnelle, ou encore des effets de piégeage (cf. respectivement [59], [60], [61]).

État stationnaire $m = n = 1$

Dans ITER, on prévoit que la période des dents de scie sera beaucoup plus longue que le temps de confinement de l'énergie ($\tau_{dds} \approx 50$ s). Le temps de relaxation de la pression après chaque crash devrait, lui, être de l'ordre de la seconde. Pour essayer de se rapprocher de ces conditions, on a fait tourner un cas avec les mêmes paramètres que celui où on a observé une reconnection accélérée : $S = 10^7$, $\beta_{p1} = 0.22$, profil high shear, effets diamagnétiques, mais en prenant $\chi_{\perp} S = 300$ et non la valeur standard, 30. On diminue ainsi le temps de diffusion perpendiculaire, et donc le temps de relaxation de la pression après chaque crash. Après rescaling avec les paramètres ITER, on obtient un plasma avec un temps de confinement de 3 s et un temps résistif $\tau_{\eta} = 10^3$ s. La simulation aboutit à un état stationnaire avec hélicité $m = n = 1$ dans le cœur du plasma. Ce résultat est bien sûr préliminaire.

3.5 Conclusion

Après avoir exposé les principales caractéristiques du code bi-fluide non-linéaire dynamique XTOR-2F, nous avons présenté une étude numérique de la dynamique du kink interne dans des plasmas de tokamaks ohmiques. Nous avons montré les premières simulations, à notre connaissance, en régime $\eta\chi$ MHD, puis en régime diamagnétique, sur des temps suffisamment longs pour que la dynamique des oscillations soit stable ($\sim 10^6 \tau_A$). Notre résultat principal est d'avoir mis à jour deux régimes stables séparés : un régime d'oscillations périodiques et un régime stationnaire hélicoïdal saturé. Le système est gouverné par une compétition entre termes de sources et croissance d'un kink interne low shear, qui évacue la pression vers l'extérieur. Les facteurs accélérant l'instabilité (*e.g.*, une augmentation de $\chi_{\perp}$ ou de β_{p1}) décalent le système vers l'état saturé. En revanche, les effets stabilisants (ω^*) favorisent le régime oscillant, ce qui nous a permis de retrouver un régime proche des dents de scie pour des paramètres similaires à ceux d'une décharge ohmique dans

JET.

Le régime oscillant est atteint quand les termes de sources forcent le plasma à repasser sous le seuil $q = 1$ plus vite que le kink ne se développe. La dynamique du régime oscillant en régime diamagnétique devient semblable à celle des dents de scie, avec notamment des oscillations post- et précurseurs, et un temps de crash sensiblement réduit, à cause d'une brusque accélération de la croissance du kink juste avant le crash. Le régime hélicoïdal saturé est associé à une cellule de convection et une couche de reconnection, qui évacuent en permanence la pression et le courant introduits au cœur du plasma par les sources. Le profil de courant n'a jamais le temps de relaxer, et q reste proche de 1 tout au long du cycle. Des états similaires ont été observés, notamment sur NSTX et TCV (cf. [63] et [62]). Il serait intéressant de mener une comparaison expérimentale poussée, pour déterminer s'il s'agit bien de l'état prédit par nos simulations.

De tels résultats, avec les effets diamagnétiques, permettent d'envisager une étude avec un code hybride, qui devrait montrer les effets connus des particules chaudes sur les dents de scie : transition en fishbones, stabilisation de la rampe et apparition de dents de scie monstres.

Deuxième partie

Contribution au développement de **XTOR-K**

4 XTOR-K : un code hybride MHD/PIC « full-f, full orbit »

Dans ce chapitre, nous définissons les caractéristiques du code hybride XTOR-K, en le comparant aux autres codes hybrides existants. La robustesse du solveur MHD implicite nous permet de choisir une représentation complète des particules : leur mouvement est intégré exactement par un code PIC. Un sous-pas de temps particulière apparaît alors, ce qui alourdit le temps de calcul, mais notre schéma de couplage fait porter l'essentiel des itérations sur le solveur Newton-Krylov. La structure du schéma permet d'effectuer une parallélisation massive de la partie cinétique.

Nous rappelons les problématiques physiques concernées (physique des ions chauds) et le modèle choisi dans XTOR-K : MHD et intégration complète des orbites des particules cinétiques en Particle-In-Cell. Ces modèles font d'XTOR-K un code hybride original, en particulier à cause de l'absence d'hypothèse de type δf sur la distribution des particules. Ensuite, nous détaillons l'algorithme de couplage utilisé, qui repose sur une méthode de Newton-Krylov/Picard, et qui présente une asymétrie temporelle, à cause de l'apparition d'un sous-pas de temps pour l'intégration des orbites des particules. Le solveur MHD Newton-Krylov implicite, qui est robuste et efficace, assure l'essentiel des itérations, tandis que l'avancée des particules, qui prend la majorité du temps de calcul, est faite au maximum 3 à 5 fois. Pour qu'une telle méthode de Picard soit viable, il faut que les itérations de Picard se comportent bien (*i.e.*, que leur facteur d'amplification soit strictement inférieur à 1) pour les modes MHD fondamentaux ; c'est ce que nous vérifions par un calcul semi-qualitatif.

4.1 Intégration d'une population cinétique dans un code MHD

Les codes hybrides MHD-cinétiques sont apparus à partir des années 1990. Ces codes sont intéressants dans le domaine des plasmas thermonucléaires comme dans celui des plasmas astrophysiques (cf. [64]). Le principe est d'explorer l'influence des échelles de temps et longueurs microscopiques sur le comportement macroscopique et au long terme du plasma. Un code hybride est donc par définition multi-échelles. Les ordres de grandeurs « microscopiques » et « macroscopiques » dépendent du type de plasma concerné.

Parce que les échelles concernées diffèrent de plusieurs ordres de grandeurs, elles sont généralement traitées par des modèles physiques indépendants. Le but d'un code numérique hybride est de réaliser le couplage de ces modèles. Un tel code contient donc deux parties largement indépendantes, qui doivent converger de façon self-consistante.

4.1.1 Physique des ions rapides dans un tokamak

Dans le cas des plasmas de fusion, la MHD donne une bonne description des phénomènes au-dessus du temps caractéristique τ_A et de l'échelle de longueur a . L'échelle microscopique est celle du mouvement cyclotronique des particules ; les effets associés sont dits « de rayon de Larmor fini ». Leur influence sur la MHD est vaste et bien documentée, particulièrement pour les particules chaudes. Ces particules peuvent être des ions chauffés par résonance cyclotronique ionique ($T \sim 100$ keV - 1 MeV), des neutres injectés pour chauffer le plasma ($T \sim 1$ MeV), ou des α issus de la fusion ($T = 3.5$ MeV).

Le confinement de ces particules doit être contrôlé pour assurer qu'elles aient le temps de transférer leur énergie au plasma ; de plus, la perte de ces particules peut causer des dommages sur la paroi. Au-delà du simple chauffage, ces particules peuvent servir pour contrôler les profils de courant, température, etc du plasma, et éviter le développement d'instabilités. La maîtrise de leur transport est donc une problématique clé pour le développement des tokamaks. L'importance des particules chaudes est appelée à grandir au fur et à mesure que les α de fusion sont plus nombreux dans les machines ; dans ITER, le chauffage dû aux α est censé être la principale source d'énergie du plasma.

Sur le plan théorique comme expérimental, la physique des particules chaudes dans les tokamaks a fait l'objet de nombreux travaux, allant de l'extraction des cendres aux instabilités d'Alfvén collectives causées par ces particules, les résonances précessionnelles avec les kink internes (entraînant l'apparition des fishbones), à l'allongement de la période des dents de scie. Ces modes peuvent créer un transport anormal de particules chaudes et conduire à des pertes inacceptables pour un réacteur à fusion. La richesse des interactions entre particules chaudes et phénomènes MHD a conduit plusieurs équipes à faire évoluer les codes MHD existant vers des codes hybrides.

Le principe de tout code hybride est de diviser le plasma en deux fractions (ou populations). La première, qui contient les particules dites thermiques, est modélisée par la MHD, selon l'algorithme pré-existant. La seconde (qu'on appellera fraction particulaire ou cinétique) représente les particules simulées par un modèle non-MHD. La structure du code hybride est déterminée par la façon dont les deux fractions sont couplées et par le choix du modèle physique pour la partie particulaire. Des codes hybrides ont été développés aussi bien à partir de codes d'équilibre que de codes dynamiques.

Les champs électromagnétiques fournis par la partie MHD servent de sources à la partie particulaire pour déterminer l'évolution de la fonction de distribution des particules chaudes. Pour obtenir un modèle self-consistant, il faut alors choisir une rétro-action des particules chaudes sur le plasma thermique. Le couplage peut être fait par deux moments de la distribution de rapides : le courant ou la pression. Dans

le modèle MHD, le courant n'apparaît pas comme variable principale (il est lié à $\mathbf{B}$ et $\mathbf{v}$ par les équations de Maxwell); on choisit donc le tenseur de pression des particules, Π_k . Il apparaît comme un terme de source dans l'équation MHD sur la vitesse :

$$\rho \frac{d\mathbf{v}}{dt} = \mathbf{J} \times \mathbf{B} - \nabla p - \nabla \cdot \Pi_k$$

La pression p est due à la partie thermique du plasma. Le tenseur de pression particulaire est issu du modèle cinétique.

4.1.2 Modèles cinétiques et hypothèse δf

Les codes hybrides existant utilisent des modèles cinétiques variés pour faire évoluer la fonction de distribution des particules chaudes. On peut utiliser l'équation de Vlasov, une équation gyro-cinétique, une équation drift-cinétique, ou l'équation de Lorentz complète si on choisit un algorithme de type PIC. Citons les codes NIM-ROD (drift-cinétique ou Lorentz), cf. [65], ou M3D-K (Vlasov ou gyrocinétique), cf. [66]. Ces codes reposent sur l'hypothèse : $n_h \ll n_c$, $\beta_h \sim \beta_c$, ce qui préserve la quasi-neutralité. Les indices h se rapportent à la fraction cinétique, et c au plasma thermique. Cette hypothèse permet l'utilisation d'une méthode de type δf : la fonction de distribution des particules chaudes est décomposée en

$$f(t) = f_0 + \delta f$$

Où f_0 est supposée connue. L'évolution de la fonction de distribution se fait uniquement sur $\delta f \ll f_0$. Cette méthode permet de réduire fortement le bruit sur le tenseur de pression cinétique, puisque la partie principale du tenseur est constante. Le tenseur de pression est supposé isotrope, et le champ électrique parallèle est pris nul.

Cette démarche ne serait pas cohérente avec le code XTOR-2F : on ne fait nulle part d'hypothèse δf , les champs sont avancés « en entier » à chaque pas de temps. De plus, l'utilisation d'un modèle gyro-cinétique ou Vlasov rend le calcul du tenseur de pression compliqué.

Au lieu de ça, on choisit un code PIC (Particle-In-Cell) pour la partie hybride, et on intègre exactement la trajectoire de toutes les particules de la population cinétique à partir de l'équation de Newton-Lorentz. L'intégration numérique de cette équation est un problème connu et analytiquement simple. Le calcul du tenseur de pression se ramène à une somme des contributions des particules. On ne fait aucune hypothèse sur la taille de population particulaire, la forme du tenseur de pression, ou le champ électrique. De plus, on prend automatiquement en compte tous les effets de type rayon de Larmor fini, qui sont particulièrement importants à haute température, puisqu'on résout exactement le mouvement.

La richesse de la physique prise en compte dans un tel modèle « full-f, full orbits » a des contreparties numériques négatives. Le temps de calcul de la partie particulaire devient dominant devant la partie MHD, nécessitant une parallélisation des particules. L'algorithme de couplage doit être choisi en tenant compte de ce fait.

4.2 Algorithme de couplage pour XTOR-K

4.2.1 Schéma Newton-Krylov/Picard et asymétrie temporelle

On choisit une méthode de Newton-Krylov/Picard pour obtenir une solution self-consistante à chaque pas de temps fluide. Au début du pas de temps t^n , on dispose des champs fluides $\mathbf{v}^n, \mathbf{B}^n$ et p^n (regroupés dans le vecteur $\mathbf{x}^n$), ainsi que du tenseur de pression particulaire $\mathbf{\Pi}^n$. La méthode de Picard est une méthode itérative ; la i^{eme} itération se déroule de la façon suivante :

- on cherche les champs x^{n+1} , et l'avancée des particules dans un champ $(\mathbf{E}, \mathbf{B})$ interpolé entre t^n et t^{n+1} . Le tenseur de pression qui couple les deux est interpolé entre t^n et t^{n+1} .
- on fait avancer les champs fluides avec le solveur Newton-Krylov implicite, selon le modèle MHD. La pression particulaire $\mathbf{\Pi}^{n+1;i}$ apparaît comme un terme de source par l'intermédiaire de sa divergence dans l'équation sur $\mathbf{v}$. Lors de la première itération, $\mathbf{\Pi}^{n+1;0}$ est initialisé par extrapolation à partir de $\mathbf{\Pi}^n$ et $\mathbf{\Pi}^{n-1}$. On obtient les champs fluides $\mathbf{x}^{n+1;i}$, représentant les champs au temps t^{n+1} . L'index i est celui de l'itération de Picard, l'index n celui du pas de temps. L'index n seul marque les solutions obtenues après convergence de l'algorithme au temps t^n .
- on utilise ensuite les champs fluides obtenus, $\mathbf{B}^{n+1;i}$ et $\mathbf{E}^{n+1;i}$, pour intégrer les trajectoires des particules sur un pas de temps fluide. Comme nous avons choisi une représentation en orbites complètes, l'intégration de la trajectoire des particules se fait sur un pas de temps particulaire δt inférieur au pas de temps fluide Δt ($\delta t / \Delta t \sim \tau_c / \tau_A \sim 100 - 1000$). L'avancée des particules est donc constituée d'itérations successives sur un grand nombre de sous-pas de temps égaux à δt . Les champs fluides sont interpolés à chaque sous-pas de temps à partir des champs $\mathbf{x}^n$ et $\mathbf{x}^{n+1;i}$. Ce point est une particularité du code XTOR-K : les autres codes hybrides utilisent généralement des méthodes cinétiques, et ont des pas de temps égaux pour les deux modèles, MHD et cinétique. Les conditions initiales sont fournies par les vitesses et positions des particules au début du pas de temps : $\mathbf{r}_k^n$ et $\mathbf{v}_k^n$.

- les positions et vitesses des particules obtenus permettent de calculer le tenseur

de pression particulière $\Pi^{n+1;i+1}$, qui représente ce champ au temps t^{n+1} . On passe à l'itération $i + 1$ en reprenant depuis la première étape, le terme de source étant maintenant $\Pi^{n+1;i+1}$.

On considère que la méthode converge si l'incrément des champs fluides $\|\mathbf{x}^i - \mathbf{x}^i\| \equiv \|\delta\mathbf{x}^i\|$ est en-dessous d'un seuil à choisir après l'avancée MHD.

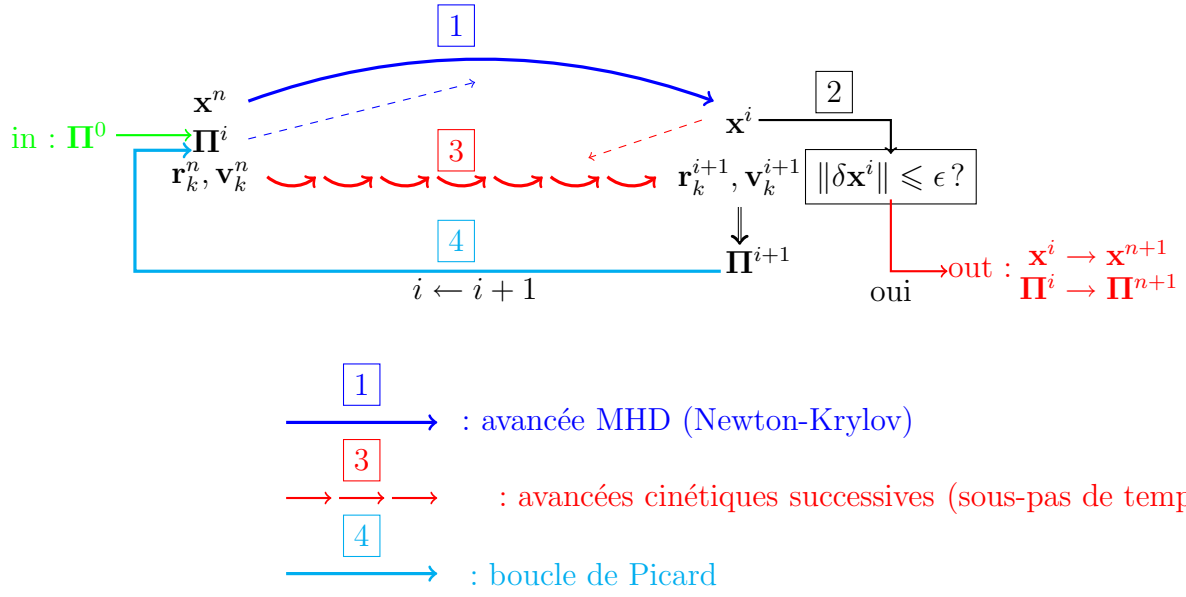
Le schéma d'un pas de temps fluide s'écrit donc symboliquement :

$$\begin{aligned}\frac{\mathbf{v}^j - \mathbf{v}^n}{\Delta t} &= \mathbf{F}_v \left(\frac{\mathbf{v}^j + \mathbf{v}^n}{2}, \frac{\mathbf{B}^j + \mathbf{B}^n}{2}, \frac{p^j + p^n}{2} \right) + \tilde{\mathbf{F}}_{vp} \left(\frac{\Pi^j + \Pi^n}{2} \right) \\ \frac{\mathbf{B}^j - \mathbf{B}^n}{\Delta t} &= \mathbf{F}_B \left(\frac{\mathbf{v}^j + \mathbf{v}^n}{2} \right) \\ \frac{p^j - p^n}{\Delta t} &= F_p \left(\frac{\mathbf{v}^j + \mathbf{v}^n}{2} \right)\end{aligned}$$

$$\Pi^j = \mathbf{K}(\mathbf{x}^{j-1}, \mathbf{x}^n; \{\mathbf{r}_k^n, \mathbf{v}_k^n\})$$

Où pour plus de simplicité, on a omis l'indice $n + 1$: on a écrit $\mathbf{x}^j$ au lieu de $\mathbf{x}^{n+1;j}$. La variable p^j représente la pression d'une population fluide, en particulier, celle des électrons. L'opérateur $\mathbf{K}$ dépend du modèle cinétique choisi. Dans notre cas, à chaque pas de temps fluide, cet opérateur intègre le mouvement des particules sur des sous-pas de temps successifs, en interpolant les champs à chaque sous-pas de temps $\mathbf{x}$ à partir de $\mathbf{x}^n$, la valeur en t^n , et $\mathbf{x}^{j-1}$, la valeur en t^{n+1} calculée à l'itération précédente $j - 1$. Les variables $\mathbf{r}_k^n, \mathbf{v}_k^n$ sont les conditions initiales de cette intégration. Les opérateurs F sont ceux qui ont été utilisés dans les chapitres précédents. L'opérateur $\tilde{\mathbf{F}}_v^p$ est différent de $\mathbf{F}_v^p$ (cf. chapitre 3) parce qu'il prend la divergence d'un tenseur et non le gradient d'un scalaire.

On a résumé le déroulement d'un pas de temps sur le schéma suivant :



Les particularités de cette méthode méritent d'être notées. Le but de l'algorithme est de faire un nombre faible d'itérations de Picard (typiquement 3) ; on autorise un nombre plus grand d'itérations de Newton-Krylov (plusieurs dizaines). L'étape de l'avancée des particules prenant un temps proportionnel au nombre de particules chargées, il est logique de chercher à minimiser les itérations de Picard. Pour fonctionner de cette façon, il faut que le schéma de Picard se comporte bien (*i.e.*, converge rapidement) : il faut que l'incrément des champs à chaque avancée fluide diminue de façon monotone d'itération en itération de Picard. Autrement dit, il faut montrer que le problème résolu par Picard soit contractant. C'est ce qu'on montre par un calcul semi-qualitatif dans la section suivante. On peut alors utiliser une condition de sortie brutale pour passer au pas de temps suivant avec la certitude de faire à chaque pas de temps peu d'itérations de Picard. Au contraire, le solveur Newton-Krylov converge de façon non nécessairement monotone, mais même de nombreuses itérations ne prennent pas un temps de calcul prohibitif.

Cette forte asymétrie numérique entre partie cinétique et partie MHD est particulière au code XTOR-K. Dans les autres codes hybrides que nous avons cité, elle se manifeste par le choix d'une méthode δf , ce qui implique que la majeure partie de la population cinétique reste proche d'un équilibre. Dans le cas d'XTOR-K, on a depuis le début rejeté le principe d'une méthode δf ; le prix à payer est l'apparition d'un sous-pas de temps particulière, qui démultiplie les opérations à effectuer pour l'avancée des particules. Le principe de l'algorithme de couplage est alors de limiter au maximum les itérations impliquant cette avancée, de façon à garder des temps de calculs corrects.

Les codes hybrides qui n'utilisent pas de sous-pas de temps particulaire effectuent un plus grand nombre d'itérations pour converger de façon self-consistante. Les champs MHD et le tenseur de pression particulaire doivent être transmis à chaque itération entre la partie MHD et la partie cinétique du code. Si on envisage une parallélisation de la partie cinétique, pour pouvoir augmenter le nombre de particules, l'étape de transfert peut alors devenir la plus importante en termes de temps de calcul. Inversement, dans notre cas, les parties fluides et cinétiques du code n'échangent d'informations qu'au début et à la fin de chaque itération de Picard. Par conséquent, on peut paralléliser massivement l'avancée des particules sur l'ensemble des pas de temps particuliers, une fois les champs fluides transmis au début de l'itération.

4.2.2 Condition suffisante de stabilité linéaire

Pour étudier la stabilité linéaire du schéma hybride, on commence par simuler la partie particulaire de la pression par l'opérateur L_p , c'est-à-dire qu'on traite en fait les particules comme un fluide magnétohydrodynamique, en traitant la partie magnétique par le solveur Newton-Krylov, et la pression par Picard, comme explicité ci-dessous. Cette manipulation n'a pas d'intérêt algorithmique, mais elle permet d'obtenir une bonne idée du comportement du schéma. À chaque itération de Picard, les champs de vitesse et magnétique sont d'abord traités ensemble, implicitement, puis la pression est avancée explicitement.

Expression du facteur d'amplification du schéma

À la i^{eme} itération de Picard, le schéma numérique linéarisé s'écrit :

$$\begin{aligned}\frac{\mathbf{v}^j - \mathbf{v}^n}{\Delta t} &= L_v^B \left(\frac{\mathbf{B}^j + \mathbf{B}^n}{2} \right) + L_v^p \left(\frac{p^j + p^n}{2} \right) \\ \frac{\mathbf{B}^j - \mathbf{B}^n}{\Delta t} &= L_B \left(\frac{\mathbf{v}^j + \mathbf{v}^n}{2} \right) \\ \frac{p^j + p^n}{\Delta t} &= L_p \left(\frac{\mathbf{v}^{j-1} + \mathbf{v}^n}{2} \right)\end{aligned}$$

L'indice n correspond aux champs au temps t^n , et reste constant au cours des itérations de Picard, dénotées par l'indice j . Les champs $\mathbf{x}^j$ représentent les champs au temps t^{n+1} calculées à la j^{eme} itération de Picard. Lorsque la condition de convergence est atteinte, *i.e.* :

$$\|\mathbf{x}^j - \mathbf{x}^{j-1}\| < \epsilon,$$

le champs $\mathbf{x}^j$ devient le $\mathbf{x}^{n+1}$, et on passe au pas de temps suivant. On a laissé de côté la viscosité, et on considère que toute la pression est issue de la partie particulière du code ; dans la pratique, une fraction de la pression est issue du code MHD.

Contrairement au schéma implicite, les champs ne sont plus tous traités de la même façon, par conséquent, on doit travailler avec le système d'équations complet. De plus, le système n'est plus diagonal pour les modes MHD, ce qui oblige à considérer les modes séparément (shear Alfvén, magnéto-sonore lente et rapide).

Les équations sur $\mathbf{B}$ et p peuvent s'écrire :

$$\begin{aligned}\mathbf{B}^j &= \mathbf{B}^n + \Delta t L_B \left(\frac{\mathbf{v}^j + \mathbf{v}^n}{2} \right) \\ p^j &= p^n + \Delta t L_p \left(\frac{\mathbf{v}^{j-1} + \mathbf{v}^n}{2} \right)\end{aligned}$$

En substituant dans l'équation sur $\mathbf{v}$, on obtient :

$$\frac{\mathbf{v}^j - \mathbf{v}^n}{\Delta t} = L_v^B \left(\mathbf{B}^n + \Delta t L_B \left(\frac{\mathbf{v}^j + \mathbf{v}^n}{4} \right) \right) + L_v^p \left(p^n + \Delta t L_p \left(\frac{\mathbf{v}^{j-1} + \mathbf{v}^n}{4} \right) \right)$$

En regroupant les termes, on arrive à une forme plus compacte :

$$\left(1 - \frac{\Delta t^2}{4} L_v^B L_B \right) (\mathbf{v}^j - \mathbf{v}^n) = S + \frac{\Delta t^2}{4} L_v^p L_p (\mathbf{v}^{j-1} - \mathbf{v}^n)$$

Avec :

$$S = \Delta t L_v^B \left(\mathbf{B}^n + \frac{\Delta t}{2} L_B (\mathbf{v}^n) \right) + \Delta t L_v^p \left(p^n + \frac{\Delta t}{2} L_p (\mathbf{v}^n) \right)$$

Le terme S est constant d'une itération de Picard à la suivante. Par conséquent, lorsqu'on soustrait l'équation à l'itération i à l'équation à $i + 1$, ce terme disparaît, ainsi que tous les termes en $\mathbf{x}^n$, et on obtient :

$$\left(1 - \frac{\Delta t^2}{4} L_v^B L_B \right) (\mathbf{v}^{j+1} - \mathbf{v}^j) = \frac{\Delta t^2}{4} L_v^p L_p (\mathbf{v}^j - \mathbf{v}^{j-1}) \quad (34)$$

Autrement dit, le facteur d'amplification du schéma est égal à la norme du rapport des opérateurs $\Delta t^2/4 L_v^p L_p$ et $(1 - \Delta t^2 L_v^B L_B/4)$ appliqués aux modes physiques.

L'opérateur $L_v^B L_B$ est l'opérateur MHD auto-adjoint du deuxième ordre³² ; les modes physiques sont donc ses vecteurs propres, et les valeurs propres associées sont connues (ce sont les fréquences fournies par les relations de dispersion classiques, cf *e.g.* [1], p. 235). L'action de l'opérateur $L_v^p L_p$ est moins simple. L_v^p contenant un gradient et L_p contenant une divergence, l'opérateur $L_v^p L_p$ peut agir comme un laplacien, qui menacerait de faire diverger le schéma.

Estimation du facteur d'amplification par décomposition spectrale de $L_v^B L_B$

Soit $|\mathbf{v}_\alpha\rangle$ un vecteur propre normé³³ de cet opérateur, associé à la valeur propre³⁴ ω_α^2 . On note le transposé de son conjugué $\langle \mathbf{v}_\alpha|$, avec $\langle \mathbf{v}_\alpha|\mathbf{v}_\alpha\rangle = 1$. Avec ces notations, on écrit formellement le rapport des opérateurs $(1 - \Delta t^2 L_v^B L_B/4)^{-1} \Delta t^2/4 L_v^p L_p$, appliqué au mode $|\mathbf{v}_\alpha\rangle$. En utilisant le fait que $(L_v^B L_B)$ est auto-adjoint, on trouve que :

$$\langle \mathbf{v}_\alpha|(1 - \frac{\Delta t^2}{4} L_v^B L_B)^{-1} \frac{\Delta t^2}{4} L_v^p L_p|\mathbf{v}_\alpha\rangle = \frac{\Delta t^2}{4} \frac{\langle \mathbf{v}_\alpha|L_v^p L_p|\mathbf{v}_\alpha\rangle}{1 + \frac{1}{4}\omega_\alpha^2 \Delta t^2} \quad (35)$$

Et la condition suffisante de stabilité devient :

$$\frac{\Delta t^2}{4} \frac{\langle \mathbf{v}_\alpha|L_v^p L_p|\mathbf{v}_\alpha\rangle}{1 + \frac{1}{4}\omega_\alpha^2 \Delta t^2} < 1 \quad \text{pour tout vecteur propre } \mathbf{v}_\alpha. \quad (36)$$

4.2.3 Comportement pour les modes MHD fondamentaux

Le schéma numérique n'étant pas diagonal pour les modes, le rapport doit être évalué séparément pour les trois modes MHD de base.

Ondes magnétosonores rapides

Les ondes magnétosonores rapides se traduisent par une compression du champ magnétique et de la densité (*i.e.*, $\nabla \cdot \mathbf{v}$ est non nul), avec $\omega_\alpha > \omega_A$ (ω_A est la fréquence d'Alfvén). À faible β , l'onde se réduit à une onde compressionnelle transverse, et la pulsation tend vers ω_A .

Dans ce cas, les opérateurs $L_v^p L_p$ et $L_v^B L_B$ se ramènent à des laplaciens :

$$\frac{\Delta t^2}{4} L_v^p L_p \sim \frac{\Delta t^2}{4} \Gamma p_0 \Delta \quad \text{et} \quad \frac{\Delta t^2}{4} L_v^B L_B \sim B_0^2 \frac{\Delta t^2}{4} \Delta$$

³²Pour simplifier, on a omis la pression dans l'opérateur MHD.

³³Pour la norme usuelle, $\|\mathbf{v}\| = \langle \mathbf{v}|\mathbf{v}\rangle \equiv \int \mathbf{v}^* \cdot \mathbf{v} d\mathbf{r}$.

³⁴Notons que le spectre de l'opérateur MHD est continu dans le domaine des shear Alfvén, mais il est discret pour l'opérateur numérique, ce qui nous permet de faire le calcul avec les notations habituelles.

De sorte que le facteur d'amplification du schéma peut s'écrire :

$$\langle \mathbf{v}_\alpha | (1 - \frac{\Delta t^2}{4} L_v^B L_B)^{-1} \frac{\Delta t^2}{4} L_v^p L_p | \mathbf{v}_\alpha \rangle \sim \left| \frac{\frac{\Delta t^2}{4} \Gamma p_0 k^2}{1 + B_0^2 k^2 \frac{\Delta t^2}{4}} \right| \sim \beta$$

Comme $\beta \ll 1$ dans nos cas, on est bien stable.

Dans ce cas, le laplacien contenu dans $L_v^p L_p$ est compensé par celui qui sort de $L_v^B L_B$. Cela reflète le fait qu'à faible β , la compression du champ magnétique est dominante devant celle de la pression. Le schéma est donc très contractant sur les ondes magnéto-sonores rapides.

Ondes magnéto-sonores lentes

Dans ce cas, on a à nouveau compression de la densité et du champ, mais avec $\omega_\alpha < \omega_A$. À faible β , l'onde se ramène à une onde sonore dans la direction parallèle à $\mathbf{B}$: la compression de la densité est dominante. Dans ce cas, l'opérateur $L_v^B L_B$ appliqué au mode tend vers 0, puisque la perturbation magnétique devient négligeable. On doit alors considérer un schéma numérique légèrement différent, avec la pression divisée en deux parties, une traitée en implicite avec $\mathbf{v}$ et $\mathbf{B}$ et l'autre traitée à part pour représenter la pression particulière. La fraction de pression particulière p_k est notée δ_k . et la schéma complet s'écrit :

$$\begin{aligned} \frac{\mathbf{v}^j - \mathbf{v}^n}{\Delta t} &= L_v^B \left(\frac{\mathbf{B}^j + \mathbf{B}^n}{2} \right) + (1 - \delta_k) L_v^p \left(\frac{p^j + p^n}{2} \right) + \delta_k L_v^p \left(\frac{p_k^j + p_k^n}{2} \right) \\ \frac{\mathbf{B}^j - \mathbf{B}^n}{\Delta t} &= L_B \left(\frac{\mathbf{v}^j + \mathbf{v}^n}{2} \right) \\ \frac{p^j + p^n}{\Delta t} &= L_p \left(\frac{\mathbf{v}^j + \mathbf{v}^n}{2} \right) \\ \frac{p_k^j + p_k^n}{\Delta t} &= L_p \left(\frac{\mathbf{v}^{j-1} + \mathbf{v}^n}{2} \right) \end{aligned}$$

L'équation 34 devient :

$$\left(1 - \frac{\Delta t^2}{4} (L_v^B L_B + (1 - \delta_k) L_v^p L_p) \right) (\mathbf{v}^{j+1} - \mathbf{v}^j) = \delta_k \frac{\Delta t^2}{4} L_v^p L_p (\mathbf{v}^j - \mathbf{v}^{j-1}) \quad (37)$$

On voit que le terme de gauche est en fait une somme d'opérateurs ; l'opérateur dû à la pression est bien nul quand $\delta_k = 1$, ce qui correspond au cas étudié jusqu'ici. Quand l'opérateur $L_v^B L_B$ s'annule, il devient nécessaire d'inclure le second terme.

Les deux côtés de l'équation contiennent le même opérateur $L_v^p L_p$; appliqué à une onde sonore, il fait apparaître la valeur propre $\omega_\alpha^2 = (k_\parallel v_s)^2$, où $v_s \equiv (\Gamma p_0)^{1/2}$ est la vitesse du son (on pose $\rho_0 = 1$). Le facteur d'amplification est donc :

$$\lambda = \frac{\delta_k \Gamma p_0 k_\parallel^2 \Delta t^2 / 4}{1 + (1 - \delta_k) \Gamma p_0 k_\parallel^2 \Delta t^2 / 4} = \frac{\delta_k x}{1 + (1 - \delta_k) x}$$

Où on a posé $x = \Gamma p_0 k_\parallel^2 \Delta t^2 / 4$. On a alors :

$$\frac{\delta_k x}{1 + (1 - \delta_k) x} < 1 \iff x(2\delta_k - 1) < 1$$

Par conséquent, le schéma est stable si $\delta_k < 1/2$ (car $x > 0$), et sinon, l'inégalité

$$(2\delta_k - 1) \Gamma p_0 k_\parallel^2 \frac{\Delta t^2}{4} < 1$$

est une condition suffisante de stabilité.

On a vérifié qualitativement cette condition en modifiant XTOR-K selon le schéma numérique montré sur la page précédente, et en étudiant sa convergence dans l'espace de paramètres $(\delta_k, \Delta t)$. Les résultats sont présentés sur la figure 17.

Pour obtenir une condition numérique sur Δt , on maximise les paramètres : $\delta_{k,max} = 1, k_{\parallel,max} = mmax/R$, où $mmax$ est le nombre de modes poloïdaux simulés et ϵ la taille du maillage. Cette égalité vient du fait que : $k_\parallel = \mathbf{k} \cdot \mathbf{B}/B \sim (B_\phi n/R + B_\theta m/r)/B \sim (n + m/q)/R$, en prenant $q \sim 1$ et $mmax \gg nmax$. Notons que dans le domaine non-linéaire, une onde sonore pourrait avoir une limite plus faible en $k_\parallel$, ce qui rendrait le schéma plus stable.

On obtient finalement :

$$\Delta t < \frac{2R}{mmax} \frac{1}{(\Gamma p_0)^{1/2}}$$

Dans nos unités, $p_0/R^2 \sim p_0/A^2 \sim \beta$, et on écrit finalement la condition :

$$\Delta t \lesssim \frac{2}{mmax} \frac{1}{\sqrt{\beta}} \sim 1$$

La figure 17 montre un bon accord entre le comportement prévu et le test numérique, pour $\delta_k \sim 1/2$ et quand Δt tend asymptotiquement vers 1. Le schéma est toujours stable si $\delta_k < 1/2$ et peut être stable si $\delta_k > 1/2$ et Δt n'est pas trop important (la valeur critique est de l'ordre de 1, ce qui correspond au pas de temps utilisé par XTOR-2F en régime faiblement non-linéaire).

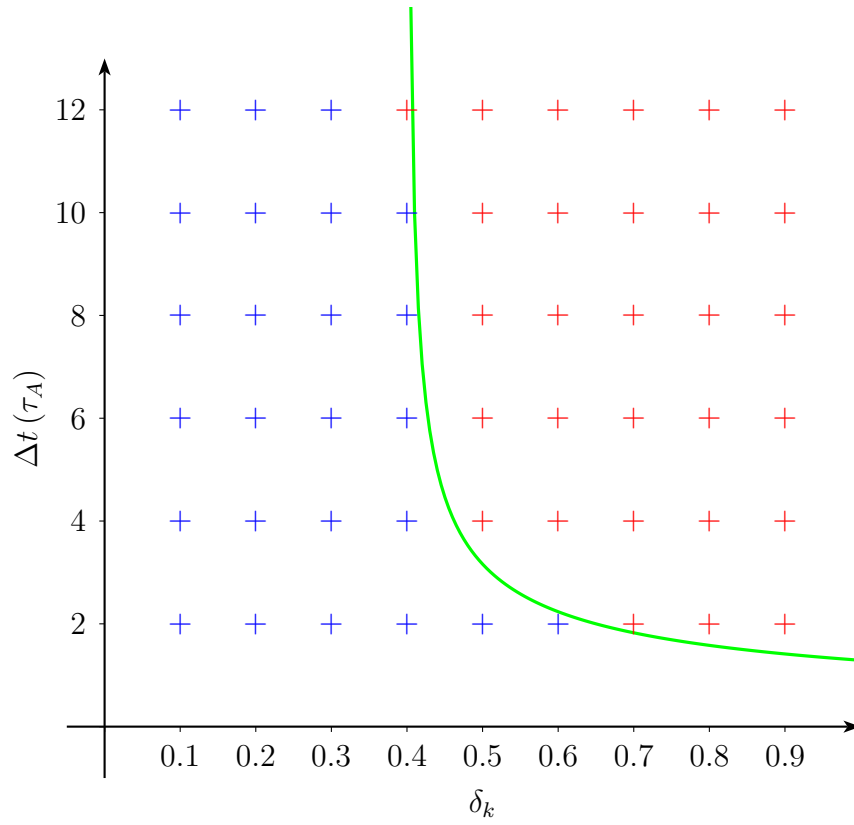


FIG. 17 – Test de convergence du schéma hybride en $\delta_k, \Delta t$. Les cas convergents sont marqués en bleu, les cas divergents en rouge. La courbe verte est proportionnelle à $(\delta_k - 0.5)^{-1/2}$.

Ondes shear Alfvén

Dans le cas des shear Alfvén, la branche de la relation de dispersion concernée permet d'avoir $\omega_\alpha \rightarrow 0$, c'est-à-dire $\langle \mathbf{v}_\alpha | L_v^B L_B | \mathbf{v}_\alpha \rangle \rightarrow 1$. Pour avoir la condition la plus stricte, on considère la limite $L_v^B L_B = 1$. Reste à évaluer $\langle \mathbf{v}_\alpha | L_v^p L_p | \mathbf{v}_\alpha \rangle$.

Pour une onde shear Alfvén, on peut utiliser la représentation suivante pour les vitesses :

$$\mathbf{v}_\alpha = \mathbf{b} \times \nabla \Phi$$

Où $\mathbf{b} \equiv \mathbf{B}/B$. On a alors :

$$\begin{aligned} (L_v^p L_p) | \mathbf{v}_\alpha \rangle &= \nabla (\mathbf{v}_\alpha \cdot \nabla p_0 + \Gamma p_0 \nabla \cdot \mathbf{v}_\alpha) \\ &= \nabla (\mathbf{v}_\alpha \cdot \nabla p_0 + \Gamma p_0 \nabla \Phi \cdot (\nabla \times \mathbf{b})) \\ &= \nabla (\mathbf{v}_\alpha \cdot \nabla p_0 + \Gamma p_0 \nabla \Phi \cdot \mathbf{K}) \end{aligned}$$

Où $\mathbf{K} = \nabla \times \mathbf{b}$ est lié à la courbure du champ, et est de l'ordre de $1/R$.

On écrit alors³⁵ :

$$\langle \mathbf{v}_\alpha | L_v^p L_p | \mathbf{v}_\alpha \rangle = \int \mathbf{v}_\alpha^* \cdot \nabla (\mathbf{v}_\alpha \cdot \nabla p_0 + \Gamma p_0 \mathbf{K} \cdot \nabla \Phi)$$

Par intégration par parties, en éliminant le terme de bord qui est nul pour des raisons de continuité, on obtient :

$$\begin{aligned} \langle \mathbf{v}_\alpha | L_v^p L_p | \mathbf{v}_\alpha \rangle &= \int (\nabla \cdot \mathbf{v}_\alpha^*) (\mathbf{v} \cdot \nabla p_0 + \Gamma p_0 \mathbf{K} \cdot \nabla \Phi) \\ &= \int (\mathbf{K} \cdot \nabla \Phi) (\mathbf{v} \cdot \nabla p_0 + \Gamma p_0 \mathbf{K} \cdot \nabla \Phi) \\ &\leq (\Gamma K_{max}^2 p_{0,max} + K_{max} (\nabla p_0)_{max}) \int \|\nabla \Phi\| = (\Gamma K_{max}^2 p_{0,max} + K_{max} (\nabla p_0)_{max}) \end{aligned}$$

En notant qualitativement que $K \sim 1/R$ et que $p_0 \sim \beta_p/r$, aux constantes près, on constate que la relation ci-dessus laisse une certaine marge de manœuvre pour vérifier la condition de stabilité qu'on avait énoncée. En pratique une relation du type $\Delta t \lesssim \beta_p^{-1/2}$ suffit à avoir une bonne stabilité.

Cette estimation est qualitative. On pourrait essayer de trouver des valeurs numériques plus précises, mais notre démarche ici est plutôt d'établir quelles sont les conditions de stabilité qui sont dominantes. Numériquement, un test de convergence analogue

³⁵On suppose le vecteur $|\mathbf{v}_\alpha\rangle$ normé, pour plus de simplicité.

au précédent en $\delta_k, \Delta t$ ($0.1 \leq \delta_k \leq 0.9, 2 \leq \Delta t \leq 12$ mené en posant $\Gamma = 0$ a montré que le schéma était toujours stable dans cette plage de paramètres. Poser $\Gamma = 0$ revient à éliminer les ondes sonores sans toucher aux shear Alfvén (puisque $p_1 = 0$ pour les shear Alfvén de toute façon). La stabilité du schéma prouve que ce sont les ondes sonores qui fournissent le critère de stabilité limitant.

Les résultats obtenus sont donc encourageants. De plus, dans le cas où la partie particulaire est traitée avec un vrai modèle cinétique (et non fluide comme pour ces tests), la compressibilité (le terme en Γ) ne se comporte pas comme un opérateur raide. La dynamique des particules étant libre dans la direction parallèle au champ, il ne s'agit pas d'une dynamique de diffusion collisionnelle, et la pression particulaire ne se comporte pas comme un $\Gamma \nabla \cdot \mathbf{v}$ habituel. Au lieu de ça, le mouvement parallèle est amorti par mélange de phase.

4.3 Conclusion

Les particules rapides jouent un rôle majeur dans les tokamaks : chauffage, contrôle des profils de courant et de la stabilité du plasma. Elles sont créées par le chauffage par ondes, l'injection de neutres, ou par la réaction de fusion. Elles donnent lieu à de nouvelles instabilités (fishbones, TAEs) qui peuvent créer lieu à un transport anormal.

XTOR-K utilise un algorithme PIC pour résoudre exactement le mouvement des particules chaudes en intégrant l'équation de Newton-Lorentz sur toute la distribution des particules chaudes. L'absence d'hypothèse de type drift ou gyro-cinétique et de type δf , rend cette méthode plus coûteuse que dans les autres codes hybrides existants, mais elle permet de prendre en compte tous les effets de rayon de Larmor fini et de traiter une fraction importante du plasma en cinétique. L'algorithme de couplage choisi, reposant sur la méthode de Newton-Krylov/Picard, permet de transférer la majorité des itérations sur le solveur MHD Newton-Krylov, qui est très robuste. Les transferts de données entre les deux parties du code sont minimisés et la partie particulaire se parallélise naturellement. L'étude de la stabilité linéaire du schéma assure qu'une condition de convergence simple sera atteinte en peu d'itérations de Picard. Dans une étude fluide du schéma Newton-Krylov/Picard, la condition de stabilité la plus restrictive (sur Δt et δ_k) est due aux ondes sonores, à cause du terme en $\Gamma \nabla \cdot \mathbf{v}$, qui introduit un laplacien. Cependant, la dynamique parallèle des particules n'est pas gouvernée par les collisions mais par le mélange de phase ; l'opérateur obtenu est moins raide et le schéma plus stable.

5 Intégration numérique des trajectoires

Dans ce chapitre nous présentons l'algorithme utilisé pour intégrer l'équation du mouvement pour chaque particule : la méthode de Boris. Cet algorithme (le « pousseur ») sera appelé un très grand nombre de fois pendant une simulation ; il doit donc être optimisé pour réduire au maximum le temps de calcul.

On commence par présenter la méthode de Boris, en toute généralité, puis adaptée en géométrie torique. Il s'agit d'un schéma leap-frog, ce qui assure des bonnes propriétés de stabilité. Nous détaillons ensuite l'implémentation dans le code : normalisations, paramètres, maillage, interpolations des champs importés depuis la partie MHD. Nous faisons le compte des opérations élémentaires nécessaires au pousseur ; le faible nombre de fonctions transcendentes (interpolations de Fourier dans la direction toroïdale) assure un temps de calcul minimal ; nous étudions également la répartition du temps de calcul dans chaque étape du pousseur. Enfin, nous caractérisons le fonctionnement du pousseur en donnant les écart-types des invariants du mouvement, E_{cin} , μ et P_ϕ . Le pousseur a une très bonne stabilité, mais pour assurer une bonne conservation des invariants sur des temps longs (plusieurs milliers de tours de tore), on se ménage une marge de sécurité sur le pas de temps (on fera en général au moins 8 pas de temps par rotation cyclotronique).

5.1 Schéma leap-frog et méthode de Boris

5.1.1 Algorithme de Boris

Discretisation de l'équation du mouvement

L'équation du mouvement que le pousseur résout s'écrit :

$$m \frac{d\mathbf{v}}{dt} = q(\mathbf{E} + \mathbf{v} \times \mathbf{B}) \quad (38)$$

La nature du mouvement est bien connue : il s'agit d'une accélération rectiligne dans la direction de $\mathbf{E}$ et d'une rotation autour de $\mathbf{B}$. La rotation se fait avec une pulsation $\omega_c = qB/m$, la pulsation cyclotronique, et un rayon de giration $\rho_L = v_\perp/\omega_c$, le rayon de Larmor.

Plusieurs méthodes d'intégration numérique pour cette équation ont été développées dans les années 70, notamment par Buneman et Boris. Elles sont décrites par exemple dans [67], p. 58, et dans [68], p. 111. Les plus efficaces reposent sur un schéma de type leap-frog :

$$m \frac{\mathbf{v}^{n+1/2} - \mathbf{v}^{n-1/2}}{\delta t} = q(\mathbf{E}^n + \frac{\mathbf{v}^{n+1/2} + \mathbf{v}^{n-1/2}}{2} \times \mathbf{B}^n) \quad (39)$$

$$\frac{\mathbf{r}^{n+1} - \mathbf{r}^n}{\delta t} = \mathbf{v}^{n+1/2} \quad (40)$$

Avec $\mathbf{x}^{n\pm 1/2} = \mathbf{x}(t_n \pm \delta t/2)$. On note le pas de temps particulière δt pour bien le différencier du pas de temps MHD, Δt . Ce schéma est centré en temps (*i.e.* réversible). Dans la direction perpendiculaire à $\mathbf{B}^n$, il est inconditionnellement stable à $\mathbf{E}$ nul. La partie parallèle a la stabilité d'un schéma leap-frog classique. Il est exact jusqu'à l'ordre 2 en δt . À titre de comparaison, une méthode Runge-Kutta, même à l'ordre 4, entraîne des dérives numériques dues à une accumulation d'erreurs d'un pas de temps à l'autre, que l'on évite avec le leap-frog.

On vérifie qu'en régime stationnaire ($\delta t \rightarrow \infty$), la vitesse moyenne tend vers la dérive de champs croisés :

$$\frac{\mathbf{v}^{n+1/2} + \mathbf{v}^{n-1/2}}{2} \xrightarrow{\delta t \rightarrow \infty} \frac{\mathbf{E}^n \times \mathbf{B}^n}{(B^n)^2} \quad (41)$$

Schéma de Boris

L'algorithme le plus utilisé pour calculer 39 est celui proposé par Boris, en 1970 (cf. [69], [70]). Cet algorithme sépare les forces électrique et magnétique, puis permet d'effectuer la rotation du vecteur $\mathbf{v}$ autour de $\mathbf{B}$ sans décomposer les vecteurs en composantes, ce qui évite de devoir projeter les vecteurs sur des axes fixes à chaque pas de temps, et donc de calculer des sinus et des cosinus. Il traite exactement le produit vectoriel (*i.e.*, $\mathbf{v}$ et $\mathbf{B}$ sont perpendiculaires par construction).

Commençons par poser :

$$\begin{aligned} \mathbf{v}^- &= \mathbf{v}^{n-1/2} + \frac{q}{m} \frac{\delta t}{2} \mathbf{E}^n \\ \mathbf{v}^+ &= \mathbf{v}^{n+1/2} - \frac{q}{m} \frac{\delta t}{2} \mathbf{E}^n \end{aligned}$$

Cela donne en remplaçant dans l'équation 39 :

$$\frac{\mathbf{v}^+ - \mathbf{v}^-}{\delta t} = \frac{q}{2m} ((\mathbf{v}^+ + \mathbf{v}^-) \times \mathbf{B}^n) \quad (42)$$

On peut voir facilement que la norme de $\mathbf{v}$ est conservée, *i.e.*, $v^+ = v^-$. Le vecteur $\mathbf{v}^+$ est donc l'image de $\mathbf{v}^-$ par une rotation autour de l'axe $\mathbf{B}^n$. En posant $\boldsymbol{\Omega} = (q/m)\mathbf{B}^n$, et $\Omega = \omega_c = \|\boldsymbol{\Omega}\|$, on peut exprimer l'angle θ de la rotation comme³⁶ :

$$\tan\left(\frac{\theta}{2}\right) = \Omega \frac{\delta t}{2}$$

³⁶On choisit le signe de q pour avoir θ positif et alléger les notations

Soit :

$$\theta = 2 \arctan \left(\frac{\Omega \delta t}{2} \right) = \Omega \delta t \left(1 - \frac{(\Omega \delta t)^2}{12} + o(\delta t^2) \right)$$

La discrétisation entraîne donc une erreur $\epsilon \propto \delta t^2$ sur l'angle de rotation cyclotronique ($\epsilon < 1\%$ si $\Omega \delta t < 0.35$). Le rayon cyclotronique est lui aussi faussé : $R_c = \rho_L / \cos(\Omega \delta t / 2)$. Notons cependant qu'il est possible de rectifier cette erreur en introduisant le facteur correctif λ dans l'équation 42 :

$$\frac{\mathbf{v}^+ - \mathbf{v}^-}{\delta t} = \lambda \frac{q}{2m} ((\mathbf{v}^+ + \mathbf{v}^-) \times \mathbf{B}^n)$$

L'angle de rotation modifié vaut alors $\tan(\tilde{\theta}/2) = \lambda \Omega \delta t / 2$, et il suffit de poser :

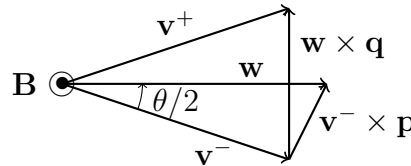
$$\lambda \equiv \frac{\tan(\Omega \delta t / 2)}{(\Omega \delta t / 2)}$$

pour obtenir un angle de rotation corrigé $\tilde{\theta} = \Omega \delta t$, ce qui correspond à la rotation cyclotronique exacte.

Pour un champ $\mathbf{B}$ arbitraire, la méthode de Boris permet de réaliser une rotation d'angle θ uniquement avec des opérations vectorielles. Cette rotation est exécutée en deux étapes :

$$\begin{aligned} \mathbf{w} &= \mathbf{v}^- + \mathbf{v}^- \times \mathbf{p} & \text{où } \mathbf{p} &= \tan \left(\frac{\theta}{2} \right) \frac{\mathbf{B}^n}{B^n} \\ \mathbf{v}^+ &= \mathbf{v}^- + \mathbf{w} \times \mathbf{q} & \text{où } \mathbf{q} &= \frac{2}{1 + p^2} \mathbf{p} \end{aligned}$$

Le schéma ci-dessous illustre l'algorithme de rotation.



Remarquons pour finir que l'introduction du facteur $\lambda = \tan(\Omega\delta t/2)/(\Omega\delta t/2)$ corrige la rotation cyclotronique, mais fausse la dérive des particules, qui diffère alors de la dérive de champs croisés. L'équation 41 n'est plus vraie. Cela est inacceptable dans le cas d'un code hybride, puisque c'est cette dérive qui domine la dynamique perpendiculaire de la partie MHD du plasma. Il est cependant possible de corriger la dérive des particules en introduisant λ en facteur devant tout le terme de droite de l'équation 39. Dans ce cas, le mouvement cyclotronique et la dérive de champs croisés sont conservés ; mais le problème a été déplacé sur la dynamique parallèle. En effet, on a introduit un facteur λ devant le champ électrique parallèle de façon arbitraire.

Pour éviter cela, on préfère ne pas introduire le facteur λ , et maintenir δt petit de sorte que l'erreur sur le mouvement cyclotronique soit négligeable. Notons que dans ce cas, il n'est pas nécessaire d'évaluer de tangente pour effectuer la rotation, puisque la seule grandeur à évaluer est $\tan(\theta/2)$, qui est comme on l'a vu égale à $\Omega\delta t/2$.

5.1.2 Adaptation en géométrie torique

Système de coordonnées

L'algorithme de Boris est utilisé en géométrie cartésienne. L'utiliser dans une base locale non-orthogonale, comme celle de XTOR, n'est pas cohérent. Cela oblige à calculer des sinus et des cosinus pour les produits scalaires et vectoriels, alors que l'avantage de Boris est justement de n'utiliser que des opérations arithmétiques. D'un autre côté, un maillage cartésien complet est clairement inadapté à la géométrie tokamak.

De plus, le nombre de points de maillage dans la direction ϕ est généralement faible (~ 12) dans XTOR. Les mailles ont donc une grande largeur angulaire dans la direction toroïdale : interpoler les champs à partir des valeurs aux sommets d'une maille serait abusif. Il est plus logique de conserver la coordonnée angulaire ϕ , et d'interpoler les champs en sommant leurs composantes de Fourier en ϕ . Cela permet d'utiliser le maximum d'informations. Dans les autres directions, le maillage d'XTOR est plus resserré, et une interpolation bilinéaire est utilisée actuellement. On peut raffiner cette interpolation si besoin est.

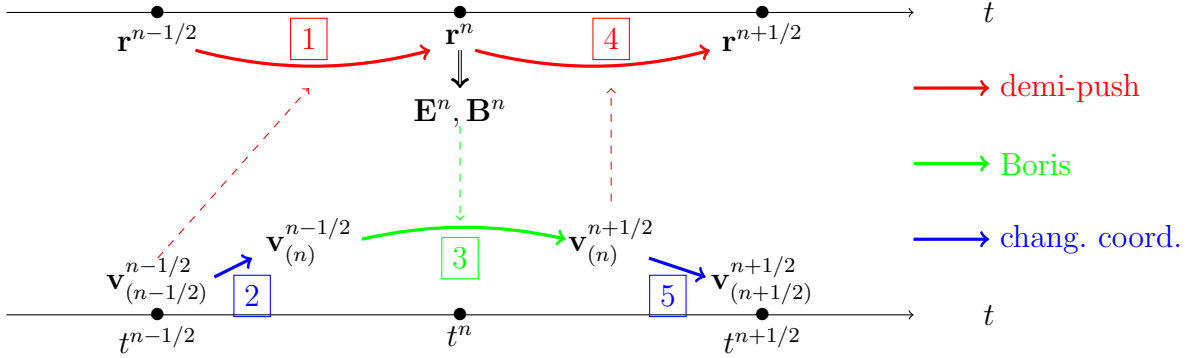
Par conséquent, on choisit un maillage cartésien du plan poloidal, et l'angle toroïdal ϕ comme troisième coordonnée. On obtient les coordonnées cylindriques classiques. La base locale $(\mathbf{e}_R, \mathbf{e}_\phi, \mathbf{e}_Z)$ est bien orthonormale³⁷. La simplicité du pousseur est alors en partie compensée par la nécessité de transporter les grandeurs fluides du maillage XTOR au nouveau maillage ; mais cette opération n'est faite qu'à chaque pas de temps fluide, et non à chaque pas de temps particulière.

³⁷Remarquons qu'XTOR utilise la base $(\mathbf{e}_\sigma, \mathbf{e}_\theta, \mathbf{e}_\phi)$, ce qui oblige à inverser le signe de la composante toroïdale du champ avant de l'utiliser dans le pousseur.

Pr  servation du sch  ma

Les champs $\mathbf{E}$ et $\mathbf{B}$ et la vitesse $\mathbf{v}$ doivent   tre exprim  s dans la m  me base $(\mathbf{e}_R, \mathbf{e}_\phi, \mathbf{e}_Z)$ lorsqu'on applique l'algorithme de Boris pour actualiser $\mathbf{v}$. Cette base d  pend de la position toroidale ϕ de la particule. Or la vitesse est calcul  e sur les pas de temps demi-entiers ;    cause du caract  re centr   du leap-frog, la position, elle, est connue sur les pas de temps entiers. Les champs ne sont donc pas exprim  s dans la m  me base que la vitesse. Ignorer cela reviendrait    exprimer les champs au temps $t^{n+1/2}$ dans la base locale de la particule au temps t^n . La pr  cision du sch  ma serait d  grad  e de fa  on importante. Le sch  ma deviendrait d'ordre 1 en pr  cision, ce qui est inacceptable. Il serait possible de corriger le probl  me en utilisant une m  thode pr  dicteur-correcteur, comme dans [71], mais il faudrait alors interpoler les champs deux fois plus souvent.

Le probl  me est r  solu en divisant l'avanc  e des particules (l'  quation 40) en deux parties. Cela permet de conna  tre la position des particules au pas de temps demi-entiers. Les champs et la vitesse sont alors naturellement exprim  s dans la m  me base locale. Pour conserver un sch  ma centr  , on effectue un « demi-push » (*i.e.*, on int  gre la vitesse pour faire avancer la particule sur un demi pas de temps) avant l'actualisation de la vitesse, et l'autre apr  s. Un pas de temps se d  compose alors en 5   tapes : demi-push, changement de coordonn  es de $\mathbf{v}^{n-1/2}$, r  cup  ration des champs, Boris, demi-push, changement de coordonn  es de $\mathbf{v}^{n+1/2}$. Le sch  ma suivant r  sume la situation³⁸ :



Un sch  ma similaire a   t   propos   par Cobb et Leboeuf dans [72] en 1994 pour int  grer les orbites des particules dans un r  acteur    plasma cylindrique, mais il utilisait un syst  me de coordonn  es cart  siennes sous-jacent ;    chaque pas de temps, tous les vecteurs   taient enti  rement convertis en cart  sien, puis re-convertis en

³⁸ $\mathbf{x}_{(j)}^i$ d  note un vecteur calcul   au temps t^i exprim   dans la base locale de la particule au temps t^j

cylindrique après Boris et les demi-push. En fait, il n'est pas nécessaire de passer en cartésien pour appliquer Boris, la base cylindrique locale convient car elle est orthonormale. De même, on peut écrire les demi-push localement. Cela permet d'économiser du temps de calcul.

Le schéma ci-dessus met en valeur le centrage de l'algorithme. On remarque cependant qu'il peut être simplifié en rassemblant le demi-push qui achève le pas de temps n et celui qui commence le pas de temps $n+1$ en un seul push, soit à la fin du pas de temps n , soit au début du suivant. On économise alors un demi-push et un changement de coordonnées par pas de temps. Le schéma obtenu est strictement équivalent, et on ne perd pas la stabilité. Cependant, il faut être bien conscient que vitesse et position restent exprimés à des temps différents. Cela aura notamment son importance quand il s'agira d'extraire la contribution d'une particule au tenseur de pression : il faudra s'assurer d'exprimer vitesse et position au même temps³⁹.

Le push, effectué en début de pas de temps, s'écrit (en introduisant les variables intermédiaires x, y , et z) :

$$\begin{aligned} x &= R^{n-1} + v_R^{n-1/2} \delta t \\ y &= v_\phi^{n-1/2} \delta t \\ z &= Z^{n-1} + v_Z^{n-1/2} \delta t \\ \\ R^n &= \sqrt{x^2 + y^2} \\ \phi^n &= \phi^{n-1} + \text{asin}(y/R^n) \\ Z^n &= z \end{aligned}$$

La vitesse $\mathbf{v}^{n-1/2}$ a été calculée au pas de temps précédent ; elle est exprimée dans la base locale de la particule au temps t^{n-1} .

Changement de coordonnées

Le changement de coordonnées de la vitesse effectué avant Boris, qui exprime ses composantes dans la base locale de la particule au temps t^n , s'écrit⁴⁰ :

$$\begin{aligned} \tilde{v}_R^{n-1/2} &= \frac{xv_R^{n-1/2} + yv_\phi^{n-1/2}}{R^{n+1/2}} \\ \tilde{v}_\phi^{n-1/2} &= \frac{xv_\phi^{n-1/2} - yv_R^{n-1/2}}{R^{n+1/2}} \end{aligned}$$

³⁹L'algorithme de Boris sépare naturellement l'intégration de la vitesse en deux ; il suffit d'extraire la vitesse au milieu de l'algorithme après normalisation.

⁴⁰le tilde marque les nouvelles composantes

5.1.3 Stabilité et précision

Du point de vue de la stabilité, la partie parallèle du schéma, en présence d'un champ électrique, se ramène à un schéma leap-frog classique avec une force indépendante de la vitesse de la particule. Dans la direction perpendiculaire, le schéma est inconditionnellement stable en absence de champ électrique.

Du point de vue de la précision, le pas de temps est limité par une condition du type :

$$\omega_c dt \lesssim \frac{2\pi}{N}$$

Si $N = 4$, cela exprime simplement le fait que quatre points sont nécessaires pour modéliser une rotation cyclotronique. La grandeur N sera déterminée à partir de tests sur le poussoir dans la suite de ce chapitre.

On préfère souvent un algorithme de type Boris à une méthode de Runge-Kutta, même si cette dernière est souvent utilisée à l'ordre 4, avec une meilleure précision. D'une part, cela est dû à la simplicité du schéma de Boris. D'autre part, dans ce schéma, l'erreur sur les invariants du mouvement ne s'accumule pas au cours du temps. Physiquement, cela signifie que les orbites modélisées n'auront pas de dérive numérique (conservation de P_ϕ), et qu'elles ne passeront pas de piégées à circulantes (ou inversement) numériquement (conservation de μ).

La bonne conservation des invariants du mouvement est le critère qui nous permettra d'évaluer le comportement du poussoir. En effet, celui-ci étant stable (en l'absence de champ $\mathbf{E}$), et les grandeurs caractéristiques des orbites étant petites (rayon de Larmor...), il est difficile de juger autrement de l'exactitude d'une orbite. La conservation doit être vérifiée sur des temps longs, puisque nos particules devront tourner pendant les temps suffisant pour qu'une instabilité résistive se développe, soit quelques $10^5 \tau_a$ avec les valeurs de S standard de XTOR. Avec des valeurs de plasma typiques de JET, cela donne un temps de simulation de quelques $4 \times 10^6 \omega_c^{-1}$, soit au minimum $2.4 \times 10^7 \delta t$ pour des alphas. Il faut donc que les invariants soient stables sur au moins 10^7 pas de temps (avec un pas de temps de l'ordre de 1/8 de la période cyclotronique).

5.2 Implémentation pour XTOR

5.2.1 Unités et normalisation

On note les grandeurs adimensionnées (normalisées) avec une tilde. Après ce paragraphe, nous travaillerons toujours en unités normalisées, en omettant les tildes.

Le poussoir étant interfacé avec XTOR, on veut rester le plus proche possible du système d'unités XTOR. L'unité de longueur sera donc le petit rayon a .

$$\tilde{l} = \frac{l}{a}$$

La normalisation naturelle de l'équation du mouvement des particules est :

$$\tilde{m} \frac{d\tilde{\mathbf{v}}}{d\tilde{t}} = \tilde{q}(\tilde{\mathbf{E}} + \tilde{\mathbf{v}} \times \tilde{\mathbf{B}}) \quad (43)$$

avec :

$$\begin{aligned} \tilde{m} &= \frac{m}{m_p} && \text{où } m_p, \text{ la masse du proton, est l'unité de masse} \\ \tilde{q} &= \frac{q}{e} && \text{où } e, \text{ la charge du proton, est l'unité de charge} \\ \tilde{\mathbf{B}} &= \frac{\mathbf{B}}{B_u} && \text{où } B_u, \text{ une grandeur caractéristique, est l'unité de champ magnétique} \\ \tilde{t} &= \frac{t}{\tau_c} && \text{où } \tau_c = \omega_c^{-1} = \frac{m_p}{eB_u} \text{ est l'unité de temps} \\ \tilde{E} &= \frac{E}{E_u} && \text{où } E_u = \frac{aB_u}{\tau_c} = \frac{eaB_u^2}{m_p} \text{ est l'unité de champ électrique} \\ \tilde{\mathbf{v}} &= \mathbf{v} \frac{\tau_c}{a} && \text{où } \frac{a}{\tau_c} \text{ est l'unité de vitesse} \end{aligned}$$

Unité de B

Dans les faits, le champ B fourni en entrée de l'équation est en unités XTOR, c'est-à-dire tel que $B_{XT}(A) = A$. On a vu plus haut que cela correspondait à une unité $B_u = B_0/A$, où B_0 , en Teslas, n'est pas précisé dans la partie MHD du code. En remplaçant donc B_u par B_0/A , on peut injecter B_{XT} directement dans l'équation. La définition de τ_c devient :

$$\tau_c = \frac{m_p A}{e B_0}$$

Autrement dit, on joue sur la définition de l'échelle de temps pour prendre en compte l'amplitude de champ magnétique voulue.

Unité de E

De même, le champ E est fourni par XTOR; son unité est donc $v_A B_0/A$, où $v_A = a/\tau_a$. On rappelle que $v_A = B_0/A\sqrt{\mu_0\rho_0}$. Il faut donc renormaliser $\mathbf{E}$ d'un facteur $\alpha = \tau_c/\tau_a$ avant de l'introduire dans l'équation. En effet :

$$\tilde{\mathbf{E}} = \frac{\mathbf{E}_{\text{SI}}}{E_u} = \frac{\mathbf{E}_{\text{XT}}}{E_u} \frac{v_A B_0}{A} = \mathbf{E}_{\text{XT}} \frac{\tau_c}{\tau_a}$$

5.2.2 Paramètres, entrées et sorties.

Le tableau ci-dessous regroupe les paramètres d'entrée physiques du pousseur.

Nom	Symbole	Type	Description
rialpha	α	R	Rapport entre τ_c et τ_a .
b0exp	B_0	R	Facteur de scaling du champ magnétique (T).
mi	m_i	I	Masse normalisée des particules.
dt	δt	R	Pas de temps normalisé
qi	q_i	I	Charge normalisée des particules.
r0exp	R_0	R	Grand rayon normalisé. Doit être égal au paramètre r0 dans XTOR.
Ti	T_i	R	Température des particules (K).

À partir de ces paramètres, on calcule les grandeurs utilisées directement par le pousseur : τ_c , noté tau, notre unité de temps, puis v_{th} , noté vthi, la vitesse thermique normalisée.

$$\tau_c = \frac{m_p A}{e B_0}$$

$$v_{th} = \sqrt{\frac{k_b T_i}{m_i m_p}} \tau_c$$

Dans la pratique, le pas de temps est déterminé différemment selon qu'on utilise le pousseur seul (pour faire des tests avec des particules-test), ou couplé avec XTOR. Dans l'option pousseur seul, le paramètre d'entrée est n_{cy} , le rapport entre la période cyclotronique et le pas de temps : $\delta t = 2\pi m_i / (q_i A n_{cy})$. Dans le cas du code hybride, il est plus pratique de partir du rapport entre pas de temps fluide et pas de temps particulaire. De plus, on utilise des paramètres de normalisation pour pouvoir passer d'un pas de temps centre-guide, plus important, au pas de temps complet.

Paramètres libres dans XTOR

Dans la partie MHD résistive de XTOR, aucune hypothèse n'est faite sur la valeur de B_0 (en Teslas) ou de α (sans dimension).

Fixer α revient en fait à fixer la densité sur l'axe ρ_0 . En effet, le temps d'Alfvén est défini par :

$$\tau_a = \frac{a}{v_A} = (\mu_0 \rho_0)^{1/2} \frac{R_0}{B_0}$$

et par conséquent :

$$\alpha = \frac{m_p}{e}(\mu_0 \rho_0)^{-1/2} = 9.31 \times 10^{-6} \rho_0^{-1/2}$$

avec ρ_0 en kg.m^{-3} .

Remarquons cependant que fixer la valeur de α est nécessaire dans XTOR-2F pour introduire les effets de dérive diamagnétiques.

L'introduction d'une population de particules rapides modélisées par un système d'équations indépendant oblige logiquement à fixer toutes les paramètres libres. On doit donc choisir des valeurs pour B_0 et pour ρ_0 (ou α). Cela revient à fixer les valeurs de τ_a et de τ_c .

Dans la pratique, le code a besoin de la valeur de B_0 pour calculer τ_c . La valeur de τ_c est nécessaire pour normaliser la vitesse fluide, qui est calculée en SI à partir de la température. D'autre part, une valeur de α est nécessaire puisque α intervient directement dans l'équation du mouvement comme facteur devant le champ E .

Les choix de B_0 et α sont guidés par des valeurs expérimentales. Typiquement, pour se placer dans une configuration proche de celle de JET, on prendra :

$$B_0 = 3\text{T} \quad ; \quad \alpha = 0.029$$

Ou de façon équivalente :

$$\tau_a \approx 3.33 \times 10^{-7}\text{s} \quad ; \quad \tau_c \approx 9.66 \times 10^{-9}\text{s}$$

Cela correspond à une densité $n_0 = 3 \times 10^{19} \text{ m}^{-3}$ (c'est-à-dire $\rho_0 = 1.0 \times 10^{-7} \text{ kg.m}^{-3}$ pour un plasma de deutérium).

Il est important de noter que α peut être reajusté quand on l'utilise devant les termes diamagnétiques, dans les équations MHD. Cela est dû au fait qu'XTOR ne tourne en général pas avec une valeur de S typique des grandes machines, pour des raisons de temps de calcul. Ainsi, si on reprend le cas ci-dessus d'une configuration de type JET, on a un temps résistif $\tau_r \sim 100\text{s}$, et $\tau_a \sim 3.33 \times 10^{-7}\text{s}$, ce qui donne $S \sim 3.3 \times 10^8$. Une telle valeur de S entraînerait des temps de calculs trop longs ; par conséquent, on choisit de tourner avec une valeur de S plus faible. Pour bien modéliser la stabilisation du kink par les effets diamagnétiques, il faut par exemple augmenter α .

Autres paramètres

Le pousseur prend de nombreux autres paramètres en entrée. Les principaux sont :

Nom	Symbole	Type	Description
nbpart	N_{part}	I	Nombre de particules.
nrbox	nx	I	Nombre de points de maillages en R . Doit être égal au INRBOX de CHEASE.
nzbox	nx	I	Nombre de points de maillages en Z . Doit être égal au INZBOX de CHEASE.

5.2.3 Mapping, changement de base, et interpolations des champs

Au début de l'étape Boris, il faut passer d'un champ⁴¹ $\mathbf{B} = (B_s, B_\phi, B_\theta)$ défini sur le maillage (s, ϕ, θ) fourni par XTOR à un champ $\mathbf{B} = (B_R, B_\phi, B_Z)$ défini sur le maillage (R, ϕ, Z) , utilisé par le pousseur (en fait, le pousseur se sert directement de la transformée de Fourier en ϕ du champ). Le changement de maillage (mapping), le changement de coordonnées, et la transformée de Fourier ne sont effectués qu'une fois par pas de temps fluide. On prédit donc qu'ils auront peu d'influence sur le temps de calcul global. En revanche, à chaque pas de temps pour chaque particule (*i.e.*, à chaque appel de Boris), le champ doit être interpolé à la position de la particule. Plus l'ordre de l'interpolation est élevé, plus la précision est bonne, mais plus on exécute d'opérations par appel de Boris.

Mapping

Pour chaque point (R_i, ϕ_n, Z_j) du maillage du pousseur, on écrit⁴²

$$B(R_i, \phi_n, Z_j) = f_1(s_\star, \theta_\star)B(s_l, \phi_n, \theta_m) + f_2(s_\star, \theta_\star)B(s_{l+1}, \phi_n, \theta_m) \\ f_3(s_\star, \theta_\star)B(s_l, \phi_n, \theta_{m+1}) + f_4(s_\star, \theta_\star)B(s_{l+1}, \phi_n, \theta_{m+1})$$

Avec les notations suivantes :

- (R_i, ϕ_n, Z_j) sont les coordonnées du point du maillage où on veut évaluer le champ, donc $1 \leq i \leq nr$, $1 \leq j \leq nz$, $0 \leq n \leq nmax + 1$. nr et nz sont choisis dans le pousseur, mais $nmax$ est imposé par XTOR.
- (s_l, ϕ_n, θ_m) sont les coordonnées des sommets de la maille du maillage XTOR qui contient le point (R_i, ϕ_n, Z_j) . Le schéma suivant illustre ces définitions :
- Les f_n sont les fonctions de bases utilisées pour l'interpolation (typiquement, elles sont bilinéaires). Elles sont évaluées en $(s_\star, \theta_\star)$, qui sont les coordonnées XTOR du point (R_i, Z_j) . Cela implique de connaître les valeurs de s et θ sur les points du maillage poloïdal cartésien. Ces valeurs sont calculées en sortie du code CHEASE, et écrites dans le fichier **ASTRO**.

⁴¹Dans toute la suite, on ne parle que du champ B . Le champ E subit le même traitement, sauf lorsque les composantes co- et contravariantes sont distinguées.

⁴²Pour la simplicité de l'écriture, on considère une composante arbitraire de B .

Changement de coordonnées

On introduit les champs calculés dans les expression suivantes : pour des composantes covariantes (B^r, B^θ, B^ϕ) , on écrira

$$\begin{aligned} B_R &= \sqrt{g_{rr}} \cos \theta B^r + \sqrt{g_{\theta\theta}} \cos(\theta + \chi) DB^\theta \\ B_Z &= \sqrt{g_{rr}} \sin \theta B^r + \sqrt{g_{\theta\theta}} \sin(\theta + \chi) DB^\theta \\ B_\phi &= \sqrt{g_{\phi\phi}} B^\phi \end{aligned}$$

Pour des composantes contravariantes (E_r, E_θ, E_ϕ) , on écrira :

$$\begin{aligned} E_R &= \sqrt{g^{rr}} \sin \theta E_r + \sqrt{g^{\theta\theta}} \sin(\theta + \chi) E_\theta \\ E_Z &= \sqrt{g^{rr}} \cos \theta E_r + \sqrt{g^{\theta\theta}} \cos(\theta + \chi) E_\theta \\ E_\phi &= \sqrt{g^{\phi\phi}} E_\phi \end{aligned}$$

Avec les définitions :

- les g_{ij} sont les composantes du tenseur métrique.
- θ est l'angle poloidal.
- α est défini par : $\cos \chi = g_{r\theta} / \sqrt{g_{rr} g_{\theta\theta}}$
- D est le déterminant du jacobien.⁴³

Tous les champs sont évalués au point de maillage du pousseur. De même que pour s et θ , les valeurs des champs métriques g et $\cos / \sin(\theta + \chi)$ sur le maillage cartésien sont calculés par CHEASE et écrits dans le fichier ASTRO.

Transformée de Fourier

Comme évoqué plus haut, on stocke les champs dans l'espace de Fourier pour la direction ϕ . La transformation de Fourier est sans doute l'un des problèmes numériques les plus étudiés et documentés, et les algorithmes de FFT sont très nombreux. Nous utilisons celui proposé par Intel (bibliothèque MKL) par défaut, et éventuellement la bibliothèque FFTW sur d'autres machines.

Interpolation des champs aux positions des particules

A chaque appel de Boris, on interpole les champs à la position (R_p, ϕ_p, Z_p) de la particule. L'interpolation s'écrit :

⁴³XTOR utilise $DB^{\theta,\phi}$ comme variables pour éviter une singularité sur l'axe magnétique, où $B^{\theta,\phi} \sim 1/r$.

$$\begin{aligned}
B(R_i, \phi_p, Z_j) &= \sum_{k=0}^{nmax/2} b(R_i, Z_j, k, 1) \cos(k\phi_p) + b(R_i, Z_j, k, 2) \sin(k\phi_p) \\
B(R_{i+1}, \phi_p, Z_j) &= \dots \\
B(R_i, \phi_p, Z_{j+1}) &= \dots \\
B(R_{i+1}, \phi_p, Z_{j+1}) &= \dots \\
\\
B(R_p, \phi_p, Z_p) &= f_1(R_p, Z_p)B(R_i, \phi_p, Z_j) + f_2(R_p, Z_p)B(R_{i+1}, \phi_p, Z_j) + \\
&\quad f_3(R_p, Z_p)B(R_i, \phi_p, Z_{j+1}) + f_4(R_p, Z_p)B(R_{i+1}, \phi_p, Z_{j+1})
\end{aligned}$$

avec les notations :

- (R_i, Z_j) est le coin inférieur de la maille poloïdale contenant la particule,

$$R_i \leq R_p < R_{i+1}, \text{ et } Z_i \leq Z_p < Z_{i+1},$$

- les $f_i(R, Z)$ sont les fonctions de base qui servent à l'interpolation (pour une interpolation bilinéaire, on a $f_1 = (R_{i+1} - R_p)(Z_{i+1} - Z_p)$, etc).
- les $b(R_i, Z_j, k, 1/2)$ représentent les coefficients paires/impaires de Fourier du champ, sur les points du maillage poloïdal. Notons que dans la pratique, pour accélérer les accès mémoire, on stocke en fait ces champs sous la forme $b(:, 1/2, k, R_i, Z_j)$, la première colonne représentant les trois composantes des vecteurs.

5.2.4 Détail des opérations et distribution du temps de calcul

On a récapitulé sur la page suivante les étapes d'un pas de temps particulière⁴⁴.

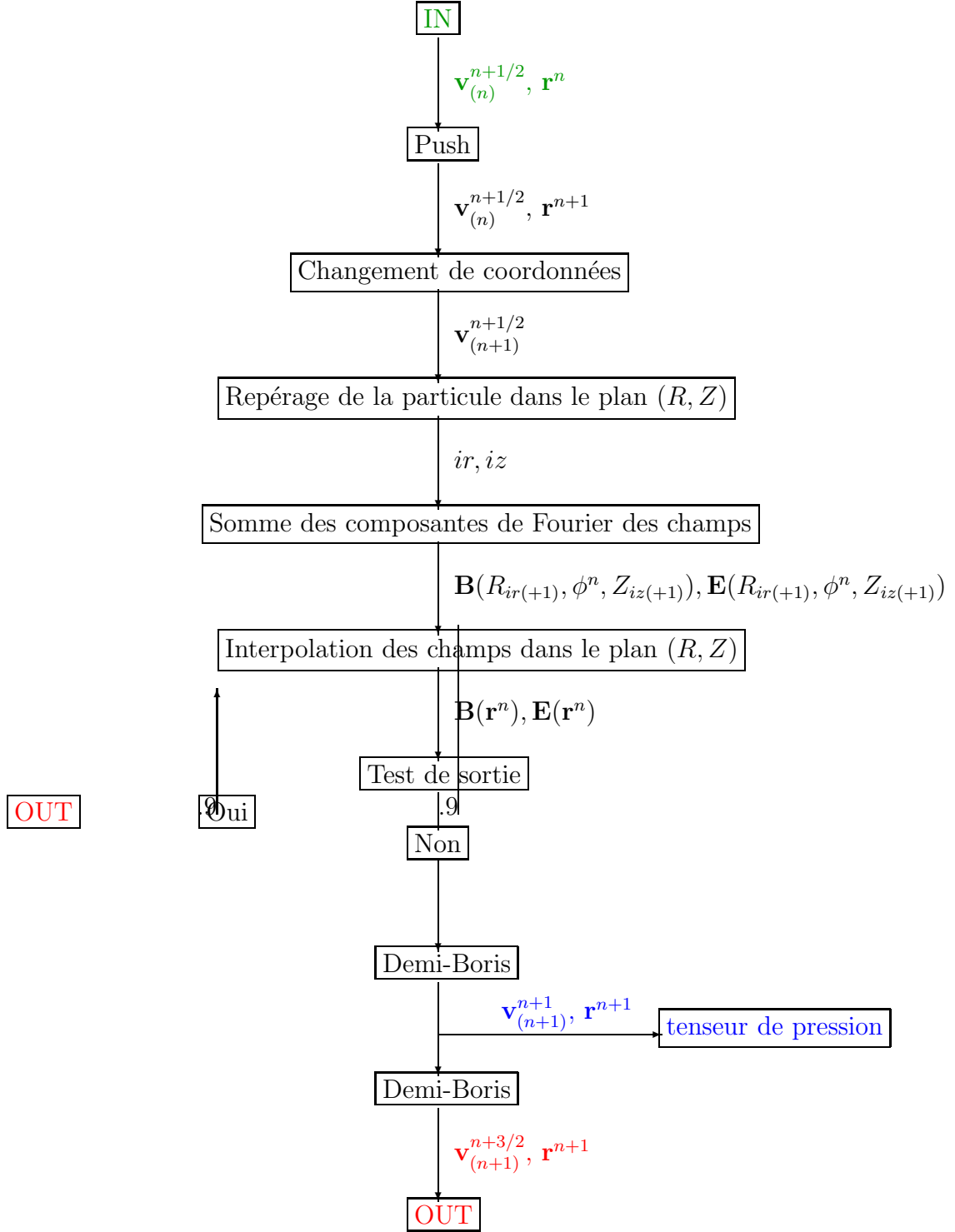
L'étape d'avancée des particules est l'étape discriminante en terme de temps d'exécution : la routine sera appelée $N_{part}\Delta t/\delta t$ fois à chaque pas de temps fluide. On récapitule ici en détail les opérations effectués par le pousseur pendant un appel.

Push et changement de coordonnées

On se réfère au détails de ces étapes donnés plus haut, et on voit que ces étapes ajoutent :

- 6 additions/soustractions
- 9 multiplications

⁴⁴ $\mathbf{x}_{(j)}^i$ dénote un vecteur calculé au temps t^i exprimé dans la base locale de la particule au temps t^j



- 3 divisions
- 1 racine
- 1 arcsinus

Les opérations les plus coûteuses sont bien sûr la racine, et surtout l'arcsinus.

Interpolation des champs

L'interpolation dans le plan poloidal ne contient que des opérations arithmétiques, et n'a pas une grande influence (dans le cas d'une interpolation bilinéaire, on aura une dizaine (11) d'addition et soustractions, et une douzaine (12, 8+4) de multiplications et divisions).

En revanche, la somme de coefficients de Fourier introduit des fonctions trigonométriques. En tout, en notant $n_{op} = n_{max}/2 + 1$, on rajoute :

- n_{op} sinus
- n_{op} cosinus
- $2 n_{op}$ additions et soustractions
- $4 n_{op}$ multiplications⁴⁵.

Et cela pour chaque composante de chaque champ, soit 6 fois en tout.

Boris

L'implémentation du schéma de Boris est décomposée suivant les étapes⁴⁶ :

1. $\Omega = q/m \cdot B$
2. $\text{normega} = \sqrt{\Omega(1)^2 + \Omega(2)^2 + \Omega(2)^2}$
3. $E = 0.5 \cdot dt \cdot q/m \cdot \text{rialpha} \cdot E$
4. $V = V + E$
5. $P = 0.5 \cdot dt \cdot \Omega$
6. $\text{normp2} = P(1)^2 + P(2)^2 + P(3)^2$
7. $Q = 2/(1 + \text{normp2}) \cdot P$
8. $W = V + V \times P$
9. $V = V + W \times Q$
10. $V = V + E$

⁴⁵À cause de la convention utilisée dans la librairie MKL, il faut multiplier tous les sinus et cosinus par 2.

⁴⁶Les vecteurs, tous de dimension 3, commencent par des majuscules

En termes d'opérations élémentaires, l'algorithme effectue :

- 17 additions et 6 soustractions
- 34 multiplications
- 2 divisions
- 1 évaluation de racine

Les compilateurs acceptent différentes options d'optimisation (inlining, loop unrolling, etc), qui restructurent le code et peuvent réduire le nombre d'opérations nécessaires. Cependant, l'algorithme se présente ici sous une forme si simple qu'il paraît a priori impossible d'économiser des opérations. Pour vérifier, on a compilé le code à l'aide du compilateur Intel ifort, avec l'option `-fcode-asm`, qui imprime la version du code compilé en assembleur. Cela permet de vérifier le nombre réel d'opérations effectuées lors de l'exécution :

- 17 additions et 6 soustractions
- 36 multiplications
- 2 divisions
- 1 évaluation de racine

L'ajout d'options d'optimisation (`-fast`, correspondant au niveau d'optimisation maximal pour les optimisations classiques, et `-axT`, correspond à des optimisations machine-dépendantes) ne réduit pas le temps de calcul, ce qui confirme notre prévision.

La seule façon d'améliorer la rapidité de Boris est de minimiser le temps d'accès mémoire. Pour cela, il faut que les coordonnées d'une même particules occupent des espaces mémoires adjacents : la forme du tableau des particules est fixée en conséquence.

Distribution du temps de calcul.

On a estimé ci-dessous le temps passé par le pousseur dans chaque partie de l'algorithme, pour une valeur standard de $n_{max} = 12$, 100 mailles dans les directions R et Z , avec un nombre de particules de 10^6 , et 100 pas de temps (ce qui correspond à l'ordre de grandeur du nombre de pas de temps entre deux pas de temps fluide.)

Étape	% temps de calcul
avancée	7.0
interpolation en ϕ	73.9
interpolation en RZ	8.5
boris	9.4
Temps total (s)	78,46

Comme prévu, la partie la plus longue est celle de l'interpolation en ϕ , à cause des calculs des sinus et cosinus. Les compilateurs proposent différents degrés de précision pour calculer les fonctions transcendentes, ce qui a une incidence sur la vitesse de calcul. Ainsi, avec l'option `-no-fast-transcendentals`, le temps de l'interpolation en ϕ est doublé, alors que les autres ne varient pas significativement.

5.3 Trajectoires et conservation des invariants

5.3.1 Trajectoires des charges dans un tokamak

Les orbites des particules chargées dans une configuration magnétique de type tokamak ont été étudiée depuis longtemps. Les particules effectuent des rotations cyclotroniques autour des lignes de champ qui s'enroulent comme des hélices sur des surfaces toriques. La courbure et le gradient du champ magnétique engendrent des dérives verticales, vers le haut ou le bas suivant le signe de la charge :

$$\mathbf{v}_G = \frac{v_\perp^2}{2\omega_c} \frac{\mathbf{b} \times \nabla B}{B} \quad ; \quad \mathbf{v}_C = \frac{v_\parallel^2}{\omega_c} \frac{\mathbf{b} \times \mathbf{n}}{R}$$

Où $\mathbf{b} = \mathbf{B}/B$ et $\mathbf{n}$ est le vecteur unitaire normal aux lignes de champ, orienté vers l'intérieur de la courbure. Ces dérives font s'éloigner les particules de leur surface de champ initiale. Cependant, cet éloignement est exactement compensé de part et d'autre du demi-plan horizontal $z = 0$, et le mouvement des particules est stable. L'écart d'une particule par rapport à sa surface de champ (et, par conséquent, son excursion radiale) reste négligeable devant sa distance moyenne à l'axe magnétique. Les particules qui ont l'excursion radiale la plus grande sont les particules piégées (le champ magnétique variant comme $1/R$, les particules peuvent rencontrer un miroir magnétique en arrivant côté fort champ) ; leur orbite vue dans le plan poloïdal rappelle la forme d'une banane, et la largeur de cette banane vaut :

$$\delta_b = q\rho_L \epsilon^{-1/2} \quad (44)$$

Les figures 18-22 présentent des orbites caractéristiques des particules thermiques (1keV) pour les tokamaks, calculées par le pousseur.

Remarquons que l'hypothèse $\delta_b \leq r$ peut s'écrire $\epsilon \geq (2q_0\rho_L/R)^{2/3}$, ce qui devient faux près de l'axe magnétique. L'orbite banane est alors déformée pour former une orbite « patate », de largeur caractéristique $\delta_p \equiv (2q_0\rho_L/R)^{2/3}/R$. Autrement dit, dans une région de rayon δ_p autour de l'axe magnétique, les orbites des particules ont des excursions radiales plus larges que dans le reste du plasma, de l'ordre de δ_p . Or, δ_p étant proportionnelle au rayon de Larmor, le lieu des orbites patates est plus

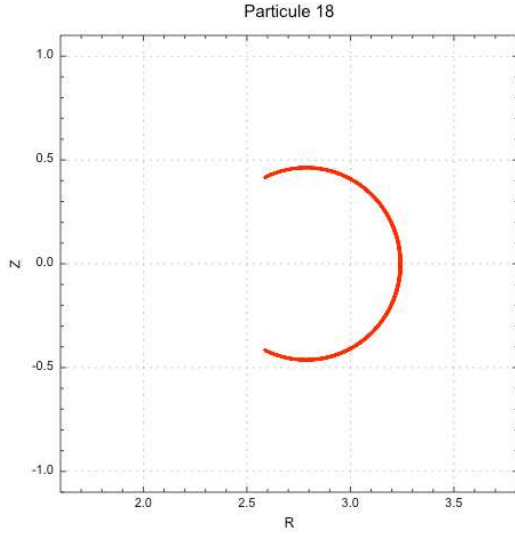


FIG. 18 – Trajectoire piégée , coupe poloïdale

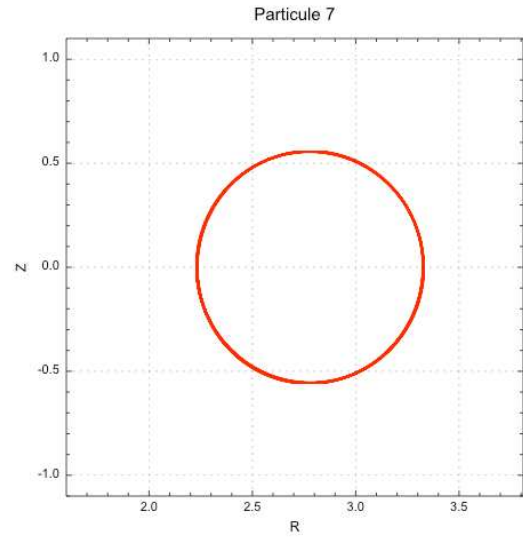


FIG. 19 – Orbite passante, coupe poloïdale

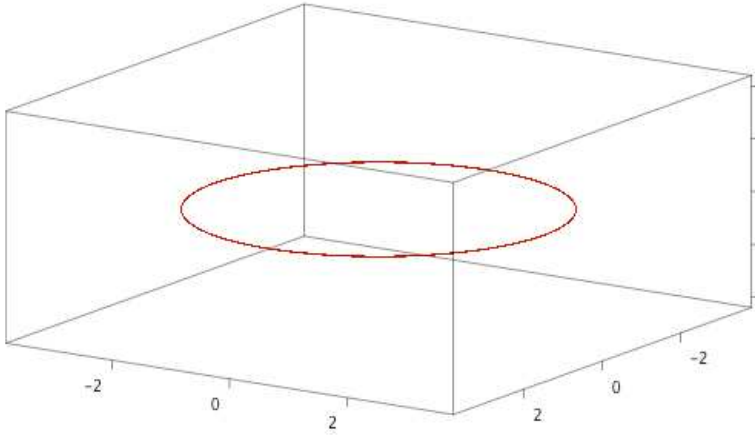


FIG. 20 – Orbite de stagnation

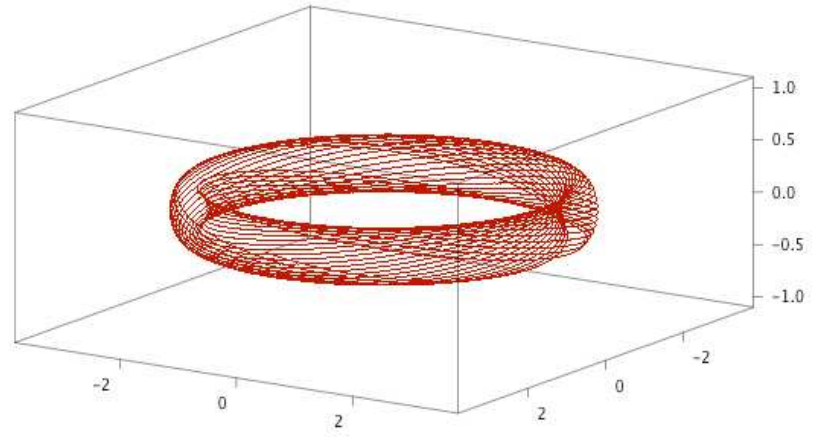


FIG. 21 – Orbite passante

important pour des particules chaudes. Cet effet est amplifié par le fait que les particules chaudes sont précisément créées au centre du plasma, que ce soit par chauffage par onde ou par fusion. Ainsi, la largeur patate est estimée à 20 - 30 % du petit rayon dans une machine à fusion comme JET. Dans cette région autour de l'axe, l'excursion

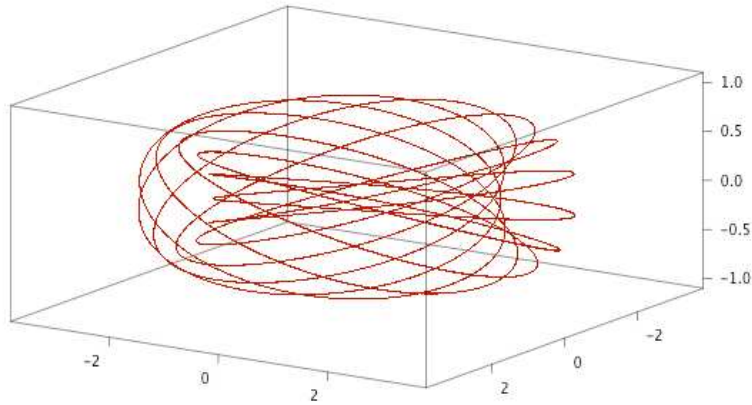


FIG. 22 – Orbite passante, surface rationnelle

radiale importante des particules chaudes homogénéise leur dépôt d'énergie, et peut avoir des conséquences importantes sur les dents de scie, les fishbones, ou des modes d'Alfvén localisés au centre du plasma. Une résolution complète des orbites des particules chaudes ne semble donc pas être un luxe, contrairement au cas des particules thermiques. On peut se référer par exemple à [76] pour une typologie complète des orbites des ions rapides ; on montre sur les figures 23-27 quelques exemples d'orbites précessionnelles de ce genre, obtenus avec notre pousseur.

5.3.2 Conservation des invariants du mouvement

Le mouvement d'une particule chargée dans le champ magnétique d'un tokamak conserve trois invariants bien connus (analytiquement, on travaille souvent dans des bases de coordonnées liées à ces invariants) ; cf. par exemple [73]. Ils sont l'énergie cinétique, le moment magnétique intégré sur une rotation cyclotronique, et le moment cinétique toroïdal. Leur conservation est le principal critère de qualification du pousseur ; c'est elle qui définit le domaine de fonctionnement numérique en fonction du pas de temps et du maillage spatial.

Énergie cinétique

Le schéma numérique traite séparément les parties parallèles et perpendiculaires de l'accélération. Dans la direction parallèle, sans champ électrique, il n'y a pas d'accélération et la composante de la vitesse reste inchangée par le pousseur. Dans la direction perpendiculaire, la norme de la vitesse est conservée par construction.

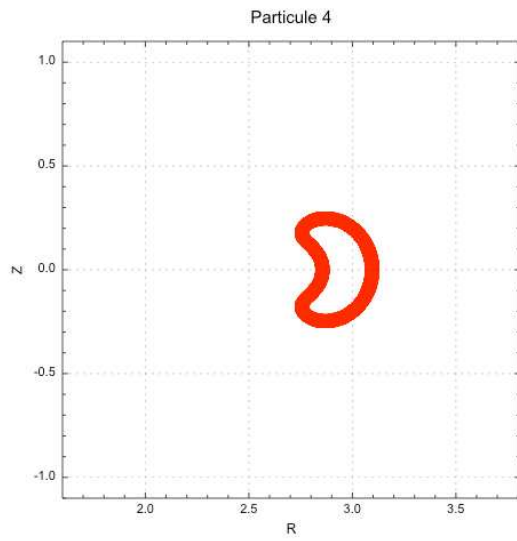


FIG. 23 – Orbite patate

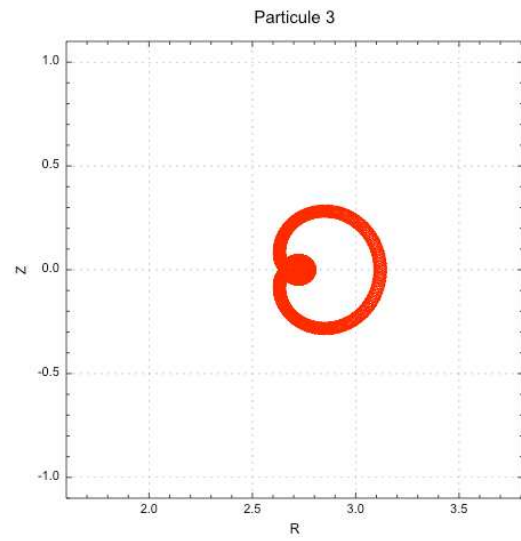


FIG. 24 – Orbite marginalement piégée

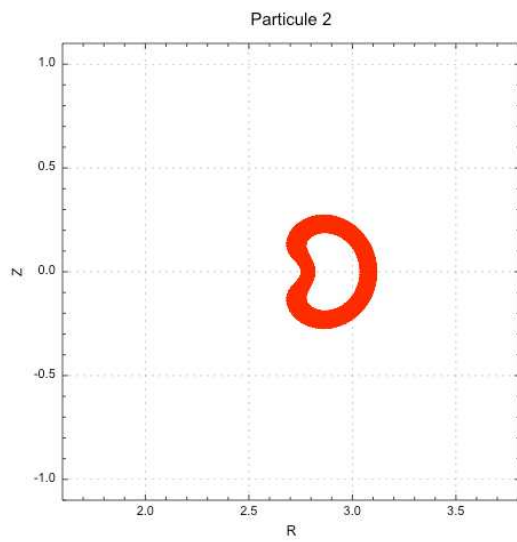


FIG. 25 – Orbite patate comprenant l'axe

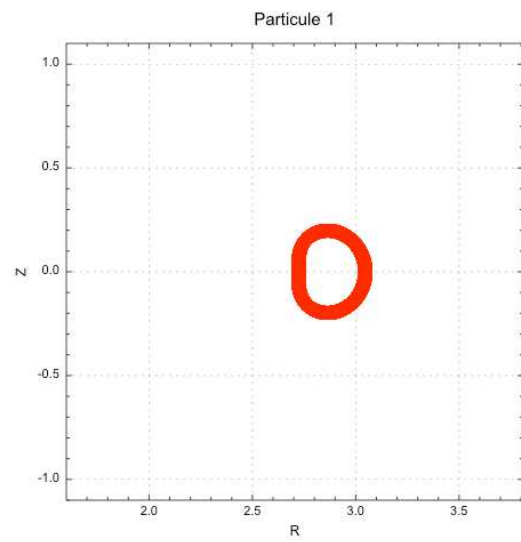


FIG. 26 – Orbite tangente à l'axe

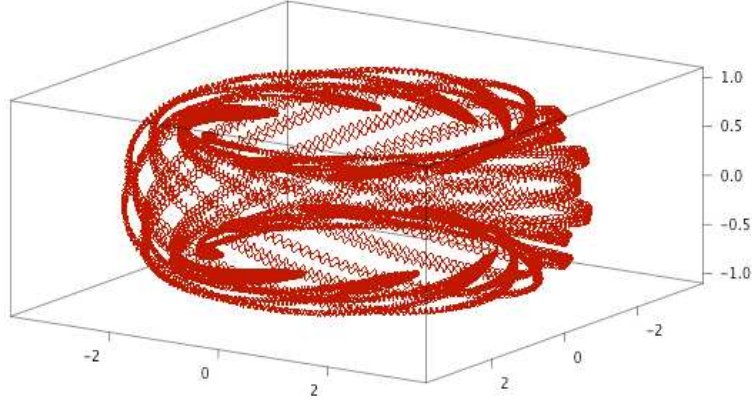


FIG. 27 – Précession toroïdale d'une orbite piégée

Par conséquent, la force de Lorentz numérique ne travaille pas, et l'énergie cinétique est exactement conservée au cours du mouvement.

Moment magnétique

Le moment magnétique μ est un invariant adiabatique, étudié notamment par Northrop dans [74], et Kruskal, dans [75]. Son invariance est liée aux hypothèses :

$$\begin{aligned} \rho_L \frac{\nabla B}{B} &\ll 1 \\ \omega_c \frac{dB}{dt} &\ll 1 \end{aligned}$$

C'est-à-dire qu'une particule voit un champ magnétique constant et homogène pendant une gyration cyclotronique. Le terme $\mu_0 \equiv mv_{\perp}^2/(2B)$, que l'on mesurera pour nos diagnostics, est le terme d'ordre 0 de l'invariant adiabatique, qui s'exprime comme un développement limité en $\epsilon = \rho_L \nabla B/B$. On précise ici les corrections du premier ordre (cf. [73]) :

$$\frac{1}{m}\mu = \mu_0 + \frac{\mu_0}{B}(\rho \cdot \nabla B) + \frac{v_{\parallel}^2}{B}(\mathbf{b} \cdot \nabla)(\mathbf{b} \cdot \rho) - \frac{\mu_0 v_{\parallel}}{2\omega_c} 5[\mathbf{b} \cdot \nabla \times \mathbf{b} - (\mathbf{a} \cdot \nabla \mathbf{b}) \cdot \mathbf{c}]$$

Où $\mathbf{b}$ est le vecteur unitaire dans la direction de $\mathbf{B}$, $\mathbf{c} \equiv \mathbf{v}_{\perp}/v_{\perp}$, $\mathbf{a} \equiv \mathbf{b} \times \mathbf{c}$, et ρ est le vecteur du rayon de gyration de la particule. L'ordre de grandeur de ces termes du premier ordre est $\rho_L \mu_0$.

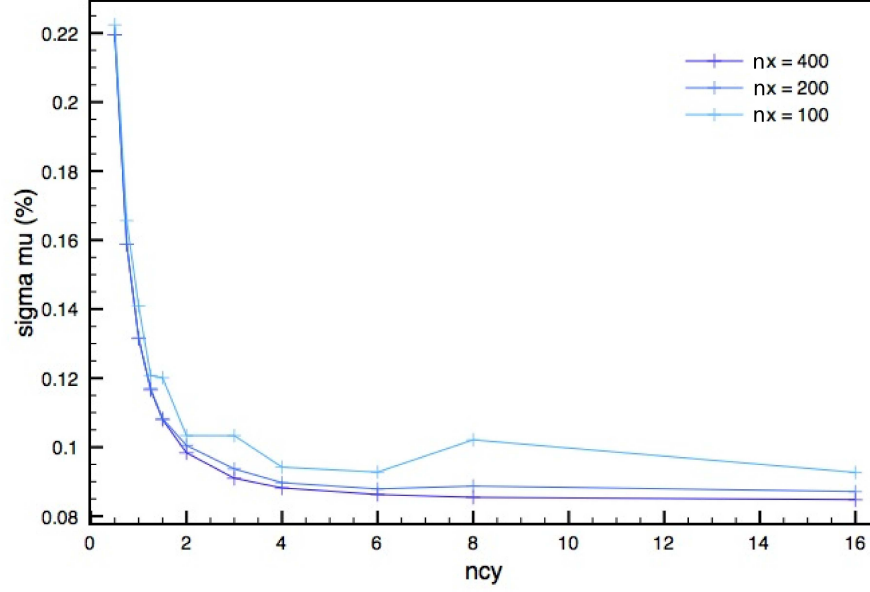


FIG. 28 – Écart-type de μ ($T_i = 1\text{keV}$) en fonction de δt pour différents nx

On présente sur la figure 28 les variations de l'écart-type de μ_0 normalisé à sa moyenne (qu'on note σ_μ pour alléger dans la suite) au cours du temps, pour une particule thermique au centre du plasma avec un pitch $\xi \equiv v_{\parallel}/v_{\perp} = 0.88$, en fonction du pas de temps et de la taille du maillage. Les graphes ci-dessous sont faits pour une particule située sur l'axe magnétique (lieu du maximum de densité des particules rapides). Les variations de σ_μ avec le petit rayon moyen de l'orbite et avec le pitch se déduisent des variations de μ en fonction de B (et donc de R au premier ordre) et de $v_{\perp}$. Comme prévu, le pousseur n'entraîne pas d'accumulation d'erreur sur l'invariant, si bien que l'écart-type est la grandeur pertinente pour le diagnostic. Les variables sont n_{cy} , le nombre de pas de temps par rotation cyclotronique (dans nos unités, $dt = 2\pi m_i/(q_i A n_{cy})$), et nx , le nombre de mailles dans les directions R et Z . La variable nx varie de 100 à 400 ; le cas $nx = 100$ est extrême, puisqu'il correspond à des mailles de 5cm sur 5cm, comparables au rayon de Larmor des particules chaudes. Dans la pratique, il est difficile de monter au-dessus de $nx = 400$, puisqu'il faut affiner le maillage polaire (σ, θ) utilisé par CHEASE et XTOR en conséquence, ce qui nous limite en termes de ressources numériques. La courbe $nx = 300$ étant confondue avec $nx = 200$ et $nx = 400$, on l'a omise. Le nombre de pas de temps correspond à 500 - 1000 tours de tore (*i.e.*, 10^7 pas de temps pour $n_{cy} = 8$).

Ce graphe appelle les interprétations suivantes :

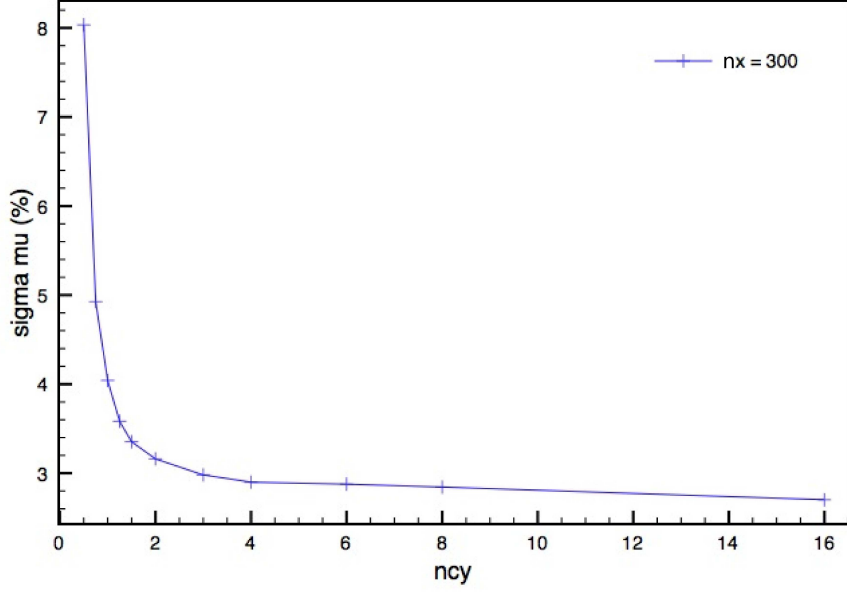


FIG. 29 – Écart-type de μ ($T_i = 1\text{MeV}$)

1. Asymptotiquement, σ_μ tend vers une constante en δt et nx . Cette constante est physique, et non numérique ; elle correspond aux termes du premier ordre de μ . Les ordres de grandeurs sont conformes aux prévisions. Autrement dit, comme on mesure en fait μ_0 et non μ , notre signal contient des composantes physiques, indépendantes de la discrétisation (ce point est vérifié dans le paragraphe suivant).
2. Du point de vue du domaine de fonctionnement, le schéma est proche de la tendance asymptotique et peu sensible à (nx, dt) à partir de $nx \geq 200$ et $n_{cy} \geq 4$. Remarquons que la définition du domaine de validité du pousueur se fait en référence aux valeurs asymptotiques. En valeur absolue, on aurait pu considérer que les cas $n_{cy} \leq 4$ étaient acceptables (σ_μ reste petit devant 1), alors que les orbites sont en fait déformées dans ces cas (le trop faible pas de temps entraîne des erreurs numériques sur la gyration cyclotronique).

On vérifie sur la figure 29 que le comportement est identique pour des particules chaudes, avec une particule à $T_i = 1\text{ MeV}$. La courbe est multipliée par $T_i^{1/2} \sim 30$, et le σ_μ asymptotique reste compatible avec les corrections du premier ordre⁴⁷.

⁴⁷Les écarts dus à nx étant numériques et non physiques, ils deviennent négligeables, et les courbes sont confondues, c'est pourquoi on n'en a montré qu'une.

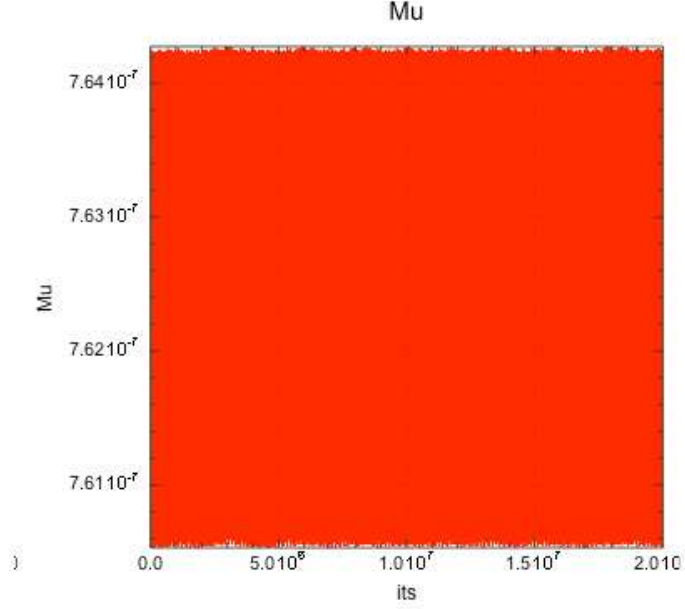


FIG. 30 – μ pour une particule circulante à 1 keV au centre du plasma ; 16 pas de temps par rotation, $nx = 400$.

La figure 30 illustre le fait que l'invariant ne dérive pas sur des temps longs ($10^7 \tau_A \sim$ quelques milliers de tours de tore).

On confirme que le signal asymptotique de μ représente bien des termes physiques en observant les comportements des différentes fréquences qui composent le signal lorsque $nx \rightarrow 0$ et $n_{cy} \rightarrow \infty$. Ces composantes ($T_i = 1$ keV , $n_{cy} = 16$, $nx = 400$, interpolation bicubique et plus bilinéaire des champs pour plus de précision), contiennent 3 fréquences principales. On les montre sur la figure 31.

1. La fréquence de transit poloïdal (battement basse fréquence sur le graphe du haut) qui correspondent aux termes en $v_{||}$ dans le développement de μ .
2. La fréquence de changement de maille (crête du signal surlignée en rouge sur le graphe du milieu), qui, elle, est numérique, et tend à disparaître.
3. La fréquence cyclotronique (graphe du bas), qui correspond aux termes en $v_{\perp}$.

Moment cinétique toroïdal

Enfin, la configuration magnétique est invariante en ϕ . Cela entraîne que le moment conjugué de ϕ dans un formalisme hamiltonien est conservé au cours du mouvement. Si on note $\mathcal{L}$ le Lagrangien associé à une particule, ce moment cinétique toroïdal, noté P_{ϕ} , est défini par :

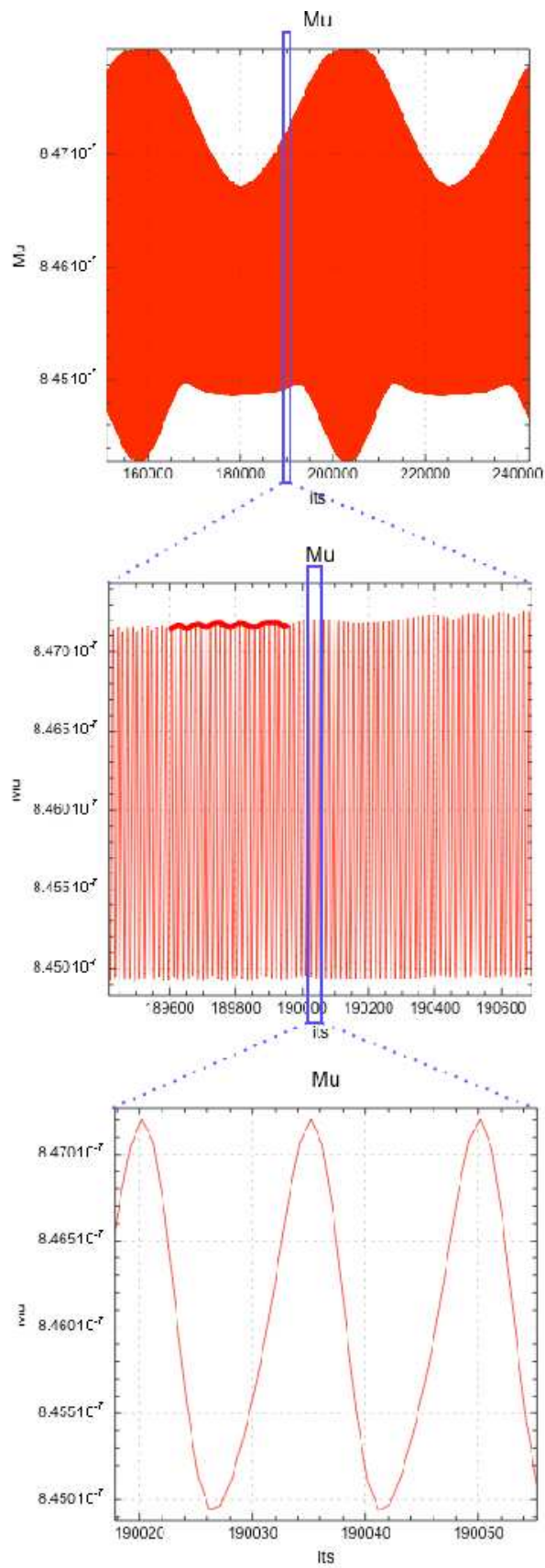


FIG. 31 – Fréquences principales du signal σ_μ . L'axe des abscisses est gradué en itérations du pousseur.

$$P_\phi = \frac{d\mathcal{L}}{d\dot{\phi}}$$

Avec $\mathcal{L} = mv^2/2 - q\Phi + q\vec{A} \cdot \vec{v}$, où Φ est le potentiel électrique et $\vec{A}$ est le potentiel vecteur, on obtient que :

$$P_\phi = mRv_\phi - q\Psi \quad (45)$$

Avec Ψ le flux du champ magnétique poloïdal⁴⁸. P_ϕ est donc un invariant géométrique ; il ne s'exprime pas comme un développement limité comme μ , mais simplement comme la somme de deux termes. Les variations de ces deux termes s'annulent exactement, sans autre hypothèse que la géométrie du système. Les variations de $q\Psi$ expriment le fait que la particule s'éloigne de sa surface de champ originelle. Comme on l'a dit plus haut, les particules chaudes ont une excursion radiale importante ; la conservation de P_ϕ est donc non triviale et c'est un bon indicateur de la qualité du poussoir.

Le premier résultat, illustré sur la figure 32, est l'absence de dérive numérique de P_ϕ , sur des temps de l'ordre de quelques milliers de tours de tore, ce qui signifie que les particules ne dérivent pas numériquement dans la direction r . On montre la variation de l'écart-type de P_ϕ , normalisé à la moyenne de $q\Psi$ ⁴⁹, noté σ_{P_ϕ} , pour une particule thermique, sur la figure 33, avec des paramètres similaires à ceux utilisés pour σ_μ .

Deux remarques s'imposent :

1. Si on atteint une asymptote en δt , ce n'est clairement pas le cas en nx . Avec nos paramètres, les variations du signal ne sont pas physiques, comme pour μ , mais bien numériques. Notons d'ailleurs que l'utilisation d'une interpolation bicubique à la place de l'interpolation bilinéaire usuelle pour les champs permet de gagner un ordre de grandeur sur σ_{P_ϕ} .
2. Par conséquent, et contrairement au cas de μ , le domaine de fonctionnement en nx ne se détermine pas par comparaison avec une valeur asymptotique, puisque celle-ci est nulle si le poussoir est bien fait. Il faut se contenter d'obtenir des valeurs très petites devant 1, ce qui est notre cas.

On explique ce comportement en rappelant que l'invariance de P_ϕ vient du fait que $\mathbf{B}$ dérive d'un potentiel vecteur, ce qui est bien sûr lié à sa représentation discrète.

⁴⁸Remarquons que dans la réalité, un tokamak n'est pas invariant en ϕ à cause du nombre fini de bobines de champ toroïdal : c'est le ripple.

⁴⁹Sachant que Ψ reste de l'ordre de $\Psi_{edge} - \Psi_{centre}$.

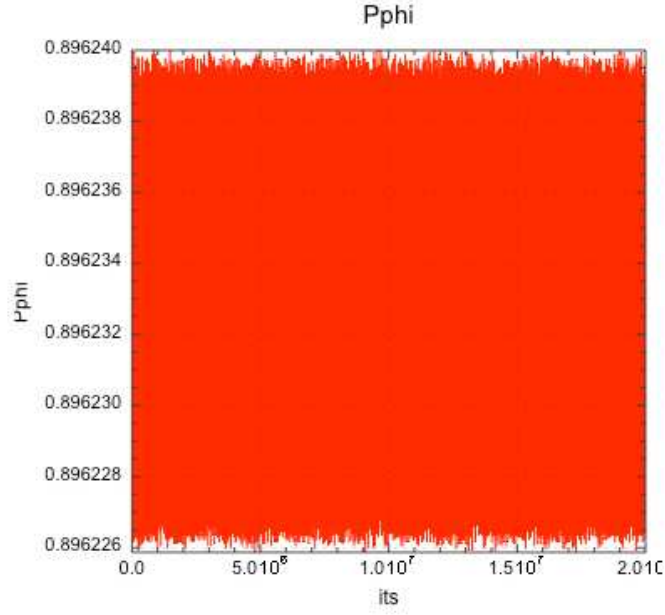


FIG. 32 – P_ϕ pour une particule circulante à 1 keV au centre du plasma ; 16 pas de temps par rotation, $nx = 400$.

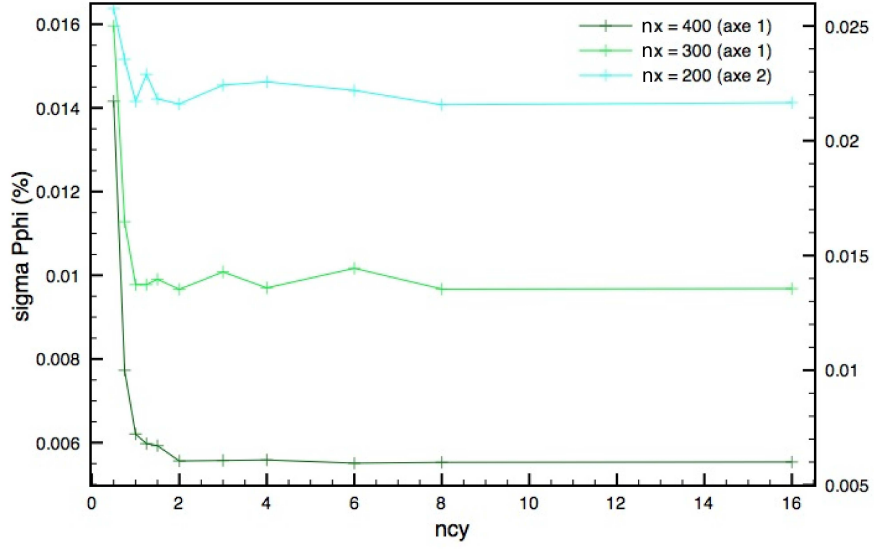


FIG. 33 – Écart-type de P_ϕ en %, $T_i = 1\text{keV}$.

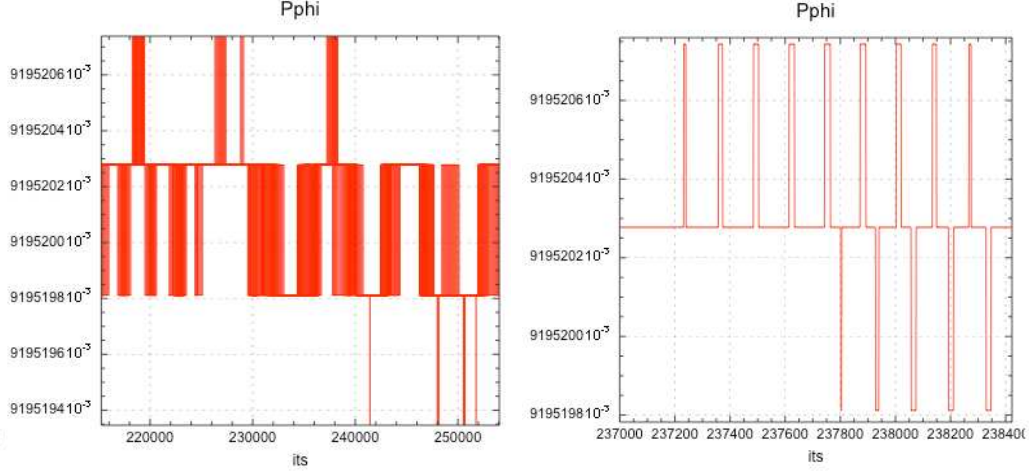


FIG. 34 – Structure temporelle du signal σ_{P_ϕ} , avec $n_{cy} = 128$ et interpolation bicubique.

Autrement dit, il n'y a pas de moyen d'assurer que le champ discret est bien à divergence nulle en tout point de l'espace, car notre code utilise $\mathbf{B}$ et non $\mathbf{A}$ comme variable. Pour aller plus loin, et vérifier cette observation, on a raffiné les interpolations des champs. Ainsi, en utilisant des interpolations bicubiques pour $\mathbf{B}$ et Ψ , pour $n_{cy} = 128$ et $T_i = 1$ keV, on a pu réduire σ_{P_ϕ} à 0 pour la majorité des particules, ce qui confirme le bon fonctionnement du pousseur. Une minorité présente une structure temporelle discrète du signal σ_{P_ϕ} , directement liée à la technique d'interpolation des champs (voir figure 34).

Le comportement de P_ϕ étant numérique et non physique, il diffère pour des particules rapides. La figure 35 illustre cette différence.

Dans ce cas, on n'atteint plus d'asymptote en δt^{50} . Cela peut s'interpréter en considérant qu'augmenter v en gardant nx constant revient numériquement à augmenter δt . Par rapport au cas 1 keV, on s'est donc décalé vers la gauche des graphes, nous éloignant de l'asymptote. Une étude rapide a montré qu'il faudrait monter à $n_{cy} \gtrsim 500$ pour atteindre l'asymptote, ce qui est clairement inenvisageable en pratique, pour des raisons de temps de calcul. Les valeurs restant petites devant 1, on régime de fonctionnement évoqué plus haut : $nx \gtrsim 200$, $n_{cy} \gtrsim 4$ n'a pas lieu d'être modifié.

⁵⁰C'est pour cette raison qu'on n'a pas montré l'influence de nx .

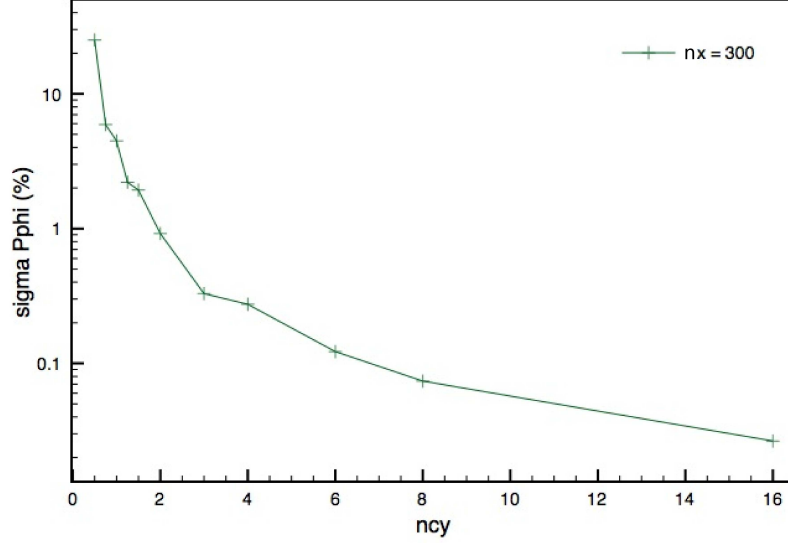


FIG. 35 – Écart-type de P_ϕ en %, $T_i = 1\text{MeV}$.

5.4 Conclusion

Dans ce chapitre, nous avons détaillé le fonctionnement du pousseur qui assure l'avancée des particules rapides.

Nous avons adapté une méthode numérique classique d'intégration de la trajectoire d'une particule dans un champ électromagnétique pour la géométrie des tokamaks. Le pousseur obtenu présente des bonnes qualités de stabilité et de précision, autorisant un pas de temps particulière de l'ordre du quart de la période cyclotronique. Il exige cependant des transformations sur les champs (changement de maillage, interpolations linéaires dans le plan poloïdal, et de Fourier dans la direction toroïdale). L'étape la plus importante en temps de calcul est l'interpolation de Fourier, à cause des calculs de sinus et cosinus. Après avoir détaillé les normalisations et paramètres utilisés par le pousseur, nous avons caractérisé son fonctionnement en étudiant les invariants du mouvement. L'énergie est exactement conservée. Le moment magnétique d'ordre 0 contient les fréquences attendues (tour de tore, changement de maille, rotation cyclotronique) qui représentent les termes de premier ordre du moment magnétique total. Le moment cinétique toroïdal est bien conservé sur des temps très longs, mais on a clairement montré l'importance de la représentation discrète de $\mathbf{B}$, qui devra être prise en compte ultérieurement. En conclusion, pour avoir un comportement correct (variance au plus de l'ordre du %) sur des temps longs, il suffit d'effectuer au moins 4 pas de temps par gyration cyclotronique, et d'avoir des mailles d'une taille d'environ 10 rayons de Larmor à 1 keV (cas $nx = 200$).

6 Stratégie de réduction du bruit sur le tenseur de pression

La discrétisation de l'espace des phases introduit un bruit numérique sur le tenseur de pression. À cause de l'absence d'hypothèse δf , ce bruit a une grande amplitude. Il est impossible de le supprimer totalement, mais il est nécessaire de le minimiser pour optimiser la convergence de l'algorithme Newton-Krylov/Picard, sans dans la mesure du possible altérer la physique en jeu.

Localement, ce sont les fonctions de bases caractérisant nos interpolations qui créent le bruit. Augmenter l'ordre de ces fonctions permet donc de le réduire. Cependant, le dépôt du tenseur de pression est l'étape la plus longue du code particulaire : pour des raisons de temps de calcul, on se limitera à des interpolations bilinéaires, ou au mieux bicubiques. Nous devons donc nous tourner vers d'autres solutions. Notre objectif n'est pas de réduire le bruit de plusieurs ordres de grandeurs, mais plutôt de le diviser par un facteur compris entre 2 et 10.

Nous commençons par rappeler la définition du tenseur de pression numérique, et par détailler l'étape de dépôt, pendant laquelle chaque particule ajoute sa contribution au tenseur. Nous expliquons ensuite une stratégie envisagée de réduction du bruit, qui consiste à rendre la distribution initiale des particules la plus homogène possible dans l'espace des phases. Les différentes dimensions de l'espace des phases ne sont pas équivalentes. Dans l'espace des vitesses, une analogie entre le calcul de la pression comme moment de la fonction de distribution et le théorème de Monte-Carlo pour l'intégration numérique permet d'espérer que la pression converge en N_{part}^{-1} (au lieu de $N_{part}^{-1/2}$ pour une distribution aléatoire) en chaque point. Dans l'espace des positions, la coordonnée toroïdale est traitée à part, du fait de la symétrie du système. Dans le reste de l'espace des phases, nous utiliserons une « séquence à faible discrétisation » (LDS) lors de l'initialisation (la séquence de Sobol) pour homogénéiser la distribution de particules. Nous présentons les résultats après initialisation, qui montrent la diminution de bruit attendue (convergence en N_{part}^{-1}). La régularité obtenue à l'initialisation sera perdue au cours du mouvement ; on s'attend en particulier à ce que la perte d'information dans la direction θ soit la plus critique.

Nous décrivons ensuite la mise en place d'un filtrage temporel du tenseur de pression pour en supprimer les fréquences indésirables. Nous étudions la structure temporelle de la pression, en mettant en évidence les fréquences clés du système, et nous décrivons le filtre numérique utilisé (filtre de Butterworth, passe-bas). On supprime ainsi les fréquences de l'ordre de ω_c , que l'algorithme voit comme un bruit blanc ; de plus, un filtrage de fréquences plus élevées pourrait permettre de compenser au moins partiellement la perte de régularité due au mélange des phases en θ . Nous présentons pour finir l'exemple d'un kink interne simulé sur quelques $10^3 \tau_A$ avec XTOR-K, où le filtrage cyclotronique permet d'obtenir un rapport signal sur bruit tolérable.

6.1 Définition et calcul du tenseur de pression

La rétroaction des particules sur le fluide MHD se fait par l'introduction d'un tenseur de pression particulaire dans le solveur MHD. Ce tenseur, symétrique et de dimension 9, s'écrit sous forme analytique :

$$\Pi_a(\mathbf{r}, t) = m_i \int_{-\infty}^{\infty} \mathbf{v} \mathbf{v} f(\mathbf{v}, \mathbf{r}, t) d^3\mathbf{v} \quad (46)$$

Où f est la fonction de distribution des particules, qui ont toutes la même masse m_i .

6.1.1 Expression du tenseur de pression numérique

Numériquement, le tenseur est donné par la somme discrète :

$$\Pi(\mathbf{r}, t) = m_i \sum_{n=1}^{N_{part}} \mathbf{v}_n \mathbf{v}_n \delta(\mathbf{r} - \mathbf{r}_n(t))$$

La somme porte sur les particules, repérées par l'indice n .

Les fonctions de Dirac sont remplacées par des fonctions de poids F qui répartissent la contribution de chaque particule sur les points de maillage avoisinants. On ajoute également des poids P_n qui peuvent être ajustés selon la fonction de distribution qu'on veut modéliser et du nombre de particules disponibles⁵¹

$$\Pi(\mathbf{r}, t) = m_i \sum_{n=1}^{N_{part}} \mathbf{v}_n \mathbf{v}_n F(\mathbf{r} - \mathbf{r}_n(t)) P_n \quad (47)$$

6.1.2 Dépôt du tenseur numérique

Ce tenseur est évalué sur les points de maillage du pousseur, c'est-à-dire sur un maillage cartésien en R et Z . On ne l'évalue pas pour des valeurs de ϕ données, mais on calcule directement ses coefficients de Fourier (i.e, il est stocké dans l'espace de Fourier en ϕ).

Les coefficients de Fourier de la composante * de Π sur le maillage (R, Z) s'écrivent :

⁵¹Comme dans tout code PIC, chaque particule numérique est en fait une super-particule qui représente de nombreuses particules physiques.

$$\begin{aligned}\Pi_{i,j,k,1}^* &= m_i \sum_{n=1}^N (\mathbf{v}_n \mathbf{v}_n)^* P_n F_{i,j}(R_n, Z_n) \cos(k\phi_n) \\ \Pi_{i,j,k,2}^* &= m_i \sum_{n=1}^N (\mathbf{v}_n \mathbf{v}_n)^* P_n F_{i,j}(R_n, Z_n) \sin(k\phi_n)\end{aligned}$$

Les indices i et j correspondent au maillage dans les directions R et Z ; l'indice k au coefficient de Fourier, et l'indice final au type de coefficient de Fourier (pair ou impair)⁵².

Une fois les contributions des particules sommées, on prend la divergence du tenseur, puis on l'évalue sur les points du maillage **XTOR**. Pour ce faire, on fait des interpolations bilinéaires dans le plan poloïdal (pour passer du maillage (R, Z) au maillage (σ, θ)), et on fait la somme de Fourier pour se placer à des valeurs de ϕ correspondant au maillage **XTOR**. Cette opération est l'inverse de la transformation qu'on a faite pour évaluer les champs E et B sur le maillage cartésien au tout début du pousseur.

La somme de Fourier (ou transformée de Fourier inverse) s'écrit :

$$\Pi_{i,j,k}^* = \sum_{l=0}^{lmax} (\Pi_{i,j,k,1}^* \cos(2\pi l \phi_k) + \Pi_{i,j,k,2}^* \sin(2\pi l \phi_k))$$

Notons que l'opération de prendre la divergence du tenseur sur le maillage cartésien s'est avérée très dommageable en termes de bruit, c'est pourquoi il existe également la possibilité de déposer directement le tenseur sur le maillage (σ, θ) .

6.2 Régularisation de la distribution initiale dans l'espace des phases

Le bruit est dû à la discrétisation spatiale du système : il est causé par le fait qu'on considère un nombre fini de particules se déplaçant dans un champ défini sur un maillage. Ce point de vue met en lumière deux axes de réduction du bruit : on peut soit travailler à affiner le maillage, en augmentant le nombre de points ou l'ordre des interpolations, soit travailler à optimiser le chargement des particules sur ce maillage. On parle ici du chargement numérique, préalable à la répartition des particules dans l'espace des phases selon la fonction de distribution choisie.

⁵²L'ordre des indices est important. En effet, le fait de placer l'indice k le plus loin possible minimise les opérations nécessaires lors des accès mémoires, ce qui se traduit par un gain de temps substantiel lors des transformations et des sommes de Fourier.

Dans cette seconde optique, on a intérêt à avoir un chargement aussi homogène que possible dans l'espace des phases. Les différentes dimensions de cet espace ne sont cependant pas équivalentes. Le tenseur de pression est une intégrale portant sur le sous-espace des vitesses, qui dépend la position (la périodicité en ϕ confère un rôle particulier à la direction toroïdale). Pour obtenir un chargement optimisé, on choisit d'utiliser une « Low Discrepancy Sequence » (LDS) : la séquence de Sobol. Cette séquence est ensuite transformée pour modéliser la fonction de distribution choisie ; nous avons fait le choix de considérer des particules de poids égaux. L'idée générale de l'utilisation d'une LDS pour l'optimisation d'une intégrale numérique est donnée dans [77], chapitre 7.

6.2.1 Anisotropie de l'initialisation dans l'espace des phases

Notre espace des phases est de la forme $[R_m, R_M] \times [0, 2\pi] \times [Z_m, Z_M] \times \mathbb{R}^3$. Les variables correspondantes sont $(R, \phi, Z, v_R, v_\phi, v_Z)$. Le choix des variables d'espace est dicté par les coordonnées utilisées par le pousseur⁵³.

Intégrale sur la vitesse et théorème de Monte-Carlo

Le tenseur de pression numérique défini par l'équation 47 est une somme discrète, analogue à celle qui intervient dans le théorème de Monte-Carlo pour l'intégration numérique. Ce théorème est rappelé :

Soit E un espace euclidien de dimension n , de mesure de Lebesgue V , g une fonction définie sur E et intégrable, et $\{\mathbf{x}_1, \dots, \mathbf{x}_N\} \in E^N$ une suite de points de E ; le théorème s'écrit alors :

$$\frac{1}{V} \int_E g(\mathbf{x}) d^n \mathbf{x} \simeq \frac{1}{N} \sum_{i=1}^N g(\mathbf{x}_i) + \sqrt{\frac{\langle g^2 \rangle - \langle g \rangle^2}{N}} \quad (48)$$

Autrement dit, on approche l'intégrale de g par la moyenne de g sur un ensemble de points $\{\mathbf{x}_1, \dots, \mathbf{x}_N\} \in E^N$. Le dernier terme donne une estimation de l'erreur. Ce n'est en aucun cas une borne supérieure. Il fournit cependant une information importante : la variation de l'erreur en fonction de N . La lenteur de la convergence en fonction du nombre de points ($\propto N^{-1/2}$) est une limitation classique de l'intégration numérique. Ce théorème est utilisé fréquemment pour calculer la surface ou le volume de sous-espaces de forme complexe, en choisissant comme fonction la fonction caractéristique du sous-espace. On voit qu'on peut également l'appliquer au calcul de moments, comme dans notre cas, en posant $g(\mathbf{v}) = \mathbf{v}^n f(\mathbf{r}, \mathbf{v}, t)$.

⁵³Remarquons que l'espace poloïdal étant cartésien (*i.e.*, rectangulaire), il est plus grand que la plus grande surface poloïdale du plasma.

Ce théorème représente l'une des utilisations les plus fréquentes d'une méthode de Monte-Carlo. Le terme de «méthode de Monte-Carlo», introduit dans les années 1940, désigne en fait tout algorithme reposant sur un processus aléatoire pour produire un résultat numérique⁵⁴. Dans notre cas, ce processus est le choix de $\{\mathbf{x}_1, \dots, \mathbf{x}_N\}$, qui représentent les coordonnées des particules dans l'espace des vitesses.

La popularité des méthodes de Monte-Carlo, et en particulier du théorème ci-dessus, a conduit à de nombreux efforts pour minimiser le terme d'erreur : on parle de réduction de variance. Il est notamment possible d'obtenir une décroissance plus forte que $N^{-1/2}$. Ainsi, en dimension 1, si les points forment un maillage uniforme de l'intervalle d'intégration ($\{x_i = i/N\}_{i=1, \dots, N}$ pour $E = [0, 1]$), le terme d'erreur varie en N^{-1} , puisqu'on a alors une simple méthode des rectangles. Il y a donc lieu de penser qu'on peut obtenir une convergence rapide de l'erreur, éventuellement jusqu'en N^{-1} , en répartissant les points idéalement dans l'espace d'intégration. Par analogie avec le cas 1D, « idéalement » signifierait « de la façon la plus uniforme possible » (avec le moins possible de « trous »). Il faut donc construire une séquence de points adaptée : on parlera de séquence quasi-random ou encore sub-random. Ces termes sont trompeurs puisqu'en réalité, la séquence est entièrement déterministe : elle est construite avec soin pour approcher au plus près (et tendre asymptotiquement vers) une séquence uniforme. C'est la méthode Quasi Monte-Carlo.

Contrairement à ce qu'on pourrait penser, une séquence homogène à n dimensions n'est pas idéale. Elle perd de son efficacité lorsqu'on la projette sur un sous-domaine de dimension moins élevée. En effet, plusieurs points vont alors se trouver alignés (c'est-à-dire que la fonction g aura la même valeur sur ces points). Alors qu'on aura utilisé N points pour mailler le domaine à n dimensions, ces points n'apporteront que $N' = N^{1/n}$ valeurs différentes par dimension. Pour optimiser la séquence, il faudrait que chacune des N particules utilisées apporte une information différente dans chacune des n dimensions : autrement dit, on veut éviter les alignements réguliers de points (« clustering ») dans toutes les directions de l'espace considéré. On voudrait donc construire une suite de vecteurs à n éléments tels que leur projection sur un sous-espace de dimension quelconque forme toujours un maillage « idéal ». De telles suites sont appelées « Low Discrepancy Sequences » (LDS), ou, au moyen d'un barbarisme inoffensif, séquences à faible discrétance.

Espace des positions

La pression est une intégrale sur la vitesse dépendant de la position. Il convient d'insister sur la différence entre les rôles de ces variables dans le calcul de la pression. Une meilleure répartition dans l'espace des vitesses va faire converger la pression plus rapidement en N_{part} en chaque position de l'espace. Une meilleure répartition

⁵⁴Les méthodes de Monte-Carlo sont particulièrement adaptées lorsqu'on étudie un système qui comporte un grand nombre de degrés de libertés ; elles sont notamment utilisées en finance, en physique statistique, en physique des fluides, etc.

en position va homogénéiser la pression sur le maillage spatial. Cela est d'autant plus important que l'on sera amené à prendre la divergence de la pression, qui fait intervenir des différences entre des points de maillage voisins. Il est clairement avantageux d'avoir la répartition la plus homogène possible dans tout l'espace des phases, mais il est important de se souvenir de cette distinction entre vitesse et position.

On va donc construire une LDS de dimension 5 dont les éléments seront des vecteurs de $\mathcal{I} \equiv [R_m, R_M] \times [Z_m, Z_M] \times \mathbb{R}^3$, chacun représentant les coordonnées d'une particule dans l'espace rassemblant les vitesses et le plan poloidal. On va ensuite transformer cette LDS pour modéliser la fonction de distribution choisie. La coordonnée ϕ est particulière, puisque la configuration tokamak est invariante en ϕ . Ainsi, les champs d'équilibre fournis par CHEASE et utilisés par XTOR comme valeurs d'origine sont invariants en ϕ . Il est donc logique d'initialiser les particules périodiquement en ϕ . On les répartit sur $nmax$ plans verticaux, définis par $\phi_k = 2k\pi/(nmax)$, avec $k = 0, \dots, nmax - 1$. La distribution des particules en (R, Z, v_R, v_ϕ, v_Z) est la même pour tous ces plans. Autrement dit, la population des particules est constituée de $nmax$ sous-populations qui ont la même fonction de distribution restreinte à $\mathcal{I}$, décalées entre elles dans la direction toroïdale d'un angle $\delta\phi = 2\pi/nmax$ à l'instant initial.

Considérons des champs invariants en ϕ . Au cours du mouvement, la distribution des particules reste constituée de $nmax$ « nuages » de particules, qui se déduisent les uns des autres par une rotation d'angle $2\pi/nmax$. La distribution des particules étant alors strictement périodique, il n'y a aucune perte d'information lors du passage dans l'espace de Fourier des moments de cette distribution (à condition bien sûr de considérer un nombre de modes suffisamment grand, ou inversement un nombre de plans suffisamment faible). Par conséquent, la discrétisation du système selon la direction ϕ n'introduit aucun bruit, ce qui est très avantageux pour l'étude de la phase linéaire des instabilités.

En réalité, au cours d'une simulation hybride, divers modes MHD se développeront ; tout mode ayant un nombre d'onde toroïdal $n \geq 1$ brisera l'invariance en ϕ . La distribution ne sera plus périodique et le tenseur de pression sera bruité. L'initialisation périodique en ϕ a donc pour effet d'éliminer initialement le bruit dû aux composantes $n \neq 0$ des modes MHD.

6.2.2 Low Discrepancy Sequences et séquence de Sobol

Mathématiquement, on veut caractériser et construire des séquences de points formant un maillage le plus uniforme possible dans un espace euclidien de dimension n . Il s'agit d'un problème de théorie de la mesure, formalisé par K.F. Roth dans le cadre de l'approximation diophantienne⁵⁵ : le problème de la faible discrédance (D). On conseille au lecteur de se référer à l'article de Niederreiter sur les LDS, cf. [78].

⁵⁵C'est-à-dire l'approximation des nombres réels par des nombres rationnels.

Définitions

Commençons par définir la discrédance d'une séquence. Nous suivons les notations de Niederreiter. Plaçons-nous dans un hypercube élémentaire de dimension n : $E = [0, 1]^n$. La discrédance D_N mesure l'éloignement maximal de la séquence par rapport à une séquence uniforme :

$$D_N = \sup_{B \in J} \left| \frac{A(B)}{N} - \lambda(B) \right|$$

où J est l'ensemble des hypercubes de la forme :

$$\prod_{i=1}^n [a_i, b_i[= \{\mathbf{x} \in \mathbb{R}^n, a_i \leq x_i < b_i\}$$

N est le nombre total de points de la séquence, B est un élément de J , $A(B)$ est le nombre d'éléments de la séquence qui sont dans B , et λ est la mesure de Lebesgue.

En minimisant D_N , on tend vers la situation où le nombre de points de la séquence dans tout sous-espace est proportionnel au volume de ce sous-espace. On tend donc bien vers un maillage uniforme de l'hypercube.

On définit D_N^* de la même façon que D_N mais en prenant la borne supérieure sur l'ensemble des hypercubes de la forme :

$$\prod_{i=1}^n [0, b_i[$$

On a la relation : $D_N^* \leq D_N \leq 2^n D_N^*$.

Bornes inférieures

Il existe des contraintes sur la qualité d'une séquence (i.e., des bornes inférieures sur D_N). En dimension 1 et 2, W.M. Schmidt a prouvé que la discrédance de toute séquence vérifie :

$$D_N^* \geq C_n \frac{(\log N)^{n-1}}{N}$$

Où C_n ne dépend pas de N . Cela équivaut à $D_N^* \geq C'_n (\log N)^n / N$. Autrement dit, la discrédance d'une séquence varie au mieux en N^{-1} .

Pour $n \geq 3$, cette conjecture n'est pas démontrée, mais Roth a donné des bornes inférieures du même type : $D_N^* \geq K_N \log N^d / N$, avec $d \leq n$.

Low Discrepancy

Puisqu'on connaît des bornes inférieures de D^* , qu'on suppose les meilleures possibles, on peut évaluer la qualité d'une séquence selon l'éloignement de D^* par rapport à ces bornes.

Le problème est qu'il est difficile de calculer D^* pour une séquence donnée. Il faut trouver des algorithmes permettant de construire des séquences et de connaître, ou du moins de savoir encadrer D^* . De nombreux algorithmes ont été inventés (Faure, Halton, Sobol, Niederreiter) qui vérifient :

$$D^*(\{\mathbf{x}_1, \dots, \mathbf{x}_N\}) \leq K_n \frac{(\log N)^n}{N} \quad (49)$$

La constante K_n dépend de n et de l'algorithme, mais pas de N . Vu les bornes inférieures indiquées ci-dessus, on peut considérer que ces séquences sont les meilleures possibles, au sens où D^* est le plus près possible de sa borne inférieure. Ce sont ces séquences qu'on appellera LDS. La seule façon d'améliorer leur qualité est de raffiner les algorithmes de façon à réduire la constante K_N au minimum.

On a donc une définition quantitative d'une LDS, au moyen d'un encadrement de D^* .

Inégalité de Koksma-Hlawka :

Une fois l'existence de LDS constatée, cette inégalité permet de conclure que :

$$\left| \int_{[0,1]^n} g(\mathbf{x}) d^n \mathbf{x} - \frac{1}{N} \sum_{i=1}^N g(\mathbf{x}_i) \right| \leq V(g) D_N^*(\{\mathbf{x}_1, \dots, \mathbf{x}_N\})$$

Cette inégalité est valable si g est à variation bornée, sa variation totale étant $V(g)$. Elle est vraie quelle que soit la séquence $\{\mathbf{x}_1, \dots, \mathbf{x}_N\}$. Le terme de gauche est égal au terme d'erreur de l'équation 48. Ce terme d'erreur est maintenant borné⁵⁶ par un produit de deux facteurs. L'un ne dépend que de la fonction à intégrer ; dans notre cas, la fonction de distribution, dont les variations sont bien bornées. L'autre est la discrétion de la séquence. En appliquant l'inégalité à une LDS, on déduit que le terme d'erreur lors de l'intégration numérique varie bien en N^{-1} .

Il est donc possible, si les particules sont réparties au mieux dans l'espace d'intégration (*i.e.*, l'espace des vitesses), d'avoir une erreur sur le tenseur de pression variant en N^{-1} . Cette variation est la meilleure qu'on puisse espérer ; il reste toutefois possible d'essayer de diminuer le facteur indépendant de N du terme d'erreur (ce qui revient à diminuer la constante K_N associée à la séquence).

⁵⁶contrairement à celui du théorème de Monte-Carlo.

Principe et construction d'une LDS

Pour des raisons pratiques, on construit une LDS de façon à ce que toute sous-séquence $\{x_1, \dots, x_i, i < N\}$ d'une LDS soit également une LDS. Cela implique que pour agrandir la séquence, il suffit de calculer un nouvel élément, et de l'ajouter à la fin de la séquence. Chaque nouvel élément vient remplir le plus grand vide laissé par les éléments précédents. Ce n'est pas le cas, par exemple, pour un maillage uniforme (en 1D sur l'intervalle $[0, 1]$, $\{x_1, \dots, x_N\}, x_k = k/N$), où il faut recalculer tous les éléments pour agrandir la séquence. Cette propriété constitue donc un avantage numérique évident.

L'idée de base est simple : pour construire une LSD, on divise un intervalle en segments de longueurs égales. On place un point par segment dans un certain ordre, puis on raffine le maillage en re-divisant chaque segment, et on itère le procédé.

La solution la plus simple pour éviter les problèmes de clustering est de choisir comme base un nombre premier différent pour chaque dimension (c'est le principe de l'algorithme de Halton). Dans la première dimension, le côté de l'hypercube est ainsi divisé en demis, puis quarts, huitièmes, etc. Dans la seconde, le côté est divisé en tiers, neuvièmes, vingt-septièmes, etc. Les premiers points de la séquence de Halton en 2D seront donc $(1/2, 1/3)$, $(1/4, 2/3)$, $(3/4, 1/9)$, et ainsi de suite.

Cette méthode a l'avantage de la simplicité, mais elle présente l'inconvénient qu'à haute dimension, la longueur des intervalles de base diminue rapidement (puisqu'elle est égale à l'inverse du nombre premier correspondant). Par conséquent, il y a de plus en plus d'intervalles à remplir, et le cycle à effectuer pour mailler un côté de l'hypercube devient de plus en plus long. Des trous apparaissent dans l'hypercube ; les points de la séquence sont corrélés. C'est pourquoi dans la pratique, cette séquence est peu utilisée à partir de 6 dimensions.

Séquence de Sobol $s_{i,j}$

L'algorithme de Sobol a été proposé par I.M. Sobol en 1967 (cf. [79]). Il repose sur une seule base pour toutes les dimensions : la base 2. Chaque côté de l'hypercube est donc divisé en 2^p intervalles de longueur 2^{-p} . L'algorithme choisit les coordonnées de chaque point de façon à éviter les problèmes de projection. Autrement dit, chaque point a pour coordonnées une combinaison de nombres de type $k 2^{-p}$, avec $k < p$. Toute la subtilité de l'algorithme repose dans la façon dont les coordonnées sont ré-ordonnées lors du calcul d'un nouveau point.

L'algorithme est complexe, et nous nous contentons ici de la décrire, faute de pouvoir l'interpréter d'une façon simple. Il est constitué d'une suite d'opérations binaires $\oplus$ (« ou » exclusif bit par bit), ce qui est logique vu que seule la base 2 est utilisée.

On part d'un ensemble de nombres appelés « direction numbers », notés $\{v_{i,j}\}_{i=1,\dots,s,j=1,\dots,n}$. Ces nombres sont calculés à partir de polynômes primitifs sur $\mathbb{Z}/2\mathbb{Z}$. Pour chaque dimension j de la séquence à calculer, on choisit un polynôme primitif différent. Les

coefficients de ce polynôme sont notés $\{a_i\}_{i=1,\dots,d_j-1}$ (les coefficients de degré 0 et d_j sont 1). Ils valent soit 0 soit 1. A partir de ces coefficients, on utilise la relation de récurrence suivante pour construire la suite de nombres $m_{k,j}$:

$$m_{k,j} = 2a_{1,j}m_{k-1,j} \oplus 2^2a_{2,j}m_{k-2,j} \oplus \dots \oplus 2^{d_j-1}a_{d_j-1,j}m_{k-d_j+1,j} \\ \oplus 2^{d_j}m_{k-d_j,j} \oplus m_{k-d_j,j}$$

On forme alors les $v_{k,j}$ avec la règle : $v_{k,j} = m_{k,j}/2^k$.

Calcul de $s_{i,j}$

La partie fractionnaire de la j^{eme} composante du i^{eme} élément de la séquence de Sobol est alors obtenue en faisant le ou exclusif ($\oplus$) d'un sous-ensemble des $v_{k,j}$ (écrits en binaire, bien sûr).

Pour déterminer le sous-ensemble, on applique le critère suivant : un $v_{k,j}$ fera partie de la somme si le k^{eme} bit de la représentation binaire de i vaut 1.

Lorsque i augmente, chaque $v_{k,j}$ rentre et sort de la somme avec une période qui dépend de k . Ainsi, $v_{1,j}$ rentre et sort une fois sur deux⁵⁷, $v_{2,j}$ deux fois sur quatre, et ainsi de suite.

Une amélioration de cet algorithme a été proposée par Antonov et Saleev (cf. [80]). Ils ont montré que la séquence conservait ses propriétés si on considérait les bits de $G(i)$ au lieu de ceux de i pour le critère de choix des $v_{k,j}$. $G(i)$ est le « Gray code » de l'entier i . Il se calcule très simplement :

$$G(i) = i \oplus E \left[\frac{i}{2} \right]$$

Où $E[n]$ représente la partie entière de n . Tous les termes de l'équations ci-dessus sont bien sûr écrits en binaire pour pouvoir effectuer l'opération $\oplus$.

Il est facile de voir que les représentations binaires de $G(i)$ et de $G(i+1)$ ne diffèrent que par un seul bit. La position de ce bit différent est la même que celle du 0 le plus à droite dans l'écriture binaire de i . Par conséquent, pour passer du i^{eme} élément de la séquence de Sobol au $i+1^{eme}$, dans une dimension j donnée, il suffira de faire l'opération $\oplus$ avec un seul $v_{k,j}$, celui tel que k est l'indice du 0 le plus à droite de i :

$$s_{i+1,j} = s_{i,j} \oplus v_{k,j}$$

⁵⁷puisque le premier bit de i en binaire vaut 1 si i est impair et 0 sinon.

Choix de l'utilisateur

La séquence de Sobol est connue pour présenter de bonnes propriétés : calcul rapide, notamment grâce à l'algorithme d'Antonov et Saleev, peu de clustering même à hautes dimension. Certains réglages peuvent néanmoins se révéler utiles. Ainsi, il est recommandé dans la littérature d'omettre les m premiers points de la séquence, m semblant arbitraire. Nous avons choisi d'omettre les 64 premiers. Cette mesure est censée améliorer le comportement de la séquence à hautes dimensions, ce dont nous n'avons a priori pas besoin, c'est en quelque sorte une mesure de sécurité. Les tests montrent que la vitesse de convergence de Π dépend peu de m .

Notons que bien que la séquence soit entièrement déterministe, il y a deux choix à faire avant de lancer le calcul. Tout d'abord, il convient de choisir les polynômes primitifs dont les coefficients serviront au calcul des v_k ; il faut en choisir un par dimension. Sobol, dans son premier article, indique qu'il faut choisir le polynôme de plus bas degré possible pour minimiser la discrédance.

Un autre choix reste cependant, celui des direction numbers. Plus précisément, pour chaque dimension, et donc pour chaque polynôme, l'utilisateur doit choisir les premiers m_k qui serviront à initialiser la récurrence. Il y a autant de m_k à choisir que le degré du polynôme considéré. Les m_k doivent être des entiers impairs inférieurs à 2^k . Le choix de ces nombres, et donc des v_k , conditionne la qualité de la séquence, c'est-à-dire la constante K_N qui intervient dans l'inégalité 49.

De nombreux choix pour les direction numbers ont été testés et rendus publics, notamment par Sobol lui-même, P. Jäckel, et Joe and Kuo. De nombreuses implémentations de l'algorithme d'Antonov et Saleev sont disponibles ; citons celle de Bratley and Fox, et celle de Guignol, que nous avons choisie. Le choix des directions numbers est un réglage fastidieux, mais il peut mener à de réelles différences de performances. C'est donc une étape nécessaire, qui devra être effectuée à l'avenir. Dans ce travail, nous nous sommes contentés de vérifier la convergence de l'erreur en N_{part}^{-1} .

6.2.3 Variation en N_{part}^{-1} du bruit après initialisation

On teste l'efficacité de la méthode Quasi Monte-Carlo en calculant les moments immédiatement après initialisation, *i.e.*, avant mouvement des particules. On vérifie que l'utilisation d'une LDS pour la répartition des particules dans l'espace des phases garantit bien une dépendance en N_{part}^{-1} du bruit sur le tenseur de pression.

Dans la figure 36, on représente le bruit en utilisant la variance de la composante parallèle de la pression scalaire, définie par :

$$p_{\parallel}(\mathbf{r}) = \sum_{n=1}^{N_{part}} v_{n,\parallel}^2 F(\mathbf{r} - \mathbf{r}_n)$$

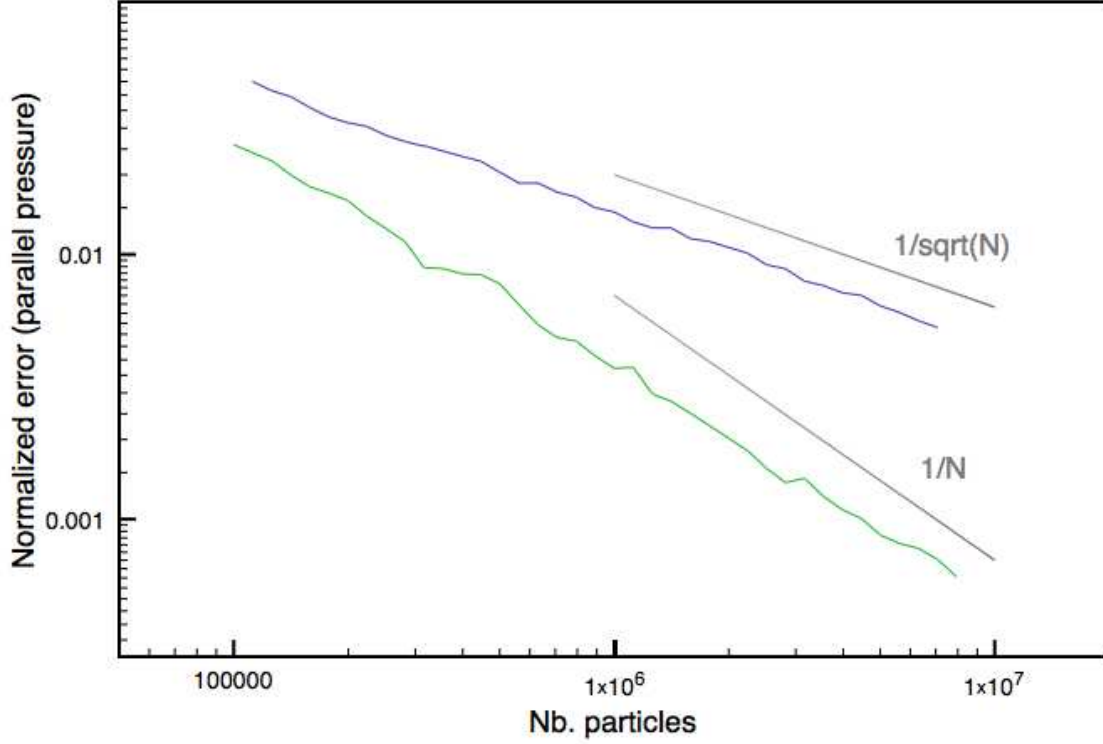


FIG. 36 – Variance de la pression parallèle en fonction du nombre de particules, Sobol (vert) ou random (bleu)

On compare deux stratégies d'initialisation : quasi-random (LDS), et random. On utilise un maillage cartésien de 100×100 . Notons que nous avons pris des poids égaux pour toutes les particules.

6.2.4 Fonction de distribution discrète à poids égaux

L'équation 47 introduit la notion de poids particulaire P_n . Ce poids pondère la contribution de chaque particule à la pression. Il exprime le fait que dans un code PIC, une particule numérique représente un grand nombre de particules physiques. Plus précisément, elle représente l'ensemble des particules qui occupent le volume élémentaire de l'espace des phases centré sur ses coordonnées. Le nombre total de particules impliquées est donc égal à la valeur de la fonction de distribution prise sur ces coordonnées, multipliée par la mesure de ce volume élémentaire. Autrement dit :

$$P_n \equiv f(\mathbf{v}_n, \mathbf{r}_n) \delta^3 \mathbf{v}_n \delta^3 \mathbf{r}_n$$

Le terme $f(\mathbf{v}_n, \mathbf{r}_n)\delta^3\mathbf{v}_n$ est l'analogue discret du $f(\mathbf{v}, \mathbf{r})d^3\mathbf{v}$ de l'équation 46. La mesure spatiale $\delta^3\mathbf{r}_n$ provient du fait que la pression est une grandeur volumique. Cette pondération permet d'ajouter des particules numériques, et donc d'améliorer la précision des calculs, sans modifier la fonction de distribution modélisée. L'ajout de particules numériques entraîne une diminution du volume qu'elles occupent dans l'espace des phases, et donc de leur poids.

Plutôt que d'associer un poids à chaque particule, on a choisi une autre méthode : la distribution des particules est transformée de façon à placer plus de particules aux endroits où le poids devrait être plus grand. Notre démarche est donc de créer une distribution numérique correspondant à une fonction de distribution donnée (dans notre cas, f) à partir d'une séquence uniforme sur $[0, 1[$ (fournie par l'algorithme de Sobol) par l'application d'une fonction mathématique à déterminer⁵⁸. On appliquera ce procédé indépendamment dans chaque dimension de l'espace des phases.

Partons d'une distribution uniforme d'une variable x sur le segment $[0, 1[$; sa densité de probabilité vaut 1. On applique à x une fonction $y(x)$, telle que la densité de probabilité de y est une fonction f donnée. On a alors que :

$$dx = f(y)dy$$

Ce qui dans notre cas peut s'interpréter comme une conservation du nombre de particules⁵⁹. Il s'ensuit que

$$f(y) = \frac{dx}{dy} \implies x(y) = \int_a^y f(u)du \equiv F(y)$$

et finalement :

$$y(x) = F^{-1}(x)$$

La fonction $y(x)$ est la transformation que nous devons appliquer à la séquence de Sobol pour obtenir une séquence dont la densité de probabilité est la fonction de distribution f : c'est la réciproque d'une primitive de f . Cette réciproque n'est pas nécessairement définie *a priori*. Cependant, dans notre cas, l'intégrale de la fonction de distribution restreinte à n'importe quelle dimension de l'espace des phases est une fonction strictement croissante, bornée et continue, et par conséquent sa réciproque existe toujours.

⁵⁸Ce procédé est analogue à l'*importance sampling*, qui est couramment utilisé pour modéliser des fonctions variant brutalement à des endroits précis : on densifie alors le maillage pour améliorer la précision.

⁵⁹Dans un sens mathématique, et non physique.

Pour les variables de vitesse on prend typiquement une fonction de distribution gaussienne ; dans ce cas, la transformation est connue, puisqu'il s'agit de la fonction ierf. Dans la direction R , il convient de commencer par prendre en compte la courbure (le volume élémentaire du tore est proportionnel à R). Le choix exact des paramétrisations de la fonction initiale, pour représenter les différents types de particules chaudes, est en cours de finalisation. Il ne sera donc pas discuté plus avant dans ce travail.

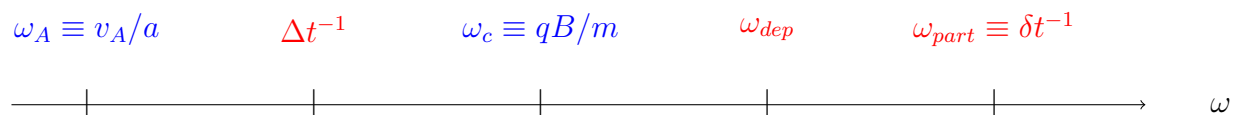
6.3 Aspect temporel et filtrage du tenseur de pression

Dans cette section, on considère la séquence Π_n (où n dénote le pas de temps) comme un signal numérique représentant une grandeur physique qui évolue au cours du temps. On choisit et on implémente un filtrage temporel de ce signal pour éliminer les fréquences de bruit. Notons que le tenseur de pression dépend également de l'espace, et qu'il est possible de régulariser les variations spatiales du tenseur en appliquant un filtre de lissage spatial. Cependant, cette manipulation n'a pas d'autre justification que de faciliter la convergence du schéma, et se heurte à la précision requise dans les couches limites MHD.

Le filtrage temporel, en revanche, découle logiquement de la structure du schéma, c'est pourquoi on l'implémente en priorité. Le tenseur étant calculé à partir des grandeurs particulières, il contient des fréquences supérieures à τ_A^{-1} . Typiquement, la fréquence la plus élevée est de l'ordre de la fréquence cyclotronique. Notons que pour gagner en temps de calcul, le tenseur est déposé tous les n_{dep} pas de temps particulières, avec $n_{dep} \sim 10$. Il est ensuite introduit en entrée du solveur Newton-Krylov à chaque pas de temps fluide (*i.e.*, tous les τ_A environ). Le signal Π_n vu par le solveur est donc un échantillonnage (sampling) du signal initial. La difficulté fondamentale est que les composantes de fréquence plus élevée que la fréquence de sampling τ_A^{-1} ne sont pas résolues et se traduisent par un bruit blanc à l'entrée du solveur, qui contient donc toutes les fréquences, y compris les fréquences basses que l'on veut étudier. On filtre Π_n pour supprimer ces composantes.

6.3.1 Structure temporelle du signal Π_n

Les échelles de temps du système sont récapitulées sur le schéma suivant (sans échelle, les fréquences numériques en rouge et les fréquences physiques en rouge) :



On va utiliser un filtre passe-bas pour éliminer les fréquences de l'ordre de la fréquence cyclotronique. La seule fréquence que nous n'avons pas représentée ici est ω_{rand} . Cette fréquence est caractéristique du transit poloïdal des particules. Grâce à la stratégie d'initialisation, on sait que les irrégularités de l'espace des phases seront créées dynamiquement : la répartition initiale dans l'espace des phases étant idéale, c'est le mouvement des particules qui crée ces irrégularités. En filtrant temporellement le signal sur cette fréquence, on espère éliminer en partie le bruit dû à ces fluctuations. On n'a pas encore mené de tests concluants dans ce sens.

Perte de la régularité initiale de l'espace des phases

Le bénéfice de l'initialisation est lié à l'intégration sur l'espace des phases. La situation serait simple si les moments pouvaient être calculés en intégrant sur les invariants du mouvement des particules (puisque la répartition dans leur espace est constante au cours du temps). Malheureusement, ce n'est pas le cas : à cause de la dépendance en R de P_ϕ et de celle en B de μ , il n'y a pas de relation entre (E_{cin}, μ, P_ϕ) et (v_R, v_ϕ, v_Z) . L'existence d'invariants exprime des contraintes sur la dynamique de chaque particule, mais ces contraintes sont insuffisantes pour conserver une bonne répartition dans toutes les dimensions de l'espace des phases. Le théorème de Liouville assure que le volume d'un élément de l'espace des phases reste constant le long de sa trajectoire. Le problème est que cet élément est néanmoins déformé dans certaines directions.

On s'attend à ce que le mélange des phases dans l'espace réel soit l'effet le plus dommageable du point de vue du bruit. La coordonnée r est directement liée à Ψ , le flux magnétique poloïdal. Les particules ne s'écartent de leur surface $\Psi = cte$ que sur une longueur de l'ordre de la largeur banane δ_b . Cet éloignement est négligeable pour la majorité des particules (surtout pour les particules thermiques, moins pour les particules chaudes), et il est périodique. On s'attend à ce que la perte de régularité majeure vienne du mélange de phases en θ , puisque le mouvement des particules est libre dans cette direction, piégeage mis à part. Au cours du mouvement, des « trous » et des zones d'accumulation apparaissent dans la direction θ . À terme, la répartition des particules en θ devient aléatoire : il y a « randomisation » totale de la distribution. Le temps caractéristique associé est égal au temps moyen de transit poloïdal des particules, c'est-à-dire à $\tau_{rand} = 2\pi R/v_\parallel$.

Le filtrage des fréquences supérieures à τ_{rand} est une option intéressante, mais le choix du filtre est conditionné par le gyrofiltrage⁶⁰, dont l'effet est plus certain. Le tenseur de pression ne contient que peu de fréquences entre la plage à préserver (de l'ordre des fréquences fluides, ω_A) et la plage à couper (la fréquence cyclotronique). On n'a donc pas besoin d'un roll-off très élevé (*i.e.*, d'une pente très brutale dans le graphe de gain du filtre). Cela permet de privilégier une bonne restitution du

⁶⁰le filtrage des composantes de fréquence proche de ω_c

signal dans la bande passante. Le choix s'oriente alors naturellement vers un filtre de Butterworth.

6.3.2 Filtre de Butterworth

Un filtre de Butterworth est conçu pour avoir un gain le plus plat possible dans la bande passante. Cela assure une restitution du signal optimale dans la bande passante. En contrepartie, ce type de filtre a un roll-off assez faible, ce qui explique pourquoi on préfère généralement des filtres elliptiques ou des filtres de Chebychev. On a vu cependant que dans notre cas, le roll-off n'est pas le critère le plus important. Contrairement à ces filtres, le gain d'un filtre de Butterworth reste toujours monotone dans la bande passante lorsque l'ordre du filtre augmente. Les filtres de Butterworth ont été proposés en 1930 par S. Butterworth (cf. [81]) ; ils sont présentés en détail dans [82]. Les calculs des coefficients numériques sont développés sur [83].

Fonction de transfert

Le carré de la fonction de gain d'un filtre de Butterworth d'ordre N est définie par :

$$|H(i\omega)|^2 = \frac{G_0^2}{1 + (\frac{\omega}{\omega_c})^{2N}}$$

$H(i\omega)$ est la fonction de transfert du filtre, ω_c la fréquence de coupure, et G_0 est le gain à fréquence nulle (nous le prenons égal à 1).

De la forme de la fonction de gain, on déduit les caractéristiques suivantes :

- lorsque N augmente, la pente logarithmique du gain augmente linéairement (-20dB par décade pour $N = 1$, -40dB pour $N = 2$, etc.).
- lorsque N tend vers $+\infty$, la fonction de gain tend vers une marche, avec une fonction constante pour $\omega < \omega_c$ et $\omega > \omega_c$ et une discontinuité en ω_c .
- la dérivée du gain est strictement négative quel que soient N et ω , i.e., le gain est toujours monotone (décroissant), sans « ripple ».
- toutes les dérivées du gain jusqu'à la $(2N - 1)^{\text{ème}}$ sont nulles en $\omega = 0$, i.e., le gain est le plus plat possible dans la bande passante.
- le gain vaut $1/\sqrt{2}$ en $\omega = \omega_c$ quel que soit l'ordre du filtre.

Factorisation

Pour simplifier la discrétisation du système, il est utile de factoriser la fonction de transfert du filtre. La fonction étant donnée sous la forme d'une fraction rationnelle, il faut déterminer ses pôles. Soit s un variable complexe, on pose $s = i\omega$, et on écrit que

$$|H(i\omega)|^2 = H(s)H(-s) = \frac{G_0^2}{1 + (\frac{s}{i\omega_c})^{2N}}$$

Les pôles de cette fraction sont les solutions de $1 + (\frac{s}{i\omega_c})^{2N} = 0$: ce sont les racines $2N^{\text{èmes}}$ de -1 multipliées par $i\omega_c$:

$$\tilde{s}_k = i\omega_c e^{i\pi \frac{(2k-1)}{2N}} = \omega_c e^{i\pi \frac{(2k-1+N)}{2N}}, \text{ avec } k = 0, \dots, 2N-1$$

La fraction étant le carré de la fonction de transfert, la factorisation de la fonction de transfert elle-même n'utilise que la moitié des pôles. Dans le plan complexe, ces pôles sont symétriques par rapport à l'axe vertical. On choisit pour la factorisation les pôles situés à gauche de cet axe (partie réelle négative). Après simplification de l'écriture, on obtient les pôles :

$$\left\{ s_k = \omega_c e^{i\pi \frac{(2k+1)}{2N}}, k = 0, \dots, N-1 \right\}$$

Finalement, la fonction de transfert factorisée s'écrit :

$$H(s) = \frac{\omega_c^N}{\prod_{k=0}^{N-1} (s - s_k)}$$

Le dénominateur est un polynôme de Butterworth.

Décomposition en filtres d'ordre 2

Ainsi, un filtre d'ordre 2 s'écrit :

$$\begin{aligned} H(s) &= \frac{1}{(\frac{s}{\omega_c} - e^{\frac{3i\pi}{4}})} \cdot \frac{1}{(\frac{s}{\omega_c} - e^{\frac{5i\pi}{4}})} \\ &= \frac{\omega_c^2}{s^2 + 2\omega_c \cos(\pi/4)s + \omega_c^2} \end{aligned} \quad (50)$$

A partir de l'expression générale, pour N pair, on peut coupler les facteurs deux à deux : l'indice k avec l'indice $N-1-k$ (*i.e.*, on couple chaque racine avec son symétrique par rapport à l'axe des réels dans le plan complexe). On peut alors écrire tout filtre d'ordre N pair comme le produit de $N/2$ filtres d'ordre 2 :

$$H(s) = \prod_{k=0}^{N/2-1} \frac{\omega_c^2}{s^2 + 2s \cos\left(\frac{2k+1}{2N}\pi\right) + \omega_c^2}$$

Ainsi, pour un filtre d'ordre pair, le polynôme de Butterworth s'écrit sous la forme d'un produit de trinômes. Un filtre d'ordre impair donne lieu à la même factorisation avec un facteur $(s + 1)^{-1}$ supplémentaire. Cette écriture factorisée est très intéressante, puisqu'elle va permettre une implémentation en cascade de tout filtre d'ordre arbitrairement élevé. La fonction de transfert exposée plus haut correspond à un filtre analogique, associé à une variable de temps continue. Il nous faut maintenant convertir ce système en un système discret, associé à une variable de temps discrète, pour l'implémenter numériquement.

Transformée en Z et équation aux différences.

En analyse du signal, il est pratique de représenter un signal discret x_n (dans notre cas, Π_n) par sa transformée en z (« z -transform »). Cette transformée est une généralisation de la transformée de Fourier discrète ; il s'agit d'une fonction complexe X d'une variable complexe z , définie par :

$$X(z) = \sum_{n=0}^{\infty} x_n z^{-n}$$

L'intérêt de cette représentation est que la relation entre le signal d'entrée et le signal de sortie du filtre dans l'espace des fréquences (espace de z) correspond à la transformée en Z d'une simple équation aux différences, linéaire, dans l'espace temporel (espace de x_n). En effet, cette relation s'écrit à l'aide la fonction de transfert :

$$H(z) = \frac{Y(z)}{X(z)}$$

Où $Y(z)$ est le signal de sortie du système et $X(z)$ est le signal d'entrée, dans l'espace des fréquences (i.e., X et Y sont les transformées en Z des signaux x_n et y_n , y_n étant le signal filtré).

Comme on a exprimé H sous la forme d'une fraction rationnelle, on peut écrire que :

$$\begin{aligned} \frac{Y(z)}{X(z)} &= \frac{\sum_{i=0}^I z^{-i} \alpha_i}{\sum_{j=0}^J z^{-j} \beta_j}, \text{ soit :} \\ Y(z) \sum_{j=0}^J z^{-j} \beta_j &= X(z) \sum_{i=0}^I z^{-i} \alpha_i \end{aligned}$$

Cette dernière expression est la transformée en Z de l'équation aux différences :

$$\sum_{j=0}^J \beta_j y_{n-j} = \sum_{i=0}^I \alpha_i x_{n-i}$$

Si le coefficient β_0 est non nul, cette équation donne l'expression de y_n en fonction des $J - 1$ éléments précédents de y et des I éléments précédents de x (x_n inclus). Autrement dit, une fois les coefficients β_j et α_i du système calculés, une simple expression algébrique permet de trouver les éléments successifs du signal de sortie y . Le nombre de termes de chaque côté de cette expression est égal au degré des polynômes dans l'expression de H , qui est égal à l'ordre du filtre. Remarquons qu'il sera nécessaire de faire une hypothèse sur les N premiers éléments du signal de sortie, puisque l'équation aux différences ne peut pas les fournir.

« Bilinear transform ».

Pour pouvoir exprimer H sous la forme écrite au-dessus, il est nécessaire de faire un changement de variable pour passer de s à z . En effet, $s = i\omega$, peut avoir un module qui tend vers ∞ . À cause de la définition de la transformée en Z , cela est inacceptable pour z . Pour garantir l'existence de la transformée, on veut projeter l'axe imaginaire du plan complexe de s sur le cercle unité du plan complexe de z . Pour cela, la transformation naturelle est le logarithme :

$$z = e^{sT}$$

Où T est le temps d'échantillonnage du système (x_n et x_{n+1} représentent une grandeur physique en t et en $t + T$). Dans notre cas, ce sera le pas de temps de dépôt du tenseur de pression. Ce temps intervient pour conserver l'homogénéité de l'équation, z étant sans dimension.

En prenant une approximation au premier ordre, et en introduisant une constante multiplicative c (telle que $s = e^{sT/c}$) on obtient :

$$z \approx \frac{1 + sT/(2c)}{1 - sT/(2c)}$$

La relation inverse est :

$$s \approx c \frac{2}{T} \frac{z - 1}{z + 1}$$

En remplaçant dans la fonction de transfert, on obtient :

$$H_d(z) \approx H_a\left(c \frac{2}{T} \frac{z-1}{z+1}\right) \quad (51)$$

H_d représente la fonction de transfert du système discret, H_a celle du système continu.

La bilinear transform projette le demi-plan gauche du plan complexe de s à l'intérieur du cercle unité du plan complexe de z . Cette propriété implique que les pôles du système continu se trouvent dans le demi-plan gauche (condition suffisante de stabilité pour le système continu), alors les pôles du système discret seront dans le cercle unité (condition suffisante de stabilité pour le système discret). Autrement dit, cette transformation conserve la stabilité du système (un système est instable si les valeurs du signal de sortie divergent).

Le même raisonnement appliqué aux zéros du système assure que la minimisation de la phase est également conservée.

« Frequency warping »

À la relation entre z et s correspond une relation entre fréquence continue ω_a (définie par $s = i\omega_a$) et fréquence discrète ω_d (définie par $z = e^{i\omega_d T}$). Comme la relation $z = e^{sT/c}$ a été approchée à l'ordre 1, la relation entre les fréquences est non triviale :

$$\omega_a = c \frac{2}{T} \tan\left(\frac{\omega_d T}{2}\right) \quad (52)$$

L'intervalle de fréquence continue $]-\infty; \infty[$ est projeté sur l'intervalle de fréquence discrète $]-\pi/T; \pi/T[$. Les fréquences ne sont dans une relation linéaire que loin de la fréquence d'échantillonnage, c'est-à-dire si :

$$\omega_d \ll \frac{2}{T} \equiv 2f_s$$

La relation 52 est une bijection : à toute fréquence continue correspond une unique fréquence discrète, et réciproquement⁶¹. Autrement dit, le système discret reproduit bien sur l'intervalle $]-\pi/T; \pi/T[$ toutes les variations du système continu sur $]-\infty; \infty[$. Les caractéristiques (gain et phase) du filtre continu sont conservées lors du passage au filtre discret, mais elles sont décalées non-linéairement dans l'espace des fréquences. À haute fréquence, les intervalles de fréquences continues sont projetées sur des intervalles de fréquences discrètes de plus en plus petits. Une autre façon de

⁶¹Remarquons au passage que cela évite tout problème d'aliasing.

voir les choses serait de dire que le système discret se comporte au point ω comme le système continu se comporte au point $c \frac{2}{T} \tan\left(\frac{\omega T}{2}\right)$. Ce phénomène s'appelle le « frequency warping »

C'est ici que l'intérêt de la constante c apparaît : elle fournit un degré de liberté qui permet de choisir le projeté sur $] - \pi/T; \pi/T[$ d'une fréquence particulière de l'axe $] - \infty; \infty[$. Pour un filtre passe-bas, qui ne contient qu'une seule fréquence particulière, la fréquence de coupure ω_c , le choix s'impose de lui-même : on va régler c de façon à ce que fréquence continue et fréquence discrète soient égales en ω_c . Le warping de l'axe des fréquences se fera de part et d'autre de cette fréquence centrale. Ainsi les déformations seront importantes le plus loin possible de la fréquence de coupure, c'est-à-dire dans des domaines où le gain sera, de toute façon, plat (qu'il soit égal à 0 ou à 1).

Cette opération est appelée « pre-warping ». Dans notre cas, il faut poser :

$$c = \omega_c \frac{T}{2} \tan^{-1} \left(\frac{\omega_c T}{2} \right) \quad (53)$$

On a bien alors $\omega_a = \omega_d$ en ω_c .

Coefficients d'un filtre d'ordre pair

En remplaçant l'expression 53 dans 51, puis dans 50, on obtient la fonction de transfert du filtre numérique d'ordre 2 :

$$H(z) = \frac{\Omega(1 + 2z^{-1} + z^{-2})}{(1 + 2\cos(\frac{\pi}{4})\Omega + \Omega^2) + 2z^{-1}(\Omega^2 - 1) + z^{-2}(1 - 2\cos(\frac{\pi}{4})\Omega + \Omega^2)}$$

Où on a posé $\Omega = \tan(\omega_c T)/2$.

Plus généralement, tout système d'ordre N pair s'écrit comme un produit de systèmes d'ordre 2 :

$$H(z) = \prod_{k=0}^{N/2-1} \frac{\alpha_{k0} + \alpha_{k1}z^{-1} + \alpha_{k2}z^{-2}}{1 + \beta_{k1}z^{-1} + \beta_{k2}z^{-2}}$$

Avec :

$$\begin{aligned}
\alpha_{k0} &= \alpha_{k2} = \Omega^2/c_k \\
\alpha_{k1} &= 2\Omega^2/c_k \\
\beta_{k1} &= \frac{2(\Omega^2 - 1)}{c_k} \\
\beta_{k2} &= \frac{1 - 2 \cos((2k + 1)\pi/2N)\Omega + \Omega}{c_k} \\
c_k &= 1 + 2 \cos((2k + 1)\pi/2N)\Omega + \Omega^2
\end{aligned}$$

Conditions initiales

Une fois établis les coefficients des $N/2$ filtres élémentaires d'ordre 2 qui composent le filtre d'ordre N , l'implémentation du filtre est très simple puisqu'il s'agit d'une suite d'équations algébriques. Rappelons cependant que les N premiers termes du signal filtré doivent être choisis arbitrairement. De plus, les premiers éléments du signal filtré sont faussés ; un pic numérique se propage avant d'être amorti, en général sur les ~ 100 premiers éléments du filtre.

Dans la pratique, ces considérations n'auront pas d'importance pour nous, puisque le début du signal correspond au début d'une simulation, pendant lequel les modes n'ont pas encore eu le temps de croître et de sortir du bruit numérique.

6.4 Résultats préliminaires : kink interne avec particules chaudes

On l'a dit, les dents de scie sont un candidat particulièrement intéressant pour la simulation avec un code hybride : l'effet des particules chaudes est à la fois important et bien documenté. Une simulation hybride d'un kink interne s'inscrirait dans la continuation des simulations que nous avons présentées en première partie, et constituerait la suite logique de cette thèse. On montre ci-dessus des résultats préliminaires, utilisant les outils hybrides que nous avons présentés pour simuler un kink interne sur un temps limité (quelques $10^3 \tau_A$).

Dans cette simulation, on a utilisé 8×10^5 particules α à 3.5 MeV, pour une pression cinétique de $\beta_\alpha = 0.01$ (un ordre de grandeur en-dessous des grandeurs réalistes). La pression du « bulk » est de $\beta_p = 0.3$, c'est-à-dire juste en-dessous du seuil idéal pour le profil parabolique de q présenté dans la première partie, qui permettait d'observer des dents de scie grâce aux effets diamagnétiques. La pression particulière est trop faible pour stabiliser complètement le kink : on s'attend à voir croître un kink résistif. La simulation a tourné sur une dizaine d'heure sur quatre processeurs ; il est donc facilement envisageable de multiplier le nombre de particules

par un facteur 10 à condition de disposer de plus de cœurs. Le pas de temps fluide est $0.25\tau_A$, le pas de temps de dépôt est de $0.025\tau_A$, et le pas de temps particulaire est de $0.0125\tau_A$, soit environ 10 points par période cyclotronique. La figure 37 montre l'allure de la composante perpendiculaire de la pression ($\Pi_{22} + \Pi_{33}$, dans une base locale liée à B), au centre du plasma, avant et après gyrofiltrage.

On constate que le filtrage des composantes cyclotroniques élimine efficacement la partie majeure du bruit du tenseur de pression. L'axe des abscisse est gradué en pas de temps de dépôt. On distingue une fréquence rémanente après filtrage, correspondant à environ 200 pas de temps, soit $5\tau_A$. Cette fréquence est proche de la fréquence de transit des particules ω_{rand} que nous avons évoquée; un filtrage de cette fréquence, non montré ici, pourrait éventuellement encore améliorer le signal.

Les figures 38, 39 montrent l'énergie cinétique associée aux modes $n = 0$ et $n = 1$ (le kink), avec et sans filtre numérique. Les 100 premiers τ_A sont consacrés à la relaxation par mélange de phase de la distribution des α , initialisée comme une maxwellienne isotrope. Après ces 100 τ_A , les particules sont couplées à la partie fluide.

On voit clairement que le cas filtré permet de voir le kink croître en-dehors du bruit numérique, contrairement au cas non filtré. La majorité du bruit restant, associé au transit toroïdal et poloïdal des particules, excite des ondes d'Alfvén aux fréquences correspondantes, de l'ordre de $0.1\tau_A^{-1}$. Les valeurs des vitesses bruitées, pour la figure 38, sont de l'ordre de $10^{-4}v_A$ (l'échelle verticale est trompeuse), et les valeurs du bruit sur le champ magnétique sont de 10^{-5} à 10^{-4} en unités XTOR-K. Ces résultats, bien que préliminaires, sont très encourageants, puisqu'ils démontrent la faisabilité de simulations hybrides, tant en termes de temps de calcul qu'en termes de bruit.

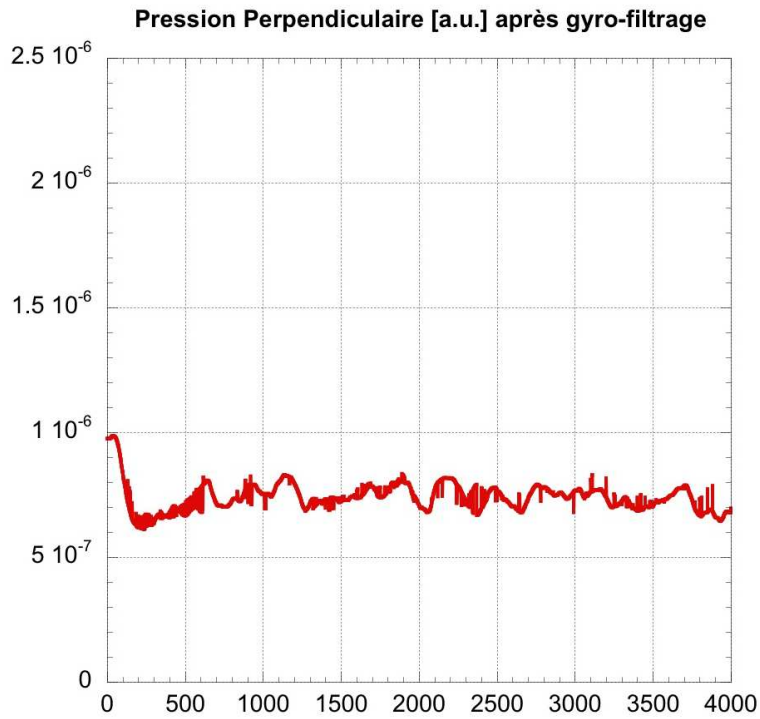
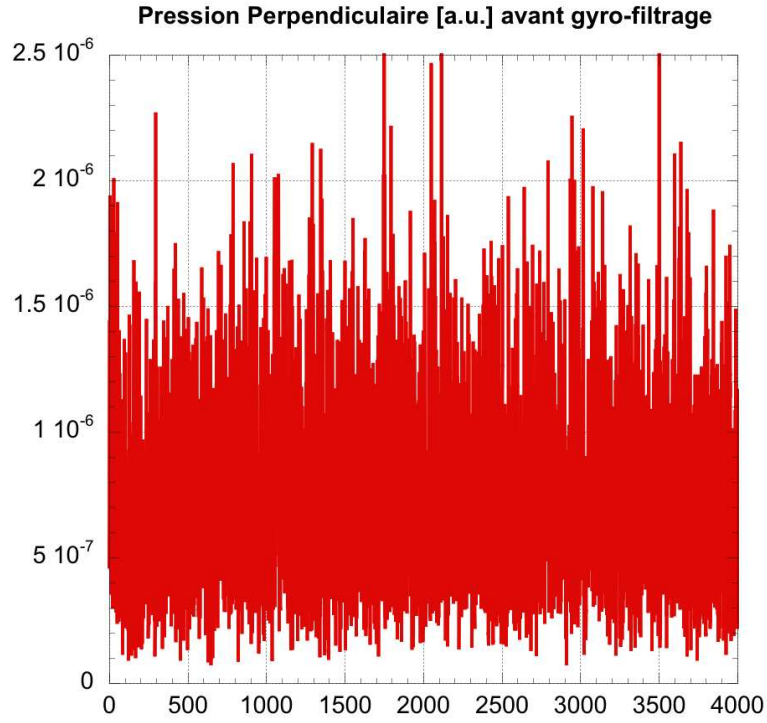


FIG. 37 – Pression perpendiculaire (unité arbitraire) au cours du temps, sans (haut) et avec (bas) gyrofiltrage

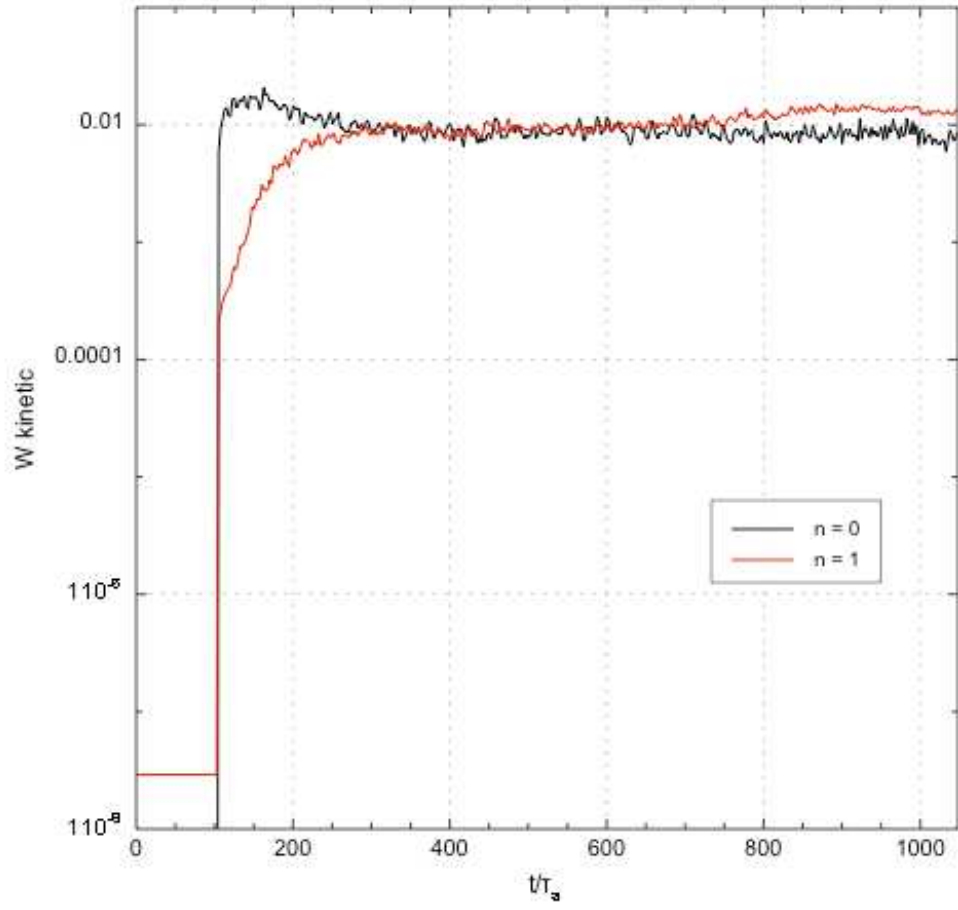


FIG. 38 – Énergie cinétique du kink interne modélisé par XTOR-K, sans filtrage

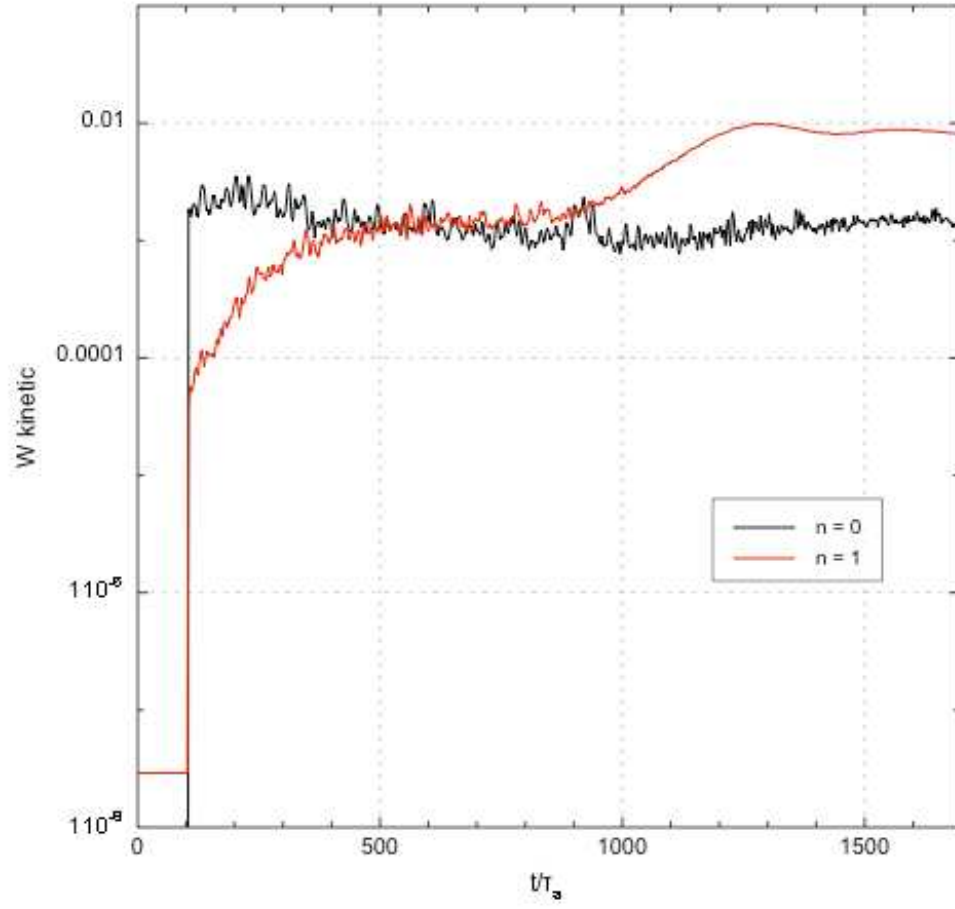


FIG. 39 – Énergie cinétique du kink interne modélisé par XTOR-K, avec gyrofiltrage

7 Conclusion Générale

Le travail effectué lors de cette thèse a été présenté en deux temps : exploitation du code modulaire XTOR-2F, et développement d'outils nécessaires à sa transition vers un code hybride MHD/cinétique.

D'une part, on a effectué une étude numérique des régimes d'oscillation du kink interne sur des temps longs à l'aide du code bi-fluide XTOR-2F. Notre but était d'explorer la dynamique de la rampe des dents de scie, qui sont un phénomène universel des plasmas de tokamaks. Il n'existait pas jusque-là, à notre connaissance, d'étude numérique similaire : simulation de plusieurs dizaines de cycles, incluant des effets MHD et diamagnétiques, en géométrie toroïdale complète.

Nous avons d'abord établi un cadre de référence pour la dynamique du kink en fonction des paramètres $(\chi_{\perp}, \beta_p)$, en $\eta\chi$ MHD. Nous avons constaté l'existence de deux régimes : l'un, à faible $(\chi_{\perp}, \beta_p)$, est un régime d'oscillations régulières de la pression centrale ; l'autre, atteint lorsque les facteurs favorisant le kink sont augmentés, est un régime stationnaire dans lequel le cœur du plasma est bloqué dans une configuration 3D d'hélicité $m = n = 1$. Ce régime est accompagné de l'apparition d'une cellule de convection et d'une couche de reconnection permanentes autour de la surface $q = 1$. Nous avons ensuite concentré notre attention sur une fenêtre de paramètres réduite, et nous avons vu que l'ajout d'effets diamagnétiques décalait le régime oscillant vers des valeurs de pression plus hautes, ce qui permet de retrouver des oscillations pour des paramètres proches de ceux de JET, pour une décharge ohmique. Les temps caractéristiques du régime oscillant se rapprochent de ceux de dents de scie expérimentales. Nous avons noté qu'un régime stationnaire similaire à celui que nous simulons a été observé dans certaines machines ; une comparaison avec les données expérimentales permettrait d'en savoir plus sur ce régime et ses conditions d'existence.

D'autre part, nous avons présenté les choix effectués et certains des outils développés dans le cadre de la transition du code vers le code hybride MHD/cinétique XTOR-K. Cette transition est encore en cours ; elle permettra à terme d'inclure des populations sortant du cadre de la MHD, en particulier les particules rapides, pour arriver à une modélisation plus intégrée des plasmas de tokamaks. Les particules chaudes, déjà présentes dans les machines à cause des chauffages par ondes ou injection de neutres, et utilisées pour contrôler les profils de courant et de chauffage du plasma, sont vouées à acquérir une importance encore accrue dans les futures machines comme ITER, où les α issus de la fusion assureront la majorité du chauffage.

Contrairement aux autres codes hybrides en cours de développement, XTOR-K est un code « full-orbit, full-f ». Il intègre exactement le mouvement de la totalité des particules formant la population cinétique. Les effets de rayons de Larmor fini sont inclus naturellement, aucune hypothèse n'est faite sur la dynamique de la fonction de distribution des particules, et le tenseur de pression particulière s'évalue simple-

ment. Numériquement, la partie cinétique (le « pousseur ») utilise un sous-pas de temps particulaire, et prend la majorité du temps de calcul. L’algorithme choisi, de type Newton-Krylov/Picard, concentre la plus grande partie des itérations sur le solveur implicite Newton-Krylov (qui résout les équations fluides), permettant une parallélisation massive et efficace des itérations de Picard (typiquement 2 ou 3 par pas de temps pour gérer la partie cinétique) : une fois les champs fluides transmis au code particulaire, chaque processeur avance ses particules sur un grand nombre de sous-pas de temps indépendamment. Le schéma permet de modéliser des fractions importantes de la pression avec le modèle cinétique, et n’exige pas de réduction importante du pas de temps fluide. Le pousseur, adapté de la méthode de Boris pour une géométrie torique, présente des bonnes qualités de stabilité et de précision, qu’on a vérifiées en analysant la conservation des invariants du mouvement. Le sous-pas de temps doit être inférieur ou égal à $1/8$ de la période cyclotronique.

Une modélisation intégrée du plasma, comme celle que fournirait un code hybride, serait particulièrement intéressante dans le cas des dents de scie. En effet, leur grande variété et la multiplicité des effets physiques à prendre en compte rend improbable la construction d’un modèle analytique simple. Dernièrement, la communauté scientifique s’est plutôt orientée vers une approche semi-empirique visant à prédire les scaling des grandeurs caractéristiques (période, amplitude). À l’inverse, un code hybride tel que **XTOR-K** permettrait de mener des simulations incluant la majorité des effets considérés comme importants pour les dents de scie (effets bi-fluides, influence des particules rapides, etc), dans une démarche plus analytique. Les deux parties du travail de thèse se rejoignent dans cet objectif. L’état d’avancement actuel du code hybride permet d’envisager une telle simulation dans les mois qui viennent.

À l’heure actuelle, le code hybride a déjà été utilisé pour simuler un kink interne à partir de la configuration de départ utilisée dans le chapitre 2, avec une pression proche du seuil de l’instabilité ; les graphes ont été présentés à la fin du dernier chapitre. Les simulations sont suffisamment longues (quelques $10^3 \tau_A$) pour voir clairement le kink sortir du bruit, après gyrofiltrage ; grâce à l’utilisation du filtre temporel, on a obtenu un rapport signal sur bruit correct. Les simulations durent une dizaine d’heures sur station (quatre cœurs), avec 800 000 particules. Ces paramètres sont encourageants, puisqu’ils permettent d’envisager des simulations tournant en un temps raisonnable, avec un bruit faible, en répartissant un nombre plus élevé de particules ($10^8 - 10^9$) sur de nombreux processeurs.

Références

- [0] A. Leblond. Poétique du roman-fleuve, de Jean-Christophe à Maumort. Université Paris-III, Paris (2010).
- [1] J. P. Freidberg. *Ideal Magnetohydrodynamics*, ed. Plenum Press, New York, USA (1987).
- [2] R. D. Hazeltine and J. D. Meiss. *Plasma Confinement*, ed. Dover-Publications, Inc., New York, USA (2003).
- [3] S. I. Braginskii. Transport processes in a plasma. *Reviews of Plasma Physics*, **vol.1**, ed. M. A. Leontovich, Consultants Bureau, New York, p. 206 (1965).
- [4] S. von Goeler, W. Stodiek, and N. Sauthoff. Studies of internal disruptions and $m = 1$ oscillations in tokamak discharges with soft x-ray techniques. *Physical Review Letters*, **33(20)**, p. 1201 (1974).
- [5] R. J. Hastie. Sawtooth Instability in Tokamak Plasmas. *Astrophysics and Space Science*, **256**, p. 177 (1997).
- [6] T. C. Hender *et al.* and the ITPA MHD, Disruption and Magnetic Control Topical Group 2007. Chapter 3 : MHD stability, operational limits and disruptions. *Nuclear Fusion*, **47**, p. S128 (2007).
- [7] F. Porcelli. Fast particles stabilisation. *Plasma Physics and Controlled Fusion*, **33**, p. 1601 (1991).
- [8] V. D. Shafranov. Hydromagnetic stability of a current-carrying pinch in a strong longitudinal magnetic field. *Soviet Physics - Technical Physics*, **15**, p. 175 (1970).
- [9] M. N. Bussac, R. Pellat, D. Edery, and J. L. Soulé. Internal kink modes in toroidal plasmas with circular cross sections. *Physical Review Letters*, **35(24)**, p. 1638 (1975).
- [10] M. N. Rosenbluth, R. Y. Dagazian, and P. H. Rutherford. Nonlinear properties of the internal $m = 1$ kink instability in the cylindrical tokamak. *Physics of Fluids*, **16**, p. 1894 (1973).
- [11] G. Ara, B. Basu, B. Coppi, G. Laval, M. N. Rosenbluth, and B. V. Waddell. Magnetic reconnection and $m = 1$ oscillations in current carrying plasmas. *Annals of Physics*, **112**, p. 443 (1978).
- [12] H. P. Furth, J. Killeen, and M. N. Rosenbluth. Finite Resistivity Instabilities of a Sheet Pinch. *Physics of Fluids*, **6**, p. 459 (1963).
- [13] B. Coppi, R. Galvão, R. Pellat, M. N. Rosenbluth, and P. Rutherford. *Fizika Plazmy*, **6**, p. 961 (1976).
- [14] B. B. Kadomtsev. *Soviet Journal of Plasma Physics*, **1**, p. 389 (1975).
- [15] J. A. Wesson. *Tokamaks (Third Edition)*, ed. Clarendon Press, Oxford, UK (2003).

- [16] A. Sykes and J. A. Wesson. Relaxation Instability in Tokamaks. *Physical Review Letters*, **37**, p.140 (1976).
- [17] R. E. Denton, J. F. Drake, and R. G. Kleva. The $m = 1$ convection cell and sawteeth in tokamaks. *Physics of Fluids*, **30**, p. 1447 (1987).
- [18] A. Y. Aydemir, J. C. Wiley, and D. W. Ross. Toroidal studies of sawtooth oscillations in tokamaks. *Physics of Fluids B*, **1**, p. 774 (1989).
- [19] W. Park and D. A. Monticello. Sawtooth oscillations in tokamaks. *Nuclear Fusion*, **30**, p. 2413 (1990).
- [20] M. Dubois, A. L. Pecquet, and C. Reverdin. Internal disruptions in the TFR tokamak : a phenomenological analysis. *Nuclear Fusion*, **23**, p.147 (1983).
- [21] J. A. Wesson. Sawtooth reconnection. *Nuclear Fusion*, **30**, p. 2545 (1990).
- [22] L. E. Zakharov and B. N. Rogers. Two-fluid magnetohydrodynamic description of the internal kink mode in tokamaks. *Physics of Fluids B*, **4**, p. 3285 (1992).
- [23] X. Wang and A. Bhattacharjee. Nonlinear dynamics of the $m = 1$ kink tearing instability in a modified magnetohydrodynamic model. *Physics of Plasmas*, **2**, p. 171 (1995).
- [24] A. Y. Aydemir. Magnetohydrodynamic modes driven by anomalous electron viscosity and their role in fast sawtooth crashes. *Physics of Plasmas*, **2**, 2135 (1990).
- [25] J. A. Wesson. Sawtooth oscillations. *Plasma Physics and Controlled Fusion*, **28**, p. 243 (1986).
- [26] A. Y. Aydemir. Mechanism for rapid sawtooth crashes in tokamaks. *Physical Review Letters*, **59**, p. 649 (1987).
- [27] H. Soltwisch. Measurement of current density changes during sawtooth activity in a tokamak by far-infrared polarimetry. *Review of Scientific Instruments*, **59**(8), p. 1599 (1988).
- [28] M. Yamada, F. M. Levinton, N. Pomphrey, R. Budny, J. Manickam, and Y. Nagayama. Investigation of magnetic reconnection during a sawtooth crash in a high-temperature plasma. *Plasma Physics*, **1**(10), p.3269 (1994).
- [29] J. Blum, J. Lazzaro, J. O'Rourke, B. Keegan, and Y. Stefan. Problems and methods of self-consistent reconstruction of tokamak equilibrium profiles from magnetic and polarimetric measurements. *Nuclear Fusion*, **30**(8), p. 1475 (1990).
- [30] K. McCormick, F. X. Söüldner, D. Eckhardt, F. Leuterer, H. Murmann, H. Derfler, A. Eberhagen, O. Gehre, J. Gernhardt, G. v. Gierke, O. Gruber, M. Keilhacker, O. Kluüber, K. Lackner, D. Meisel, V. Mertens, H. Roühr, K.-H. Schmitter, K.-H. Steuer, and F. Wagner. Temporal behavior of the plasma current distribution in the ASDEX tokamak during lower-hybrid current drive. *Physical Review Letters*, **58**, p. 487 (1987).

- [31] L. K. Huang, M. Finkenthal, D. Wróblewski, H. W. Moos, W. P. West, P. E. Phillips, and D. C. Sing. Safety factor measurements on the magnetic axis of the Texas experimental tokamak plasma in ohmic and electron-cyclotron-resonance heated discharges. *Physical Review A*, **45**(2), p. 1089 (1992).
- [32] H. Weisen, G. Borg, B. Joye, A. J. Knight and, J. B. Lister. Measurements of the tokamak safety factor profile by means of driven resonant Alfvén waves. *Physical Review Letters*, **62**, p. 434 (1989).
- [33] A. Weller, A. D. Cheetham, A. W. Edwards, R. D. Gill, A. Gondhalekar, R. S. Granetz, J. Snipes, and J. A. Wesson. Persistent density perturbations at rational-q surfaces following pellet injection in the Joint European Torus. *Physical Review Letters*, **59**, p. 2303 (1987).
- [34] C. Angioni, A. Pochelon, N. N. Gorelenkov, K. G. McClements, O. Sauter, R. V. Budny, P. C. de Vries, D. F. Howell, M. Mantsinen, M. F. F. Nave, S. E. Sharapov and contributors to the EFDA-JET Workprogram. Neutral beam stabilization of sawtooth oscillations in JET. *Plasma Physics and Controlled Fusion*, **44**, p. 205 (2002).
- [35] H. Reimerdes, A. Pochelon, O. Sauter, T. P. Goodman, M. A. Henderson, and An. Martynov. Effect of triangular and elongated plasma shape on the sawtooth stability. *Plasma Physics and Controlled Fusion*, **42**, p. 629 (2000).
- [36] C. Angonia, T.P. Goodman, M.A. Henderson and O. Sauter. Effects of localized electron heating and current drive on the sawtooth period. *Nuclear Fusion*, **43**, p. 455 (2003).
- [37] J. P. Graves. Influence of asymmetric energetic ion distributions on sawtooth stabilization. *Physical Review Letters*, **92** (18), 185003 (2004).
- [38] A. Mück, T. P. Goodman, M. Maraschek, G. Pereverzev, F. Ryter, H. Zohm, and ASDEX Upgrade Team. Sawtooth control experiments on ASDEX Upgrade. *Plasma Physics and Controlled Fusion*, **47** p. 1633 (2005).
- [39] F. Porcelli, D. Boucher, and M. N. Rosenbluth. Model for the sawtooth period and amplitude. *Plasma Physics and Controlled Fusion*, **38**, p. 2163 (1996).
- [40] H. Lütjens and J.-F. Luciani. The XTOR code for nonlinear 3D simulations of MHD instabilities in tokamak plasmas. *Journal of Computational Physics*, **227** (14), p. 6944 (2008).
- [41] H. Lütjens, J.-F. Luciani, and X. Garbet. Nonlinear three-dimensional MHD simulations of tearing modes in tokamak plasmas. *Plasma Physics and Controlled Fusion*, (Suppl. 12A), **43** A339 (2001).
- [42] H. Lütjens and J.-F. Luciani. Saturation levels of neoclassical tearing modes in ITER plasmas. *Physics of Plasmas*, **12**, 080703 (2005).
- [43] H. Lütjens and J.-F. Luciani. Stochastic couplings of neoclassical tearing modes in ITER plasmas. *Physics of Plasmas*, **13**, 112501 (2006).

- [44] P. Maget, H. Lütjens, G. Huysmans, Ph. Moreau, B. Schunke, J.-L. Ségui, X. Garbet, E. Joffrin and J.F. Luciani. MHD stability of (2,1) tearing mode : an issue for the preforming phase of Tore Supra non-inductive discharges. *Nuclear Fusion*, **47**, p. 233 (2007).
- [45] P. Maget, G. T. A. Huysmans, X. Garbet, M. Ottaviani, H. Lütjens, and J.-F. Luciani. Nonlinear magnetohydrodynamic simulation of Tore Supra hollow current profile discharges. *Physics of Plasmas*, **14**, 052509 (2007).
- [46] P. Maget, G. T. A. Huysmans, H. Lütjens, M. Ottaviani, Ph. Moreau and J.-L. Ségui. From MHD regime to quiescent non-inductive discharges in Tore Supra : experimental observations and MHD modelling. *Plasma Physics and Controlled Fusion*, **51**, 065005 (2009).
- [47] H. Lütjens and J.-F. Luciani. XTOR-2F : a fully implicit Newton-Krylov solver applied to nonlinear 3D extended MHD in tokamaks. *Journal of Computational Physics*, **229**, p. 8130 (2010).
- [48] M. Pernice and H. F. Walker. NITSOL : a Newton Iterative Solver for nonlinear systems. *SIAM Journal on Scientific Computing, Special issue on iterative methods*, **19**, p. 302 (1998).
- [49] F. Halpern, D. Leblond, H. Lütjens, and J.-F. Luciani. Oscillation regimes of the internal kink mode in a tokamak plasma. *Plasma Physics and Controlled Fusion*, **53**, 15011 (2011).
- [50] J. M. Rax. *Physique des Plasmas*, ed. Dunod, Paris, France (2005).
- [51] H. Lütjens, A. Bondeson, and O. Sauter. The CHEASE code for toroidal MHD equilibria. *Computer Physics Communications*, **97(3)**, p. 219 (1996).
- [52] D. Leblond, H. Lütjens, and J.-F. Luciani. *Proc. 36th EPS Conference on Plasma Physics and Controlled Fusion (Sofia, Bulgaria, 2009)*, **P1.127**, (2009).
- [53] F. L. Waelbroeck and R. D. Hazeltine. Stability of low-shear tokamaks. *Physics of Fluids*, **31(5)**, p. 1217 (1988).
- [54] H. Baty, J.-F. Luciani, and M. N. Bussac. Field line stochasticity during the m=1 reconnection. *Physics of Fluids B : Plasma Physics*, **5(4)**, p. 1213 (1993).
- [55] A. Y. Aydemir. Nonlinear studies of m=1 modes in high-temperature plasmas. *Physics of Fluids B : Plasma Physics*, **4(11)**, p. 3469 (1992).
- [56] L. Zakharov, B. Rogers, and S. Migliuolo. The theory of the early nonlinear stage of m=1 reconnection in tokamaks. *Physics of Fluids B : Plasma Physics*, **5(7)**, p. 2498 (1993).
- [57] X. Wang and A. Bhattacharjee. Nonlinear dynamics of the m=1 instability and fast sawtooth collapse in high-temperature plasmas. *Physical Review Letters*, **70(11)**, p. 1627 (1993).

- [58] F. Romanelli and R. Kamendje. Overview of JET results. *Nuclear Fusion*, **49(10)**, p. 104006 (2009).
- [59] S. C. Cowley, R. M. Kulsrud, and T. S. Hahm. Linear stability of tearing modes. *Physics of Fluids*, **29(10)**, p. 3230 (1986).
- [60] J. F. Drake, T. M. Antonsen, Jr., A. B. Hassam, and N. T. Gladd. Stabilization of the tearing mode in a high-temperature plasma. *Physics of Fluids*, **26(9)**, p. 2509 (1983).
- [61] R. Fitzpatrick. The effect of trapped particles on the linear stability of long wavelength resistive modes. *Physics of Fluids B : Plasma Physics*, **2(11)**, p. 2636 (1990).
- [62] I. Furno, C. Angioni, F. Porcelli, H. Weisen, R. Behn, T. P. Goodman, M. A. Henderson, Z. A. Pietrzyk, A. Pochelon, H. Reimerdes, and E. Rossi. Understanding sawtooth activity during intense electron cyclotron heating experiments on TCV. *Nuclear Fusion*, **41(4)**, p. 403 (2001).
- [63] J. E. Menard, R. E. Bell, E. D. Fredrickson, D. A. Gates, S. M. Kaye, B. P. LeBlanc, R. Maingi, S. S. Medley, W. Park, S. A. Sabbagh, A. Sontag, D. Stutman, K. Tritz, W. Zhu, and the NSTX research team. Internal kink mode dynamics in high- β NSTX plasmas. *Nuclear Fusion*, **45(7)**, p. 539 (2005).
- [64] C. Z. Cheng. A Kinetic-Magnetohydrodynamic Model for Low-Frequency Phenomena. *Journal of Geophysical Research*, **96, A12**, p.159 (1991).
- [65] C. C. Kim, C. R. Sovinec, S. E. Parker, and The NIMROD Team. Hybrid kinetic-MHD simulations in general geometry. *Computer Physics Communications*, **163**, p.448 (2004).
- [66] W. Park, S. Parker, H. Biglari, M. Chance, L. Chen, C. Z. Cheng, T. S. Hahm, W. W. Lee, R. Kulsrud, D. Monticello, L. Sugiyama, and R. White. Three-dimensional hybrid gyrokinetic-magnetohydrodynamics simulation. *Physics of Fluids B*, **4**, p. 2033 (1992).
- [67] C. K. Birdsall and A. B. Langdon. *Plasma Physics via Computer Simulation*, ed. McGraw-Hill, New York, USA (1985).
- [68] R.W. Hockney and J.W. Eastwood. *Computer Simulation Using Particles*, ed. McGraw-Hill, New York, USA, (1981).
- [69] J. Boris, *Proc. 4th Int. Conf. on Numerical Simulation of Plasmas* (NRL, 1970).
- [70] D. C. Barnes and R. D. Milroy, *Physics of Fluids B*, **3**, 2609 (1991).
- [71] J. M. Wallace, J. U. Nrackbill, and D. W. Forslund, *Journal of Computational Physics*, **63**, 434 (1986)
- [72] J. W. Cobb and J. N. Leboeuf, *Proc. 15th Int. Conf. on Numerical Simulation of Plasmas*, (Valley Forge, 1994).

- [73] R. G. Littlejohn. Variational principles of guiding centre motion. *Journal of Plasma Physics*, **29**, p. 111 (1983).
- [74] T. G. Northrop, *The Adiabatic Motion of Charged Particles*, ed. Interscience, New York, USA (1963).
- [75] M. Kruskal, *Advanced Plasma Theory, course 25*, Proc. Int. School of Phys. Enrico Fermi, ed. Academic Press, New York and London (1964).
- [76] L.-G. Eriksson and F. Porcelli. Dynamics of energetic ions orbits in magnetically confined plasmas. *Plasma Physics and Controlled Fusion* **43**, 145 (2001).
- [77] W. H. Press, S. A. Teukolsky, W. T. Vetterling, and B. P. Flannery. *Numerical Recipes in Fortran 77 : The Art of Scientific Computing, 4th ed.* Cambridge University Press, Cambridge, U.K. (1992).
- [78] H. Niederreiter. *Random Number Generation and Quasi-Monte Carlo Methods*, ed. Society for Industrial and Applied Mathematics, Philadelphia, USA (1992).
- [79] I. M. Sobol. Distribution of points in a cube and approximate evaluation of integrals. *U.S.S.R Computational Mathematics and Mathematical Physics*, **7**, p. 86 (1967).
- [80] I. A. Antonov, V. M. and Saleev. An economic method of computing LP τ -sequences. *U.S.S.R Computational Mathematics and Mathematical Physics*, **19**, p. 252 (1979).
- [81] S. Butterworth. On the theory of filter amplifiers, *Wireless Engineer*, **7**, p. 536 (1930).
- [82] A. V. Oppenheim and R. W. Schaffer, *Digital Signal Processing*, ed. Prentice-Hall, London, UK (1999).
- [83] Y.-H. Kwon. Butterworth Digital Filter, *page web personnelle* : <http://www.kwon3d.com/theory/filtering/butt.html>, consultée au 21/05/2011.