



Etude du système couplé Boltzmann sans collisions-Poisson pour la gravitation. Simulations numériques de la formation des systèmes auto-gravitants

Fabrice Roy

► To cite this version:

Fabrice Roy. Etude du système couplé Boltzmann sans collisions-Poisson pour la gravitation. Simulations numériques de la formation des systèmes auto-gravitants. Astrophysique [astro-ph]. Université de Versailles-Saint Quentin en Yvelines, 2004. Français. NNT : . tel-00007424

HAL Id: tel-00007424

<https://theses.hal.science/tel-00007424>

Submitted on 16 Nov 2004

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

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

**Thèse de Doctorat de l'Université de Versailles –
S^t Quentin-en-Yvelines**

Spécialité :

Mathématiques, Simulation et Modélisation

présenté par

Fabrice Roy

Pour obtenir le titre de Docteur de l'Université de Versailles – S^t Quentin-en-Yvelines

Sujet de la thèse :

**Étude du système couplé Boltzmann sans
collisions–Poisson pour la gravitation.**

**Simulations numériques de la formation des systèmes
auto-gravitants.**

Soutenu le 8 juillet 2004 devant le jury composé de

Mme Evangelie	ATHANASSOULA
M. Patrick	CIARLET
Mme Françoise	COMBES
M. Hervé	DE FERAUDY
M. Gary	MAMON
M. Jérôme	PEREZ

Thèse réalisée au Laboratoire de Mathématiques Appliquées de l'ENSTA



À ma famille et mes amis...

REMERCIEMENTS

La rédaction de ce manuscrit constitua l'une des étapes les plus longues et douloureuses de mes quatre années de thèse. Et cette page fut sans aucun doute celle qui aura nécessité le plus d'efforts et d'attention de ma part, tant je tenais à remercier comme il se doit les personnes ayant participé, directement ou non, à l'achèvement de ce travail.

Mes premiers remerciements vont à Patrick Ciarlet et Jérôme Perez, qui ont dirigé et encadré mes travaux. Merci à Patrick pour ses précieux conseils et sa disponibilité, et à Jérôme pour son indéfectible enthousiasme et sa passion, qu'il sut me communiquer. Je n'aurais pu imaginer meilleurs guides pour accompagner mes premiers pas dans le monde de la recherche scientifique.

Mes plus sincères remerciements vont également à Evangélie Athanassoula et Gary Mamon, qui m'ont fait l'honneur d'être rapporteurs de cette thèse, ainsi qu'à Françoise Combes et Hervé de Feraudy pour avoir accepté de faire partie du jury.

Je remercie chaleureusement tous les membres du Laboratoire de Mathématiques Appliquées, qu'ils soient permanents, thésards ou stagiaires, que j'ai eu beaucoup de plaisir à côtoyer durant ces quatre dernières années.

Je remercie particulièrement Jean-Luc Commeau, pour sa gentillesse et son travail inestimable sur l'installation et la maintenance de la machine parallèle du laboratoire. Sans son assistance, je ne serais probablement jamais parvenu aux mêmes résultats.

Je tiens à exprimer ma reconnaissance à Annie Marchal et Chrystelle Scafarto pour m'avoir aidé à traverser sans encombres les méandres administratifs qui accompagnent chaque étude doctorale.

Merci à Benoît, Christophe et Franck, tous trois jeunes ou futurs docteurs rencontrés sur les bancs de l'université de Nantes, avec qui j'ai pu partager mes petits bonheurs et malheurs de thésard.

Merci à Céline et Franck, et Delphine et Gaël, qui m'ont permis d'avoir un toit et ainsi de terminer mes travaux dans les meilleures conditions, et qui ne sauraient imaginer combien ce témoignage d'amitié compta à mes yeux.

Je remercie infiniment ces amis à qui je n'ai jamais su dire toute l'affection que j'ai pour eux, qui furent présents à mes côtés lorsque j'en avais besoin. Merci donc à Boris, Colin, Émilie, François, Guillaume, Nathalie et Éric, et Mathieu.

Merci enfin à ma famille, pour son soutien et ses encouragements, pour le havre qu'a représenté le domicile familial en ces moments de doute que connaît tout doctorant. Merci pour tout !

TABLE DES MATIÈRES

Introduction	3
1 Modélisation et étude analytique des systèmes auto-gravitants	7
1 Amas globulaires et galaxies elliptiques : une tentative de définition	7
2 Modélisation des systèmes auto-gravitants	8
2.1 Temps caractéristiques	8
2.2 Descriptions discrète et continue	10
2.3 L'équation de Poisson	11
2.4 L'équation de Boltzmann sans collisions	12
2.5 Théorème du viriel	15
3 Solutions analytiques et étude de stabilité	18
3.1 Premiers résultats sur les solutions de Boltzmann sans collisions	18
3.2 Polytropes et modèle de Plummer	20
3.3 Sphère isotherme	22
3.4 Étude complémentaire sur la sphère isotherme et instabilité d'Antonov	26
3.5 Méthode d'énergie et instabilité d'orbite radiale	33
2 Méthodes numériques	41
1 Le treecode	41
1.1 Quelques méthodes numériques et leurs caractéristiques	41
1.2 Description du treecode	44
1.3 Calcul parallèle et treecode	49
2 Conditions initiales	51
2.1 Géométrie du système	51
2.2 Taille du système R_{ini}	51
2.3 Profil de densité	52
2.4 Rapport du viriel initial $\mathcal{V}_{ini}$	53
2.5 Spectre de masse	54
2.6 Distribution des vitesses	55
2.7 Rotation	56
2.8 Nomenclature	56
3 Calcul des observables	57
3.1 Énergie totale	57
3.2 Rapport du viriel	58

3.3	Centre de densité et rayon de densité	59
3.4	Rapports d'axes	61
3.5	Rayons contenant 10, 50 et 90% de la masse	61
3.6	Potentiel et densité	62
3.7	Dispersion de vitesse	62
3.8	Température	62
3.9	Fonctions de distribution	62
4	Système d'unités	63
3	Formation des amas globulaires et des galaxies elliptiques	65
1	Effondrement gravitationnel	65
1.1	Mécanisme	65
2	Études préparatoires	66
2.1	Influence du paramètre d'adoucissement du potentiel	66
2.2	Influence du nombre de particules	68
3	Influence des conditions initiales sur la géométrie du système à l'équilibre	70
4	Taille du cœur et ségrégation	72
5	Fonction de distribution à l'équilibre	75
6	Analyse thermodynamique	77
	Conclusion	81
A	Éléments de calcul scientifique	83
1	Génération de nombres pseudo-aléatoires : méthode de transformation inverse	83
2	Schéma saute-mouton (ou leap-frog)	84
3	Méthode des moindres carrés	86
B	Notions de calcul parallèle et distribué – Introduction à la bibliothèque MPI	87
1	Calcul parallèle	87
1.1	Présentation du calcul parallèle	87
1.2	Machines à mémoire partagée, machines à mémoire distribuée : les différences pratiques	89
1.3	La solution logicielle et matérielle utilisée	90
2	MPI – Message Passing Interface	91
2.1	Le passage de message et le standard MPI	91
2.2	Nœuds et processus	92
2.3	Communicateurs et communications point à point	93
2.4	Communications collectives et communicateurs (compléments)	95
2.5	Types composés et autres fonctions	97
C	Tableaux	99
	Bibliographie	105

INTRODUCTION

Les systèmes auto-gravitants représentent un excellent sujet d'étude pour qui s'intéresse à l'interaction gravitationnelle. Ces objets astronomiques sont en effet gouvernés par la gravitation, seule interaction à l'œuvre sur les échelles de distance sur lesquelles ils s'étendent. Les systèmes auto-gravitants sont de plusieurs types et tailles, et leur étude s'étend de la mécanique céleste, pour les systèmes composés de quelques corps, à la cosmologie, pour les grandes structures de l'Univers. L'astrophysique étant exclue des sciences expérimentales par le fait qu'aucune expérience reproductible n'est réalisable dans ce domaine de la physique – n'est pas démiurge qui le souhaite – elle repose sur les trois autres instruments laissés à la disposition des scientifiques, à savoir l'observation, la modélisation et l'analyse mathématique, et la simulation informatique.

Le travail présenté dans ce manuscrit exploite ces deux derniers outils, et délaisse le premier, essentiellement par manque de temps. Son ambition, limitée en regard des questions encore en attente de réponses, consiste à effectuer une synthèse des principaux résultats analytiques concernant la formation par effondrement et la dynamique des systèmes à N corps gravitationnels non cosmologiques et non collisionnels – les amas globulaires et les galaxies elliptiques – et à exploiter les ressources informatiques grandissantes mises à disposition des scientifiques afin de confronter ces résultats analytiques avec les seules expériences qui nous sont accessibles : les simulations numériques.

De nombreux astrophysiciens se sont penchés sur cette questions depuis des décennies, les principaux résultats analytiques exposés par la suite remontant pour la plupart aux années 1950 ou 1960, voire même plus tôt (travaux de J.H.Jeans ou de A.Schuster). Les premières simulations numériques auxquelles nous faisons référence ont quant à elles été réalisées il y a vingt ans environ (T.S.van Albada [54], 1982).

Il apparaît cependant qu'aucune étude exhaustive n'ait été menée concernant les simulations numériques de formation de systèmes auto-gravitants par effondrement. Nous ne prétendons pas réaliser ici cette étude exhaustive mais du moins compléter les études précédentes et mettre en relation les résultats obtenus avec les développements analytiques correspondants. Nous avons également tenté d'obtenir des résultats sous une forme telle qu'ils pourront être confronté à des observations.

La modélisation mathématique de la dynamique des systèmes à N corps gravitationnels non collisionnels repose sur le système couplé des deux équations de Boltzmann sans collisions et de Poisson. L'équation de Boltzmann sans collisions, aussi appelée équation de Vlasov, mais originellement écrite par L.Boltzmann en 1872, fut introduite dans le cadre de l'astrophysique par J.H.Jeans en 1915. L'équation de Poisson est quant à elle due à S.-D.Poisson, et date de 1813. Un autre résultat analytique se révèle indispensable à l'étude de ces systèmes : le théorème du viriel, qui établit une condition d'équilibre que nous avons utilisée lors de nos simulations (R.Clausius, 1870).

L'étude du système Boltzmann sans collisions–Poisson permet d'obtenir plusieurs résultats concernant leurs solutions à l'équilibre, parmi lesquels le théorème de Jeans sur la dépendance de ces solutions

stationnaires aux intégrales du mouvement. Il est également possible de dériver des résultats plus précis, notamment deux solutions analytiques bien connues, qui sont le modèle de Plummer et la sphère isotherme.

En approfondissant l'étude concernant le modèle de la sphère isotherme, comme l'ont fait V.A. Antonov [4] puis T. Padmanabhan [45], nous pouvons mettre à jour une instabilité, nommée instabilité d'Antonov. L'utilisation de la méthode d'énergie nous permet ensuite d'obtenir de nouveaux critères de stabilité des solutions du système couplé Boltzmann sans collisions–Poisson. L'instabilité d'orbite radiale se trouve également révélée par cette méthode.

Afin d'étudier numériquement la question que nous nous sommes posée, nous avons utilisé une version parallèle du treecode. L'idée du treecode revient à J. Barnes et P. Hut [6], J. Barnes ayant par ailleurs écrit la version séquentielle du code dont nous nous sommes servi. D. Pfenniger a par la suite modifié et parallélisé le treecode. Nous lui avons ensuite apporté les modifications nécessaires à son utilisation dans le cadre de notre travail.

Nous avons également développé un certain nombre d'outils permettant la création de la gamme de conditions initiales que nous désirions explorer, ainsi que l'analyse des résultats fournis par la simulation. Certains de ces codes développés par nos soins sont parallèles, d'autres séquentiels.

Le calcul parallèle a de ce fait tenu un rôle important dans le travail que nous avons effectué. L'apprentissage de ses bases, puis de certaines de ses spécificités, ainsi que la mise en place et l'administration de la grappe d'ordinateurs, en collaboration avec l'ingénieur système du laboratoire, furent des étapes importantes et indispensables à la réalisation de notre projet.

La première de nos tâches fut de déterminer des valeurs acceptables du paramètre d'adoucissement du potentiel et du nombre de particules utilisé. Le paramètre d'adoucissement, ou de lissage, utilisé dans une majorité de codes particules–particules, permet de masquer la singularité du potentiel gravitationnel. Sa valeur détermine également la résolution spatiale du code. Puisque nous ne désirions pas perdre inutilement du temps de calcul en utilisant un nombre aussi grand que possible de particules, nous avons effectué une étude préliminaire heuristique afin de définir un nombre minimal de particules au-delà duquel les observables que nous souhaitions mesurer demeurent inchangées. Sachant que les temps de calcul cumulés de l'ensemble des simulations réalisées s'élèvent à six mois environ, cette étude fut salutaire.

Notre second objectif consistait à vérifier les résultats précédemment obtenus concernant la dépendance des rapports d'axes finaux au rapport du viriel initial. À des degrés divers, L. Aguilar et D. Merritt [3], K.W. Min et C.S. Choi [40] et J.K. Cannizzo et T.C. Hollister [14] ont étudié ce sujet. Cependant, en raison des nombres de particules utilisés ou de l'objectif très ciblé de ces travaux, il nous est apparu intéressant de nous pencher sur ce problème parfois considéré comme résolu et de le mettre en relation avec l'instabilité d'orbites radiales (comme l'avaient déjà fait L. Aguilar et D. Merritt).

Nous avons ensuite remarqué un phénomène à notre connaissance inconnu jusqu'ici. Les systèmes formés à partir de matériel homogène et ceux formés à partir de matériel inhomogène possèdent des rayons caractéristiques (rayon du cœur et rayons contenant 10 et 50% de leur masse) très différents. Cette ségrégation semble liée à l'instabilité d'Antonov, présentée dans les développements analytiques. Nous avons également comparé les couples potentiel–densité des systèmes à l'équilibre avec ceux des deux modèles théoriques exposés préalablement. La modélisation mathématique prédit qu'il existe de nombreuses fonctions de distribution solutions à l'équilibre du système Boltzmann sans collisions–Poisson, mais nous voulions vérifier si nos solutions ne s'approchaient pas de ces deux modèles théoriques, et en particulier du modèle de Plummer puisque celui-ci représente correctement les amas globulaires.

Enfin, nous avons étudié la dépendance entre l'énergie totale et la température de nos systèmes. L'énergie totale des systèmes isolés doit être conservée aussi précisément que possible malgré l'utilisation d'approximations dans les méthodes numériques employées. La température évolue par contre différemment selon le type de conditions initiales.

Les analyses de tous les résultats obtenus permettent de parvenir à des conclusions très intéres-

santes concernant l'importance des instabilités d'Antonov et d'orbites radiales sur les états d'équilibre ainsi que les facteurs menant à l'apparition de ces instabilités.

Enfin, certains aspects de notre travail se révèlent porteurs d'interrogations, et soulèvent des questions auxquelles seules de nouvelles simulations ou des observations pourraient répondre.

MODÉLISATION ET ÉTUDE ANALYTIQUE DES SYSTÈMES AUTO-GRAVITANTS

1 Amas globulaires et galaxies elliptiques : une tentative de définition

Les définitions des amas globulaires rencontrées dans la littérature sont nombreuses, mais s'accordent toutes sur un certain nombre de points. Ainsi, un amas globulaire est un amas d'étoiles sphérique (ou quasi-sphérique) contenant de 10^3 à 10^6 étoiles. Son rayon de marée est de l'ordre de quelques dizaines de pc, et sa masse est comprise entre 10^3 et $5 \cdot 10^6 M_\odot$. La densité centrale d'un amas globulaire (typiquement $10^4 M_\odot \text{pc}^{-3}$) est supérieure de plusieurs ordres de grandeur à celle du voisinage du Soleil (environ $0.05 M_\odot \text{pc}^{-3}$). Un amas globulaire ne contient apparemment ni gaz, ni poussière. Il ne semble de plus pas nécessaire de faire appel à la présence de matière noire pour expliquer leur dynamique.

Il apparaît donc que les amas globulaires sont des systèmes à N corps gravitationnels purs. Ils constituent dès lors un objet d'étude privilégié pour la dynamique des systèmes à N corps.

Les observations menées sur les amas de la Galaxie nous permettent d'avoir quelques données chiffrées plus précises concernant ces objets.

La Galaxie contient environ 150 amas globulaires (47 Tuc, ω Cen, M15 parmi les plus connus), qu'on observe aussi bien au voisinage du centre galactique que dans le halo.

Une étude concernant 99 amas de la Galaxie réalisée par R.E.White et S.J.Shaw [56] portant sur la forme des amas globulaires arrive aux conclusions suivantes : le rapport moyen entre leur petit axe et leur grand axe est égal à 0.93. La valeur minimale de ce rapport est environ 0.73, mais seulement 5% des amas ont un rapport d'axes inférieur à 0.8. Les amas globulaires peuvent donc être considérés comme sphériques dans la plupart des cas.

L'amas le plus brillant de la Galaxie possède une magnitude absolue $\mathcal{M}_v = -10.4$, pour un nombre d'étoiles de l'ordre de 10^6 , alors que la magnitude des moins lumineux est $\mathcal{M}_v \sim -3.0$, et leur nombre d'étoiles environ 10^3 .

L'âge des amas de la Voie Lactée varie selon la méthode utilisée pour le déterminer. Les résultats obtenus sont compris entre $\sim 13 \cdot 10^9$ et $\sim 15 \cdot 10^9$ années. Il existe donc une grande incertitude concernant leur âge absolu, environ 2 milliards d'années, mais les mesures d'âge relatif (d'un amas par rapport à un autre) sont plus précises.

Les galaxies elliptiques représentent environ 13% des galaxies observées. Elles possèdent quelques caractéristiques identiques à celles des amas globulaires. Elles ne contiennent par exemple que peu de gaz et de poussière, et sont sphériques ou proches de la sphéricité. Elles sont composées de 10^9 à 10^{13}

étoiles, leur rayon est compris entre 1 et 200 kpc et leur masse entre $10^6 M_\odot$ et $10^{13} M_\odot$.

Le rapport du grand axe sur le petit axe d'une galaxie elliptique varie de un à trois environ. E.Hubble a classifié les galaxies elliptiques selon leur ellipticité, égale à $1 - b/a$. Une galaxie de type En est ainsi une galaxie dont l'ellipticité vaut $n/10$. Les galaxies les plus elliptiques sont de type E7, les galaxies E0 étant sphériques.

Certaines caractéristiques dynamiques des galaxies elliptiques ne peuvent s'expliquer que par la présence de matière noire.

2 Modélisation des systèmes auto-gravitants

La modélisation des systèmes que nous avons décrits dans la section précédente est nécessaire à leur étude analytique, ainsi qu'à la préparation et la compréhension de leur étude numérique.

Le premier point que nous abordons concerne les échelles de temps de la dynamique des systèmes auto-gravitants. La compréhension de ces échelles de temps est indispensable à une modélisation physique pertinente de ces systèmes, et, bien entendu, à la mise en œuvre de méthodes numériques adaptées à leur dynamique.

Ensuite, nous présentons rapidement deux types de formalismes employés pour décrire les systèmes à N corps, l'un discret et l'autre continu. Nous exprimons également les principales grandeurs utilisées lors de nos études.

La troisième et la quatrième partie de cette section sont respectivement consacrées à la dérivation des équations de Poisson et de Boltzmann sans collisions. L'équation de Poisson permet de relier le potentiel gravitationnel engendré par une assemblée de particules à la densité de cette assemblée. L'équation de Boltzmann sans collisions décrit quant à elle l'évolution temporelle de la fonction de distribution d'un ensemble de particules. Le système couplé formé de ces deux équations décrit complètement la dynamique d'un système purement gravitationnel et non collisionnel. L'étude analytique des systèmes auto-gravitants auxquels nous nous intéressons repose sur ce système couplé de deux équations.

Enfin, nous présentons le théorème du viriel, dont l'application revêt un rôle important lors de nos travaux numériques. En effet, ce théorème relie la valeur moyenne de l'énergie cinétique et celle de l'énergie potentielle d'un système à l'équilibre. Ce théorème est l'outil qui nous permet d'identifier les états d'équilibre produits par nos simulations.

2.1 Temps caractéristiques

La dynamique des systèmes auto-gravitants possède plusieurs échelles de temps. Nous utiliserons principalement deux de ces échelles lors de nos travaux. La première est caractérisée par le temps dynamique t_d , qui correspond au temps mis par une particule pour traverser le système considéré. La seconde est définie par le temps de relaxation par collisions t_r , qui désigne le temps nécessaire aux collisions à deux corps pour modifier de manière significative la trajectoire d'une particule test traversant le système considéré.

La définition du temps dynamique peut être illustrée par son calcul dans le cas simple d'un système sphérique homogène de densité ρ_h , constitué de N particules. D'après le principe fondamental de la dynamique, le module v de la vitesse $\mathbf{v}$ d'une particule en orbite circulaire à une distance r du centre du système est donnée par

$$v^2 = \frac{GM(r)}{r},$$

où $M(r)$ désigne la masse totale contenue dans la sphère de rayon r centrée sur le centre du système et G la constante de la gravitation. Notre système étant homogène, nous avons

$$v = \sqrt{\frac{4\pi G \rho_h}{3}} r.$$

La période de révolution de cette particule autour du centre du système est donc égale à

$$t_0 = \frac{2\pi r}{v} = \sqrt{\frac{3\pi}{G\rho_h}}.$$

Cette période est indépendante du rayon de l'orbite circulaire de notre particule.

D'autre part, si une particule test est placée au repos à une distance r du centre de notre système, l'équation de son mouvement est

$$\frac{d^2r}{dt^2} = -\frac{GM(r)}{r^2} = -\frac{4\pi G\rho_h}{3}r = -\left(\frac{2\pi}{t_0}\right)^2 r.$$

Cette équation est celle d'un oscillateur harmonique de fréquence $2\pi/t_0$. Le temps t_d mis par cette particule pour se déplacer du bord au centre du système est donc égal à

$$t_d = \frac{t_0}{4} = \sqrt{\frac{3\pi}{16G\rho_h}}. \quad (1.1)$$

L'ordre de grandeur du temps dynamique d'un amas globulaire et d'une galaxie elliptique peut être calculé en utilisant (1.1) et en approchant ρ_h par la moyenne de la densité sur l'ensemble du système ρ_m , qui nous est donnée par

$$\rho_m = \frac{3M}{4\pi R^3},$$

M étant la masse totale du système. Nous obtenons ainsi un temps dynamique égal à $7.4 \cdot 10^5$ années pour un amas globulaire de rayon $R = 10$ pc et de masse $M = 10^6 M_\odot$, et égal à $7.4 \cdot 10^7$ années pour une galaxie elliptique de rayon $R = 10$ kpc et de masse $M = 10^{11} M_\odot$.

En pratique, il n'est pas possible de calculer ce temps puisque la densité des systèmes que nous simulons n'est pas analytique. Nous avons donc approché le temps dynamique par le temps de croisement t_c , qui est égal au temps mis par une particule de vitesse moyenne pour traverser notre système. Nous avons calculé le temps de croisement en divisant la distance moyenne séparant les particules du centre du système par la vitesse moyenne de ces mêmes particules :

$$t_c = \frac{\langle r \rangle}{\langle v \rangle},$$

où $\langle . \rangle$ désigne la moyenne sur toutes les particules et v est le module de la vitesse. $\langle r \rangle$ se révèle être un estimateur acceptable de la taille R du système, et a le mérite d'être très simple et rapide à calculer. Le calcul du temps de relaxation par collisions est un peu plus long mais assez simple (voir par exemple [9], p.187-190). Il suffit, pour le réaliser, de considérer le parcours d'une particule test de masse m et de vitesse initiale $\mathbf{v}$ traversant un système auto-gravitant constitué de N particules identiques de masse m . En faisant quelques hypothèses supplémentaires, nous trouvons que le carré de la vitesse de la particule test varie, à chaque traversée du système, de δv^2 , avec

$$\frac{\delta v^2}{v^2} \approx \frac{8}{3N} \log \frac{3N}{4\pi}.$$

Le temps de relaxation par collisions peut être défini comme le temps nécessaire pour que le rapport $\delta v^2/v^2$ soit égal à 1. Ce temps est donc

$$t_r = \frac{3N t_c}{8 \log \frac{3N}{4\pi}}.$$

Nous avons dit précédemment que le temps de croisement t_c et le temps dynamique t_d sont du même ordre de grandeur. Il est donc raisonnable de considérer que le temps de relaxation par collisions d'un

système auto-gravitant est supérieur d'un facteur $N/\log N$ à son temps dynamique.

Notons enfin que l'ordre de grandeur du temps de relaxation par collisions d'un gros amas globulaire est 10^{11} années, et celui d'une galaxie elliptique 10^{17} années. Les collisions à deux corps sont donc susceptibles de jouer un rôle dans la dynamique des amas globulaires si l'on considère l'intégralité de leur évolution dans une galaxie pendant plusieurs milliards d'années, et ce particulièrement dans leur cœur. Par contre, ces collisions n'ont aucune influence sur la dynamique des galaxies elliptiques.

2.2 Descriptions discrète et continue

Nous pouvons décrire les systèmes à N corps de deux façons distinctes. La première, la plus intuitive, est une description que nous qualifierons de discrète – chaque particule i est caractérisée par sa position $\mathbf{r}_i$ et son impulsion $\mathbf{p}_i$ – pour laquelle le système correspond à l'ensemble des corps le composant. Ainsi, le potentiel $\psi(\mathbf{r}_i)$ exercé par le système sur une particule i s'exprime comme la somme des composantes provenant de chacune des autres particules, c'est-à-dire

$$\psi(\mathbf{r}_i) = -G \sum_{\substack{j=1 \\ j \neq i}}^N \frac{m_j}{|\mathbf{r}_i - \mathbf{r}_j|} . \quad (1.2)$$

L'énergie cinétique K , l'énergie potentielle U et l'énergie totale $\mathcal{H}$ du système s'écrivent alors

$$K = \sum_{i=1}^N \frac{\mathbf{p}_i^2}{2m_i} , \quad (1.3)$$

$$U = -G \sum_{i=1}^N \sum_{j>i}^N \frac{m_i m_j}{|\mathbf{r}_i - \mathbf{r}_j|} = \frac{1}{2} \sum_{i=1}^N m_i \psi(\mathbf{r}_i) , \quad (1.4)$$

$$\mathcal{H} = K + U = \sum_{i=1}^N \frac{\mathbf{p}_i^2}{2m_i} - G \sum_{i=1}^N \sum_{j>i}^N \frac{m_i m_j}{|\mathbf{r}_i - \mathbf{r}_j|} . \quad (1.5)$$

La deuxième formulation provient de la mécanique statistique et repose sur l'utilisation d'une fonction de distribution de $6N + 1$ variables $f(\mathbf{r}_1, \mathbf{r}_2, \dots, \mathbf{p}_1, \mathbf{p}_2, \dots, \mathbf{p}_N, t)$. Cette fonction est définie de telle manière que

$$f d^3\mathbf{r}_1 d^3\mathbf{p}_1 \dots d^3\mathbf{r}_N d^3\mathbf{p}_N = f d\Gamma_1 \dots d\Gamma_N$$

représente la probabilité pour qu'à l'instant t la particule 1 se trouve dans le volume de l'espace des phases $d\Gamma_1$ centré autour de $(\mathbf{r}_1, \mathbf{p}_1)$, la particule 2 dans le volume de l'espace des phases $d\Gamma_2$ centré autour de $(\mathbf{r}_2, \mathbf{p}_2)$, et ainsi de suite. f est une densité de probabilité, donc non négative et de norme égale à 1 (voir [13], p.107 et suivantes).

Nous verrons par la suite (voir section 2.4) que nous pouvons, sous certaines hypothèses, décrire un système par une fonction de distribution de la forme $f(\mathbf{r}, \mathbf{p}, t)$ à partir de laquelle il est possible de définir la densité $\rho(\mathbf{r}, t)$ du système, le potentiel gravitationnel $\psi(\mathbf{r}, t)$ au point $\mathbf{r}$, ainsi que les énergies cinétique, potentielle et totale du système par

$$\rho(\mathbf{r}, t) = \int m f(\mathbf{r}, \mathbf{p}, t) d\mathbf{p} , \quad (1.6)$$

$$\psi(\mathbf{r}, t) = -G \int \frac{\rho(\mathbf{r}', t)}{|\mathbf{r} - \mathbf{r}'|} d^3\mathbf{r}' , \quad (1.7)$$

$$K(t) = \frac{1}{2} \int \frac{\mathbf{p}^2}{m} f(\mathbf{r}, \mathbf{p}, t) d\Gamma , \quad (1.8)$$

$$U(t) = -\frac{G}{2} \int \int mm' \frac{f(\mathbf{r}, \mathbf{p}, t) f(\mathbf{r}', \mathbf{p}', t)}{|\mathbf{r} - \mathbf{r}'|} d\mathbf{\Gamma} d\mathbf{\Gamma}' \quad (1.9)$$

$$\mathcal{H}(t) = K(t) + U(t) = \frac{1}{2} \int \frac{\mathbf{p}^2}{m} f(\mathbf{r}, \mathbf{p}, t) d\mathbf{\Gamma} - \frac{G}{2} \int \int mm' \frac{f(\mathbf{r}, \mathbf{p}, t) f(\mathbf{r}', \mathbf{p}', t)}{|\mathbf{r} - \mathbf{r}'|} d\mathbf{\Gamma} d\mathbf{\Gamma}' . \quad (1.10)$$

Remarque : La fonction de distribution sera par la suite parfois exprimée comme une fonction des positions et des vitesses (dans la section 3.3), et non des positions et des impulsions. Il convient alors d'écrire les différentes quantités définies ci-dessus en fonction des vitesses, ce qui ne pose aucune difficulté.

Nous utiliserons par la suite ces deux types de formulations en fonction des aspects théoriques que nous chercherons à mettre en valeur.

2.3 L'équation de Poisson

Considérons un système stellaire tel qu'un amas globulaire ou une galaxie elliptique. Un tel système est composé de 10^5 à 10^{10} étoiles, et ne contient que peu de gaz interstellaire. Sur des échelles de taille telles que celle des amas globulaires (environ 10 pc de rayon) ou des galaxies elliptiques (10^4 pc), et si l'on néglige le rôle du gaz interstellaire, la seule interaction intervenant dans la dynamique de notre système est l'interaction gravitationnelle.

Nous considérons donc un système $\Omega \subset \mathbb{R}^3$ dont la masse est distribuée selon une certaine densité

$$\rho = \begin{cases} \rho(\mathbf{r}') & \text{si } \mathbf{r}' \in \Omega \\ 0 & \text{si } \mathbf{r}' \notin \Omega . \end{cases}$$

La loi de la gravitation de Newton nous indique alors que la force $\mathbf{F}(\mathbf{r})$ que crée le système en tout point $\mathbf{r} \in \mathbb{R}^3$, dans l'hypothèse d'une distribution continue de matière, est

$$\mathbf{F}(\mathbf{r}) = G \int \frac{\mathbf{r} - \mathbf{r}'}{|\mathbf{r} - \mathbf{r}'|^3} \rho(\mathbf{r}') d^3\mathbf{r}' . \quad (1.11)$$

Considérons maintenant le potentiel gravitationnel engendré par cette distribution de particules, dont nous rappelons la forme :

$$\psi(\mathbf{r}) = -G \int \frac{\rho(\mathbf{r}')}{|\mathbf{r} - \mathbf{r}'|} d^3\mathbf{r}' , \quad (1.12)$$

Ce potentiel peut être écrit sous la forme d'un produit de convolution. Le produit de convolution est défini par

$$(A * B)(x) = \int A(x - x') B(x) dx' ,$$

et possède les propriétés suivantes :

$$A(x) * B(x) = B(x) * A(x) ,$$

et

$$\frac{\partial}{\partial x} (A(x) * B(x)) = \left(\frac{\partial}{\partial x} A(x) \right) * B(x) = A(x) * \left(\frac{\partial}{\partial x} B(x) \right) . \quad (1.13)$$

Nous avons donc

$$\psi(\mathbf{r}) = -G\rho(\mathbf{r}) * \frac{1}{|\mathbf{r}|} . \quad (1.14)$$

Le produit de convolution (1.14) a un sens car la densité ρ est à support borné (voir [51] pour plus de précisions concernant le produit de convolution).

D'autre part, on peut vérifier assez facilement, grâce aux distributions, que $1/|\mathbf{r}| = 1/r$ est, à une constante près, la fonction de Green du Laplacien. En effet, si l'on considère une fonction test $u(r)$, C^∞ à support compact¹, on a, en dérivant au sens des distributions,

$$\left\langle \Delta \left(\frac{1}{r} \right), u(r) \right\rangle = \left\langle \frac{1}{r}, \Delta(u(r)) \right\rangle .$$

De plus, puisque nous ne considérons que des fonctions de r , seule partie radiale du Laplacien doit être prise en compte. Nous avons donc

$$\begin{aligned} \left\langle \frac{1}{r}, \Delta(u(r)) \right\rangle &= \int \frac{1}{r} \Delta(u(r)) d^3\mathbf{r} \\ &= 4\pi \int_0^\infty \frac{1}{r} \frac{1}{r^2} \frac{d}{dr} \left(r^2 \frac{du}{dr} \right) r^2 dr . \end{aligned}$$

En intégrant par parties, on obtient

$$\left\langle \frac{1}{r}, \Delta(u(r)) \right\rangle = 4\pi \left[r \frac{du}{dr} \right]_0^\infty - 4\pi \int_0^\infty \frac{d}{dr} \left(\frac{1}{r} \right) r^2 \frac{du}{dr} dr .$$

Le premier terme s'annule : rdu/dr s'annule à l'infini parce que u est à support compact, et s'annule en 0 du fait de r . Ce qui nous donne finalement

$$\left\langle \frac{1}{r}, \Delta u \right\rangle = 4\pi \int_0^\infty \frac{du}{dr} dr = 4\pi [u]_0^\infty = -4\pi u(0) = \langle -4\pi \delta_D, u \rangle ,$$

où δ_D est la masse de Dirac en 0. On a donc $\Delta(1/r) = -4\pi \delta_D$ au sens des distributions. Toujours au sens des distributions, nous pouvons écrire

$$\Delta\psi = -G\Delta \left(\rho * \frac{1}{|\mathbf{r}|} \right) = -G\rho * \left(\Delta \frac{1}{r} \right) = 4\pi G\rho * \delta_D .$$

δ_D étant l'élément neutre du produit de convolution, ceci nous mène à l'équation de Poisson

$$\Delta\psi(\mathbf{r}, t) = 4\pi G\rho(\mathbf{r}, t) . \quad (1.15)$$

La connaissance de la densité de matière et l'application de cette équation nous permettent donc d'obtenir le potentiel gravitationnel et la force résultant de ce potentiel.

2.4 L'équation de Boltzmann sans collisions

Considérons un système de N étoiles, représentées chacune par une masse ponctuelle m . Ces étoiles interagissent seulement par la gravitation newtonienne.

Le système est décrit par sa fonction de distribution $f(\mathbf{r}_1, \dots, \mathbf{r}_N, \mathbf{p}_1, \dots, \mathbf{p}_N, t)$. La fonction f vérifie l'équation de Liouville

$$\frac{\partial f}{\partial t} + \sum_{i=1}^N \frac{\partial}{\partial \mathbf{r}_i} \left(\frac{d\mathbf{r}_i}{dt} f \right) + \sum_{i=1}^N \frac{\partial}{\partial \mathbf{p}_i} \left(\frac{d\mathbf{p}_i}{dt} f \right) = 0 \quad (1.16)$$

qui exprime la conservation du nombre de particules dans le système au cours de son évolution. Les dérivées de $\mathbf{r}_i$ et $\mathbf{p}_i$ par rapport au temps nous sont données par les équations du mouvement du système, c'est-à-dire

$$\frac{d\mathbf{r}_i}{dt} = \frac{\mathbf{p}_i}{m} \quad (1.17)$$

¹Une fonction C^∞ à support compact de $\mathbb{R}^3$ est non nulle sur un domaine fermé et borné et est indéfiniment dérivable. De plus, cette fonction et toutes ses dérivées sont nulles au bord.

et

$$\frac{d\mathbf{p}_i}{dt} = \sum_{j \neq i} \frac{\partial}{\partial \mathbf{r}_i} \frac{Gm^2}{|\mathbf{r}_i - \mathbf{r}_j|} . \quad (1.18)$$

Ces équations du mouvement dérivent du Hamiltonien $\mathcal{H}$ du système (équation (1.5)) et des équations canoniques de Hamilton

$$\begin{cases} \frac{d\mathbf{r}_i}{dt} = \frac{\partial \mathcal{H}}{\partial \mathbf{p}_i} , \\ \frac{d\mathbf{p}_i}{dt} = -\frac{\partial \mathcal{H}}{\partial \mathbf{r}_i} . \end{cases} \quad (1.19)$$

Si nous injectons (1.19) dans l'équation de Liouville (1.16), nous obtenons une nouvelle équation d'évolution

$$\begin{aligned} 0 &= \frac{\partial f}{\partial t} + \sum_{i=1}^N \frac{\partial}{\partial \mathbf{r}_i} \left(\frac{\partial \mathcal{H}}{\partial \mathbf{p}_i} f \right) - \sum_{i=1}^N \frac{\partial}{\partial \mathbf{p}_i} \left(\frac{\partial \mathcal{H}}{\partial \mathbf{r}_i} f \right) \\ &= \frac{\partial f}{\partial t} + \sum_{i=1}^N \left[\frac{\partial \mathcal{H}}{\partial \mathbf{p}_i} \frac{\partial f}{\partial \mathbf{r}_i} + f \frac{\partial}{\partial \mathbf{r}_i} \left(\frac{\partial \mathcal{H}}{\partial \mathbf{p}_i} \right) - \frac{\partial \mathcal{H}}{\partial \mathbf{r}_i} \frac{\partial f}{\partial \mathbf{p}_i} - f \frac{\partial}{\partial \mathbf{p}_i} \left(\frac{\partial \mathcal{H}}{\partial \mathbf{r}_i} \right) \right] \\ &= \frac{\partial f}{\partial t} + \sum_{i=1}^N \frac{d\mathbf{r}_i}{dt} \frac{\partial f}{\partial \mathbf{r}_i} + \sum_{i=1}^N \frac{d\mathbf{p}_i}{dt} \frac{\partial f}{\partial \mathbf{p}_i} = \frac{df}{dt} . \end{aligned} \quad (1.20)$$

Revenons maintenant à la fonction de distribution f . Il est possible de définir une fonction de distribution f_1 de 1 particule

$$f_1(\mathbf{r}_1, \mathbf{p}_1, t) = \int f(\mathbf{r}_1, \mathbf{p}_1, \dots, \mathbf{r}_N, \mathbf{p}_N, t) d\mathbf{\Gamma}_2 \dots d\mathbf{\Gamma}_N , \quad (1.21)$$

et une fonction de distribution $f_{1,2}$ de 2 particules

$$f_{1,2}(\mathbf{r}_1, \mathbf{p}_1, \mathbf{r}_2, \mathbf{p}_2, t) = \int f(\mathbf{r}_1, \mathbf{p}_1, \dots, \mathbf{r}_N, \mathbf{p}_N, t) d\mathbf{\Gamma}_3 \dots d\mathbf{\Gamma}_N . \quad (1.22)$$

$f_1 d\mathbf{\Gamma}_1$ représente la probabilité de trouver, à l'instant t , la particule 1 dans le volume de l'espace des phases $d\mathbf{\Gamma}_1$ centré autour de $(\mathbf{r}_1, \mathbf{p}_1)$, quelles que soient les positions et les impulsions des autres particules. $f_{1,2} d\mathbf{\Gamma}_1 d\mathbf{\Gamma}_2$ s'interprète de la même façon, mais pour les particules 1 et 2.

Nous pouvons obtenir une équation d'évolution pour f_1 en intégrant l'équation d'évolution (1.20) sur $d\mathbf{\Gamma}_2 \dots d\mathbf{\Gamma}_N$

$$\int \left[\frac{\partial f}{\partial t} + \sum_{i=1}^N \frac{d\mathbf{r}_i}{dt} \frac{\partial f}{\partial \mathbf{r}_i} + \sum_{i=1}^N \frac{d\mathbf{p}_i}{dt} \frac{\partial f}{\partial \mathbf{p}_i} \right] d\mathbf{\Gamma}_{2,N} = 0 , \quad (1.23)$$

où $d\mathbf{\Gamma}_{2,N}$ est le produit $d\mathbf{\Gamma}_2 \dots d\mathbf{\Gamma}_N$. Nous pouvons inverser l'ordre des intégrations et de la dérivation en temps puisque le domaine d'intégration est indépendant du temps. Le premier terme est donc égal à

$$\int \frac{\partial f}{\partial t} d\mathbf{\Gamma}_{2,N} = \frac{\partial}{\partial t} \int f d\mathbf{\Gamma}_{2,N} = \frac{\partial f_1}{\partial t} .$$

Puisque f est une densité de probabilité, *i.e.* $f \geq 0$ et $\int f d\mathbf{\Gamma}_{1,N} = 1$, nous savons que

$$\lim_{|\mathbf{r}_i| \rightarrow \infty} f = 0 \quad \text{et} \quad \lim_{|\mathbf{p}_i| \rightarrow \infty} f = 0 .$$

En intégrant le deuxième terme de (1.23) par parties et en utilisant notre première hypothèse, nous obtenons

$$\frac{\partial}{\partial \mathbf{r}_1} \int f \frac{\partial \mathcal{H}}{\partial \mathbf{p}_1} d\mathbf{\Gamma}_{2,N} ,$$

et finalement, en utilisant (1.5)

$$\frac{\partial}{\partial \mathbf{r}_1} \int f \frac{\partial \mathcal{H}}{\partial \mathbf{p}_1} d\mathbf{\Gamma}_{2,N} = \frac{\mathbf{p}_1}{m} \frac{\partial f_1}{\partial \mathbf{r}_1} .$$

En intégrant le troisième terme (1.23) par parties et en utilisant notre deuxième hypothèse, nous obtenons

$$-\frac{\partial}{\partial \mathbf{p}_1} \int f \frac{\partial \mathcal{H}}{\partial \mathbf{r}_1} d\mathbf{\Gamma}_{2,N} .$$

Si nous explicitons $\partial \mathcal{H} / \partial \mathbf{r}_1$, nous avons

$$\frac{\partial \mathcal{H}}{\partial \mathbf{r}_1} = - \sum_{j \neq 1} \frac{\partial}{\partial \mathbf{r}_1} \left(\frac{Gm^2}{|\mathbf{r}_1 - \mathbf{r}_j|} \right) = \sum_{j \neq 1} \frac{\partial}{\partial \mathbf{r}_1} U_2(1, j) \quad \text{où} \quad U_2(i, j) = - \frac{Gm^2}{|\mathbf{r}_i - \mathbf{r}_j|} ,$$

ce qui nous donne

$$\begin{aligned} - \frac{\partial}{\partial \mathbf{p}_1} \int f \frac{\partial \mathcal{H}}{\partial \mathbf{r}_1} d\mathbf{\Gamma}_{2,N} &= - \frac{\partial}{\partial \mathbf{p}_1} \int f \sum_{j \neq 1} \frac{\partial}{\partial \mathbf{r}_1} U_2(1, j) d\mathbf{\Gamma}_{2,N} \\ &= - \frac{\partial}{\partial \mathbf{p}_1} \sum_{j \neq 1} \int \frac{\partial U_2(1, j)}{\partial \mathbf{r}_1} \left[\int f d\mathbf{\Gamma}_{2,j-1} d\mathbf{\Gamma}_{j+1,N} \right] d\mathbf{\Gamma}_j \\ &= - \frac{\partial}{\partial \mathbf{p}_1} \sum_{j \neq 1} \int \frac{\partial U_2(1, j)}{\partial \mathbf{r}_1} f_{1,j} d\mathbf{\Gamma}_j . \end{aligned} \tag{1.24}$$

Nous devons maintenant faire une hypothèse forte sur la nature de nos particules afin de pouvoir simplifier cette expression : nous considérons que nos étoiles sont identiques et indiscernables. Le terme (1.24) peut alors se réduire au produit

$$-(N-1) \frac{\partial}{\partial \mathbf{p}_1} \int \frac{\partial U_2(1, 2)}{\partial \mathbf{r}_1} f_{1,2} d\mathbf{\Gamma}_2 .$$

L'équation d'évolution (1.23) de f_1 devient donc

$$\frac{\partial f_1}{\partial t} + \frac{\mathbf{p}_1}{m} \frac{\partial f_1}{\partial \mathbf{r}_1} = (N-1) \int \frac{\partial U_2(1, 2)}{\partial \mathbf{r}_1} \frac{\partial f_{1,2}}{\partial \mathbf{p}_1} d\mathbf{\Gamma}_2 . \tag{1.25}$$

L'équation d'évolution d'une fonction de distribution d'une particule dépend ainsi de la fonction de distribution de deux particules. Il est possible de dériver de la même manière l'équation d'évolution d'une fonction de distribution de n particules identiques et indiscernables, et de voir qu'elle dépend de la fonction de distribution de $(n+1)$ particules. Cette cascade d'équations, qui part de la fonction à N particules et mène jusqu'à la fonction à une particule, est connue sous le nom de hiérarchie BBGKY (pour Born, Bogoliubov, Green, Kirkwood et Yvon, voir par exemple [17], p.9, pour plus de détails). Supposons que $f_{1,2}$ soit de la forme

$$f_{1,2} = f_1 f_2 + f_{int}(\mathbf{r}_1, \mathbf{p}_1, \mathbf{r}_2, \mathbf{p}_2, t) ,$$

où f_{int} représente les interactions entre les deux particules, alors que la fonction de distribution de deux particules serait égale au produit $f_1 f_2$ s'il n'y avait aucune interaction.

Comme les particules sont identiques indiscernables, il nous est possible de définir une fonction $f(\mathbf{r}, \mathbf{p}, t)$, décrivant le système tout entier, par

$$f(\mathbf{r}, \mathbf{p}, t) = N f_1(\mathbf{r}, \mathbf{p}, t) .$$

Cette fonction a donc pour norme N . Si nous réécrivons l'équation (1.25) pour f , nous avons

$$\frac{1}{N} \left(\frac{\partial f}{\partial t} + \frac{\mathbf{p}}{m} \frac{\partial f}{\partial \mathbf{r}} \right) = (N-1) \int \frac{\partial U_2(1,2)}{\partial \mathbf{r}} \frac{\partial}{\partial \mathbf{p}} \left(\frac{f(\mathbf{r}, \mathbf{p}, t)}{N} \frac{f(\mathbf{r}', \mathbf{p}', t)}{N} + f_{int} \right) d\mathbf{\Gamma}' ,$$

soit

$$\begin{aligned} \frac{\partial f}{\partial t} + \frac{\mathbf{p}}{m} \frac{\partial f}{\partial \mathbf{r}} &= \frac{(N-1)}{N} \frac{\partial f(\mathbf{r}, \mathbf{p}, t)}{\partial \mathbf{p}} \frac{\partial}{\partial \mathbf{r}} \int U_2(1,2) f(\mathbf{r}', \mathbf{p}', t) d\mathbf{\Gamma}' \\ &+ N(N-1) \int \frac{\partial U_2(1,2)}{\partial \mathbf{r}} \frac{\partial f_{int}}{\partial \mathbf{p}} d\mathbf{\Gamma}' . \end{aligned} \quad (1.26)$$

Or, d'après la définition de $U_2(1,2)$ et celle de la densité de masse ρ (équation (1.6)), on sait que

$$\int U_2(1,2) f(\mathbf{r}', \mathbf{p}', t) d\mathbf{\Gamma}' = -Gm \int \frac{\rho(\mathbf{r}', t)}{|\mathbf{r} - \mathbf{r}'|} d^3\mathbf{r}' ,$$

soit, d'après l'équation (1.7)

$$\int U_2(1,2) f(\mathbf{r}', \mathbf{p}', t) d\mathbf{\Gamma}' = m\psi(\mathbf{r}, t) .$$

En nous plaçant dans le cas limite des très grands N , l'équation (1.26) devient

$$\frac{\partial f}{\partial t} + \frac{\mathbf{p}}{m} \frac{\partial f}{\partial \mathbf{r}} - m \frac{\partial f}{\partial \mathbf{p}} \frac{\partial \psi}{\partial \mathbf{r}} = GN^2 m^2 \int \frac{\partial f_{int}(\mathbf{r}, \mathbf{p}, \mathbf{r}', \mathbf{p}', t)}{\partial \mathbf{p}} \frac{(\mathbf{r} - \mathbf{r}')}{|\mathbf{r} - \mathbf{r}'|^3} d\mathbf{\Gamma}' . \quad (1.27)$$

Dans le cas où les interactions à deux particules, représentées par le terme de droite de (1.27), sont négligeables, cette équation devient l'équation de Boltzmann sans collisions (ou BSC)

$$\frac{\partial f}{\partial t} + \frac{\mathbf{p}}{m} \frac{\partial f}{\partial \mathbf{r}} - m \frac{\partial f}{\partial \mathbf{p}} \frac{\partial \psi}{\partial \mathbf{r}} = 0 . \quad (1.28)$$

L'équation de Boltzmann sans collisions décrit l'évolution temporelle de la fonction de distribution du assemblée de particules identiques évoluant dans un potentiel.

Nous avons vu dans la section 2.1 que le temps de relaxation par collisions est environ égal à $t_d N / \log N$. Un amas globulaire contient typiquement 10^6 particules, et une galaxie elliptique 10^{10} particules. Ceci implique que si l'on considère l'évolution d'un tel système pendant une durée ne dépassant pas un nombre réduit² de temps dynamiques, les collisions peuvent être négligées puisqu'elles n'influent pas de façon notable sur la dynamique du système.

2.5 Théorème du viriel

Considérons un système de N particules de masses m_i situées aux positions $\mathbf{r}_i$. L'énergie cinétique K de ce système est égale à

$$K(\dot{\mathbf{r}}_1, \dots, \dot{\mathbf{r}}_N) = \frac{1}{2} \sum_{i=1}^N m_i \dot{\mathbf{r}}_i^2 , \quad (1.29)$$

²Il est délicat de déterminer une limite acceptable pour ce "nombre réduit", mais il est certain que le système peut être considéré comme non-collisionnel sur une centaine de temps dynamiques.

où $\dot{\mathbf{r}}$ est la dérivée totale par rapport au temps de $\mathbf{r}$.

K est une fonction homogène de degré 2. En effet, pour tout réel λ on a

$$K(\lambda \dot{\mathbf{r}}_1, \dots, \lambda \dot{\mathbf{r}}_N) = \lambda^2 K(\dot{\mathbf{r}}_1, \dots, \dot{\mathbf{r}}_N) .$$

D'après le théorème d'Euler des fonction homogènes, nous savons que

$$\sum_{i=1}^N \dot{\mathbf{r}}_i \text{grad}_{\dot{\mathbf{r}}_i} K = 2K . \quad (1.30)$$

De plus, si l'énergie potentielle de notre système ne dépend que des coordonnées généralisées, ce qui est le cas pour un système purement gravitationnel, nous savons également, d'après la mécanique hamiltonienne, que les impulsions généralisées sont définies par :

$$\mathbf{p}_i = \frac{\partial K}{\partial \dot{\mathbf{r}}_i} .$$

La relation (1.30) s'écrit donc

$$\sum_{i=1}^N \dot{\mathbf{r}}_i \mathbf{p}_i = 2K \Leftrightarrow \frac{d}{dt} \left(\sum_{i=1}^N \mathbf{r}_i \mathbf{p}_i \right) - \sum_{i=1}^N \dot{\mathbf{p}}_i \mathbf{r}_i = 2K . \quad (1.31)$$

La valeur moyenne temporelle d'une observable A de notre système est définie par

$$\bar{A} = \lim_{\tau \rightarrow \infty} \frac{1}{\tau} \int_0^\tau A dt .$$

La valeur moyenne temporelle de K est donc donnée par

$$2\bar{K} = \overline{\frac{d}{dt} \left(\sum_{i=1}^N \mathbf{r}_i \mathbf{p}_i \right)} - \overline{\sum_{i=1}^N \dot{\mathbf{p}}_i \mathbf{r}_i}$$

La valeur moyenne temporelle de la dérivée d'une fonction du temps bornée est nulle. Donc, si nous faisons l'hypothèse, très raisonnable, que les vitesses et les impulsions de notre système sont bornées, nous obtenons

$$2\bar{K} = - \overline{\sum_{i=1}^N \dot{\mathbf{p}}_i \mathbf{r}_i} . \quad (1.32)$$

Si nous restons dans le cadre d'un système dont l'énergie potentielle ne dépend que des coordonnées généralisées, nous savons également que

$$\dot{\mathbf{p}}_i = - \frac{\partial U}{\partial \mathbf{r}_i}$$

d'après la deuxième équation canonique de Hamilton (1.19). L'équation (1.32) s'écrit donc

$$2\bar{K} = \overline{\sum_{i=1}^N \mathbf{r}_i \text{grad}_{\mathbf{r}_i} U} . \quad (1.33)$$

Considérons maintenant un système purement gravitationnel. Nous avons

$$U(\mathbf{r}_1, \dots, \mathbf{r}_N) = -G \sum_{i=1}^N \sum_{j>i}^N \frac{m_i m_j}{|\mathbf{r}_i - \mathbf{r}_j|} . \quad (1.34)$$

Nous pouvons remarquer que U est une fonction homogène de degré -1 des coordonnées. Le théorème d'Euler nous donne donc

$$-U = \sum_{i=1}^N \mathbf{r}_i \frac{\partial U}{\partial \mathbf{r}_i} ,$$

la relation (1.33) donne donc

$$2\bar{K} + \bar{U} = 0 ,$$

ce qui constitue le théorème du viriel dans le cadre d'un système gravitationnel. Ce théorème donne une relation entre les valeurs moyennes temporelles des énergies cinétique et potentielle d'un système mécanique.

Ce théorème peut également être dérivé dans le cadre de la mécanique statistique (voir [9], p.211-213). On obtient alors le théorème du viriel tensoriel

$$\frac{1}{2}m \frac{d^2 \mathbf{I}_{jk}}{dt^2} = \mathbf{U}_{jk} + 2\mathbf{C}_{jk} + \mathbf{K}_{jk} , \quad (1.35)$$

$\mathbf{I}$ est le tenseur d'inertie

$$\mathbf{I}_{jk} = \int \nu(\mathbf{r}) r_j r_k d^3\mathbf{r} ,$$

où

$$\nu(\mathbf{r}) = \int f(\mathbf{r}, \mathbf{p}) d^3\mathbf{p}$$

est la densité de nombre du système³, $\mathbf{U}$ le tenseur potentiel

$$\mathbf{U}_{jk} = - \int \rho(\mathbf{r}) r_k \frac{\partial \psi}{\partial r_j} d^3\mathbf{r}$$

et $\mathbf{C}$ et $\mathbf{K}$ sont deux tenseurs symétriques définis respectivement par

$$\mathbf{C}_{jk} = \frac{1}{2m} \int \rho \langle p_j \rangle \langle p_k \rangle d^3\mathbf{r}$$

et

$$\mathbf{K}_{jk} = \frac{1}{m} \int \rho \varsigma_{jk}^2 d^3\mathbf{r} ,$$

avec

$$\varsigma_{ij}^2 = \langle p_i p_j \rangle - \langle p_i \rangle \langle p_j \rangle$$

et

$$\langle p_i \rangle = \frac{1}{\rho} \int p_i f(\mathbf{p}, \mathbf{r}) d^3\mathbf{p} .$$

En prenant la trace de l'équation (1.35), on obtient

$$\begin{aligned} \text{Tr}(2\mathbf{C}_{jk} + \mathbf{K}_{jk}) &= 2K , \\ \text{Tr}(\mathbf{U}_{jk}) &= U , \end{aligned}$$

d'où

$$\frac{1}{2}m \frac{d^2}{dt^2} \text{Tr}(\mathbf{I}_{jk}) = U + 2K .$$

Pour un système à l'équilibre, le terme de gauche est identiquement nul, ce qui nous redonne le résultat de la mécanique classique. Un système à l'équilibre est donc tel que son rapport du viriel $2K/U$ est égal à -1.

³ $\nu(\mathbf{r}) = \rho(\mathbf{r})/m$ pour un système où toutes les particules ont une masse égale à m .

3 Solutions analytiques du système couplé Boltzmann sans collisions – Poisson et étude de stabilité

Nous savons que la dynamique des systèmes purement gravitationnels, non collisionnels et formés de N particules identiques est régie par le système d'équations BSC – Poisson

$$\begin{cases} \frac{\partial f}{\partial t} + \frac{\mathbf{p}}{m} \frac{\partial f}{\partial \mathbf{r}} - m \frac{\partial f}{\partial \mathbf{p}} \frac{\partial \psi}{\partial \mathbf{r}} = 0, \\ \Delta \psi(\mathbf{r}, t) = 4\pi G \rho(\mathbf{r}, t). \end{cases} \quad (1.36)$$

Notre travail numérique a pour objectif l'étude des états d'équilibre de tels systèmes. Il est donc naturel d'étudier au préalable les solutions stationnaires de BSC – Poisson.

Il existe des résultats très généraux, que nous présentons dans la section suivante, concernant ces solutions stationnaires.

Nous voulons également mettre l'accent sur deux solutions analytiques particulières BSC – Poisson, les polytropes et la sphère isotherme.

Enfin, une étude de la stabilité de ces solutions analytiques nous paraît nécessaire, puisqu'il semble que les objets que nous voulons modéliser sont eux-mêmes stables. Les observations montrent en effet que ces objets – amas globulaires et galaxies elliptiques – présentent des profils de densité et des distributions de vitesses semblables et réguliers (voir par exemple [53]), ce que nous ne pourrions observer si ces objets étaient majoritairement instables.

3.1 Premiers résultats sur les solutions de Boltzmann sans collisions

Afin de pouvoir énoncer un premier résultat important concernant les solutions de BSC, il est nécessaire d'introduire la notion d'intégrale du mouvement dans un potentiel $\psi(\mathbf{r})$. Une fonction $I_m(\mathbf{r}, \mathbf{p})$ est une intégrale du mouvement si et seulement si

$$\frac{d}{dt} I_m(\mathbf{r}, \mathbf{p}) = 0 \quad (1.37)$$

le long de chaque orbite dans l'espace des phases, l'orbite étant l'ensemble des coordonnées de l'espace des phases occupées au cours du temps : $\{\mathbf{r}(t), \mathbf{p}(t) / t \geq 0\}$. Nous pouvons exprimer la dérivée par rapport au temps de l'équation (1.37), ce qui nous donne

$$\frac{dI_m}{dt} = \frac{\partial I_m}{\partial \mathbf{r}} \cdot \frac{d\mathbf{r}}{dt} + \frac{\partial I_m}{\partial \mathbf{p}} \cdot \frac{d\mathbf{p}}{dt} = 0,$$

soit

$$\frac{\partial I_m}{\partial \mathbf{r}} \cdot \mathbf{p} - \frac{\partial I_m}{\partial \mathbf{p}} \cdot \frac{\partial \psi}{\partial \mathbf{r}} = 0.$$

Cette équation est identique à l'équation de Boltzmann sans collision pour une fonction ne dépendant pas explicitement du temps. Cette constatation amène directement le théorème suivant

THÉORÈME 1 *Toute solution stationnaire de l'équation de Boltzmann sans collisions ne dépend des coordonnées de l'espace des phases qu'à travers les intégrales du mouvement dans le potentiel gravitationnel. Réciproquement, toute fonction des intégrales du mouvement est une solution stationnaire de l'équation de Boltzmann sans collisions.*

Ce théorème, connu sous le nom de premier théorème de Jeans [24], revêt une grande importance pour l'étude des états d'équilibre des systèmes auto-gravitants. Cependant, la première partie de

ce théorème n'est pas exploitable, puisqu'il n'est généralement pas possible de déterminer toutes les intégrales du mouvement d'un système. Mais la deuxième partie de ce théorème nous permet d'étudier certaines classes de fonctions de distributions solutions stationnaires de BSC.

Considérons ainsi l'énergie par unité de masse $E_m = v^2/2 + \psi(\mathbf{r})$, qui est une intégrale du mouvement. D'après le premier théorème de Jeans, toute fonction $f(E_m)$ est une solution stationnaire de BSC. Il est intéressant de s'intéresser aux caractéristiques que peuvent posséder les systèmes décrits par des fonctions de distribution ne dépendant que de E_m . Cette question, beaucoup plus délicate qu'elle ne le paraît, a trouvé une réponse grâce à l'application du théorème de Gidas-Ni-Nirenberg⁴ :

THÉORÈME 2 *Si u est une fonction C^2 de $\mathbb{R}^3$ dans $\mathbb{R}$ telle que*

$$\forall \mathbf{r} \in \mathbb{R}^3 \quad u(\mathbf{r}) < 0 \quad \text{et} \quad \lim_{\mathbf{r} \rightarrow \infty} u(\mathbf{r}) = 0$$

et s'il existe une fonction h continue de $\mathbb{R}$ dans $\mathbb{R}$ telle que

$$\Delta u = h(u) \quad ,$$

alors pour un choix convenable de l'origine, u est une fonction symétrique croissante de $\mathbf{r}$, c'est-à-dire

$$\forall \mathbf{r} \in \mathbb{R}^3 \quad u(\mathbf{r}) = u(|\mathbf{r}|) \quad \text{et} \quad \frac{\partial u}{\partial |\mathbf{r}|} > 0 \quad .$$

La densité d'un système dont la fonction de distribution ne dépend que de E_m est donnée par

$$\rho = \int f(E_m) d^3\mathbf{v} = \int f\left(\frac{v^2}{2} + \psi\right) d^3\mathbf{v} = \rho(\psi) \quad ,$$

et ne dépend donc que du potentiel ψ . En considérant l'équation de Poisson (1.15) que l'on peut ainsi réécrire de la manière suivante

$$\Delta\psi = 4\pi G\rho(\psi) \quad , \quad (1.38)$$

et sachant que $\psi(\mathbf{r}) < 0$ et que $\lim_{\mathbf{r} \rightarrow \infty} \psi(\mathbf{r}) = 0$ puisque notre système est isolé, nous pouvons donc en déduire, à la condition que $\psi \in C^2(\mathbb{R}^3)$, que

$$\psi(\mathbf{r}) = \psi(r) \quad \text{et} \quad \frac{\partial \psi}{\partial r} > 0 \quad .$$

Une fonction de distribution $f(E_m)$ décrit donc un système à symétrie sphérique. Ces systèmes sont également isotropes dans l'espace des vitesses. En effet, nous avons

$$\forall \nu \in \{x, y, z\} \quad , \quad \bar{v}_\nu^2 = \int v_\nu^2 f(E_m) d^3\mathbf{v} = \int v_\nu^2 f((v_x^2 + v_y^2 + v_z^2)/2 + \psi) d^3\mathbf{v} \quad , \quad (1.39)$$

avec $\psi = \psi(\mathbf{r})$. Il vient donc immédiatement $\bar{v}_x^2 = \bar{v}_y^2 = \bar{v}_z^2$.

Nous pouvons nous demander également quelles sont les caractéristiques d'un système dont la fonction de distribution dépend seulement de l'énergie par unité de masse E_m et du carré du moment cinétique J^2 , également intégrale du mouvement. Une extension du théorème de Gidas-Ni-Nirenberg permet de montrer qu'un tel système est à symétrie sphérique si sa densité est décroissante en fonction du rayon. Les systèmes de ce type sont de plus anisotropes, puisque $J^2 = r^2(v_\varphi^2 + v_\theta^2)$ en coordonnées sphériques, ce qui nous donne

$$\forall \nu \in \{x, y, z\} \quad , \quad \bar{v}_\nu^2 = \int v_\nu^2 f(E_m, r^2(v_\varphi^2 + v_\theta^2)) d^3\mathbf{v} \quad ,$$

soit

$$\bar{v}_r^2 \neq \bar{v}_\varphi^2 = \bar{v}_\theta^2 \quad .$$

⁴Pour plus de détails, le lecteur peut se référer à [19].

3.2 Polytropes et modèle de Plummer

Les polytropes sont des modèles inspirés des gaz polytropiques qui représentent des systèmes à symétrie sphérique.

Pour des raisons pratiques liées à l'étude des polytropes, nous définissons le potentiel relatif ϕ et l'énergie relative $\mathcal{E}$ par

$$\phi = -\psi + \psi_0 \quad \text{et} \quad \mathcal{E} = -E_m + \psi_0 = \phi - \frac{1}{2}v^2.$$

ψ_0 est choisi pour que la fonction de distribution f des polytropes soit positive pour des énergies relatives positives, et nulle pour des énergies relatives négatives.

Les polytropes sont alors définis par une fonction de distribution f de la forme

$$f(\mathcal{E}) = \begin{cases} \alpha \mathcal{E}^{n-3/2} & (\mathcal{E} > 0) \\ 0 & (\mathcal{E} \leq 0) \end{cases} \quad (1.40)$$

L'intégration de cette fonction de distribution sur les vitesses nous fournit la densité associée

$$\begin{aligned} \rho &= \int f(\mathcal{E}) d^3\mathbf{v} \\ \rho &= 4\pi \int_0^\infty f(\phi - v^2/2) v^2 dv \\ \rho &= 4\pi\alpha \int_0^{\sqrt{2\phi}} \left[\phi - \frac{1}{2}v^2 \right]^{(n-3/2)} v^2 dv. \end{aligned} \quad (1.41)$$

En opérant un nouveau changement de variable, $v^2 = 2\phi \cos^2 \theta$, c'est-à-dire $\mathcal{E} = \phi \sin^2 \theta$, l'équation (1.41) devient

$$\rho = \begin{cases} c_n \phi^n & (\phi > 0) \\ 0 & (\phi \leq 0) \end{cases} \quad (1.42)$$

avec

$$c_n = \frac{(2\pi)^{3/2} (n - \frac{3}{2})! \alpha}{n!}.$$

La fonction factorielle est définie, pour un réel, par

$$x! = \int_0^\infty t^x e^{-t} dt.$$

Cette intégrale n'est finie que pour $x > -1$. Cette contrainte nous impose donc $n > 1/2$.

Injectons maintenant (1.42) dans l'équation de Poisson (1.15). Nous obtenons

$$\begin{cases} \frac{1}{r^2} \frac{d}{dr} \left(r^2 \frac{d\phi}{dr} \right) + 4\pi G c_n \phi^n = 0 & (\phi > 0) \\ \frac{1}{r^2} \frac{d}{dr} \left(r^2 \frac{d\phi}{dr} \right) = 0 & (\phi \leq 0) \end{cases} \quad (1.43)$$

Nous pouvons adimensionner cette équation en utilisant les variables $q = r/b$ et $\Phi = \phi/\phi_0$ où

$$\phi_0 = \phi(0)$$

et

$$b = (4\pi G \phi_0^{n-1} c_n)^{-1/2}.$$

L'équation (1.43) devient alors

$$\frac{1}{q^2} \frac{d}{dq} \left(q^2 \frac{d\Phi}{dq} \right) = \begin{cases} -\Phi^n & \Phi > 0 \\ 0 & \Phi \leq 0 \end{cases} \quad (1.44)$$

avec les conditions aux limites

$$\lim_{q \rightarrow 0} \Phi = 1 \text{ et } \lim_{q \rightarrow 0} \frac{d\Phi}{dq} = 0 .$$

La première condition aux limites découle directement de la définition de Φ . La seconde est une conséquence du fait que le référentiel du centre d'inertie du système est galiléen, donc en translation rectiligne uniforme, *i.e.* $d\Phi/dq = \mathbf{F}$ est nulle en 0.

Cette équation, nommée équation de Lane-Emden, possède une solution analytique simple dans les cas $n = 0$, $n = 1$ et $n = 5$. Le cas $n = 5$ nous intéresse plus particulièrement. En posant

$$q = e^{-t} \text{ et } \Phi(q) = \frac{e^{t/2}}{\sqrt{2}} u(t) ,$$

nous obtenons

$$\frac{d^2 u}{dt^2} = \frac{1}{4} u - \frac{1}{4} u^5 ,$$

avec les conditions aux limites suivantes

$$\lim_{t \rightarrow \infty} u = \lim_{t \rightarrow \infty} \frac{du}{dt} = 0 .$$

La solution de cette équation différentielle nous mène à l'équation

$$t = \frac{1}{2} \ln \left(k_1 \frac{1 - \sqrt{1 - u^4/3}}{1 + \sqrt{1 - u^4/3}} \right) ,$$

où k_1 est une constante d'intégration. Après quelques manipulations, nous obtenons finalement

$$\Phi = \frac{k_2^{1/4}}{\sqrt{1 + \frac{k_2}{3} q^2}} ,$$

où $k_2 = -3k_1$. La limite en 0 de Φ nous permet de fixer notre constante : $k_2 = 1$. Ce qui nous donne finalement

$$\Phi = \frac{1}{\sqrt{1 + \frac{1}{3} \left(\frac{r}{b} \right)^2}} . \quad (1.45)$$

Le potentiel (1.45) est donc la solution de l'équation de Lane-Emden (1.44) pour $n = 5$. La densité associée à ce potentiel est égale à

$$\rho = \frac{c_5 \phi_0^5}{\left(1 + \frac{1}{3} \left(\frac{r}{b} \right)^2 \right)^{5/2}} . \quad (1.46)$$

Ce polytrophe d'indice $n = 5$ est appelé modèle de Plummer. Ce modèle a été découvert par Schuster (1883), à qui nous devons la dérivation présentée ci-dessus, mais tient son nom de Plummer, qui a montré en 1911 que la densité (1.46) correspond assez bien à celle de certains amas globulaires.

Nous pouvons remarquer que cette densité ne s'annule jamais, mais que la masse du système est finie. En effet, considérons l'équation de Poisson pour notre système sphérique

$$\frac{1}{r^2} \frac{d}{dr} \left(r^2 \frac{d\psi}{dr} \right) = 4\pi G \rho .$$

En intégrant cette équation, nous avons

$$\frac{r^2}{G} \frac{d\psi}{dr} = M(r) .$$

La masse totale M du système est la limite de $M(r)$ lorsque r tend vers $+\infty$, soit

$$M = \lim_{r \rightarrow \infty} \left(\frac{r^2}{G} \frac{d\psi}{dr} \right) = - \lim_{q \rightarrow \infty} \frac{b}{G} \left(q^2 \frac{d\phi}{dq} \right) = \lim_{q \rightarrow \infty} \frac{b\phi_0}{3G} \frac{1}{\left(\frac{1}{q^2} + \frac{1}{3} \right)^{3/2}} = \frac{\sqrt{3}\phi_0 b}{G} .$$

La masse nous permet d'exprimer le potentiel au centre ϕ_0 . En posant $a = \sqrt{3}b$, nous obtenons

$$\phi_0 = \frac{GM}{a}$$

et donc

$$\phi = \frac{GM}{(r^2 + a^2)^{1/2}} .$$

De plus, en utilisant ρ_0 donné par l'équation (1.42), nous obtenons

$$\rho = \frac{\rho_0}{(1 + r^2/a^2)^{5/2}} .$$

On peut également dériver la relation (1.42) de l'équation d'état des gaz polytropiques

$$P = \alpha_p \rho^\gamma , \tag{1.47}$$

où P est la pression et γ l'indice polytropique, en considérant une sphère de gaz polytropique en équilibre hydrostatique. Cet état d'équilibre est décrit par la relation

$$\frac{dP}{dr} = -\rho \frac{d\psi}{dr} . \tag{1.48}$$

Si l'on injecte (1.47) dans (1.48), puis que l'on change de variable au profit du potentiel relatif, on obtient, après intégration et en choisissant bien ψ_0 (voir [9] p.224)

$$\rho^{\gamma-1} = \frac{\gamma-1}{\alpha_p \gamma} \phi .$$

Cette relation est équivalente à la relation (1.42), avec

$$n = \frac{1}{\gamma-1} \text{ et } c_n = \left(\frac{1}{\alpha_p \gamma n} \right)^n .$$

3.3 Sphère isotherme

Considérons un ensemble de N masses ponctuelles, chacune de masse m . Ces N particules sont confinées dans un volume sphérique de rayon R . Les particules rebondissent (chocs élastiques) sur les parois de ce volume et interagissent seulement gravitationnellement. Le système formé de ces N particules est par conséquent isolé.

Ce système peut être décrit par une fonction de distribution de masse $f(\mathbf{r}, \mathbf{v}, t)$. L'entropie du système est donnée par

$$S = - \int f \ln f d^3\mathbf{r} d^3\mathbf{v} . \tag{1.49}$$

L'évolution de f est régie par l'équation (1.27). Le terme de collisions de cette équation satisfait les deux conditions suivantes : l'énergie et la masse du système sont conservées, et l'entropie définie par (1.49) ne décroît pas au cours de l'évolution du système.

Il est possible de déterminer un état d'équilibre en cherchant la fonction f pour laquelle l'entropie est maximale, pour une énergie $\mathcal{H}$ et une masse M données. Il a été montré qu'un maximum global de l'entropie, dans notre cadre, n'existe pas (voir par exemple [9], p.268). Mais il peut exister des maxima locaux.

Nous allons dans un premier temps chercher la fonction de distribution f qui rend extrémale l'entropie de notre système, en imposant les contraintes suivantes : la densité ρ et l'énergie cinétique totale K doivent être fixées.

Considérons donc l'ensemble des fonctions de distribution $f(\mathbf{r}, \mathbf{v})$, avec $\rho(\mathbf{r})$ et K fixées. Si nous fixons ρ , la masse totale M , le potentiel gravitationnel $\psi(\mathbf{r})$ et l'énergie potentielle totale U sont fixés. Puisque K et U sont fixées, l'énergie totale $\mathcal{H} = K + U$ est également fixée. Nous sommes donc dans le cadre défini précédemment. Rappelons que ρ et K sont respectivement données par

$$\rho(\mathbf{r}) = \int f(\mathbf{r}, \mathbf{v}) d^3\mathbf{v} \quad (1.50)$$

et

$$K = \int \frac{1}{2} v^2 f(\mathbf{r}, \mathbf{v}) d\mathbf{r}^3 d^3\mathbf{v} . \quad (1.51)$$

Nous cherchons donc

$$\max_{f \in \sigma_{(\rho, K)}} S(f)$$

où

$$\sigma_{(\rho, K)} = \left\{ f ; \int f d^3\mathbf{v} = \rho , \int \frac{1}{2} v^2 f d\mathbf{r}^3 d^3\mathbf{v} = K \right\} ,$$

ce qui est équivalent à chercher

$$\max_f \mathcal{S}_1(f, \Lambda, \beta)$$

$\mathcal{S}_1$ étant extrémal en Λ et en β , et étant égal à

$$\mathcal{S}_1 = S(f) + \int \Lambda(\mathbf{r}) \left[\int f d^3\mathbf{v} - \rho(\mathbf{r}) \right] d^3\mathbf{r} + \beta \left[K - \frac{1}{2} \int v^2 f d^3\mathbf{r} d^3\mathbf{v} \right] .$$

Puisque $d(f \ln f) = \ln(ef) df$, nous avons

$$\delta \mathcal{S}_1 = - \int \left[\ln(ef) \delta f + \Lambda(\mathbf{r}) \delta f + \frac{1}{2} \beta v^2 \delta f \right] d\mathbf{r}^3 d^3\mathbf{v} .$$

Pour que la variation du premier ordre $\delta \mathcal{S}_1$ soit nulle quelle que soit la variation δf , il faut que le facteur multiplicateur de δf soit nul, c'est-à-dire

$$\ln(ef) + \Lambda(\mathbf{r}) + \frac{1}{2} \beta v^2 = 0 .$$

Nous avons donc

$$f = \exp \left(- [\Lambda(\mathbf{r}) + 1] - \frac{1}{2} \beta v^2 \right) .$$

Nous pouvons relier les multiplicateurs de Lagrange Λ et β respectivement à ρ et K . En effet, d'après (1.50), nous avons

$$\rho(\mathbf{r}) = \int f d^3\mathbf{v} = \exp(-[\Lambda(\mathbf{r}) + 1]) \int e^{-1/2\beta v^2} d^3\mathbf{v} = \left(\frac{2\pi}{\beta} \right)^{3/2} e^{-\Lambda(\mathbf{r})-1} .$$

De même, d'après (1.51), nous avons

$$\begin{aligned}
 K &= \int \frac{v^2}{2} f d^3 \mathbf{r} d^3 \mathbf{v} \\
 &= \frac{1}{2} \int e^{-\Lambda(\mathbf{r})-1} d^3 \mathbf{r} \int v^2 e^{-1/2\beta v^2} d^3 \mathbf{v} \\
 &= \frac{1}{2} \int \left(\frac{\beta}{2\pi} \right)^{3/2} \rho(\mathbf{r}) d^3 \mathbf{r} \int 4\pi v^4 e^{-1/2\beta v^2} d^3 \mathbf{v} \\
 &= \frac{1}{2} M \left(\frac{\beta}{2\pi} \right)^{3/2} \frac{3}{2} \pi^{3/2} \left(\frac{2}{\beta} \right)^{5/2}
 \end{aligned}$$

et finalement

$$K = \frac{3M}{2\beta}. \quad (1.52)$$

Si nous posons $T = 1/\beta$, nous obtenons finalement

$$f(\mathbf{r}, \mathbf{v}) = (2\pi T)^{-3/2} \rho(\mathbf{r}) e^{-v^2/(2T)}. \quad (1.53)$$

À ρ et K fixées, l'entropie est donc extrémale⁵ pour cette fonction de distribution f . Qui plus est, le signe de la variation du deuxième ordre $\delta^2 \mathcal{S}_1$ de $\mathcal{S}_1$ est celui de $-(\delta f)^2/(2f)$, et nous permet de savoir si cette extremum est un maximum ou un minimum. f étant une fonction positive, $\delta^2 \mathcal{S}_1$ est négative, et f représente donc un maximum local de l'entropie.

Remarque : Le signe de la perturbation de deuxième ordre de l'entropie dans le cas d'une perturbation de la densité est l'objet de la section 3.4.

Cet extremum local de l'entropie a pour valeur

$$S(\rho(\mathbf{r}), K) = \frac{3}{2} M (1 + \ln(2\pi T)) - \int \rho \ln \rho d^3 \mathbf{r},$$

soit

$$S(\rho(\mathbf{r}), K) = \frac{3}{2} M \ln T - \int \rho \ln \rho d^3 \mathbf{r} + \text{Constante}. \quad (1.54)$$

Nous avons formellement fixé ρ et K , mais nous ne savons pas encore quelle est la densité correspondant à l'extremum de l'entropie. Afin de calculer cette densité, nous allons cette fois rechercher la fonction de distribution qui rend extrémale l'entropie en fixant $\mathcal{H}$ et M et faisant varier K et $\rho(\mathbf{r})$.

Nous allons utiliser la variation de l'énergie $\mathcal{H}$ tronquée au deuxième ordre, et qui doit être nulle puisque $\mathcal{H}$ est fixée, afin de trouver une expression de la variation de K utilisable pour nos calculs ultérieurs. Rappelons que U est donnée par

$$U = \frac{1}{2} \int \rho(\mathbf{r}) \psi(\mathbf{r}) d^3 \mathbf{r}.$$

Cette variation de $\mathcal{H}$ est égale à

$$0 = \delta \mathcal{H} = \delta K + \delta U = \frac{3M}{2} \delta T + \frac{1}{2} \int (\psi \delta \rho + \rho \delta \psi) d^3 \mathbf{r} + \frac{1}{2} \int \delta \rho \delta \psi d^3 \mathbf{r}. \quad (1.55)$$

Nous savons de plus que

$$\int \rho \delta \psi d^3 \mathbf{r} = \int \psi \delta \rho d^3 \mathbf{r}.$$

⁵ Il est nécessaire que $\delta \mathcal{S}_1$ soit nulle pour que f représente un extremum de l'entropie, mais ce n'est pas une condition suffisante. Il faut également que $\mathcal{S}_1$ soit convexe.

En effet,

$$\int \psi(\mathbf{r}) \delta \rho(\mathbf{r}) d^3 \mathbf{r} = -G \int \frac{f(\mathbf{r}', \mathbf{v}')}{|\mathbf{r} - \mathbf{r}'|} df(\mathbf{r}, \mathbf{v}) d^3 \mathbf{r}' d^3 \mathbf{v}' d^3 \mathbf{v} d^3 \mathbf{r} ,$$

soit, en échangeant les variables d'intégration $\mathbf{r}$ et $\mathbf{r}'$,

$$\int \psi(\mathbf{r}) \delta \rho(\mathbf{r}) d^3 \mathbf{r} = -G \int \frac{df(\mathbf{r}', \mathbf{v}')}{|\mathbf{r} - \mathbf{r}'|} f(\mathbf{r}, \mathbf{v}) d^3 \mathbf{r} d^3 \mathbf{v} d^3 \mathbf{v}' d^3 \mathbf{r}' .$$

Si nous inversons maintenant l'ordre d'intégration entre $\mathbf{r}$ et $\mathbf{r}'$, nous obtenons

$$\int \psi(\mathbf{r}) \delta \rho(\mathbf{r}) d^3 \mathbf{r} = \int \delta \psi(\mathbf{r}) \rho(\mathbf{r}) d^3 \mathbf{r} .$$

L'équation (1.55) nous donne donc

$$\delta T = -\frac{2}{3M} \int \left(\psi \delta \rho + \frac{1}{2} \delta \rho \delta \psi \right) d^3 \mathbf{r} . \quad (1.56)$$

Nous allons maintenant chercher

$$\max_{(\rho, T) \in \sigma(\mathcal{H}, M)} S(\rho, T) ,$$

avec

$$\sigma(\mathcal{H}, M) = \left\{ (\rho, T) ; \frac{3}{2} MT + \frac{1}{2} \int \rho \psi d^3 \mathbf{r} = \mathcal{H} , \int \rho d^3 \mathbf{r} = M \right\} ,$$

ce qui est équivalent, si nous utilisons un multiplicateur de Lagrange α pour prendre en compte la contrainte de M fixée, à chercher

$$\max_{(\rho, T)} \mathcal{S}_2 ,$$

où $\mathcal{S}_2$ est extrémal en M et est égal à

$$\mathcal{S}_2 = S(\rho, T) + \alpha \left[\int \rho d^3 \mathbf{r} - M \right] .$$

Si nous considérons la variation du premier ordre de $\mathcal{S}_2$, S étant donné par (1.54), en utilisant (1.56), nous avons

$$\delta \mathcal{S}_2 = - \int \delta \rho \left(\frac{\psi}{T} + \ln(e\rho) - \alpha \right) d^3 \mathbf{r} .$$

Pour que $\delta \mathcal{S}_2$ soit nulle, il faut donc que

$$\frac{\psi}{T} + \ln(e\rho) - \alpha = 0 ,$$

c'est-à-dire que

$$\rho = A e^{-\beta \psi} \text{ avec } A = e^{\alpha-1} .$$

En posant

$$\rho_c = \rho(0) = A e^{-\beta \psi(0)} ,$$

la densité est finalement égale à

$$\rho(\mathbf{r}) = \rho_c e^{-\beta(\psi(\mathbf{r}) - \psi(0))} . \quad (1.57)$$

Remarque : Nous pouvons maintenant nous assurer que la fonction de distribution (1.53) associée à (1.57) correspond à un maximum de l'entropie. En effet, il n'est pas assuré que nous trouvions un

maximum puisque nous avons cherché la solution d'un nouveau problème de maximisation. Pour cela, nous devons déterminer le signe de la variation du deuxième ordre de $\mathcal{S}_2$. Cette variation est égale à

$$\delta^2 \mathcal{S}_2 = \frac{1}{2} \left[-\frac{3}{2} M \frac{(\delta T)^2}{T^2} - \int \frac{(\delta \rho)^2}{\rho} d^3 \mathbf{r} \right] .$$

Si nous injectons (1.56) dans cette équation et que nous omettons les termes de premier ordre et d'ordre supérieur ou égal à 3, nous obtenons

$$\delta^2 \mathcal{S}_2 = - \int \left[\frac{\delta \rho \delta \psi}{2T} + \frac{(\delta \rho)^2}{2\rho} \right] d^3 \mathbf{r} - \frac{1}{3MT^2} \left(\int \psi \delta \rho d^3 \mathbf{r} \right)^2 . \quad (1.58)$$

Le signe de $\delta^2 \mathcal{S}_2$ nous permet de savoir si l'extremum atteint est un maximum ($\delta^2 \mathcal{S}_2 < 0$) ou un minimum ($\delta^2 \mathcal{S}_2 > 0$). Nous verrons par la suite que, sous certaines conditions, cet extremum est un maximum (voir la section 3.4).

Le système décrit par la fonction de distribution ainsi obtenue

$$f = (2\pi T)^{-3/2} \rho_c e^{-(\psi(\mathbf{r}) - \psi(0))/T} e^{-v^2/(2T)}$$

est qualifié d'isotherme, par analogie avec un système de gaz isotherme⁶. Cette fonction ne dépend que de l'énergie par unité de masse E_m

$$f = f(E_m) \quad \text{avec} \quad E_m = \frac{v^2}{2} + \psi(\mathbf{r}) ,$$

ce système est donc à symétrie sphérique (voir la section 3.1). Il est appelé sphère isotherme. Cette dérivation de la fonction de distribution de la sphère isotherme peut être trouvée dans [45] ou [16].

3.4 Étude complémentaire sur la sphère isotherme et instabilité d'Antonov

Considérons de nouveau un système composé de N particules ponctuelles de masse m , confiné dans un volume sphérique (de rayon R), sur les parois duquel les particules rebondissent de façon élastique. Comme nous l'avons brièvement signalé dans le paragraphe 3.3, il n'existe pas de fonction de distribution f permettant de maximiser globalement l'entropie du système. Par contre, nous avons vu dans ce même paragraphe qu'il était possible de trouver un extremum local de l'entropie, qui correspond à la fonction de distribution d'une sphère isotherme.

Il est possible de montrer que cet extremum local, la sphère isotherme, ne peut pas être atteint pour certaines valeurs de l'énergie totale, de la masse et du rayon. Dans ce cas, il n'existe pas d'extremum local de l'entropie.

Nous dériverons enfin l'instabilité d'Antonov d'une étude complémentaire permettant de déterminer le signe de la variation de deuxième ordre de l'énergie, étude que nous avons laissée en suspens précédemment.

⁶La fonction de distribution d'un système de gaz isotherme est donnée par

$$f = (2\pi T)^{-3/2} \rho_c e^{\psi(0)/T} e^{-E_m/T} .$$

Condition d'existence de la sphère isotherme

Rappelons que la fonction permettant de rendre extrémale l'entropie obtenue précédemment est la suivante

$$f = \left(\frac{2\pi}{\beta} \right)^{-3/2} \rho_c e^{-\beta(\psi(r) - \psi(0))} e^{-\beta v^2/2} .$$

Rappelons également que ce système étant à symétrie sphérique, le potentiel et la densité sont donc des fonctions de r et non de $\mathbf{r}$. De plus, nous savons que le potentiel au bord du système est égal à

$$\psi(R) = -\frac{GM}{R} ,$$

où $M = Nm$, et que la force dérivant du potentiel est nulle au centre du système, puisque celui-ci est à symétrie sphérique, soit

$$\left. \frac{d\psi}{dr} \right|_{r=0} = 0 .$$

Nous pouvons nous demander s'il est possible d'obtenir n'importe quelle énergie totale $-\infty < \mathcal{H} < +\infty$, n'importe quelle masse totale $0 < M < +\infty$ et n'importe quel rayon $0 < R < +\infty$ en choisissant correctement ρ_c et β . Pour répondre à cette question, nous allons tout d'abord définir quelques constantes

$$\begin{aligned} L_0 &= (4\pi G \rho_c \beta)^{-1/2} , \\ M_0 &= 4\pi \rho_c L_0^3 , \\ \Phi_0 &= \beta^{-1} = \frac{GM_0}{L_0} , \end{aligned}$$

qui nous permettent de définir des variables sans dimensions

$$\begin{aligned} \hat{r} &= \frac{r}{L_0} , \\ \hat{\rho} &= \frac{\rho}{\rho_c} , \\ \hat{M} &= \frac{M(r)}{M_0} , \\ \hat{\psi} &= \beta [\psi(r) - \psi(0)] . \end{aligned}$$

Remarque : L_0 est égal à $r_K/3$, où r_K est le rayon de King (voir [9], p.228 et p.502).

Nous pouvons vérifier assez simplement que ces variables vérifient les équations suivantes

$$\hat{M}' = \hat{\rho} \hat{r}^2 , \tag{1.59}$$

$$\hat{\psi}' = \frac{\hat{M}}{\hat{r}^2} , \tag{1.60}$$

$$\hat{\rho}' = -\frac{\hat{\rho} \hat{M}}{\hat{r}^2} , \tag{1.61}$$

où le prime dénote une dérivée par rapport à $\hat{r}$.

En effet, nous savons que

$$\frac{d\hat{M}}{d\hat{r}} = \frac{L_0}{M_0} \frac{dM}{dr} = \frac{4\pi r^2 L_0}{M_0} \rho = \frac{4\pi L_0^3 \rho_c}{M_0} \hat{r}^2 \hat{\rho} ,$$

soit (1.59).

Ensuite, considérons l'équation de Poisson pour un système à symétrie sphérique

$$\frac{1}{r^2} \frac{d}{dr} \left(r^2 \frac{d}{dr} \psi(r) \right) = 4\pi G \rho(r) . \quad (1.62)$$

Sachant que

$$M(0) = 0 \quad \text{et} \quad \left. \frac{d\psi}{dr} \right|_{r=0} = 0$$

et que

$$\frac{dM}{dr} = 4\pi r^2 \rho ,$$

nous obtenons, en intégrant (1.62) après l'avoir multipliée par r^2

$$r^2 \frac{d}{dr} \psi(r) = GM(r) .$$

En multipliant cette équation par β , et après quelques changements de variables, nous arrivons à (1.60).

Enfin, en remarquant, à partir de (1.57), que

$$\frac{d\rho}{dr} = \frac{d\rho}{d\psi} \frac{d\psi}{dr} = -\beta \rho \frac{d\psi}{dr} ,$$

et en utilisant (1.60), nous obtenons immédiatement (1.61).

Le système formé des équations (1.59), (1.60) et (1.61) implique l'équation suivante

$$\frac{1}{\hat{r}^2} \frac{d}{d\hat{r}} \left(\hat{r}^2 \frac{d\hat{\psi}}{d\hat{r}} \right) = e^{-\hat{\psi}} , \quad (1.63)$$

avec les conditions aux limites

$$\hat{\psi}(0) = \hat{\psi}'(0) = 0 ,$$

obtenue en partant de (1.62) et en opérant quelques changements de variables.

Le degré de cette équation peut être réduit si nous utilisons les variables suivantes, introduites par Milne (voir [38] et [39])

$$\begin{aligned} u_1 &= \frac{\hat{M}}{\hat{r}} , \\ u_2 &= \frac{\hat{\rho} \hat{r}^3}{\hat{M}} = \frac{\hat{\rho} \hat{r}^2}{u_1} , \end{aligned}$$

où u_1 représente le carré d'une vitesse angulaire et u_2 le rapport entre une densité locale et une densité moyenne. Nous obtenons en effet

$$\begin{aligned} \frac{1}{\hat{r}^2} \frac{d}{d\hat{r}} \left(\hat{r}^2 \frac{d\hat{\psi}}{d\hat{r}} \right) = e^{-\hat{\psi}} &\Leftrightarrow \frac{\hat{\rho}}{u_1 u_2} \frac{d}{d\hat{r}} (u_1 \hat{r}) = \hat{\rho} \\ &\Leftrightarrow \frac{\hat{r}}{u_1 u_2} \frac{du_1}{d\hat{r}} + \frac{1}{u_2} = 1 \\ &\Leftrightarrow \frac{\hat{r}}{u_1} \frac{du_1}{du_2} \frac{du_2}{d\hat{r}} = u_2 - 1 . \end{aligned}$$

Or, la dérivée de u_2 par rapport à $\hat{r}$ est égale à

$$\frac{du_2}{d\hat{r}} = \frac{\hat{r}^3}{\hat{M}} \frac{d\hat{\rho}}{d\hat{r}} + \frac{3\hat{r}^2\hat{\rho}}{\hat{M}} - \frac{\hat{\rho}\hat{r}^3}{\hat{M}^2} \frac{d\hat{M}}{d\hat{r}} ,$$

soit

$$\hat{r} \frac{du_2}{d\hat{r}} = -u_2 (u_2 + u_1 - 3) .$$

Nous avons donc finalement

$$\frac{u_2}{u_1} \frac{du_1}{du_2} = -\frac{u_2 - 1}{u_2 + u_1 - 3} . \quad (1.64)$$

Les conditions aux limites sont obtenues en effectuant un développement polynomial de $\hat{\psi}$ au voisinage de $\hat{\psi} = 0$. Si $\hat{\psi}$ est de la forme $\sum \alpha_i \hat{r}^{2i}$ (voir [16], p.156), alors, en utilisant (1.63), nous pouvons déterminer les coefficients α_i du polynôme. Les premiers coefficients sont par exemple

$$\alpha_1 = \frac{1}{6} , \alpha_2 = -\frac{1}{120} , \alpha_3 = \frac{1}{1890} .$$

De ce développement, nous pouvons déduire

$$u_1 = \hat{\psi}'\hat{r} = \frac{1}{3}\hat{r}^2 - \frac{1}{30}\hat{r}^4 + o(\hat{r}^6)$$

et

$$u_2 = \frac{e^{-\hat{\psi}}\hat{r}}{\hat{\psi}'} = \frac{\hat{r} - \hat{r}^3/6 + o(\hat{r}^5)}{\hat{r}/3 - \hat{r}^3/30 + o(\hat{r}^5)} = 3 - \frac{1}{5}\hat{r}^2 + o(\hat{r}^4) .$$

De ce fait, lorsque $\hat{r}$ tend vers 0, nous avons

$$u_2 = 3 , u_1 = 0 , \frac{du_1}{du_2} = -\frac{5}{3} .$$

À partir de là, il est possible d'intégrer numériquement u_1 et u_2 , ce que nous avons fait à l'aide d'un schéma Runge-Kutta d'ordre 4, en partant de $\hat{r} \sim 0$. En effet, nous savons que

$$\frac{du_2}{d\hat{r}} = \frac{-u_2 (u_1 + u_2 - 3)}{\hat{r}} \quad (1.65)$$

et que

$$\frac{du_1}{d\hat{r}} = \frac{u_1 (u_2 - 1)}{\hat{r}} . \quad (1.66)$$

La seule difficulté résulte dans la singularité en 0, mais nous pouvons la lever en utilisant les développements polynomiaux de u_1 et u_2 en 0, et en nous plaçant au voisinage de 0. Le résultat de cette intégration est représenté sur la figure 1.1. Une sphère isotherme doit obligatoirement se trouver sur la courbe, qui représente les solutions de l'équation (1.64).

Nous pouvons par ailleurs calculer l'énergie totale d'une sphère isotherme. Cette énergie est égale à la somme de l'énergie cinétique K et de l'énergie potentielle U avec, d'après (1.52) et en posant $\hat{r}_0 = R/L_0$,

$$K = \frac{3M}{2\beta} = \frac{3GM_0^2}{2L_0} \hat{M}(\hat{r}_0) = \frac{3GM_0^2}{2L_0} \int_0^{\hat{r}_0} \frac{d\hat{M}}{d\hat{r}} d\hat{r} = \frac{3GM_0^2}{2L_0} \int_0^{\hat{r}_0} \hat{\rho} \hat{r}^2 d\hat{r}$$

et

$$U = - \int_0^R \frac{GM(r)}{r} \frac{dM}{dr} dr = - \frac{GM_0^2}{L_0} \int_0^{\hat{r}_0} \hat{M} \hat{\rho} \hat{r} d\hat{r} .$$

Nous avons donc

$$\mathcal{H} = \frac{GM_0^2}{2L_0} \int_0^{\hat{r}_0} (3\hat{\rho}\hat{r}^2 - 2\hat{M}\hat{\rho}\hat{r}) d\hat{r} = \frac{GM_0^2}{2L_0} \int_0^{\hat{r}_0} \frac{d}{d\hat{r}} (2\hat{\rho}\hat{r}^3 - 3\hat{M}) d\hat{r} ,$$

soit

$$\mathcal{H} = \frac{GM_0^2}{L_0} \left(\hat{\rho}_0 \hat{r}_0^3 - \frac{3}{2} \hat{M}_0 \right) , \quad (1.67)$$

où $\hat{\rho}_0 = \hat{\rho}(\hat{r}_0)$ et $\hat{M}_0 = \hat{M}(\hat{r}_0)$. Définissons maintenant λ par

$$\lambda = \frac{R\mathcal{H}}{GM^2} .$$

λ est une quantité sans dimension. D'après (1.67), nous avons

$$\lambda = \frac{1}{u_1(\hat{r}_0)} \left(u_2(\hat{r}_0) - \frac{3}{2} \right) . \quad (1.68)$$

λ ne dépend donc que de u_1 et u_2 . Si nous caractérisons une sphère isotherme en fixant son rayon R ,

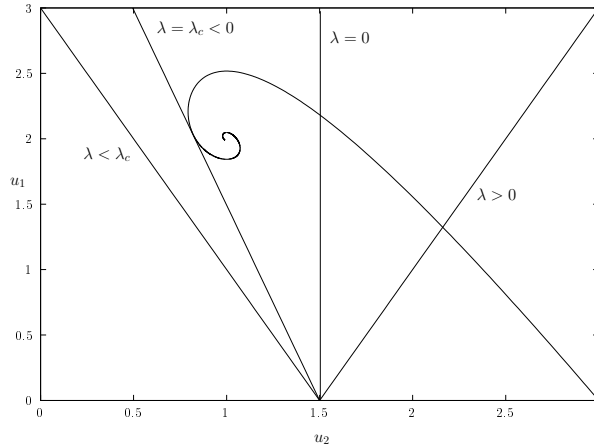


FIG. 1.1 – La courbe représente les sphères isothermes dans le plan (u_2, u_1) . Les droites représentent les valeurs prises par λ dans ce même plan (u_2, u_1) .

son énergie $\mathcal{H}$ et sa masse M , nous savons que cette sphère doit se trouver sur la droite

$$u_1 = \frac{1}{\lambda} \left(u_2 - \frac{3}{2} \right) \quad (1.69)$$

dans le plan (u_2, u_1) . Comme d'autre part, une sphère isotherme doit se trouver sur la courbe donnée par l'équation (1.64), elle doit nécessairement être à l'intersection de cette courbe et de la droite (1.69). Or, nous pouvons voir sur la figure 1.1 qu'une intersection n'est possible que si $\lambda > \lambda_c$. La valeur de λ_c que nous avons déterminé numériquement est égale à -0.335 , et cette valeur est identique à celle donnée par Padmanabhan [45]. Ceci signifie donc qu'un système tel que nous l'avons décrit au début de ce paragraphe pour lequel $\lambda < -0.335$ ne peut pas évoluer vers la sphère isotherme. Or nous avons montré que lorsqu'il existe un extremum de l'entropie, cet extremum correspond à une sphère isotherme. Un tel système ne peut donc pas atteindre d'état représentant un extremum local de l'entropie.

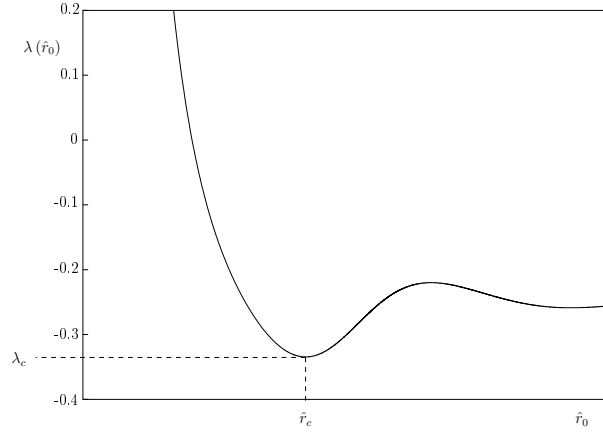


FIG. 1.2 – Représentation de λ en fonction de $\hat{r}_0$. Le minimum λ_c de λ est atteint en $\hat{r}_c$.

La figure 1.2 représente l'évolution⁷ de λ en fonction de $\hat{r}_0$. Lorsque $\hat{r}_0$ tend vers 0, λ tend vers $+\infty$. Puis, quand R , donc $\hat{r}_0$, augmente, λ diminue jusqu'à atteindre son minimum λ_c en $\hat{r}_c$. Ensuite λ augmente et oscille indéfiniment, en tendant vers -0.25.

Instabilité d'Antonov

Plaçons nous dans le cas où $\lambda > \lambda_c$. Dans ce cas, il existe un extremum local de l'entropie, la sphère isotherme. Mais cet extremum est-il un maximum de l'entropie, et représente-t-il un état d'équilibre stable ?

Afin de répondre à ces questions, nous allons étudier le signe de la variation de deuxième ordre de l'entropie (cette question avait été laissée en suspens dans le § 3.3).

Considérons une sphère isotherme de densité ρ_s , d'énergie $\mathcal{H}$, de masse M et de rayon R . Nous savons que $\lambda = R\mathcal{H}/GM^2 > -0.335$. Considérons également un système dont la densité est égale à $\rho_s + \delta\rho$, avec $\delta\rho = \epsilon u$ et $\epsilon \ll 1$, mais avec la même énergie, la même masse et le même rayon que le système précédent. La différence d'entropie entre ces deux systèmes est égale à

$$S(\rho_s + \delta\rho) - S(\rho_s) = \delta^2 S + \text{termes d'ordres supérieurs.}$$

En effet, le terme du premier ordre, δS , est nul puisque ρ_s est la densité d'une sphère isotherme (voir le § 3.3). Le terme de deuxième ordre est quant à lui égal à

$$\delta^2 S = -\epsilon^2 \int \left[\frac{u\psi_u}{2T} + \frac{u^2}{2\rho_s} \right] dr - \frac{\epsilon^2}{3MT^2} \left(\int \psi_s u dr \right)^2,$$

avec

$$\psi_s = -G \int \frac{\rho_s(\mathbf{r}')}{|\mathbf{r} - \mathbf{r}'|} d^3\mathbf{r}' \quad \text{et} \quad \psi_u = -G \int \frac{u(\mathbf{r}')}{|\mathbf{r} - \mathbf{r}'|} d^3\mathbf{r}'$$

d'après l'équation (1.58).

Si $\delta^2 S$ est négative, alors la sphère isotherme représente un maximum local de l'entropie. Mais si $\delta^2 S$ est positive, cela signifie qu'il existe des configurations voisines de ρ_s possédant une entropie supérieure à celle de ρ_s . Dans ce cas, la sphère isotherme représente un minimum local de l'entropie, et correspond à un état d'équilibre instable.

V.A.Antonov [4] a montré qu'il existe un rayon R_c pour lequel $\delta^2 S$ est négatif si $R < R_c$, s'annule en

⁷Ces résultats ont été obtenus en intégrant les équations (1.65) et (1.66) par la méthode Runge-Kutta d'ordre 4.

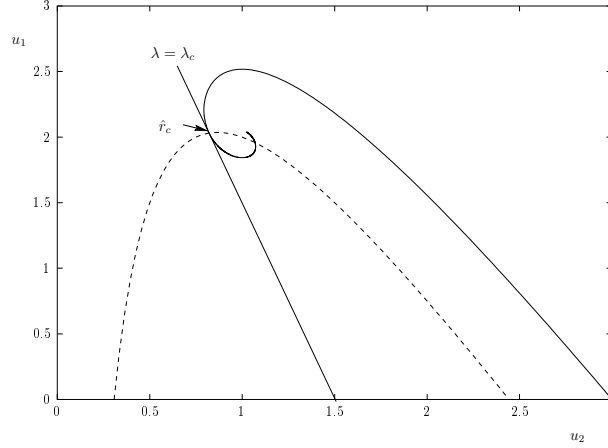


FIG. 1.3 – Représentation des sphères isothermes dans le plan (u_2, u_1) (courbe continue) et des solutions de $4u_2^2 + 2u_1u_2 - 11u_2 + 3 = 0$ (courbe discontinue).

R_c puis change de signe lorsque R devient supérieur à R_c . Les sphères isothermes avec $R < R_c$ sont donc des maxima locaux de l'entropie, et sont des états d'équilibre stable de l'équation de Vlasov.

V.A.Antonov a également montré que toutes les sphères isothermes avec $R > R_c$ sont instables, indépendamment du signe de $\delta^2 S$ – même si $\delta^2 S$ redevient négative pour des rayons $R > R_{c'} > R_c$, le système demeure instable. D'autre part, ce rayon critique R_c correspond au plus petit rayon pour lequel $d\lambda/d\hat{r} = 0$, c'est-à-dire $R_c = \hat{r}_c L_0$ (voir par exemple [45]).

Il est facile de trouver numériquement la valeur de $\hat{r}_c$. En effet, d'après (1.68), nous avons

$$\frac{d\lambda}{d\hat{r}} = \frac{1}{u_1} \frac{du_2}{d\hat{r}} + \frac{u_2 - 3/2}{u_1^2} \frac{du_1}{d\hat{r}} = -\frac{4u_2^2 + 2u_1u_2 - 11u_2 + 3}{2u_1\hat{r}}.$$

Si nous cherchons l'intersection entre la courbe représentant les sphères isothermes dans le plan (u_2, u_1) et la courbe $4u_2^2 + 2u_1u_2 - 11u_2 + 3 = 0$, nous trouvons que le plus petit $\hat{r}$ pour lequel ces deux courbes se croisent, c'est-à-dire $\hat{r}_c$ est environ égal à 34.2 (voir la figure 1.3). Cette valeur de $\hat{r}$ correspond à $\hat{\rho} = 1/709$, soit $\rho_c/\rho(R) = 709$.

Un sphère isotherme telle que $\rho_c/\rho(R) > 709$ est donc instable, et cette instabilité est appelée instabilité d'Antonov ou encore catastrophe gravothermale. Lorsque cette instabilité se déclenche, le cœur du système se réchauffe et sa densité augmente, ce qui permet au système d'atteindre des états de plus grande entropie. J.Katz [28] explique cette instabilité en considérant un système auto-gravitant composé de deux sous-systèmes distincts : le cœur et le halo. Le cœur est pratiquement purement auto-gravitant, ce qui implique, d'après le théorème du viriel, que $\mathcal{H}_c \simeq -K_c \propto -T_c$, où $\mathcal{H}_c$ est l'énergie totale du cœur, K_c son énergie cinétique et T_c sa température. La chaleur spécifique du cœur $C_c = d\mathcal{H}_c/dT_c$ est donc négative. Le halo se comporte quant à lui comme un gaz parfait dans le champ créé par le cœur, et sa chaleur spécifique C_h est donc positive. À l'équilibre, les températures du cœur et du halo sont identiques. Si, sous l'effet d'une perturbation, T_c devient supérieure à T_h , le cœur va perdre de l'énergie, puisque $C_c < 0$, et le halo va gagner cette énergie car le système est isolé. La température du halo va donc croître, puisque $C_h > 0$. D'autre part, la variation de température du cœur est égale à $\Delta T_c = \Delta \mathcal{H}_c/C_c$ et celle du halo $\Delta T_h = \Delta \mathcal{H}_h/C_h$. Donc si $C_h < |C_c|$, le halo se réchauffera plus vite que le cœur, la différence de température diminuera, et l'équilibre se rétablira de lui-même. Mais si $C_h > |C_c|$, alors l'écart de température ne cessera d'augmenter. Ce qui conduit à la catastrophe gravothermale (voir [36] et [9] p.500).

3.5 Méthode d'énergie et instabilité d'orbite radiale

La stabilité des solutions de l'équation de Boltzmann sans collisions (BSC) est une question d'importance. En effet, les systèmes auto-gravitants considérés dans cette étude (amas globulaires et galaxies elliptiques) sont des objets très anciens. Il ne fait donc aucun doute qu'ils représentent des solutions stables de l'équation de Boltzmann, à la réserve près que cette équation modélise correctement ces objets.

La stabilité des solutions de BSC peut être reliée au signe de la variation de deuxième ordre de l'énergie du système. Les outils nécessaires à cette étude seront introduits dans un premier temps, puis les principaux résultats de stabilité seront présentés.

Méthode d'énergie

Considérons les équations canoniques de Hamilton, que nous rappelons :

$$\frac{d\mathbf{r}}{dt} = \frac{\partial E}{\partial \mathbf{p}} \text{ et } \frac{d\mathbf{p}}{dt} = -\frac{\partial E}{\partial \mathbf{r}} ,$$

où E est le hamiltonien de la particule de position $\mathbf{r}$ et d'impulsion $\mathbf{p}$. Nous pouvons exprimer la dérivée totale de $\mathbf{r}$ et de $\mathbf{p}$ par rapport au temps grâce au crochet de Poisson

$$\{a, b\} = \frac{\partial a}{\partial \mathbf{r}} \frac{\partial b}{\partial \mathbf{p}} - \frac{\partial a}{\partial \mathbf{p}} \frac{\partial b}{\partial \mathbf{r}} .$$

Ainsi, nous obtenons

$$\frac{d\mathbf{r}}{dt} = \frac{\partial \mathbf{r}}{\partial \mathbf{r}} \frac{\partial E}{\partial \mathbf{p}} - \frac{\partial \mathbf{r}}{\partial \mathbf{p}} \frac{\partial E}{\partial \mathbf{r}} = \{\mathbf{r}, E\} ,$$

et, de la même façon,

$$\frac{d\mathbf{p}}{dt} = \{\mathbf{p}, E\} .$$

Il est même possible, et aisé, de montrer que, pour toute fonction χ ne dépendant que de $\mathbf{r}$ et $\mathbf{p}$, nous avons

$$\frac{d\chi}{dt} = \{\chi, E\} = -\{E, \cdot\} \chi ,$$

cette dernière notation, $-\{E, \cdot\} \chi$, exprimant l'application de l'opérateur

$$-\{E, \cdot\} = -\frac{\partial E}{\partial \mathbf{r}} \frac{\partial}{\partial \mathbf{p}} + \frac{\partial E}{\partial \mathbf{p}} \frac{\partial}{\partial \mathbf{r}}$$

à la fonction χ . Plus généralement, nous pouvons généraliser ceci à chacun des couples (Ξ, ξ) intervenant dans le théorème Noether⁸ (voir [48])

$$\forall \chi(\mathbf{r}, \mathbf{p}), \quad \frac{d\chi}{d\xi} = \{\chi, \Xi\} .$$

Notons que le crochet de Poisson a les quatre propriétés suivantes :

1. il est bilinéaire;
2. il est anticommutatif : $\{a, b\} = -\{b, a\}$;
3. il vérifie l'identité de Jacobi : $\{a, \{b, c\}\} + \{c, \{a, b\}\} + \{b, \{c, a\}\} = 0$;

⁸Le théorème de Noether relie les symétries du système avec des constantes du mouvement. Ce théorème exprime, par exemple, la conservation de l'énergie par une translation dans le temps, celle de la quantité de mouvement par une translation dans l'espace, et celle du moment cinétique par une rotation.

4. il satisfait la règle de Leibniz : $\{ab, c\} = a\{b, c\} + b\{a, c\}$.

Nous pouvons définir une algèbre de Lie grâce à ce crochet.

DÉFINITION 1 *On appelle algèbre de Lie toute algèbre dont l'application bilinéaire est anticommutative et vérifie l'identité de Jacobi.*

Nous savons de plus que E ne dépend pas du temps dans le cadre de notre étude, et que l'opérateur $\{E, \cdot\}$ est linéaire et indépendant du temps. La théorie des équations différentielles (voir par exemple [31], chapitre 2) nous permet donc d'exprimer l'évolution temporelle de $\chi(\mathbf{r}, \mathbf{p})$ par

$$\chi_{t_0+\Delta t} = e^{-\{E, \cdot\}\Delta t} \chi_{t_0} , \quad (1.70)$$

soit

$$\chi_{t_0+\Delta t} = \chi_{t_0} - \{E, \chi_{t_0}\} \Delta t + \frac{1}{2} \{E, \{E, \chi_{t_0}\}\} (\Delta t)^2 - \frac{1}{3!} \{E, \{E, \{E, \chi_{t_0}\}\}\} (\Delta t)^3 + \dots$$

Considérons maintenant une fonction de distribution f d'un système auto-gravitant dépendant de $\mathbf{r}$, $\mathbf{p}$ et t . Si nous calculons la dérivée totale de cette fonction par rapport au temps, nous obtenons

$$\frac{df}{dt} = \frac{\partial f}{\partial t} + \frac{\partial f}{\partial \mathbf{r}} \frac{d\mathbf{r}}{dt} + \frac{\partial f}{\partial \mathbf{p}} \frac{d\mathbf{p}}{dt} = \frac{\partial f}{\partial t} + \{f, E\} .$$

Si notre fonction dépend explicitement du temps t , il n'est pas possible d'écrire sa dérivée totale par rapport au temps à l'aide du crochet de Poisson seul. L'évolution temporelle de cette fonction de distribution dépendant du temps ne peut donc pas être exprimée grâce à (1.70). Mais d'après l'équation de Liouville (1.16), sa dérivée totale par rapport au temps est nulle. Par conséquent, nous savons que

$$\frac{\partial f}{\partial t} = \{E, f\} .$$

Nous pouvons également examiner le cas d'une fonctionnelle $\mathcal{F}$ de notre fonction de distribution f de la forme suivante

$$\mathcal{F}[f] = \int \mathcal{C}(f, \mathbf{r}, \mathbf{p}) d\mathbf{r} d\mathbf{p} = \int \mathcal{C}(f, \mathbf{r}, \mathbf{p}) d\mathbf{\Gamma} , \quad (1.71)$$

où $\mathcal{C}(f, \mathbf{r}, \mathbf{p})$ est une fonction de f , de $\mathbf{r}$ et de $\mathbf{p}$. Cette fonctionnelle $\mathcal{F}$ est une fonction de t , les autres variables étant intégrées sur tout l'espace des phases. Par définition, le domaine d'intégration ne dépend pas du temps. Nous pouvons donc écrire

$$\frac{d\mathcal{F}}{dt} = \frac{\partial \mathcal{F}}{\partial t} = \int \frac{\partial \mathcal{C}}{\partial f} \frac{\partial f}{\partial t} d\mathbf{\Gamma} .$$

Si nous introduisons la dérivée fonctionnelle⁹ $\delta\mathcal{F}/\delta f$ et le crochet de Poisson, nous obtenons

$$\frac{d\mathcal{F}}{dt} = \int \frac{\delta\mathcal{F}}{\delta f} \{E, f\} d\mathbf{\Gamma} . \quad (1.72)$$

Remarquons également que le crochet de Poisson vérifie la propriété suivante

$$\{h_1(a), h_2(a)\} = \frac{\partial h_1}{\partial a} \{a, a\} \frac{\partial h_2}{\partial a} = 0 ,$$

⁹Rappelons que la dérivée fonctionnelle $\delta X/\delta f$ d'une fonction $X(f)$ est donnée par

$$X(f + \delta f) - X(f) = \int \left(\delta f \frac{\delta X}{\delta f} + o(\delta f^2) \right) d\mathbf{\Gamma} .$$

et que, pour des fonctions décroissant suffisamment rapidement lorsque $|\mathbf{r}|$ et $|\mathbf{p}|$ tendent vers l'infini,

$$\int A \{B, C\} d\Gamma = \int C \{A, B\} d\Gamma \quad (1.73)$$

(voir [46], notamment l'annexe C, pour plus de précisions concernant ces deux propriétés). Considérons par ailleurs le Hamiltonien $\mathcal{H}$ d'un système auto-gravitant (voir l'équation (1.10))

$$\mathcal{H} = \int \frac{\mathbf{p}^2}{2m} f d\Gamma - \frac{G}{2} \int \frac{f f'}{|\mathbf{r} - \mathbf{r}'|} m m' d\Gamma d\Gamma'.$$

La dérivée fonctionnelle de $\mathcal{H}$ par rapport à f est égale à

$$\frac{\delta \mathcal{H}}{\delta f} = \frac{\mathbf{p}^2}{2m} - Gm \int \frac{m' f'(\mathbf{r}', \mathbf{p}')}{|\mathbf{r} - \mathbf{r}'|} d\Gamma' = E. \quad (1.74)$$

En injectant (1.74) dans (1.72), et en nous souvenant de (1.73), nous pouvons écrire

$$\frac{d\mathcal{F}}{dt} = \int \frac{\delta \mathcal{F}}{\delta f} \left\{ \frac{\delta \mathcal{H}}{\delta f}, f \right\} d\Gamma = \int f \left\{ \frac{\delta \mathcal{F}}{\delta f}, \frac{\delta \mathcal{H}}{\delta f} \right\} d\Gamma =: (\mathcal{F}, \mathcal{H}).$$

Le crochet $(.,.)$ a été introduit par Morrison [41] dans le cadre de la physique des plasmas. Ce crochet, de même que le crochet de Poisson, est bilinéaire et anticommutatif, et vérifie l'identité de Jacobi. Il permet donc de définir une algèbre de Lie. De plus, $\mathcal{H}$ étant indépendant du temps dans notre cadre, l'opérateur $(\mathcal{H}, .)$ l'est également. Nous pouvons finalement exprimer l'évolution temporelle de notre fonctionnelle $\mathcal{F}$ par

$$\mathcal{F}_{t+\Delta t} = e^{-(\mathcal{H},.)\Delta t} \mathcal{F}_t. \quad (1.75)$$

Ces outils vont nous permettre de trouver un critère de stabilité pour les solutions stationnaires de BSC.

Stabilité et variation de l'énergie

Considérons l'équation de Boltzmann sans collisions (1.28). Nous pouvons la formuler à l'aide du crochet de Poisson :

$$\frac{\partial f}{\partial t} + \{f, E\} = 0. \quad (1.76)$$

Cette équation peut également s'écrire grâce au crochet de Morrison. En effet, nous avons

$$(\mathcal{H}, f) = \int f(\Gamma') \left\{ \frac{\delta \mathcal{H}[f(\Gamma)]}{\delta f(\Gamma')}, \frac{\delta f(\Gamma)}{\delta f(\Gamma')} \right\} d\Gamma' = \left\{ f, \frac{\delta \mathcal{H}}{\delta f} \right\} = \{f, E\},$$

d'où

$$\frac{\partial f}{\partial t} + (\mathcal{H}, f) = 0.$$

Nous allons étudier les conditions de stabilité des solutions stationnaires de cette équation. Pour cela, considérons une fonction de distribution f_0 perturbée de la manière suivante

$$f = f_0 + df = f_0 + f_1 \epsilon + \frac{1}{2!} f_2 \epsilon^2 + \frac{1}{3!} f_3 \epsilon^3 + \dots, \quad (1.77)$$

où f_1 est le premier ordre de la perturbation, f_2 son deuxième ordre, et ainsi de suite.

En généralisant (1.75) (voir la section 3.2 de [46]), nous pouvons écrire

$$\frac{d\mathcal{F}}{d\epsilon} = (\mathcal{F}, \mathcal{G}),$$

où $\mathcal{G}$ est appelé le générateur de la perturbation paramétrée par ϵ . Il suit immédiatement

$$\mathcal{F}[f] = \mathcal{F}[f_0] - (\mathcal{G}, \mathcal{F}[f_0]) \epsilon + \frac{1}{2} (\mathcal{G}, (\mathcal{G}, \mathcal{F}[f_0])) \epsilon^2 - \frac{1}{3!} (\mathcal{G}, (\mathcal{G}, (\mathcal{G}, \mathcal{F}[f_0]))) \epsilon^3 + \dots \quad (1.78)$$

Considérons la fonctionnelle $\mathcal{F}$ particulière

$$\mathcal{F}[f] = \int f(\mathbf{\Gamma}') \delta_D(\mathbf{\Gamma}' - \mathbf{\Gamma}) d\mathbf{\Gamma}' = f(\mathbf{\Gamma}) \quad (1.79)$$

Le premier terme de (1.78) est donc égal à

$$\mathcal{F}[f_0] = f_0 \quad .$$

La dérivée fonctionnelle de $\mathcal{F}$ est égale à

$$\frac{\delta \mathcal{F}}{\delta f} = \delta_D(\mathbf{\Gamma}' - \mathbf{\Gamma}) \quad .$$

Nous pouvons calculer $(\mathcal{G}, \mathcal{F}[f_0])$ pour la fonctionnelle (1.79), ce qui nous donne

$$\begin{aligned} (\mathcal{G}, \mathcal{F}[f_0]) &= \int f'_0 \left\{ \frac{\delta \mathcal{G}}{\delta f_0}, \frac{\delta \mathcal{F}[f'_0]}{\delta f_0} \right\} d\mathbf{\Gamma}' \\ &= \int f'_0 \{g, \delta_D(\mathbf{\Gamma}' - \mathbf{\Gamma})\} d\mathbf{\Gamma}' \\ &= - \int \delta_D(\mathbf{\Gamma}' - \mathbf{\Gamma}) \{g, f'_0\} d\mathbf{\Gamma}' \\ &= - \{g, f_0\} \quad , \end{aligned}$$

où les primes désignent les fonctions de $\mathbf{\Gamma}'$ et avec

$$g := \frac{\delta \mathcal{G}}{\delta f_0} \quad .$$

Pour calculer le terme suivant de (1.78), nous devons tout d'abord calculer la dérivée fonctionnelle de $(\mathcal{G}, \mathcal{F}[f_0])$. Sachant que

$$\begin{aligned} (\mathcal{G}, \mathcal{F}[f_0 + \delta f]) &= \int (f'_0 + \delta f') \{g, \delta_D(\mathbf{\Gamma}' - \mathbf{\Gamma})\} d\mathbf{\Gamma}' \\ &= - \int \delta_D(\mathbf{\Gamma}' - \mathbf{\Gamma}) \{g, (f'_0 + \delta f')\} d\mathbf{\Gamma}' \\ &= - \{g, f_0 + \delta f\} \quad , \end{aligned}$$

le calcul de la dérivée fonctionnelle donne

$$\begin{aligned} \int \left(\frac{\delta}{\delta f} (\mathcal{G}, \mathcal{F}[f_0]) \right) \delta f' d\mathbf{\Gamma}' &:= (\mathcal{G}, \mathcal{F}[f_0 + \delta f]) - (\mathcal{G}, \mathcal{F}[f_0]) \\ &= - \{g, \delta f\} \\ &= - \int \{g, \delta f'\} \delta_D(\mathbf{\Gamma}' - \mathbf{\Gamma}) d\mathbf{\Gamma}' \\ &= \int \{g, \delta_D(\mathbf{\Gamma}' - \mathbf{\Gamma})\} \delta f' d\mathbf{\Gamma}' \quad , \end{aligned}$$

soit

$$\frac{\delta}{\delta f} (\mathcal{G}, \mathcal{F}[f_0]) = \{g, \delta_D (\mathbf{\Gamma}' - \mathbf{\Gamma})\} .$$

Le troisième terme de (1.78) est donc égal à

$$\begin{aligned} (\mathcal{G}, (\mathcal{G}, \mathcal{F}[f_0])) &= \int f'_0 \{g', \{g', \delta_D (\mathbf{\Gamma}' - \mathbf{\Gamma})\}\} d\mathbf{\Gamma}' \\ &= \int \{f'_0, g'\} \{g', \delta_D (\mathbf{\Gamma}' - \mathbf{\Gamma})\} d\mathbf{\Gamma}' \\ &= \int \{\{f'_0, g'\}, g'\} \delta_D (\mathbf{\Gamma}' - \mathbf{\Gamma}) d\mathbf{\Gamma}' \\ &= \{\{f_0, g\}, g\} \\ &= \{g, \{g, f_0\}\} . \end{aligned}$$

Nous pouvons poursuivre ce type de calcul avec les termes suivants, ce qui nous permet finalement d'obtenir l'expression suivante

$$f = f_0 + \{g, f_0\} \epsilon + \frac{1}{2!} \{g, \{g, f_0\}\} \epsilon^2 + \frac{1}{3!} \{g, \{g, \{g, f_0\}\}\} \epsilon^3 + \dots$$

Si nous comparons cette expression à (1.77), nous pouvons facilement déduire la forme de df

$$df = \{g, f_0\} \epsilon + \frac{1}{2!} \{g, \{g, f_0\}\} \epsilon^2 + \frac{1}{3!} \{g, \{g, \{g, f_0\}\}\} \epsilon^3 + \dots \quad (1.80)$$

Nous savons donc quelle doit être la forme notre perturbation pour être une perturbation physiquement acceptable. Nous pouvons maintenant examiner la perturbation du hamiltonien $\mathcal{H}$ engendrée par la perturbation df d'une solution stationnaire f_0 de BSC.

Il suffit pour cela de remplacer $\mathcal{F}$ par $\mathcal{H}$ dans (1.78), ce qui nous donne

$$\mathcal{H}[f] = \mathcal{H}[f_0] - (\mathcal{G}, \mathcal{H}[f_0]) \epsilon + \frac{1}{2} (\mathcal{G}, (\mathcal{G}, \mathcal{H}[f_0])) \epsilon^2 + \dots \quad (1.81)$$

La perturbation de premier ordre peut être aisément calculée puisque

$$(\mathcal{G}, \mathcal{H}[f_0]) = \int f_0 \left\{ \frac{\delta \mathcal{G}}{\delta f}, \frac{\delta \mathcal{H}}{\delta f} \right\} d\mathbf{\Gamma} = \int f_0 \{g, E\} d\mathbf{\Gamma} . \quad (1.82)$$

Nous pouvons remarquer que la perturbation de premier ordre $\delta \mathcal{H}$ du hamiltonien est égale à

$$\delta \mathcal{H} = -(\mathcal{G}, \mathcal{H}[f_0]) \epsilon = \epsilon \int f_0 \{E, g\} d\mathbf{\Gamma} = -\epsilon \int g \{E, f_0\} d\mathbf{\Gamma} = 0 ,$$

car f_0 est une solution stationnaire de BSC (voir (1.76)). Les solutions stationnaires de BSC représentent donc des extrema de l'énergie au regard des perturbations de la forme (1.80). Il est de plus possible de montrer (voir [5]) qu'un système dont la variation d'énergie est positive suivant une perturbation est stable contre cette perturbation. L'étude du signe de la perturbation de deuxième ordre $\delta^2 \mathcal{H}$ du hamiltonien nous permettra donc de dériver des critères de stabilité, au moins dans certains cas particuliers.

La dérivée fonctionnelle de $(\mathcal{G}, \mathcal{H}[f_0])$, que nous devons connaître pour calculer $\delta^2 \mathcal{H}$, est telle que

$$\begin{aligned}
 (\mathcal{G}, \mathcal{H}[f_0 + \delta f]) - (\mathcal{G}, \mathcal{H}[f_0]) &= \int \delta f \{g, E\} d\Gamma - \int f_0 \left\{ g, Gm \int \frac{\delta f' m'}{|\mathbf{r} - \mathbf{r}'|} d\Gamma' \right\} d\Gamma \\
 &\quad - \int \delta f \left\{ g, Gm \int \frac{\delta f' m'}{|\mathbf{r} - \mathbf{r}'|} d\Gamma' \right\} d\Gamma \\
 &= \int \delta f \{g, E\} d\Gamma - \int \delta f m \int f'_0 \left\{ g', \frac{Gm'}{|\mathbf{r} - \mathbf{r}'|} \right\} d\Gamma' d\Gamma \\
 &\quad - \int \delta f \left\{ g, Gm \int \frac{\delta f' m'}{|\mathbf{r} - \mathbf{r}'|} d\Gamma' \right\} d\Gamma \\
 &= \int \delta f \left(\frac{\delta}{\delta f} (\mathcal{G}, \mathcal{H}[f_0]) \right) d\Gamma - \int \delta f \left\{ g, Gm \int \frac{\delta f' m'}{|\mathbf{r} - \mathbf{r}'|} d\Gamma' \right\} d\Gamma,
 \end{aligned}$$

soit

$$\frac{\delta}{\delta f} (\mathcal{G}, \mathcal{H}[f_0]) = \{g, E\} - m \int f'_0 \left\{ g', \frac{Gm'}{|\mathbf{r} - \mathbf{r}'|} \right\} d\Gamma'.$$

Ceci nous permet de calculer la perturbation de deuxième ordre du hamiltonien

$$\begin{aligned}
 (\mathcal{G}, (\mathcal{G}, \mathcal{H}[f_0])) &= \int f_0 \{g, \{g, E\}\} d\Gamma - \int f_0 \left\{ g, m \int f'_0 \left\{ g', \frac{Gm}{|\mathbf{r} - \mathbf{r}'|} \right\} d\Gamma' \right\} d\Gamma \\
 &= - \int \{g, f_0\} \{g, E\} d\Gamma + \int \{g, f_0\} \int f'_0 \left\{ g', \frac{Gmm'}{|\mathbf{r} - \mathbf{r}'|} \right\} d\Gamma' d\Gamma \\
 &= - \int \{g, f_0\} \{g, E\} d\Gamma - \int \{g, f_0\} \{g', f'_0\} \frac{Gmm'}{|\mathbf{r} - \mathbf{r}'|} d\Gamma' d\Gamma,
 \end{aligned}$$

ce qui nous donne finalement

$$\delta^2 \mathcal{H} = -\frac{\epsilon^2}{2} \int \{g, f_0\} \{g, E\} d\Gamma - \frac{\epsilon^2}{2} \int \{g, f_0\} \{g', f'_0\} \frac{Gmm'}{|\mathbf{r} - \mathbf{r}'|} d\Gamma' d\Gamma. \quad (1.83)$$

Nous avons réduit le problème de l'étude de la stabilité des solutions stationnaires de l'équation de Boltzmann sans collisions à celui du signe de l'expression (1.83). Pour conclure ce chapitre, nous allons maintenant exposer les principaux résultats de stabilités effectivement obtenus.

Résultats de stabilité

Le premier résultat de stabilité que nous pouvons énoncer est le suivant : si la variation $\delta^2 \mathcal{H}$, donnée par l'expression (1.83), est positive pour tous les générateurs g , alors f_0 est stable contre toutes les perturbations. Ceci découle directement du théorème d'Arnold. Mais si cette variation n'est pas toujours positive, nous ne pouvons a priori pas conclure.

Des résultats plus précis peuvent être démontrés sous certaines hypothèses.

Ainsi, si f_0 est une fonction de l'énergie E seule, f_0 est stable contre toutes les perturbations à la seule condition que

$$\frac{\partial f_0}{\partial E} < 0.$$

Ce résultat est valide dans le cas d'un système isolé, et peut-être étendu à celui d'un système soumis à un potentiel gravitationnel externe, pourvu que la densité externe associée à ce potentiel ne croisse pas avec le rayon. Ce résultat a été établi par J.Perez et J.-J.Aly [47]. De plus, lorsque f_0 ne dépend que de l'énergie, il y a "équivalence" entre le signe de $\delta^2 \mathcal{H}$ et la stabilité du système, i.e. le système

est stable si $\delta^2\mathcal{H} > 0$ et il est instable si $\delta^2\mathcal{H} < 0$. Cette équivalence provient du fait que $\delta^2\mathcal{H}$ peut s'écrire, dans ce cas, sous la forme

$$\delta^2\mathcal{H} = - \int h \mathcal{T} h d\Gamma, \quad (1.84)$$

où $\mathcal{T}$ est un opérateur symétrique et $h = \{g, f_0\}$ (pour plus de détails, voir les notes de cours de H.Kandrup¹⁰ [26] et les articles de G.Laval, C.Mercier et R.Pellat [35] et de R.M.Kulsrud et J.W.-K.Mark [33]).

Si f_0 est une fonction de l'énergie E et du carré du moment cinétique J^2 , alors, à la condition que

$$\frac{\partial f_0}{\partial E} < 0 \text{ et } \frac{\partial f_0}{\partial J^2} < 0,$$

le système est stable contre les perturbations, dites *préservantes*, engendrées par des fonctions g telles que

$$\{g, J^2\} = 0. \quad (1.85)$$

L'ensemble des perturbations vérifiant (1.85) contient toutes les perturbations radiales, et certaines

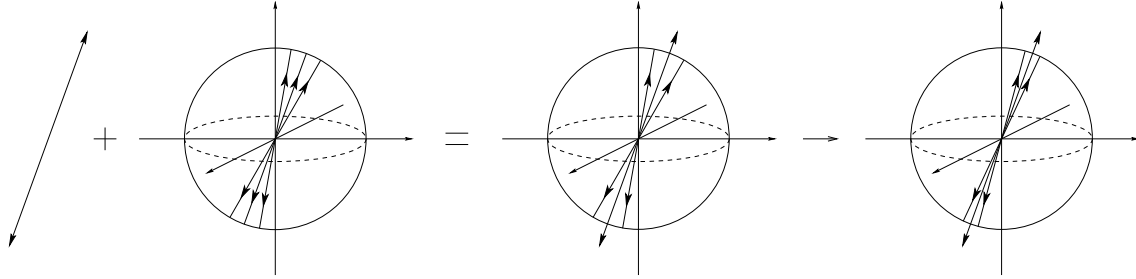


FIG. 1.4 – Mécanisme de développement de l'instabilité d'orbites radiales.

perturbations non radiales. Ce résultat, également dû à J.Perez et J.-J.Aly [47], est une généralisation du travail de H.Kandrup et J.-F.Sygnet [27]. Nous pouvons également noter que lorsque les perturbations sont radiales la variation peut être écrite sous la forme (1.84), ce qui implique de nouveau l'équivalence entre le signe de la variation et la stabilité.

Si f_0 a la forme particulière suivante $f_0 = \zeta(E) \delta_D(J^2)$, c'est-à-dire que les orbites des particules sont purement radiales¹¹, et que ζ est une fonction monotone de E , alors on peut montrer que l'ensemble des générateurs tels que $\{g, J^2\} = 0$ est vide. Il n'existe donc pas de perturbations préservantes pour ce type de fonction de distribution. Il est de plus possible de montrer qu'il existe des perturbations pour lesquelles $\delta^2\mathcal{H} < 0$. Cette instabilité est connue sous le nom d'instabilité d'orbites radiales. Elle se développe lorsque une perturbation non radiale dévie certaines orbites, créant une zone de surdensité, qui par effet d'avalanche, va piéger une à une les autres orbites et ainsi conduire à la formation d'une barre (voir la figure 1.4).

Lorsque f_0 ne dépend pas seulement de E , ou pas seulement de E et J^2 , l'étude de la stabilité est plus compliquée. Le premier résultat que nous avons énoncé – si $\forall g \delta^2\mathcal{H} > 0$, alors le système est stable – entre dans ce cadre général. Dans le cas où le signe de $\delta^2\mathcal{H}$ serait négatif, ou non défini, la situation n'est pas très claire. A.M.Bloch, P.S.Krischnaprasad, J.E.Marsden et T.S.Ratiu [10] ont montré que le système est instable si la perturbation dissipe de l'énergie. L'instabilité est alors qualifiée

¹⁰Ces notes de cours sont disponibles sur la page internet de l'Université de Floride suivante : <http://www.astro.ufl.edu/~siopis/papers/>

¹¹Rappelons que $J^2 = r^2 (v_\varphi^2 + v_\theta^2)$ en coordonnées sphériques. La seule valeur de J^2 acceptée est 0. Ainsi, la vitesse tangentielle de chaque particule est nulle. Les particules possèdent donc une vitesse purement radiale.

de séculaire. H.Kandrup [25] a étudié la stabilité de systèmes axisymétriques en rotation autour de l'axe Oz avec $f_0 = f_0(E, J_z)$, où J_z est le moment cinétique selon z . Il a pu montrer que, dans le cas où $\partial f_0 / \partial J_z$ ne s'annule pas lorsque $J_z = 0$, il existe toujours des perturbations non axisymétriques pour lesquelles $\delta^2 \mathcal{H} < 0$. Physiquement, ceci peut être également interprété comme une sorte d'instabilité d'orbite radiale.

MÉTHODES NUMÉRIQUES

1 Le treecode

Nous avons vu précédemment que l'étude théorique de la dynamique des systèmes gravitationnels à N corps nous fournit de nombreux résultats. Cependant, ces résultats ne nous permettent pas de comprendre tous les phénomènes dynamiques, certains étant trop complexes pour être étudiés correctement d'un point de vue analytique.

Nous devons donc nous en remettre à l'utilisation de simulations numériques afin de compléter nos études analytiques. Plusieurs types de méthodes numériques sont utilisés pour simuler l'évolution des systèmes gravitationnels, chacune possédant un certain nombre d'avantages et d'inconvénients. Le choix d'une méthode doit donc être conditionné par l'adéquation entre les caractéristiques de cette méthode et la nature du problème étudié.

Dans un premier temps, nous présentons succinctement différents types de méthodes numériques. Puis nous détaillons la méthode que nous avons choisi d'utiliser, le *treecode*, ainsi que les raisons de ce choix. Nous présentons enfin en quelques mots la parallélisation du treecode.

1.1 Quelques méthodes numériques et leurs caractéristiques

Les méthodes numériques présentées ci-après sont caractérisées par la façon dont elles calculent l'accélération qui s'exerce sur chaque particule du système dont l'évolution est simulée. Nous nous limitons donc à la description de l'algorithme de calcul de l'accélération, et omettons de parler de l'avancement des particules dans le temps.

Plusieurs schémas numériques d'intégration en temps différents peuvent être associés à ces algorithmes de calcul de l'accélération, le plus simple étant probablement le schéma d'ordre 2 “saute-mouton”, ou “leap-frog” (voir l'annexe A2). Le schéma “saute-mouton” est également l'un des plus largement utilisés, mais certaines méthodes sont associées à des schémas d'intégration en temps plus sophistiqués – tels que les schémas à pas de temps individuels utilisés dans les codes à N corps direct écrits par S.J.Aarseth¹ – permettant un gain de performance ou un meilleur traitement de phénomènes fins, comme les étoiles binaires.

Codes à N corps directs (P^2)

Les codes de type N corps directs – ou particule-particule (P^2) – calculent l'interaction gravitationnelle entre chacune des N particules en sommant les interactions à deux corps. Par son gradient, le

¹Voir [1] pour des explications concernant les pas de temps individuels ou [2] pour la présentation de l'un de ses codes.

potentiel permet d'obtenir l'accélération subie par chaque particule. Pour éviter que cette accélération n'atteigne des valeurs trop élevées lorsque deux particules se trouvent très proches l'une de l'autre, ce code utilise un potentiel adouci. L'expression du potentiel au point $\mathbf{r}_i$, position de la i -ème particule, est ainsi

$$\psi(\mathbf{r}_i) = - \sum_{j=1}^N \frac{Gm_j}{\left((\mathbf{r}_i - \mathbf{r}_j)^2 + \varepsilon^2 \right)^{1/2}},$$

où $\varepsilon > 0$ est le paramètre d'adoucissement du potentiel.

Le calcul de toutes les accélérations s'effectue en un nombre d'opérations proportionnel à N^2 à chaque pas de temps (si le pas de temps est identique pour toutes les particules). Ce temps de calcul représente le principal inconvénient des codes à N corps directs, mais ces codes permettent d'effectuer des simulations d'une grande précision, la seule approximation introduite étant due à ε .

Codes particule-grille (PM) et P³M

Les codes particule-grille (ou PM pour l'anglais particle-mesh) se différencient des codes particule-particule (tels que les codes à N corps directs et le treecode) en divisant l'espace à l'aide d'une grille et en calculant les interactions entre les différentes cellules de cette grille plutôt qu'entre les différentes particules constituant le système. Il n'est donc plus nécessaire de calculer la distance séparant chaque particule, ce qui représente un avantage en terme de temps de calcul.

L'espace est donc divisé en J^3 cellules cubiques, de côté a , et repérées par les indices i, j et k . La masse contenue dans une cellule est supposée se trouver en son centre. Le potentiel exercé sur la cellule (i, j, k) est ainsi égal à

$$\psi_{i,j,k} = \sum_{i'=1}^J \sum_{j'=1}^J \sum_{k'=1}^J \mathcal{D}(i, j, k, i', j', k') \mathcal{M}(i', j', k'), \quad (2.1)$$

où

$$\mathcal{D}(i, j, k, i', j', k') = - \frac{G}{a \left[(i - i')^2 + (j - j')^2 + (k - k')^2 + \varepsilon^2/a^2 \right]^{1/2}},$$

et $\mathcal{M}(i, j, k)$ la masse contenue dans la cellule d'indices (i, j, k) . Notons que cette fonction $\mathcal{D}$, qui joue le rôle de fonction de Green, ne doit être évaluée qu'une fois au début de la simulation.

L'équation (2.1) est une convolution discrète. Il est donc possible de remplacer cette convolution par un produit en travaillant dans l'espace de Fourier. En effet, la transformée de Fourier d'une convolution est égale au produit des transformées de Fourier, *i.e.*

$$\tilde{\psi} = \frac{1}{J^{3/2}} \tilde{\mathcal{D}} \tilde{\mathcal{M}}. \quad (2.2)$$

$\tilde{\mathcal{D}}$ ne doit également être évalué qu'une seule fois.

L'algorithme de transformée de Fourier rapide (ou FFT pour l'anglais Fast Fourier Transform) permet d'effectuer la transformée de Fourier discrète de N points en $N \log N$ opérations. Il est donc possible d'obtenir le potentiel exercé sur chaque grille en $J^3 \log J$ opérations à chaque pas de temps, ce qui représente un gain conséquent si on compare ce nombre aux J^6 opérations nécessaires si l'on n'utilise pas l'espace de Fourier. À ce temps de calcul, il convient d'ajouter le temps nécessaire à l'obtention de $\mathcal{M}$.

Cette méthode possède cependant un sérieux inconvénient. Le potentiel (2.1) n'est une bonne approximation que si la densité de masse dans les cellules voisines est à peu près constante. Dans le cas contraire, les contrastes de densité présents dans chaque cellule ne sont pas pris en compte dans le

calcul du potentiel, et le potentiel (2.1) ne représente plus une approximation acceptable. La densité ne peut pas être constante dans les cellules voisines lorsque le système étudié présente de fortes inhomogénéités, puisque le nombre de cellules est limité. Cette méthode ne semble donc adaptée qu'à l'étude des systèmes homogènes ou quasi homogènes.

Des évolutions de la méthode PM, les codes P³M et les grilles adaptatives, permettent de résoudre ce problème. Les codes P³M sont des hybrides de codes PM et particule-particule. Ils calculent le potentiel à l'aide d'une grille sur les grandes échelles et à l'aide des interactions particule-particule sur les petites échelles, c'est-à-dire pour les cellules voisines. L'inconvénient des codes P³M réside dans le temps de calcul parfois prohibitif de l'interaction sur les petites échelles, lorsque le système comporte une forte surdensité de matière. Les grilles adaptatives permettent quant à elles de diviser les cellules contenant beaucoup de matière, et donc de réduire le coup de calcul de la partie particule-particule d'une code P³M.

Développement en harmoniques sphériques

Cette catégorie de codes, décrite entre autres par T.S.van Albada et J.H.van Gorkom [55] repose sur le développement du potentiel en harmoniques sphériques. Si la particule i est repérée par les coordonnées sphériques $(r_i, \theta_i, \varphi_i)$, alors le potentiel créé au point de coordonnées (r, θ, φ) par cette particule est égal à (voir par exemple [9] p.62-67 pour la démonstration)

$$\psi(r, \theta, \varphi) = -4\pi G \sum_{l=0}^{\infty} \frac{r^l}{r_i^{l+1}} \sum_{m=-l}^l \tilde{\sigma}_{lm}(i) Y_l^m(\theta, \varphi) \quad \text{si } r < r_i, \quad (2.3)$$

$$\psi(r, \theta, \varphi) = -4\pi G \sum_{l=0}^{\infty} \frac{r_i^{l+2}}{r^{l+1}} \sum_{m=-l}^l \tilde{\sigma}_{lm}(i) Y_l^m(\theta, \varphi) \quad \text{si } r > r_i, \quad (2.4)$$

avec

$$\tilde{\sigma}_{lm}(i) = \frac{m_i}{4\pi(2l+1)r_i^2} Y_l^{m*}(\theta_i, \varphi_i), \quad (2.5)$$

où Y_l^m est l'harmonique sphérique et Y_l^{m*} le complexe conjugué de Y_l^m . L'algorithme permettant d'exploiter ce développement est dès lors relativement clair. Il suffit de déterminer dans un premier temps le point du système correspondant au minimum du potentiel gravitationnel – ce point est le centre du système si ce système présente une symétrie sphérique – et d'en faire l'origine de notre repère sphérique. Puis nous pouvons calculer les $L+1$ coefficients $\tilde{\sigma}_{lm}$ pour chacune des N particules, avec $l < L$. Ceci nous permet de calculer une valeur approchée du potentiel en tout point

$$\tilde{\psi}(r, \theta, \varphi) = -4\pi G \sum_{l=0}^L \sum_{m=-l}^l \left[r^l \sum_{r < r_i} \frac{\tilde{\sigma}_{lm}(i)}{r_i^{l+1}} + \frac{1}{r^{l+1}} \sum_{r_i > r} r_i^{l+2} \tilde{\sigma}_{lm}(i) \right] Y_l^m(\theta, \varphi). \quad (2.6)$$

Lorsque L tend vers l'infini, cette valeur approchée $\tilde{\psi}$ tend vers le véritable potentiel ψ . Mais il suffit d'un L relativement faible (quelques unités) pour obtenir une approximation satisfaisante du potentiel. L'utilisation optimale de cet algorithme nécessite un ordonnancement des particules par r_i croissant, mais ce tri s'effectue rapidement entre chaque pas de temps pour peu que le pas de temps soit suffisamment faible, puisque dans ce cas l'ordre des particules dans la liste ne varie que très peu. Lorsque les particules sont triées, le calcul du potentiel s'effectue en $O(N)$ opérations (voir [9] ou [22]). D'autre part, le calcul des coefficients $\tilde{\sigma}_{lm}$ s'effectue en $O(N(L+1)^2)$ opérations. Le temps de calcul pour chaque pas de temps est donc proportionnel au nombre de particules.

Ce type de codes se révèle très véloce, comparativement aux autres, pour des systèmes comprenant beaucoup de particules. Il présente également l'avantage de posséder une excellente définition au centre.

Son utilisation doit cependant être réservée à la simulation de systèmes présentant une symétrie sphérique². De plus, d'après L.Hernquist et J.E.Barnes [22], le fait que ce type de code soit “non collisionnel” n’a jamais été prouvé.

1.2 Description du treecode

Pour atteindre l’objectif que nous nous sommes fixé, nous avons besoin d’un code de calcul relativement rapide, puisque nous avons l’intention d’effectuer de nombreuses simulations et que certaines de ces simulations nécessitent de plus de nombreux pas de temps afin de décrire de façon satisfaisante la dynamique des objets que nous modélisons. Ce code doit pouvoir s’adapter à des géométries variées, de manière à pouvoir simuler correctement aussi bien la dynamique d’objets ressemblant à des galaxies elliptiques de type E7 que celle d’objets sphériques. Il doit également pouvoir gérer des types de profils de densité variés, en particuliers des densités comportant de fortes inhomogénéités. Pour l’ensemble de ces raisons, nous avons choisi d’utiliser un treecode, ou code en arbre.

Principe du treecode

Le treecode, présenté par J.E.Barnes et P.Hut [6] repose sur une idée très simple permettant d’envisager un calcul approché du potentiel subi par chacune des particules. Cette idée se comprend aisément si l’on considère l’exemple suivant. Envisageons le calcul du potentiel gravitationnel exercé sur la Terre par tous les objets visibles dans un ciel nocturne, plus, bien évidemment, le Soleil. Chacun des points lumineux observés contribuera au potentiel. Mais, alors que certains de ces points se révèlent être des planètes du système solaire ou bien des étoiles de la Voie Lactée (la Lune étant suffisamment visible pour ne pas être confondue avec quoi que ce soit d’autre), d’autres sont des galaxies extérieures à la nôtre. Cependant, leur distance par rapport à la Terre est telle qu’elles nous apparaissent peu différentes d’une étoile ou une planète. Cette simple constatation peut nous amener à faire l’approximation suivante : un ensemble d’étoiles, même s’il est de très grande taille, peut être considéré comme ponctuel s’il est suffisamment éloigné de la Terre. Le critère d’acceptation n’est alors pas la distance séparant cet ensemble de la Terre, mais le rapport entre la taille de l’ensemble et cette distance (voir la figure 2.1). Cette approximation permet, par exemple, de réduire le calcul du potentiel exercé par une galaxie très éloignée à un seul terme, en regroupant toute la masse de la galaxie en son centre de masse.

Pour pouvoir exploiter cette idée, il est nécessaire de diviser l’espace en cellules. Une cellule cubique contenant toute les particules est d’abord créée, puis cette cellule, que l’on appellera dorénavant racine, est divisée en huit sous-cellules cubiques de même taille. Chacune de ces sous-cellules est de nouveau divisée en huit sous-cellules, et ce processus se répète jusqu’à ce que chaque cellule élémentaire ne contienne au plus qu’une particule (voir la figure 2.2).

Lors du calcul du potentiel en un point de l’espace, il suffit de calculer la distance séparant chaque cellule du point considéré, puis de comparer le rapport $\gamma_T = \text{taille de la cellule} / \text{distance}$ de chaque cellule avec un critère d’acceptation Θ_T défini au préalable, en commençant par la cellule racine. Si $\gamma_T < \Theta_T$, alors le contenu de la cellule est ramené à une masse ponctuelle, égale à la somme des masses des particules contenues dans la cellule, située au centre de masse de la cellule. Si $\gamma_T \geq \Theta_T$, alors l’approximation n’est pas acceptée, et le contenu de la cellule doit être examiné plus en détail, c’est-à-dire que l’on répète l’opération “calcul du rapport-comparaison à Θ_T ” sur ses sous-cellules.

Cette approximation permet de considérablement réduire le nombre d’opérations à effectuer lors du calcul du potentiel, le faisant passer de $O(N^2)$ à $O(N \log(N))$, où N est le nombre de particules, pour une valeur de Θ_T standard ($\Theta_T \simeq 0.7$).

Le code structure les données (les positions des particules) sous la forme d’un arbre afin de gérer

²Il existe des équivalents reposant sur des expansions dans d’autres systèmes de coordonnées, cylindriques par exemple.

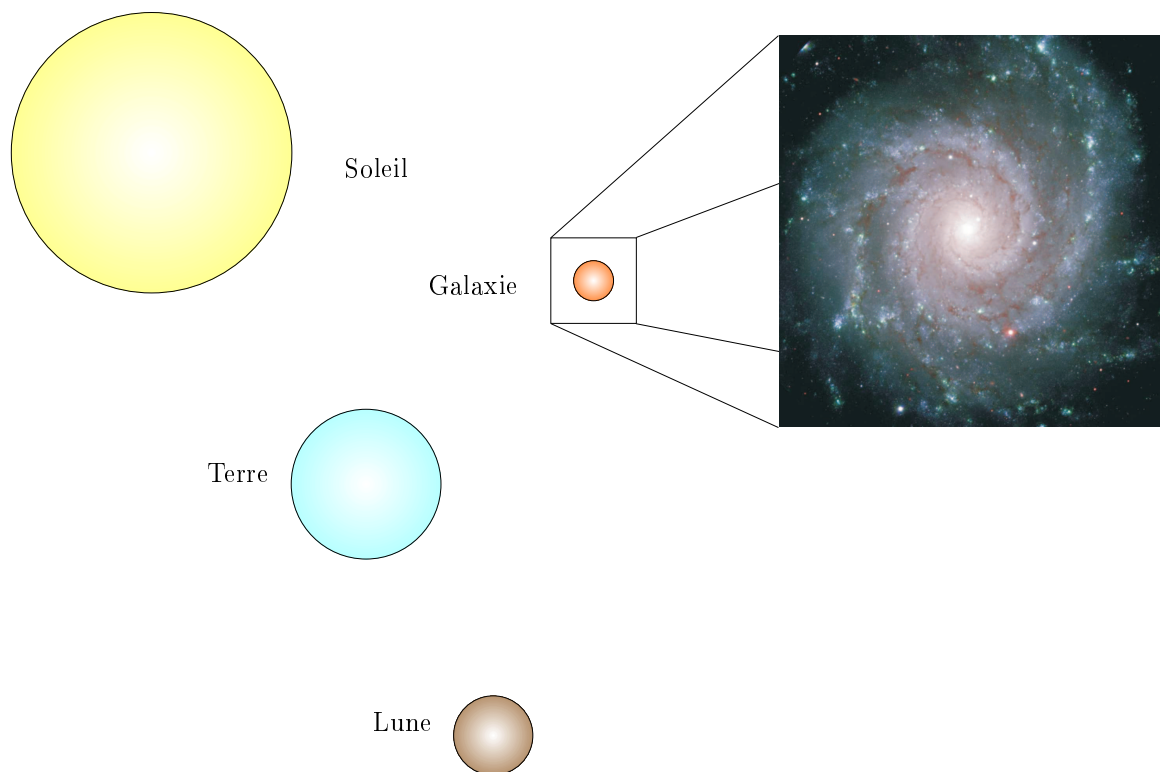


FIG. 2.1 – Principe à la base du treecode.

le plus efficacement possible la division de l'espace en cellules. La construction de cet arbre et son utilisation sont assez instinctive, même si leur mise en œuvre l'est beaucoup moins.

L'arbre et sa construction

Il est nécessaire d'examiner les cellules dans l'ordre décroissant de leur taille afin de réduire au maximum le nombre de termes à inclure dans le calcul du potentiel. Afin de faciliter cet examen, les cellules sont ordonnées et structurées sous la forme d'un arbre.

La construction de l'arbre s'opère de la manière suivante. La cellule racine est aussi la racine de l'arbre, et pour chaque sous-cellule de la racine, une branche est créée. Cette branche se termine par une feuille si la cellule correspondante ne contient qu'une particule, et elle est elle-même divisée en branches si cette cellule en contient plusieurs. Les branches ne contenant pas de particules sont coupées³. À la fin de la construction, on trouve une feuille au bout de chaque branche. Une cellule est donc un nœud de branches lorsqu'elle contient plusieurs particules et une feuille lorsqu'elle n'en contient qu'une.

Le treecode organise le parcours de cet arbre en partant de la racine, puis en descendant dans chaque branche, en ordonnant les cellules dans une liste. Référons-nous à la figure 2.2, représentant un exemple en deux dimensions et une numérotation des cellules un peu différente de celle utilisée dans le code mais plus "parlante", afin de mieux comprendre cet ordonnancement.

La première cellule est donc la racine, et puis la suivante est son premier descendant, le nœud 1. Vient ensuite le premier descendant du nœud 1, le nœud 1.1, puis son premier descendant, la feuille 1.1.1. Comme une feuille n'a pas de descendants, on passe ensuite au deuxième descendant du nœud 1.1, à

³Elles ne sont en réalité pas créées.

savoir le nœud 1.1.3. Ses descendants, la feuille 1.1.3.1 puis la feuille 1.1.3.4 sont alors les éléments suivants de la liste. Tous les descendants du nœud 1.1 ont maintenant été ordonnés, le nœud 1.3 vient ensuite, puis ses descendants, en opérant suivant le même ordre que précédemment. Le parcours de

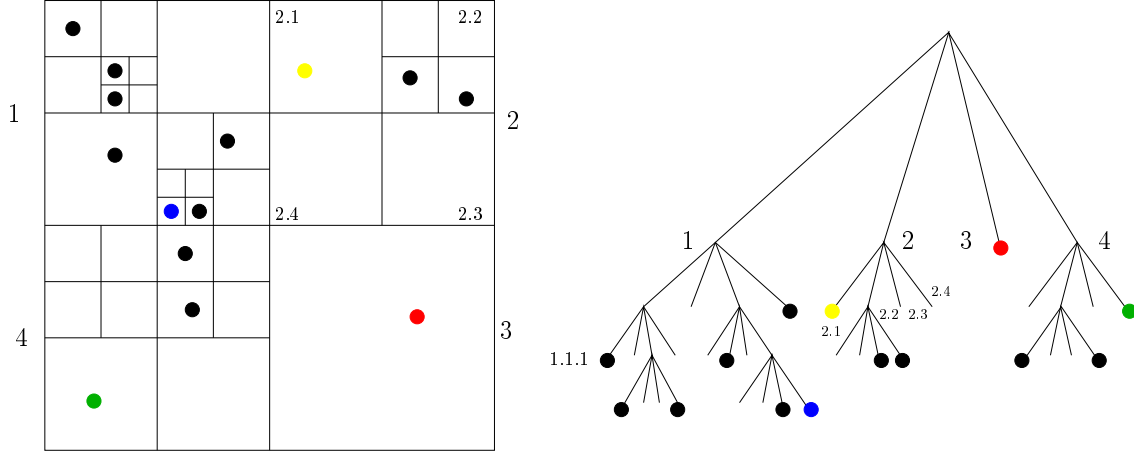


FIG. 2.2 – Construction des cellules et de l'arbre.

l'arbre selon cet ordre est permis par la structure de données. À chaque cellule correspond, pour le code, un ensemble de paramètres, dont un pointeur vers la cellule suivante dans l'ordre de parcours, et un pointeur vers chacun des descendants de la cellule. Par exemple, les données correspondant à la cellule 1.1 contiennent un pointeur vers la cellule 1.3, et un pointeur vers les données des cellules 1.1.1 et 1.1.3. Ces valeurs sont initialisées lors de la construction de l'arbre, et sont utilisées lors de son parcours, pendant le calcul du potentiel.

Deux paramètres importants : Θ_T et ε

Deux paramètres, pouvant être modifiés pour chaque simulation, influencent considérablement le calcul du potentiel. Ces paramètres sont le critère d'acceptation du treecode Θ_T et le paramètre d'adoucissement du potentiel ε .

Si Θ_T est nul, alors aucune approximation n'est faite et le calcul effectué est un calcul direct du potentiel. Le calcul est alors plus précis, mais aussi plus lent que pour une valeur de Θ_T non nulle⁴. À l'inverse, si la valeur de Θ_T est trop grande, beaucoup d'approximations seront faites dans le calcul du potentiel, donc le temps de calcul se retrouvera grandement réduit, mais l'erreur engendrée par ces approximations risque de se révéler très importante. J.E.Barnes et P.Hut [7] ont effectué une analyse des erreurs engendrées par le treecode. Cette étude présente l'influence de Θ_T sur l'erreur commise sur le calcul du potentiel et sur le nombre totales d'interactions prises en compte lors du calcul du potentiel, ce qui conditionne la rapidité du calcul. Selon la forme de la densité de masse du système étudié, l'erreur commise lors du calcul du potentiel augmente d'environ un ordre de grandeur lorsque valeur Θ_T passe de 0.3 à 1. Dans le même temps, le nombre d'interactions baisse de un demi à un ordre de grandeur. Le choix de Θ_T représente donc un choix, ou plutôt un compromis, entre la précision du calcul et sa rapidité.

Le paramètre d'adoucissement du potentiel ε est utilisé pour modéliser le potentiel gravitationnel sur de petites distances. En effet, le potentiel créé par une source en un point est proportionnel à l'inverse

⁴Utiliser le treecode pour un calcul direct est très inefficace, puisque de nombreux calculs annexes nécessaires à l'approximation mais inutiles au calcul direct sont tout de même effectués. Il existe de plus quelques codes directs très performants.

de la distance séparant la source et le point. Si cette distance est très petite, alors le potentiel est très grand. Pour éviter des valeurs de potentiel si grandes que les outils numériques ne pourraient pas les traiter correctement⁵, nous approchons le potentiel gravitationnel créé par la particule i en $\mathbf{r}$ de la façon suivante :

$$\psi(\mathbf{r}) = \begin{cases} -\frac{Gm_i}{|\mathbf{r} - \mathbf{r}_i|} & \text{si } |\mathbf{r} - \mathbf{r}_i| > \varepsilon \\ -\frac{Gm_i}{\varepsilon} & \text{si } |\mathbf{r} - \mathbf{r}_i| \leq \varepsilon \end{cases} \quad (2.7)$$

Cette approximation impose une limite supérieure au potentiel exercé par une particule sur une autre. Elle est différente d'une autre approximation assez largement employée, qui consiste à calculer le potentiel de la manière suivante :

$$\psi(\mathbf{r}) = -\frac{Gm_i}{\sqrt{|\mathbf{r} - \mathbf{r}_i|^2 + \varepsilon^2}} \quad (2.8)$$

Dans ce cas, tous les termes participant au calcul du potentiel sont inexacts, alors que dans notre cas seuls les termes faisant intervenir une distance $|\mathbf{r} - \mathbf{r}_i|$ inférieure à ε le sont.

Calcul du potentiel et moments quadripolaires

Dans un certain nombre de cas – que nous espérons le plus grand possible, afin de diminuer le temps de calcul – le treecode calcule une valeur approchée du potentiel exercé par un ensemble de n particules de masse m_j repérées par les positions $\mathbf{r}_j$, et contenues dans une cellule cubique, sur une autre particule i , située en $\mathbf{r}_i$. Cette valeur, en première approximation, est égale au potentiel qu'aurait exercé une particule cm de masse m_{cm} située à la position $\mathbf{r}_{cm}$, avec

$$m_{cm} = \sum_{j=1}^n m_j \quad (2.9)$$

et

$$\mathbf{r}_{cm} = \frac{1}{m_{cm}} \sum_{j=1}^n m_j \mathbf{r}_j . \quad (2.10)$$

Cette première approximation consiste donc à considérer que la masse totale des n particules est située au centre de masse de ces particules.

Il est souhaitable d'affiner l'approximation en utilisant un développement multipolaire du potentiel exercé par les n particules. Pour cela, plaçons nous dans un repère centré sur $\mathbf{r}_{cm}$ (voir la figure 2.3). Les n particules sont repérées par les vecteurs $\mathbf{x}_j$ et la particule i par $\mathbf{x}_i$. La valeur exacte du potentiel est égale à

$$\psi(\mathbf{x}_i) = -G \sum_{j=1}^n \frac{m_j}{|\mathbf{x}_i - \mathbf{x}_j|} . \quad (2.11)$$

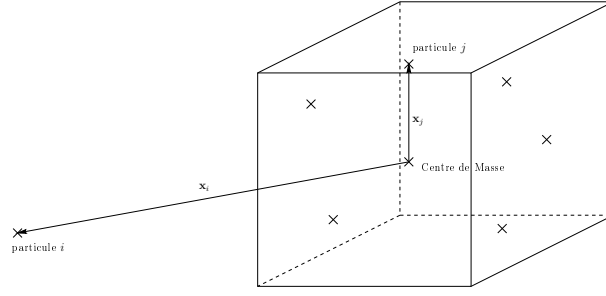
Le développement de $1/|\mathbf{x}_i - \mathbf{x}_j|$ en polynômes de Legendre (voir [23], p.85-94) donne

$$\frac{1}{|\mathbf{x}_i - \mathbf{x}_j|} = \frac{1}{x_i} \left[P_0(u) + \frac{x_j}{x_i} P_1(u) + \left(\frac{x_j}{x_i} \right)^2 P_2(u) + \dots \right] , \quad (2.12)$$

avec

$$u = \frac{\mathbf{x}_i \cdot \mathbf{x}_j}{x_i x_j} , \quad (2.13)$$

⁵L'utilisation du potentiel gravitationnel non adouci imposerait l'utilisation de pas de temps extrêmement faibles et d'un schéma numérique d'intégration en temps plus précis que le schéma "saute-mouton" que nous utilisons.

FIG. 2.3 – Exemple de calcul du potentiel exercé sur la particule i .

où $x = |\mathbf{x}|$. Si nous ne considérons que les trois premiers termes du développement, nous obtenons donc pour le potentiel l'expression suivante

$$\psi(\mathbf{x}_i) = -\frac{G}{x_i} \sum_{j=1}^n m_j \left(P_0(u) + \frac{x_j}{x_i} P_1(u) + \left(\frac{x_j}{x_i} \right)^2 P_2(u) \right). \quad (2.14)$$

Les trois premiers polynômes de Legendre ont pour valeur

$$\begin{aligned} P_0(u) &= 1, \\ P_1(u) &= u, \\ P_2(u) &= \frac{3}{2} \left(u^2 - \frac{1}{3} \right). \end{aligned}$$

Le potentiel est alors égal à

$$\psi(\mathbf{x}_i) = -\frac{G}{x_i} \sum_{j=1}^n m_j - \frac{G}{x_i^3} \mathbf{x}_i \cdot \sum_{j=1}^n m_j \mathbf{x}_j - \frac{G}{x_i^5} \sum_{j=1}^n m_j \left[\frac{3}{2} (\mathbf{x}_i \cdot \mathbf{x}_j)^2 - \frac{1}{2} x_i^2 x_j^2 \right]. \quad (2.15)$$

Le premier terme correspond à l'approximation initialement faite et consistant à placer toute la masse au centre de masse. Les deuxième et troisième termes correspondent aux termes dipolaire et quadrupolaire.

Nous pouvons remarquer que le terme dipolaire est nul. En effet, l'origine de notre repère étant situé au centre de masse, la somme des $m_j \mathbf{x}_j$ est nulle.

Le terme quadrupolaire est calculé par le *treecode*. Il permet un calcul plus précis du potentiel et une meilleure intégration du problème. Le calcul de ce terme est optionnel, mais nous avons choisi de l'effectuer pour chacune de nos simulations.

Le *treecode* en cinq étapes

Le calcul du potentiel permet à chaque pas de temps d'obtenir l'accélération de chaque particule, donc sa vitesse et sa position. L'algorithme utilisé pour l'intégration en temps est le schéma "saute-mouton". Ce schéma peut être résumé de la façon suivante :

$$\begin{aligned} \frac{\mathbf{r}_{new} - \mathbf{r}_{old}}{\Delta t} &= \mathbf{v}_{new}, \\ \frac{\mathbf{v}_{new} - \mathbf{v}_{old}}{\Delta t} &= \mathbf{f}(\mathbf{r}_{old}), \end{aligned}$$

où Δt est le pas de temps séparant deux évaluations du potentiel gravitationnel⁶.

Afin de clarifier le fonctionnement du treecode, le plus simple est certainement de présenter les différentes étapes de son exécution. Une fois les conditions initiales et les paramètres de la simulation lus, le code effectue une boucle jusqu'à ce que le temps final soit atteint ou qu'un problème d'exécution soit rencontré. Cette boucle est composée des étapes suivantes :

1. Avancement des vitesses d'un demi pas de temps, puis des positions d'un pas de temps⁶.
2. Construction de l'arbre.
3. Calcul de l'accélération.
4. Avancement des vitesses d'un demi pas de temps.
5. Calcul des observables, écriture des diagnostics d'exécution et des sauvegardes intermédiaires (cette étape n'est pas effectuée à chaque pas de temps).

Le treecode en quatre erreurs

Le treecode repose, comme nous l'avons vu, sur un calcul approché du potentiel gravitationnel. Cette approximation n'est pas la seule source d'erreurs du code, puisqu'il en existe trois autres, dont une que nous n'avons pas encore mentionnée. La deuxième source d'erreurs est l'utilisation d'un potentiel lissé par le paramètre ε , que nous donne (2.7). La troisième est également indirectement due à ε , puisque le calcul de l'accélération subie par la particule i est donné par

$$\dot{\mathbf{v}}_i = \sum_{\substack{j=1 \\ |\mathbf{r}_j - \mathbf{r}_i| > \varepsilon}}^N m_j \frac{\mathbf{r}_i - \mathbf{r}_j}{|\mathbf{r}_j - \mathbf{r}_i|^3} \quad (2.16)$$

et non par

$$\dot{\mathbf{v}}_i = \sum_{j=1}^N m_j \frac{\mathbf{r}_i - \mathbf{r}_j}{|\mathbf{r}_j - \mathbf{r}_i|^3} . \quad (2.17)$$

Tous les termes provenant des particules j telles que $|\mathbf{r}_j - \mathbf{r}_i| < \varepsilon$ sont donc omis. Cette erreur a cependant un effet positif sur l'intégration en temps, comme nous tentons de l'expliquer dans l'annexe A2. Enfin, la dernière erreur provient du schéma "saute-mouton", puisque ce schéma est d'ordre 2 en temps.

Il est important de noter que les deux erreurs induites par ε ne sont pas cumulatives, puisque les accélérations ne sont pas calculées à partir du potentiel mais parallèlement au potentiel. L'erreur due à ε dans le calcul du potentiel ne se répercute donc pas dans le calcul de l'accélération.

L'influence de ε et du pas de temps sur l'erreur commise sur le calcul de l'énergie totale au cours de la simulation fait l'objet du paragraphe 2.1 du chapitre 3. Les résultats obtenus lors de cette étude puis lors des simulations ultérieures tendent à prouver que les approximations effectuées sont acceptables.

1.3 Calcul parallèle et treecode

Calcul parallèle : pourquoi, où et comment ?

Le temps de calcul du potentiel gravitationnel, malgré l'approximation réalisée, est toujours très important. Mais il est possible de le faire baisser de façon spectaculaire en utilisant une version parallélisée du treecode⁷.

La version du treecode que nous avons utilisée est basée sur le treecode de J.E.Barnes et P.Hut et a été

⁶Le schéma "saute-mouton" est expliqué plus largement dans l'annexe A2

⁷Pour une introduction au calcul parallèle, se référer à l'annexe B.

parallélisée par D.Pfenniger (Observatoire de Genève). La parallélisation ne concerne que le calcul du potentiel. Le calcul de l'accélération subie par les N particules de la simulation (l'étape 3 de la boucle principale) est donc partagé entre les différents processeurs (le premier processeur calcule l'accélération subie par les n_1 premières particules, le deuxième l'accélération subie par les n_2 suivantes, ...). Toutes les autres étapes de calcul, telles que l'avancement des particules ou le calcul des diagnostics, sont séquentielles.

Le code de calcul parallèle a été exécuté sur un cluster d'ordinateurs composé d'environ vingt machines bi-processeurs. Ces ordinateurs sont connectés par un réseau Ethernet. La parallélisation est mise en œuvre à l'aide de la bibliothèque MPI (Message Passing Interface). Il s'agit donc ici de calcul à mémoire distribuée, et non partagée, c'est-à-dire que chaque processeur dispose de ses propres ressources mémoires et que l'échange d'information entre processeurs a lieu par l'intermédiaire de messages transitant par le réseau Ethernet.

Le code est écrit sur le modèle maître-esclave. Un programme maître, exécuté sur un seul processeur, lit les données concernant la simulation, les communique aux programmes esclaves, partage le travail de calcul des accélérations entre tous les processeurs, réunit les résultats, avance les particules et calcule les diagnostics. Le programme esclave, exécuté sur les autres processeurs, se contente de recevoir les données, de calculer l'accélération⁸ pour une partie des particules et d'envoyer ce résultat au maître.

Le parallélisme est-il efficace ?

L'accélération relative S_p (aussi appelée speed-up relatif) d'un programme parallèle est définie par le rapport entre le temps $t^{(1)}$ mis par le code pour effectuer un calcul sur un processeur et le temps $t^{(p)}$ mis pour effectuer un calcul sur p processeurs

$$S_p = \frac{t^{(1)}}{t^{(p)}}.$$

Cette accélération est dite linéaire si $S_p = p$ et sous-linéaire si $S_p < p$. Il est très difficile, voire impossible, d'obtenir une accélération linéaire, en particuliers en raison des surcoûts en calcul et en communications impliqués par la parallélisation.

Afin de prévoir la prédisposition d'un algorithme à être parallélisé efficacement, il est utile de s'intéresser à la Loi d'Amdahl. Si l'on appelle γ_{seq} ($0 \leq \gamma_{seq} \leq 1$) la fraction intrinsèquement séquentielle, et donc non parallélisable⁹, d'un algorithme, alors l'accélération S_p sera telle que

$$S_p \leq \frac{1}{\gamma_{seq} + \frac{1-\gamma_{seq}}{p}} \quad (2.18)$$

Cette loi nous dit donc que si un algorithme contient une fraction non parallélisable de 10%, alors l'accélération ne peut être supérieure à 10, et ce quel que soit le nombre de processeurs utilisés. Cette loi est relativement controversée, car basée sur un modèle très simplifié, et ses prédictions ne se révèlent pas forcément exactes. Elle fournit cependant une idée des performances auxquelles on peut s'attendre en parallélisant un algorithme.

En ce qui concerne le treecode, lorsque le nombre de particules contenues dans la simulation est grand, le calcul du potentiel (la partie parallèle) représente très nettement la fraction la plus importante, en

⁸Par calcul de l'accélération, il faut comprendre ici construction de l'arbre puis calcul de l'accélération (les étapes 2 et 3 de la boucle principale décrite dans le paragraphe 1.2). Ceci signifie qu'en pratique des arbres identiques sont construits indépendamment et simultanément (en théorie) par chaque esclave et par le maître. Bien évidemment, il est inutile de construire l'arbre de manière identique sur chaque processeur, il suffirait de le construire dans le programme maître puis de le communiquer aux programmes esclaves. Cependant, le temps occupé par le maître à la construction de l'arbre serait du temps perdu pour les esclaves, et à ceci s'ajouterait le temps consacré à la communication. Il est donc vraisemblablement préférable, en termes de performances, de construire un arbre sur chaque processeur.

⁹Par exemple, les entrées/sorties.

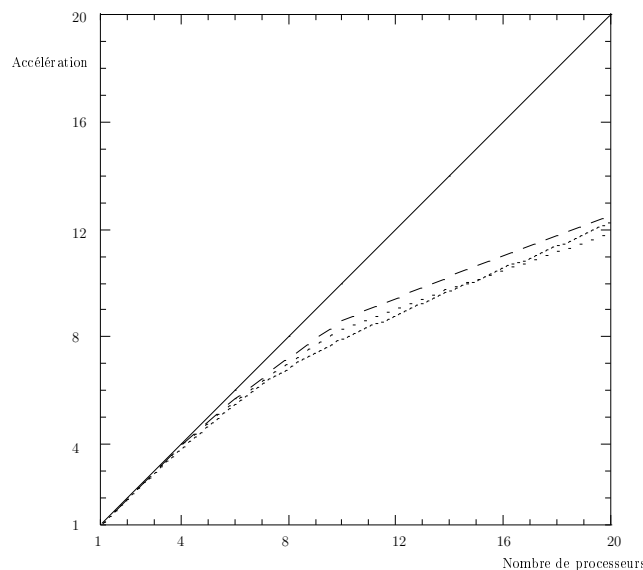


FIG. 2.4 – Accélération en fonction du nombre de processeurs pour des systèmes comprenant 5000, 50000 et 200000 particules.

général au moins 90%, du temps de calcul. L'accélération obtenue est donc satisfaisante, puisqu'elle n'est que légèrement sous-linéaire jusqu'à 10 processeurs (voir la figure 2.4). L'efficacité du parallélisme se dégrade ensuite, mais reste raisonnable ($S_p \approx 12$ pour 20 processeurs).

Ces performances se sont dégradées lorsque des routines de calcul d'observables ont été adjointes au code principal. Ces opérations ont fait augmenter la fraction de calcul séquentiel, ce qui ne pouvait que faire diminuer l'accélération maximale.

2 Conditions initiales

L'un des objectifs de notre travail consiste à analyser l'influence des conditions initiales sur les caractéristiques des systèmes produits à la suite d'un effondrement gravitationnel. Il est donc indispensable que notre générateur nous permette de créer un très large éventail de conditions initiales. Certains des paramètres déterminant la nature du système créé par notre code méritent quelques explications, à la fois concernant leur signification et la manière dont ils sont pris en compte par le générateur de conditions initiales.

2.1 Géométrie du système

La majorité des simulations ont été effectuées à partir de conditions initiales sphériques. Cependant, le code permet de générer un système ellipsoïdal ou cubique, ce qui nous a permis de tester l'influence de la géométrie initiale.

2.2 Taille du système R_{ini}

La taille du système représente la longueur du diamètre, s'il s'agit d'un système sphérique, ou de l'arête dans le cas d'un système cubique, du nuage d'étoiles initial. S'il s'agit d'un système ellipsoïdal, cette taille est la longueur du grand axe.

La taille est exprimée dans le système d'unités propre aux simulations (voir la section 4).

La taille doit être mise en rapport avec le paramètre de lissage du potentiel ε . En effet, ε représente la résolution spatiale du code. Cette résolution n'a de sens qu'en comparaison de la taille du système étudié. Le choix des valeurs de ces deux paramètres doit donc s'effectuer conjointement.

Une étude préliminaire nous a apporté les renseignements nécessaires nous permettant de faire ce choix.

2.3 Profil de densité

Nous avons considéré trois types de densité initiale différents :

- densité homogène
- densité $\rho \propto r^{-\alpha_d}$
- densité de type “clumpy” (qui contient des grumeaux), basée sur les conditions testées par Van Albada [54].

Le choix de ces profils de densité particuliers repose sur notre souci de tester des conditions initiales à la fois homogènes et inhomogènes, avec des inhomogénéités soit globales ($\rho \propto r^{-\alpha_d}$) soit locales (grumeaux).

Densité homogène

Le nombre de particules utilisées lors de nos simulations étant grand, nous aurions pu créer une densité homogène en tirant les positions de ces particules au hasard dans le volume initial. Cependant, en raison d'un problème apparu lors de l'utilisation de notre Treecode avec un grand nombre de particules, nous avons décidé de procéder différemment.

Le Treecode n'arrête la construction de son arbre que lorsque chaque particule se trouve seule dans une cellule (voir section 1.2). Le code doit donc créer des sous-cellules tant que cette condition n'est pas atteinte. Il peut donc arriver que deux positions tirées au hasard soient suffisamment proches l'une de l'autre pour que le code ne puisse pas les séparer sans dépasser le nombre limite de cellules autorisé (ce nombre est fixé par l'utilisateur en fonction des capacités de mémoire de l'ordinateur utilisé). Nous avons rencontré cette situation à de nombreuses reprises avec un nombre de particules de l'ordre de 10^5 , c'est pourquoi nous avons renoncé à tirer les positions au hasard.

La solution alternative que nous avons adoptée est la suivante : nous avons créé une grille cubique de même taille que le système initial. Puis nous avons choisi le pas Δ_x de grille de telle sorte que le nombre N_g de points de grilles soit environ égal à $6/\pi$ fois le nombre de particules souhaité. Le nombre de positions contenues dans le volume de la sphère incluse dans la grille cubique est donc environ égal au nombre de particules souhaité. Une particule est placée à chaque point de grille, puis sa position est modifiée par un bruit $\delta \mathbf{r}$, où $-\zeta \leq \delta r_i \leq \zeta$, pour $i = 1, 2, 3$, ζ étant une fraction de Δ_x . Cette méthode nous a permis de nous affranchir du problème de la construction de l'arbre en nous assurant une distance minimale entre chaque particule (cette distance dépend de la fraction de Δ_x utilisée).

Densité $\rho(r) \propto r^{-\alpha_d}$

Le profil de densité initial en $r^{-\alpha_d}$ est créé en utilisant un générateur de nombres aléatoires non uniforme. En effet, nous cherchons à obtenir une densité de type $\rho(r) \propto r^{-\alpha_d}$. Or, nous savons que

$$\rho(r) = \frac{M_d(r)}{V_d(r)}$$

où $M_d(r)$ est la masse contenue dans l'enveloppe sphérique située entre les rayons r et $r + dr$, et $V_d(r) = 4\pi r^2 dr$ le volume de cette même enveloppe. Pour ce type de condition initiale, toutes les

particules ont une masse identique m . Nous pouvons donc substituer $m n_d(r)$ à $M_d(r)$, où $n_d(r)$ est le nombre de particules contenues dans l'enveloppe. Nous obtenons donc

$$\rho(r) = \frac{m n_d(r)}{4\pi r^2} \propto n_d(r) r^{-2}$$

Pour obtenir une densité proportionnelle à $r^{-\alpha_d}$, il faut donc que le nombre de particules dont le rayon est compris entre r et $r + dr$ soit proportionnel à $r^{2-\alpha_d}$.

Pour obtenir une variable aléatoire x , comprise entre 0 et 1, et distribuée selon une fonction $f(x)$, nous avons utilisé la méthode de transformation inverse (voir l'annexe 1). Ceci nous permet d'obtenir la distance au centre, ou plutôt le rapport r/R , de chaque particule.

Une fois le rayon obtenu, nous générons deux nombres aléatoires uniformes u_2 et u_3 qui correspondront aux coordonnées cylindriques z et φ de la manière suivante :

$$(z, \varphi) = (2u_2 - 1, 2\pi u_3)$$

Nous obtenons ainsi les positions $\mathbf{r}$ de nos particules à l'aide de trois nombres aléatoires uniformes :

$$\mathbf{r} = \begin{bmatrix} r\sqrt{1-z^2} \cos \varphi \\ r\sqrt{1-z^2} \sin \varphi \\ rz \end{bmatrix}$$

Densité de type “clumpy”

Ce type de densité initiale est similaire à celui présenté par T.S.van Albada [54]. Il s'agit de systèmes sphériques homogènes auxquels ont été ajoutées des inhomogénéités. La création de ces systèmes a lieu en deux temps :

1. Création de n_g inhomogénéités
2. “Remplissage” de l'espace restant

Les inhomogénéités – que nous appellerons grumeaux – sont des sphères homogènes, de rayon égal au vingtième du rayon du système. Les n_g grumeaux sont uniformément distribués dans le système initial. Chaque grumeau contient un soixantième du nombre total de particules.

Plus concrètement, nous créons les grumeaux de la même manière que les systèmes de type homogène, puis nous les déplaçons aux positions qui leur ont été allouées. Les positions des divers grumeaux sont testées afin qu'ils ne se recouvrent pas. Des particules sont ensuite ajoutées dans l'espace laissé libre autour des inhomogénéités jusqu'à ce que le nombre total de particules atteigne le nombre voulu.

T.S.van Albada ayant obtenu des résultats intéressants à l'aide de ces conditions initiales, il nous a semblé judicieux de les utiliser. Ce type de densité représente de plus une forme d'inhomogénéité différente du gradient de densité présenté précédemment.

2.4 Rapport du viriel initial $\mathcal{V}_{ini}$

Le rapport du viriel¹⁰ est fixé après que les positions et les vitesses des particules ont été générées. Le rapport du viriel dépend à la fois des positions et des vitesses des particules, puisqu'il est égal à $2K/U$, les expressions de l'énergie cinétique K et de l'énergie potentielle U étant données par (1.29) et (1.34). Il est donc possible d'en choisir la valeur en modifiant les positions ou/et les vitesses de chacune des particules. Cependant, la dépendance du rapport du viriel par rapport aux positions est telle qu'il est évidemment préférable, pour des raisons de simplicité, de modifier les vitesses. De plus, comme nous avons souhaité fixer la taille du système ainsi que son profil de densité initial, nous ne

¹⁰Voir le paragraphe 2.5 du chapitre 1.

pouvons pas modifier les positions.

Une fois les positions et les vitesses générées, un calcul du rapport du viriel, que nous appelons $\mathcal{V}_{pro}$, est effectué. Le rapport du viriel est proportionnel à la somme des carrés des modules des vitesses des particules. Si nous multiplions chacune des composantes des vitesses des particules par $\sqrt{\mathcal{V}_{ini}/\mathcal{V}_{pro}}$, nous obtenons la valeur du rapport du viriel désiré, $\mathcal{V}_{ini}$. En effet, on a :

$$\mathcal{V}_{pro} = \alpha \sum_{i=1}^N m_i (v_{x,i}^2 + v_{y,i}^2 + v_{z,i}^2)$$

Après multiplication des composantes des vitesses par notre facteur correcteur, on obtient :

$$\mathcal{V} = \alpha \sum_{i=1}^N m_i \frac{\mathcal{V}_{ini}}{\mathcal{V}_{pro}} (v_{x,i}^2 + v_{y,i}^2 + v_{z,i}^2) = \mathcal{V}_{ini}$$

2.5 Spectre de masse

Afin de tester l'influence du spectre de masse sur nos simulations, nous avons utilisé trois types de spectres différents, en plus des systèmes où toutes les particules ont une masse identique.

Ces spectres sont définis par le nombre $n(M)dM$ de particules dont la masse est comprise entre M et $M + dM$, M variant entre $M_{inf} = 0.01M_{\odot}$ et $M_{sup} = 50.0M_{\odot}$. Les trois fonctions $n(m)$, où $m = M/M_{\odot}$, considérées sont de type loi de puissance, et sont :

(1) $n(m) \propto m^{-2.35}$

(2) $n(m) \propto m^{-\gamma}$ avec $\begin{cases} \gamma_a = 0.3 \text{ pour } m_{inf} = 0.01 \leq m \leq 0.08 = m_1 \\ \gamma_b = 1.3 \text{ pour } m_1 = 0.08 \leq m \leq 0.50 = m_2 \\ \gamma_c = 2.3 \text{ pour } m_2 = 0.50 \leq m \leq 50.0 = m_{sup} \end{cases}$

(3) $n(m) \propto m^{-1}$

Les spectres (1) et (2) correspondent respectivement aux fonctions de masses initiales introduites par E.E.Salpeter [50] et P.Kroupa [32], alors que le troisième, même s'il n'est pas forcément réaliste, nous permet de mesurer l'influence du spectre dans un cas plus extrême.

Pour créer les spectres de masse, nous procédons donc de la manière suivante. Nous divisons tout d'abord l'intervalle de masses $[m_{inf}, m_{sup}]$ en sous-intervalles $[m_i, m_{i+1}]$. Puis, pour chaque sous-intervalle, nous déterminons le nombre de particules N_i dont la masse sera comprise entre M_i et M_{i+1} . Il ne reste plus ensuite qu'à déterminer chacune de ces N_i masses de la manière suivante : $m = m_i + u \cdot dm$, où u est un nombre aléatoire uniforme compris entre 0 et 1, et dm est la longueur du sous-intervalle.

La façon dont les N_i sont calculés diffère d'un spectre à l'autre. Pour les spectres (1) et (3), nous avons :

$$N_{1 \rightarrow 2} = \int_{m_1}^{m_2} \alpha m^{-\gamma} dm \quad (2.19)$$

ce qui nous donne :

$$N_{1 \rightarrow 2} = \alpha \frac{m_2^{-1.35} - m_1^{-1.35}}{-1.35} \quad \text{pour (1)} \quad (2.20)$$

$$N_{1 \rightarrow 2} = \alpha \ln \left(\frac{m_2}{m_1} \right) \quad \text{pour (3)} \quad (2.21)$$

$N_{1 \rightarrow 2}$ est le nombre de particules dont la masse est comprise entre m_1 et m_2 . Nous pouvons donc calculer le coefficient α , dans les deux cas, en considérant l'intervalle total de masse $[m_{inf}, m_{sup}]$ (puisque $N_{1 \rightarrow 2}$ est alors égal au nombre total N de particules), puis calculer les différents N_i pour chacun des sous-intervalles, en utilisant ces mêmes formules.

Pour le spectre (2), il est tout d'abord nécessaire de déterminer les trois coefficients α_i correspondant à chacun des trois intervalles. Pour ceci, nous disposons de l'équation donnée par le nombre total de particules, à savoir :

$$\int_{m_{inf}}^{m_1} n_a^{(2)}(m) dm + \int_{m_1}^{m_2} n_b^{(2)}(m) dm + \int_{m_2}^{m_{sup}} n_c^{(2)}(m) dm = N$$

Nous disposons également de deux équations de continuité :

$$\begin{aligned} n_a^{(2)}(m_1) &= n_b^{(2)}(m_1) \Leftrightarrow \alpha_a m_1^{\gamma_a} = \alpha_b m_1^{\gamma_b} \\ n_b^{(2)}(m_2) &= n_c^{(2)}(m_2) \Leftrightarrow \alpha_b m_2^{\gamma_b} = \alpha_c m_2^{\gamma_c} \end{aligned}$$

Il est ensuite aisé de déterminer les différents N_i , puis de tirer les différentes masses dans chacun des intervalles.

Une fois les N masses déterminées, elles sont réparties aléatoirement dans le volumes du système initial. Le densité initiale est donc globalement homogène en nombre et en masse, et il n'y a pas de ségrégation de masse.

2.6 Distribution des vitesses

S.Kazantzidis, J.Magorrian et B.Moore [29] montrent qu'il est nécessaire de générer les vitesses à partir de la fonction de distribution f , et que l'approximation maxwellienne n'est pas satisfaisante, pour créer des systèmes à l'équilibre. Mais nous voulons générer des systèmes hors équilibre. Nous ne sommes donc pas tenu de suivre la méthode décrite par ces auteurs. Nous pouvons choisir des distributions de vitesses indépendantes de la densité.

Nous utilisons deux types de fonction de distribution des vitesses initiales : constante¹¹ ou gaussienne. Les distributions de vitesses constantes sont obtenues en générant les trois composantes de la vitesse de chaque particule à l'aide d'un nombre aléatoire uniforme. Plus précisément, chaque composante de la vitesse de la particule i est égale à :

$$v_{\nu,i} = 2 * u - 1 \quad \text{avec } \nu = x, y, z \quad (2.22)$$

où u est nombre aléatoire uniforme sur $[0, 1]$.

Pour la distribution que nous appelons gaussienne, nous avons généré chaque composante de la vitesse de manière aléatoire, en utilisant une variable aléatoire de densité de probabilité $f(x) = \exp(-x^2/\sigma^2)$. Nous avons une nouvelle fois utilisé la méthode de transformation inverse :

$$v_{\nu,i} = \sqrt{-\sigma^2 \ln(u_1)} \cos(2\pi\sigma^2 u_2) \quad \text{avec } \nu = x, y, z \quad (2.23)$$

u_1 et u_2 sont deux variables aléatoires uniformes sur $[0, 1]$. Dans les deux cas, les systèmes sont isotropes. En effet, la fonction de distribution des vitesses ne dépend que des vitesses, et non des positions, alors que la fonction de distribution des positions ne dépend que des positions (dans tous les cas). La fonction de distribution du système dans l'espace des phases est donc le produit de la densité $\rho(\mathbf{r})$ et d'une fonction $f_v(\mathbf{v})$. Par définition, nous avons (voir l'équation (1.39) du chapitre 1)

$$\overline{v_{\nu,i}^2} = \frac{1}{N} \int_{\Omega} v_{\nu,i}^2 \rho(\mathbf{r}) f_v(\mathbf{v}) d^3\mathbf{r} d^3\mathbf{v} . \quad (2.24)$$

Que f_v soit constante ou égale à

$$f_v(\mathbf{v}) = e^{-\frac{v_x^2}{\sigma_x^2}} e^{-\frac{v_y^2}{\sigma_y^2}} e^{-\frac{v_z^2}{\sigma_z^2}} , \quad (2.25)$$

¹¹La fonction de distribution ne dépend pas des vitesses, mais uniquement des positions.

nous retrouvons $\bar{v}_{x,i}^2 = \bar{v}_{y,i}^2 = \bar{v}_{z,i}^2$.

Seules les conditions initiales en rotation (voir ci-dessous) ont une distribution de vitesses non isotrope, puisque ce sont les seuls systèmes où la vitesse d'une particule dépend de la position de cette particule dans le système.

2.7 Rotation

Le paramètre de rotation γ_{rot} correspond à un coefficient de rotation solide appliqué à un système sphérique homogène et isotrope. Nous créons tout d'abord le système homogène et isotrope. Puis nous calculons la valeur moyenne du module de la vitesse des particules :

$$\bar{v} = \frac{1}{N} \sum_{i=1}^N |\mathbf{v}_i| \quad (2.26)$$

La vitesse de rotation ω_{rot} appliquée au système autour de l'axe Oz est égale à $\omega_{rot} = \gamma_{rot}\bar{v}/R$. Cette vitesse de rotation est ajoutée à chaque particule¹² sur la composante angulaire de la vitesse :

$$v_{\varphi,i}^{rot} = v_{\varphi,i} + \omega_{rot}r_i^c, \quad (2.27)$$

où r_i^c est la distance de la particule i à l'axe Oz . Le système tourne donc initialement sur lui-même autour de l'axe Oz . La vitesse de rotation solide étant proportionnelle à $\bar{v}$, elle dépend du rapport du viriel et du profil de densité de l'état initial. La comparaison entre divers systèmes mis en rotation est donc parfois problématique. Toutefois, imposer une vitesse angulaire fixe quel que soit le type de système nous aurait amené à créer des systèmes dans lesquels la rotation aurait représenté un très fort pourcentage de l'énergie cinétique, ce que nous voulions éviter.

La quantité de rotation introduite dans ces systèmes peut également être mesurée à l'aide du rapport μ_K entre l'énergie cinétique de rotation K_{rot} et l'énergie cinétique totale K du système. Nous avons calculé l'énergie cinétique de rotation en utilisant la définition de Navarro et White (voir [42]) :

$$K_{rot} = \frac{1}{2} \sum_{i=1}^N m_i \frac{(\mathbf{J}_i \cdot \hat{\mathbf{J}}_{tot})^2}{\left[r_i^2 - (\mathbf{r}_i \cdot \hat{\mathbf{J}}_{tot})^2 \right]} \quad (2.28)$$

$\mathbf{J}_i$ est le moment angulaire de la particule i , et $\hat{\mathbf{J}}_{tot}$ est un vecteur unitaire dans la direction du moment angulaire total du système.

La correspondance entre μ_K et le paramètre γ_{rot} est difficile à établir. Elle l'est d'autant plus que le rapport du viriel du système initial est fixé après l'ajout de la rotation solide. Or, lorsque le rapport du viriel est fixé, le facteur de multiplicateur appliqué aux vitesses dépend de l'énergie potentielle du système, et donc de son profil de densité.

Ainsi, le rapport entre le paramètre γ_{rot} et la quantité de rotation réellement introduite dans le système est modifié à la fois par le rapport du viriel initial et par le profil de densité choisi.

La valeur de μ_K de chaque système en rotation est indiquée dans le tableau C.8.

2.8 Nomenclature

Pour des raisons pratiques, nous avons décidé d'introduire des notations précises pour désigner les différents types de conditions initiales utilisées lors de nos études.

Sauf précisions de notre part, les conditions initiales ont des caractéristiques communes. Les systèmes

¹²Nous travaillons ici en coordonnées cylindriques.

initiaux sont sphériques, de rayon $R = 10$. Le paramètre de lissage du potentiel utilisé est $\varepsilon = 0.1$, et l'angle d'ouverture ou critère d'acceptation du treecode $\Theta_T = \sqrt{2}/2$ (voir la section 1.2). Le paramètre du viriel η est défini par $\eta = -100 \times \mathcal{V}_{ini}$

Les notations utilisées sont les suivantes :

- H_η : Densité homogène, distribution des vitesses uniforme, masses identiques, paramètre de viriel η
- $C_\eta^{n_g}$: Densité de type grumeaux, avec n_g grumeaux, distribution des vitesses uniforme, masses identiques, paramètre de viriel η
- G_η^σ : Densité homogène, distribution des vitesses de type gaussienne de largeur à mi-hauteur σ , masses identiques, paramètre de viriel η
- M_η^k : Densité homogène, distribution des vitesses uniforme, spectre de masse de type $k = I$ (Kroupa), II (Salpeter) ou III ($1/m$), paramètre de viriel η
- $P_\eta^{\alpha_d}$: Densité selon une loi de puissance $\rho \propto r^{-\alpha_d}$, distribution des vitesses uniforme, masses identiques, paramètre de viriel η
- $R_\eta^{\gamma_{rot}}$: Densité homogène, distribution des vitesses de type rotation avec un paramètre de rotation γ_{rot} , masses identiques, paramètre de viriel η

3 Calcul des observables

L'analyse et l'interprétation des résultats des simulations ne peuvent se faire sans le calcul de paramètres caractérisant le système, et pouvant éventuellement être comparés aux résultats d'observations d'amas globulaires. Parmi ces "observables" auxquelles nous nous sommes intéressés, les suivantes sont calculées lors de la simulations par des fonctions ajoutées au treecode :

- l'énergie totale
- le rapport du viriel
- les rapports d'axe
- les rayons contenant 10, 50 et 90% de la masse

A celles-ci, il convient d'ajouter les observables calculées à partir du résultat de la simulation et des sauvegardes intermédiaires des positions et vitesses des particules, à savoir :

- le potentiel
- la densité
- la dispersion de vitesse

Les raisons du choix de ces observables et des méthodes de calcul employées sont exposées ci-dessous.

3.1 Énergie totale

L'énergie totale du système est un indicateur de la validité de la simulation. En effet, les systèmes étudiés sont isolés et non collisionnels sur les échelles de temps considérées. Leur énergie totale est donc constante au cours du temps.

L'énergie totale est égale à

$$\mathcal{H} = \sum_{i=1}^N K_i + \frac{1}{2} \sum_{i=1}^N U_i ,$$

où K_i et U_i sont respectivement l'énergie cinétique et l'énergie potentielle de la particule i . Le potentiel gravitationnel utilisé pour le calcul de l'énergie potentielle est celui obtenu par l'intermédiaire du treecode.

Le facteur d'erreur le plus important, et celui sur lequel nous agissons pour améliorer la conservation de l'énergie, est l'intégration en temps. Le choix du schéma "saute-mouton" étant définitif, nous ne

pouvons que diminuer le pas de temps utilisé pour gagner en précision. L'erreur commise lors de l'intégration en temps est exprimée par $\Delta\mathcal{H}$, qui est égal à

$$\Delta\mathcal{H} = 100 \frac{|\mathcal{H}_{ini} - \mathcal{H}_{fin}|}{\mathcal{H}_{ini}} .$$

3.2 Rapport du viriel

Le rapport du viriel est le rapport entre l'énergie cinétique et l'énergie potentielle d'un système, ou plus exactement :

$$\mathcal{V} = \frac{2K}{U}$$

D'après le théorème du viriel, le rapport du viriel des systèmes gravitationnels est égal à -1 lorsque ces systèmes sont à l'équilibre. L'énergie cinétique et l'énergie potentielle ont déjà été calculée précédemment, pour obtenir l'énergie totale. En particulier, le code calcule l'énergie potentielle (voir l'équation (1.4) du chapitre 1) à l'aide des potentiels approchés calculés pour chaque particule. Le rapport du viriel est donc immédiatement obtenu.

R.A.Gerber [18] a analysé l'influence de l'adoucissement du potentiel sur le théorème du viriel. Son étude repose sur l'adoucissement (2.8), et présente des résultats pour trois types de fonctions de distributions dont la sphère isotherme¹³. L'application de son résultat à une sphère isotherme de masse $M = 1$ et de rayon $R = 10$ peut nous permettre d'estimer l'erreur introduite dans notre calcul du rapport du viriel par le paramètre d'adoucissement. Pour $\varepsilon = 0.1$, la valeur calculée avec l'adoucissement (2.8) diffère de 0.15% de la valeur calculée sans adoucissement. Il est de plus assez aisé de montrer que notre adoucissement (2.7) introduit dans le calcul du potentiel une erreur inférieure à celle introduite par l'adoucissement (2.8), même si (2.7) est moins régulier que (2.8). En effet, le potentiel créé par l'ensemble des N particules du système sur la particule i est égal à

$$\psi(\mathbf{r}_i) = -G \sum_{\substack{j=1 \\ j \neq i}}^N \frac{m_j}{|\mathbf{r}_j - \mathbf{r}_i|} .$$

Nous pouvons diviser cette somme de la façon suivante

$$\psi(\mathbf{r}_i) = -G \sum_{|\mathbf{r}_j - \mathbf{r}_i| > \varepsilon}^N \frac{m_j}{|\mathbf{r}_j - \mathbf{r}_i|} - G \sum_{|\mathbf{r}_j - \mathbf{r}_i| \leq \varepsilon}^N \frac{m_j}{|\mathbf{r}_j - \mathbf{r}_i|} .$$

Nous remarquons alors que la première somme est exactement calculée avec l'adoucissement (2.7), alors qu'elle est sous-évaluée (en valeur absolue) par l'adoucissement (2.8). Nous pouvons également aisément montrer que la deuxième somme est mieux approchée avec (2.8), puisque, pour ces termes, nous avons

$$0 \leq |\mathbf{r}_j - \mathbf{r}_i| < \varepsilon \text{ donc } |\mathbf{r}_j - \mathbf{r}_i| < \varepsilon \leq \sqrt{(\mathbf{r}_j - \mathbf{r}_i)^2 + \varepsilon^2} ,$$

et ainsi

$$\frac{m_j}{|\mathbf{r}_j - \mathbf{r}_i|} > \frac{m_j}{\varepsilon} \geq \frac{m_j}{\sqrt{(\mathbf{r}_j - \mathbf{r}_i)^2 + \varepsilon^2}} .$$

Le rapport du viriel calculé devrait donc être suffisamment proche du véritable rapport du viriel pour que l'application du théorème du viriel nous permette de distinguer les systèmes à l'équilibre.

¹³Nous verrons par la suite que les systèmes produits par nos simulations d'effondrement sont proches d'une sphère isotherme.

3.3 Centre de densité et rayon de densité

Le rayon de densité, introduit par S.Casertano et P.Hut (voir [15] pour les détails concernant ce qui suit), nous a permis d'obtenir une valeur approchée du rayon du cœur de nos systèmes. Ce rayon de densité est en effet relié aux rayons de cœur utilisés pour les observations.

Pour calculer le rayon de densité¹⁴ $R_{d,j}$, il convient tout d'abord de calculer ce que l'on nomme les densités locales $\rho_j^{(i)}$ d'ordre j autour de chaque particule i (voir (2.29)). Pour cela, il suffit de rechercher les j particules les plus proches de la particule i (voir la figure 2.5). Si $d_j^{(i)}$ est la distance séparant la j^{eme} plus proche particule de la particule i , et m la masse de chaque particule, alors la densité locale autour de i est égale à :

$$\rho_j^{(i)} = \frac{j-1}{V(d_j^{(i)})} m \quad (2.29)$$

où $V(d)$ est le volume de la sphère de rayon d . Ces densités locales nous permettent ensuite de calculer la position $\mathbf{r}_{cd}$ du centre de densité du système, de la manière suivante :

$$\mathbf{r}_{cd} = \frac{\sum_{i=1}^N \mathbf{r}_i \rho_j^{(i)}}{\sum_{i=1}^N \rho_j^{(i)}} \quad (2.30)$$

Le centre de densité représente le centre de l'objet que nous qualifions de lié. Il est en cela différent du centre de masse du système total, puisque le centre de masse est influencé par les particules libres s'échappant du système lié, alors que le centre de densité l'est très peu en raison de la très faible densité locale entourant ces particules libres.

Sa position sera utilisée comme origine pour tous les calculs d'observables suivants. Ainsi, le rayon d'une particule sera calculé par rapport au centre de densité de l'objet.

Le rayon de densité est égal à la moyenne des rayons des particules pondérés par la densité locale autour de ces particules. Les rayons étant calculés par rapport au centre de densité, nous avons :

$$R_{d,j} = \frac{\sum_{i=1}^N |\mathbf{r}_i - \mathbf{r}_{cd}| \rho_j^{(i)}}{\sum_{i=1}^N \rho_j^{(i)}} \quad (2.31)$$

Selon l'analyse menée par S.Casertano et P.Hut, le rayon de densité ainsi défini est un très bon estimateur du rayon du cœur observationnel, tel qu'il est défini par I.R.King [30], c'est-à-dire le rayon pour lequel la luminosité de surface est égale à la moitié de sa valeur centrale. Toujours selon S.Casertano et P.Hut, ce rayon de densité est également relié au rayon du cœur du modèle de Plummer $R_{c,p}$ de la manière suivante :

$$R_d = \frac{32}{15\pi} R_{c,p} . \quad (2.32)$$

Cette relation est obtenue en passant à la limite continue de la définition du rayon de densité, c'est-à-dire la limite $j \rightarrow \infty$ et $N/j \rightarrow \infty$. Ceci explique la disparition de l'indice j , puisque l'estimateur $\rho_j(\mathbf{r})$ de la densité locale a ici été remplacé par sa valeur exacte $\rho(r)$. Nous avons choisi d'utiliser pour j la valeur $j = 10$. Ce choix repose sur les analyses de S.Casertano et P.Hut.

Le calcul de la densité locale $\rho_{10}^{(i)}$ nécessite la connaissance des 10 plus proches voisins de la particule i . Pour cela, il est nécessaire de calculer la distance

$$d_{(i,k)} = |\mathbf{r}_k - \mathbf{r}_i|, \quad k \in \{1, \dots, N\} \setminus \{i\},$$

séparant i de chacune des autres particules k . Nous avons donc choisi de calculer les $d_{(i,k)}$ dans une boucle sur les particules. Nous avons initialisé un vecteur contenant les distances séparant i de ses 10

¹⁴La méthode exposée ici concerne les systèmes formés de N particules de masse identique m .

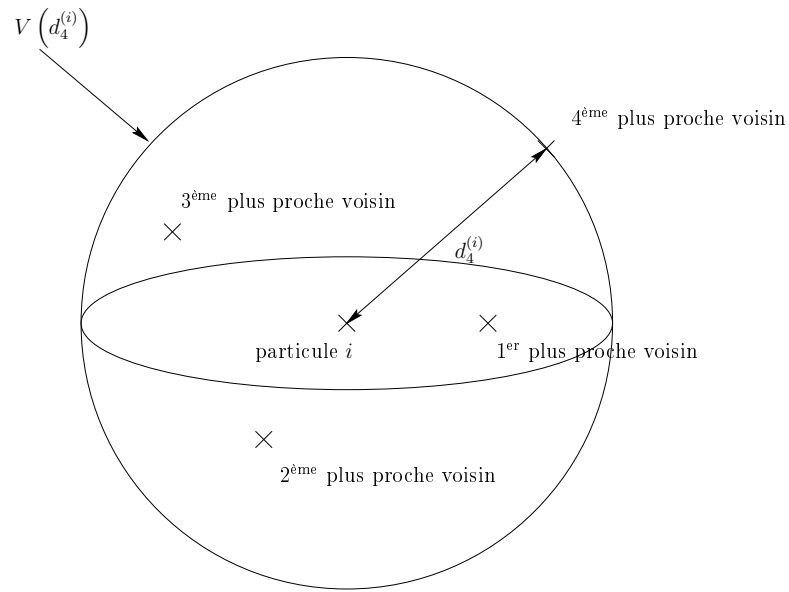


FIG. 2.5 – Exemple de calcul de la densité locale d'ordre 4 autour de la particule i .

plus proches voisins avec un réel très grand. Puis nous avons inséré chaque distance inférieure à la dixième plus proches dans ce vecteur à chaque fois que cela était nécessaire. L'algorithme se résume donc de la manière suivante :

```

pour k=1 à 10
    voisin[k] = 100000.0
fin pour
pour k=1 à N
    si k est différent de i alors
        d = |position[k]-position[i]|
        si d < voisin[10] alors
            insérer d à sa place dans voisin
        fin si
    fin si
fin pour

```

Cette algorithmme nous a permis d'obtenir chaque densité locale en un nombre de calcul proportionnel à N , le calcul de toutes les densités locales étant donc proportionnel à N^2 . Nous avons effectué le calcul des densités locales en parallèle, en distribuant l'ensemble des positions puis en attribuant à chaque processus le calcul d'une partie des densités locales, ce qui nous permis de réduire notablement le temps de calcul des densités.

3.4 Rapports d'axes

Les rapports d'axes représentent la géométrie du système. Ils sont définis par

$$\begin{aligned} a_1 &= \frac{\lambda_1}{\lambda_2} \\ a_2 &= \frac{\lambda_3}{\lambda_2} \end{aligned}$$

où $\lambda_1 \leq \lambda_2 \leq \lambda_3$ sont les valeurs propres de la matrice d'inertie I du système. Si une particule i de masse m_i est repérée par le vecteur $r_i = (x_{1,i}, x_{2,i}, x_{3,i})$, alors la matrice d'inertie est définie par (voir [34], p.136)

$$\begin{cases} I_{\mu,\nu} = - \sum_{i=1}^N m_i x_{\mu,i} x_{\nu,i} & \text{pour } \mu \neq \nu = 1, 2, 3 \\ I_{\mu,\mu} = \sum_{i=1}^N m_i (r_i^2 - x_{\mu,i}^2) & \text{pour } \mu = 1, 2, 3 \end{cases}$$

Les rapports d'axes permettent de classer les systèmes en 4 catégories :

- sphérique : a_1 et a_2 sont proches de l'unité
- aplati aux pôles¹⁵ : $a_1 < 1$ et $a_2 \approx 1$
- en forme de cigare¹⁵ : $a_1 \approx 1$ et $a_2 > 1$
- triaxial : $a_1 < 1$ et $a_2 > 1$

Il est également possible de définir une ellipticité e correspondant à l'ellipticité maximale pouvant être observée. Cette ellipticité est égale à

$$e = \frac{a_2 - a_1}{a_2}.$$

Elle peut être directement reliée à l'ellipticité des galaxies elliptiques.

3.5 Rayons contenant 10, 50 et 90% de la masse

Le calcul des rayons contenant 10, 50 et 90% de la masse du système – ces rayons sont désignés respectivement par R_{10} , R_{50} et R_{90} – nous fournit également de précieuses informations sur les caractéristiques du système. Ces rayons ne sont bien sûr exploitables que si le système est proche de la sphéricité. R_{10} nous donne une idée de la densité au cœur du système. R_{50} est une donnée pouvant être observée, elle présente donc un intérêt en termes de comparaison entre les simulations et les observations. Enfin, R_{90} est une approximation de la taille du système¹⁶.

Ainsi, même si les informations que nous pouvons en déduire sont relativement imprécises, elles présentent l'avantage de pouvoir être obtenues au cours de la simulation.

Le rayon R_{50} nous a également permis de définir un critère d'appartenance des particules au système lié. Toute particule située à une distance supérieure à $6 R_{50}$ du centre du système final est considérée comme étant libérée. Ce critère nous permet de sélectionner les particules libres plus rapidement que si nous avons calculé l'énergie de liaison de chaque particule. Nous avons comparé pour quelques cas ces deux méthodes de sélection des particules liées – énergie et distance au centre – et les résultats se sont révélés tout à fait comparables.

Les particules que nous considérons comme libres ne sont pas exclues de la simulation dynamique,

¹⁵ Les termes anglais *oblate* (aplatis aux pôles) et *prolate* (en forme de cigare) n'ont pas d'équivalent en français.

¹⁶ Cette approximation n'est valable que pour les systèmes n'ayant pas perdu beaucoup de particules lors de l'effondrement. Si un grand nombre de particules se libèrent lors de l'effondrement, elles s'éloignent du système lié et le calcul de R_{90} est alors très largement supérieur à la taille du système.

mais uniquement du calcul de certaines observables du système final, à savoir le rapport du viriel, les rapports d'axes, le potentiel et la densité, et la dispersion de vitesse.

3.6 Potentiel et densité

Le potentiel gravitationnel calculé à partir des fichiers de résultats et des sauvegardes intermédiaires est calculé de manière directe. Le paramètre d'adoucissement utilisé est le même que celui employé par le treecode. Il représente la seule source d'erreur de ce calcul du potentiel, puisqu'aucune autre approximation n'est faite. Le potentiel est exprimé en fonction du rayon. Cela sous-entend donc que ce résultat n'est pertinent que pour les systèmes à symétrie sphérique.

La méthode de calcul employée est très simple. Nous calculons le potentiel gravitationnel exercé sur chacune des particules par l'ensemble des autres. Le système est ensuite divisé en enveloppes sphériques d'épaisseur constante. Nous calculons le potentiel moyen exercé sur les particules contenues dans chaque enveloppe. Ce potentiel moyen nous donne la valeur du potentiel au rayon correspondant au milieu de l'enveloppe.

Le calcul de la densité est semblable à celui du potentiel. Nous divisons également le système en enveloppes sphériques. Nous calculons la densité moyenne dans chaque enveloppe, en divisant la masse totale contenue dans l'enveloppe par son volume.

3.7 Dispersion de vitesse

La dispersion de vitesse est calculée selon le même principe que le potentiel et la densité. Nous divisons le système en enveloppes sphériques, puis nous calculons la moyenne des carrés des vitesses radiale, $\langle v_{rad}^2 \rangle$, et tangentielle, $\langle v_{tan}^2 \rangle$, sur l'ensemble des particules situées dans chaque enveloppe. Ces moyennes sont définies par

$$\langle A \rangle = \frac{1}{N} \sum_{i=1}^N A_i . \quad (2.33)$$

Ceci nous permet d'obtenir le rapport $\kappa(r)$

$$\kappa(r) = \frac{2 \langle v_{i,rad}^2 \rangle_{r \leq r_i < r+dr}}{\langle v_{i,tan}^2 \rangle_{r \leq r_i < r+dr}} , \quad (2.34)$$

qui nous permet de mesurer le caractère radial ou tangentiel des orbites des particules dans le système.

3.8 Température

Nous avons défini une température d'équilibre du système. Une fois l'équilibre atteint, nous calculons la température T

$$T = \frac{2K}{3Nk_B} = \frac{1}{3} \langle mv^2 \rangle , \quad (2.35)$$

où K est l'énergie cinétique du système, et k_B est la constante de Boltzmann, que nous avons normalisée à 1. La température est calculée sur l'ensemble des particules, *i.e.* le système final plus les particules libérées au cours de l'effondrement.

3.9 Fonctions de distribution

Il nous a paru intéressant de comparer les résultats obtenus lors de nos simulations avec quelques résultats théoriques bien connus. Les systèmes à l'équilibre produits par nos simulations d'effondrement

ont donc été confrontés à deux modèles théoriques : le modèle polytropique et la sphère isotherme étudiés par ailleurs en détail dans le chapitre 1.

Les densités des systèmes définis par ces modèles sont respectivement

$$\rho_p = \rho_0 \psi_p^\gamma \quad (2.36)$$

et

$$\rho_s = \rho_1 e^{\psi_s/s^2} \quad (2.37)$$

Nous utilisons la méthode des moindres carrées afin d'ajuster les couples potentiel-densité que nous calculons pour nos systèmes à l'équilibre par un polytrope et une sphère isotherme. Nous passons pour cela à des quantités logarithmiques :

$$\ln(\rho_p) = \ln(\rho_0) + \gamma \ln(\psi_p) \quad (2.38)$$

et

$$\ln(\rho_s) = \ln(\rho_1) + \frac{\psi_s}{s^2} \quad (2.39)$$

Nous obtenons donc pour chaque état d'équilibre une valeur de l'indice polytropique γ et de la densité au centre ρ_0 , qui correspondent au polytrope le plus proche de l'état d'équilibre en question, ainsi qu'une valeur densité au centre ρ_1 et une dispersion de vitesse s , correspondant à la sphère isotherme la plus proche.

Ces résultats, ainsi que la méthode des moindres carrés, sont présentés en annexe. Leur analyse fait quant à elle l'objet de la section 5 du chapitre 3.

4 Système d'unités

Le système d'unités communément employé lors de simulations de dynamique gravitationnelle à N corps (voir [21]) repose sur les normalisations suivantes :

- $G = 1$,
- $M = 1$,
- $\mathcal{H} = -1/4$,

où G est la constante de la gravitation, M la masse totale du système et $\mathcal{H}$ son énergie totale.

Nous avons pour notre part choisi de fixer la valeur initiale du rapport du viriel $\mathcal{V}_{ini}$ et la taille initiale R_{ini} du système. Il nous a donc été impossible d'utiliser ce système d'unités.

Nos unités sont donc conditionnées par le choix d'utiliser les normalisations suivantes :

- $G = 1$,
- $M = 1 \, u_m$,
- $R_{ini} = 10 \, u_l$.

La valeur de $\mathcal{H}$ varie pour chaque simulation.

Afin de pouvoir comparer nos résultats numériques avec des résultats observationnels, nous avons choisi de faire correspondre nos unités de masse u_m et de longueur u_l avec les unités standards de la manière suivante

$$\begin{aligned} 1 \, u_m &= 10^6 \, \text{M}_\odot, \\ 1 \, u_l &= 1 \, \text{pc}. \end{aligned}$$

Cette correspondance nous permet de calculer la valeur de l'unité de temps u_t dans laquelle sont exprimées les durées de nos simulations. En effet, le temps dynamique¹⁷ t_d d'un système peut être

¹⁷Voir la section 2.1 du chapitre 1.

calculé de manière approchée par (equation (1.1) du chapitre 1)

$$t_d = \sqrt{\frac{3\pi}{16G\rho_m}} ,$$

où ρ_m est la moyenne de la densité sur l'ensemble du système.

Le rapport d'un temps dynamique calculé avec les unités standards sur un temps dynamique calculé avec nos unités nous donnera donc l'équivalence entre nos unités de temps et les années, soit

$$1u_t = \sqrt{\frac{R_{sta}^3 G_{sim} M_{sim}}{R_{sim}^3 G_{sta} M_{sta}}} ,$$

où les variables X_{sim} sont exprimées dans les unités des simulations et les variables X_{sta} dans les unités standards.

En prenant en compte les équivalences concernant les unités de longueur et de masse, et la valeur de G adoptée dans les simulations, nous obtenons

$$1u_t = \sqrt{\frac{(10 \cdot 3.09 \cdot 10^{18})^3 \cdot 1 \cdot 1}{10^3 \cdot 6.672 \cdot 10^{-8} \cdot 10^6 \cdot 1.989 \cdot 10^{33}}} \text{ s} \approx 4.72 \cdot 10^{11} \text{ s} = 1.49 \cdot 10^4 \text{ années} .$$

Il est intéressant de comparer cette unité de temps à la valeur approchée du temps dynamique d'un système réel dont la taille est 10 pc et la masse $10^6 M_\odot$. En effet, les amas globulaires ont des caractéristiques de ces ordres de grandeur. Pour cette taille et cette masse, nous obtenons $t_d \approx 2.34 \cdot 10^{13} \text{ s}$, c'est-à-dire $t_d \approx 50 u_t$.

Nous pouvons également choisir comme correspondance entre nos unités de masse et de longueur des ordres de grandeur plus en rapport avec ceux d'une galaxie elliptique typique, soit

$$\begin{aligned} 1u_m &= 10^{11} M_\odot , \\ 1u_l &= 10^3 \text{ pc} . \end{aligned}$$

Dans ce cas, nous obtenons la correspondance suivante

$$1u_t = \sqrt{\frac{(10^4 \cdot 3.09 \cdot 10^{18})^3 \cdot 1 \cdot 1}{10^3 \cdot 6.672 \cdot 10^{-8} \cdot 10^{11} \cdot 1.989 \cdot 10^{33}}} \text{ s} \approx 4.72 \cdot 10^{13} \text{ s} = 1.49 \cdot 10^6 \text{ années} .$$

Cette unité est simplement 100 fois supérieure à la précédente et nous obtenons la même correspondance pour le temps dynamique d'une galaxie de $10^{11} M_\odot$ de masse et de 10^4 pc de rayon, égal à environ $2.34 \cdot 10^{15} \text{ s}$, que pour l'amas globulaire, c'est-à-dire $t_d \approx 50 u_t$.

FORMATION DES AMAS GLOBULAIRES ET DES GALAXIES ELLIPTIQUES

1 Effondrement gravitationnel

1.1 Mécanisme

Nous avons décidé d'étudier un mécanisme particulier de formation des systèmes auto-gravitants : l'effondrement gravitationnel. Nous allons tenter d'expliquer de manière imagée le mécanisme de l'effondrement.

Le système que nous considérons initialement est constitué d'un nuage de particules ponctuelles. Ces particules possèdent chacune une vitesse propre. Notre nuage de particules crée un puits de potentiel gravitationnel. Les particules le constituant se trouvent à l'intérieur du puits (voir la figure 3.1). D'après le théorème du viriel (voir le paragraphe 2.5 du chapitre 1), le système n'est à l'équilibre que lorsque $2\bar{K} = -\bar{U}$, où $\bar{K}$ et $\bar{U}$ sont les moyennes temporelles de l'énergie cinétique totale et de l'énergie potentielle totale.

Si l'énergie cinétique, c'est-à-dire la vitesse, des particules est suffisamment élevée, les particules peuvent s'échapper du puits de potentiel. Dans ce cas, le nuage s'évapore. Mais si leur énergie cinétique est insuffisante, alors les particules vont tomber au fond du puits, le nuage va s'effondrer sur lui-même. L'effondrement est d'autant plus violent que le système initial est éloigné de la stabilité, caractérisée par le théorème du viriel.

Lors de l'effondrement, le système atteint une taille minimale en un temps de l'ordre du temps dynamique du nuage initial. Puis le système entre dans une phase d'expansion assez courte avant de s'effondrer de nouveau. Ses oscillations se poursuivent ensuite en s'atténuant, par un effet de type friction dynamique : les particules extérieures sont ralenties par leurs interactions avec les particules déjà agrégées. La figure 3.2 représente l'évolution du rayon R_{50} du système H_{50} au cours des 3000 unités de temps de son évolution. Les oscillations et les phases successives d'effondrement et d'expansion sont particulièrement visibles sur cette figure.

Nous allons étudier numériquement ce phénomène d'effondrement ou plus précisément les systèmes produits par l'effondrement de nuages initiaux aux propriétés diverses.

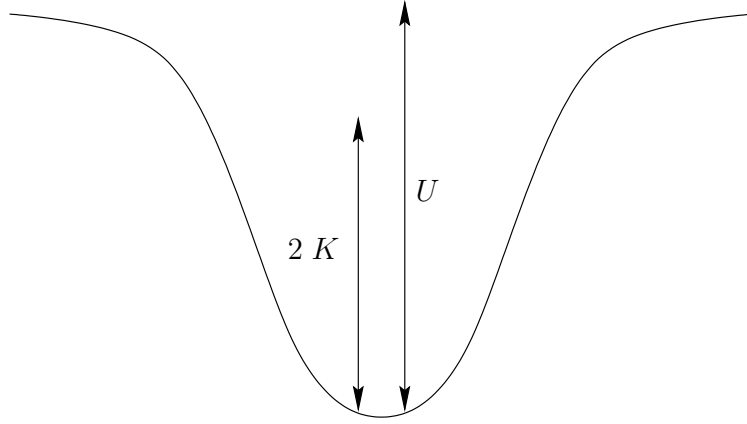


FIG. 3.1 – Si $2K$ est égal à U , le système est en équilibre, mais si $2K < U$, alors le système s’effondre.

2 Études préparatoires

2.1 Influence du paramètre d’adoucissement du potentiel

Le paramètre d’adoucissement du potentiel utilisé par le treecode est d’une importance majeure, à la fois pour les résultats des simulations et pour l’interprétation que l’on peut faire de ceux-ci. En effet, ce paramètre modifie, pour certaines particules, la valeur de l’accélération calculée par le code (voir la relation (2.7)). Si la valeur de ε est mal choisie, la dynamique globale du système s’en retrouvera donc modifiée. De plus, ce paramètre représente la résolution spatiale¹ de notre simulation. Si nous ne pouvons pas calculer avec précision le mouvement de particules qui sont séparées par une distance inférieure à ε , il n’est pas raisonnable de penser que nous pouvons observer et étudier des phénomènes se produisant sur des échelles de distance inférieures à ε .

Nous avons effectué une série de tests composée de 16 simulations afin de mesurer l’influence de ε . Les systèmes initiaux sont de type H_{50} , et comportent 20000 particules. Ils sont sphériques et de rayon égal à 1. Les simulations diffèrent par la valeur de ε utilisée et par le pas de temps Δt employés pour l’intégration. Le pas de temps utilisé dans cette série de simulations est fixe. Quatre valeurs différentes de ε et de Δt sont utilisées. Les valeurs de ε utilisées sont 0.05, 0.01, 0.005 et 0.001. Celles de Δt varient de 0.03 ($t_d/50$) à 0.003 ($t_d/500$). La durée des simulations a été fixée à 30 fois le temps dynamique des systèmes initiaux. Nous avons mesuré, à l’issue de ces effondrements, la perte d’énergie du système, le rapport du viriel, les rapports d’axes, le rayon de 50% de la masse et l’indice polytropique. Les résultats sont présentés dans le tableau C.1.

Les résultats montrent en premier lieu qu’il est nécessaire de diminuer de manière importante le pas d’intégration en temps lorsque l’on baisse la valeur de ε . Pour $\varepsilon = 0.005$, la perte d’énergie est égale à 12.13% avec $\Delta t = 0.03$, alors qu’avec le même pas de temps, mais pour $\varepsilon = 0.05$, on a $\Delta\mathcal{H} = 0.30\%$. Il est nécessaire d’abaisser le pas de temps à $\Delta t = 0.003$ pour que la perte d’énergie soit inférieure à 1% avec $\varepsilon = 0.005$. Il faut donc choisir de privilégier ou bien la résolution en espace, ou bien le temps de calcul.

Plus précisément, il semble qu’il soit nécessaire de diviser le pas de temps par un facteur λ si nous voulons diviser la résolution par ce même facteur λ comme nous pouvons le voir sur la figure 3.3. En effet, l’erreur sur l’énergie totale pour $\varepsilon = 0.05$ et $\Delta t = 0.03$ est égale à 0.30%. Si nous divisons la résolution par 10, en conservant le même pas de temps, l’erreur devient $\Delta\mathcal{H} = 12.13\%$. Si maintenant nous divisons également le pas de temps par 10, l’erreur descend à 0.14%. Un autre exemple est donnée

¹Pour être plus exact, la résolution spatiale est déterminée par le rapport entre ε et la taille du système.

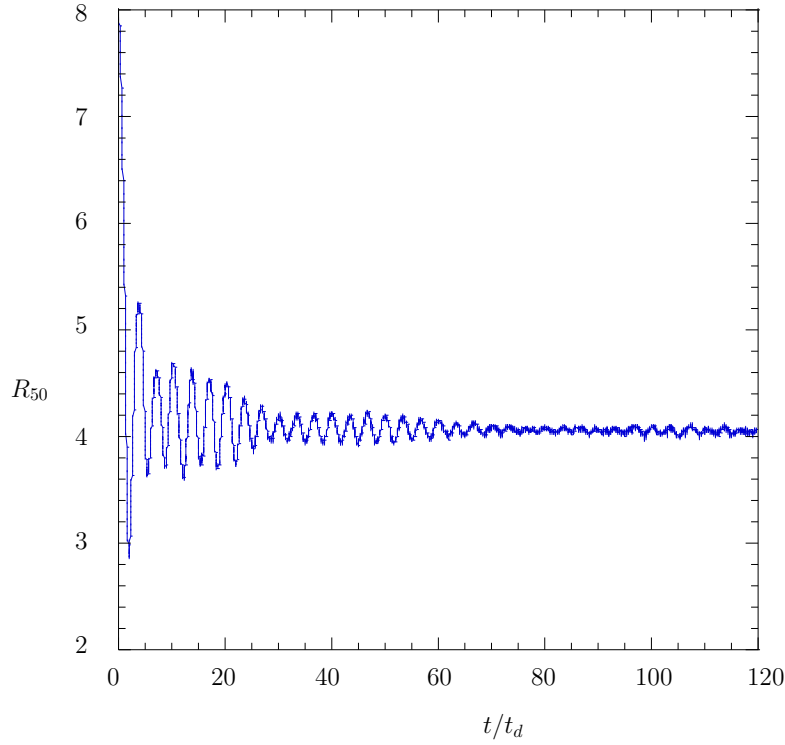


FIG. 3.2 – Évolution du rayon de 50% de la masse R_{50} du système H_{50} en fonction du temps.

par l'erreur pour $\varepsilon = 0.10$ et $\Delta t = 0.15$ (1.04%) et celle pour $\varepsilon = 0.05$ et $\Delta t = 0.08$ (1.02%). Cette constatation semble indiquer qu'il existe pour notre schéma une condition jouant un rôle analogue à la condition CFL (Courant-Friedrichs-Levy) (voir [12], p.50 et suivantes), c'est-à-dire qu'il est nécessaire que $\Delta t/\varepsilon$ soit inférieure à une constante pour que l'erreur n'explose pas. Ce comportement peut en partie s'expliquer par l'analyse du schéma d'intégration en temps de notre code que nous avons mené dans l'annexe A2.

Ces premiers résultats nous ont amené à choisir une valeur de ε égale à $R_{ini}/100$, ce qui correspond ici à $\varepsilon = 0.01$. Une valeur de ε inférieure à celle-ci nécessiterait un pas de temps très faible, puisque nous tenons à conserver l'énergie totale du système autant que possible. En effet, la perte d'énergie est toujours supérieure à 5% pour $\varepsilon = 0.001$, et la perte d'énergie est supérieure à 1% si Δt est supérieur à $t_d/500$ pour $\varepsilon = 0.005$. Les effondrements simulés lors de cette étude sont relativement "calmes", puisque le rapport du viriel initial des systèmes est égal à -0.5 . Lors d'effondrements plus violents, l'intégration en temps doit être faite avec plus de soin. Ceci implique donc que le pas de temps doit être inférieur à celui que nous utilisons pour cette expérience test. Notre objectif étant de réaliser un grand nombre de simulations, nous ne pouvons pas nous permettre d'utiliser un pas de temps trop petit, ce qui aurait pour conséquence d'allonger le temps de calcul.

Ce premier test nous montre par ailleurs que, si l'on excepte les simulations terminées avec plus de 5% de variation d'énergie totale, les autres observables mesurées ne sont pas très sensibles à la variation de ε . Les variations maximales du petit et du grand rapport d'axes sont respectivement 3.22% et 1.84%. Le rayon de 50% de masse varie de moins de 5%, et l'indice polytropique de moins de 1.61%.

C.Theis et R.Spurzem [52] ont également effectué une étude de l'influence de ε . Les conditions initiales qu'ils ont utilisées et les observables qu'ils ont décidé d'observer rendent toutefois la comparaison délicate entre leurs résultats et les nôtres. Nous pouvons toutefois noter que leurs mesures des rapports

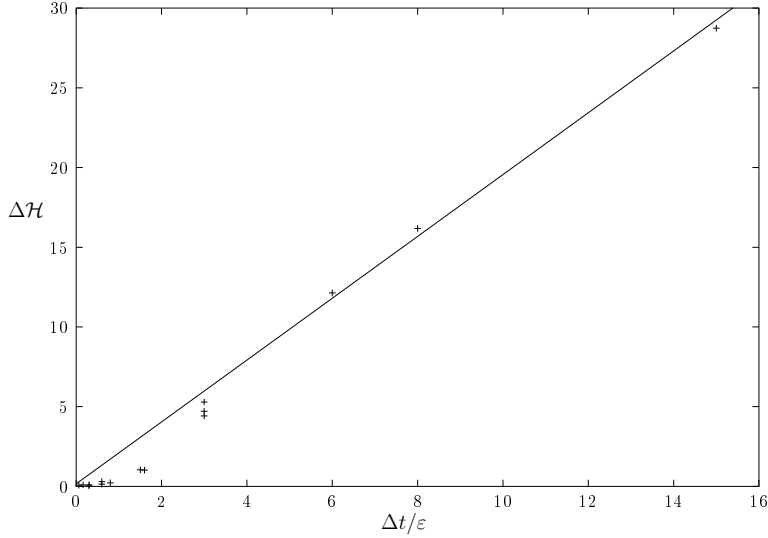


FIG. 3.3 – Erreur ΔH sur le calcul de l'énergie totale en fonction du rapport $\Delta t/\varepsilon$.

d'axes des 50% des particules les plus liées montrent que les rapports d'axes – tels que nous les avons définis – sont surestimés lorsque ε est trop grand, comme nous l'avons également observé.

2.2 Influence du nombre de particules

Il nous a paru essentiel, avant de mener notre étude, de nous poser la question suivante : “Combien faut-il utiliser de particules lors d'une simulation à N corps pour modéliser fidèlement un amas globulaire ou une galaxie elliptique?”.

C.Theis et R.Spurzem [52] ont répondu en partie à cette question, alors que van Albada [54], L.Aguilar et D.Merritt [3] ou J.K.Cannizzo et T.C.Hollister [14] ne s'y sont pas intéressés. De plus, le nombre de particules utilisé dans chacune de ces trois dernières études n'atteint que rarement 10^4 , ce qui peut paraître faible pour représenter une galaxie elliptique composée d'environ 10^{10} étoiles (Theis et Spurzem ont quant à eux utilisé jusqu'à 32768 particules).

Il est toutefois délicat de justifier l'emploi d'un nombre de particules plutôt qu'un autre, aussi longtemps que l'on n'utilise pas une particule pour représenter une étoile. Nous avons décidé d'utiliser un critère relativement simple à vérifier : si l'augmentation du nombre de particules utilisé n'induit aucun changement dans les observables calculées lors de la simulation, nous considérons avoir atteint une valeur satisfaisante du nombre de particules. Nous avons donc réalisé une série de simulations d'effondrements gravitationnels en faisant varier le nombre de particules N . Nous avons ensuite étudié l'influence de N sur les rapports d'axes a_1 et a_2 , les ajustements avec un polytrophe et une sphère isotherme et les rapports R_{10}/R_d et R_{50}/R_d . Le nombre de particules que nous avons adopté est égale à la plus petite valeur de N' telle que chacune des observables ci-dessus est à peu près constante pour $N > N'$.

Il était difficile, pour des raisons évidentes de temps de calcul, de tester tous les types de conditions initiales. C'est pourquoi nous avons utilisé trois conditions initiales de type H , différant par leur rapport du viriel initial. Les trois valeurs utilisées sont $\mathcal{V}_{ini} = -0.1, -0.5, -0.80$. Les systèmes examinés sont tous sphériques, de rayon $R_{ini} = 10$. Le nombre de particules utilisé varie quant à lui de $5 \cdot 10^3$ à $7 \cdot 10^4$.

Nous avons calculé pour chacune des simulations effectuées l'ensemble des observables décrites dans

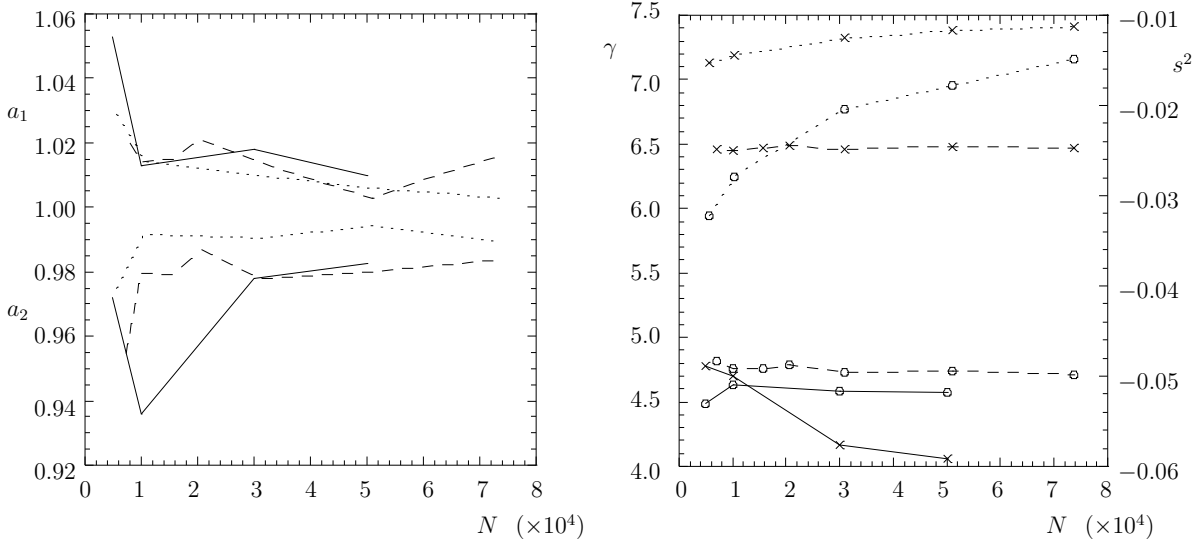


FIG. 3.4 – Figure de gauche : évolution des petit et grand rapports d’axes en fonction du nombre de particules pour trois conditions initiales : H_{10} , H_{50} et H_{80} . Figure de droite : évolution des résultats de l’ajustement par un polytrope (indice polytropique γ , représenté par un cercle) et par une sphère isotherme (dispersion de vitesse s^2 , représentée par une croix) en fonction du nombre de particules pour trois conditions initiales : H_{10} , H_{50} et H_{80} . Les lignes continues correspondent aux systèmes H_{10} , les lignes discontinues aux systèmes H_{50} et les lignes pointillées aux systèmes H_{80} .

la section 3 du chapitre 2. Cette étude préliminaire n’est donc pertinente que dans le cadre que nous avons établi, à savoir la simulation d’effondrement gravitationnel à l’aide d’un treecode. Nous ne prétendons pas déterminer avec certitude le nombre idéal de particules qu’il faut employer pour modéliser un amas globulaire ou une galaxie elliptique. Nous voulons simplement montrer que le nombre de particules employé n’est pas un paramètre négligeable, et qu’il faut le choisir en prenant quelques précautions.

Les résultats produits lors de cette étude sont représentés sur les figures 3.4 et 3.5, et repris dans le tableau C.2. Nous pouvons voir sur la figure 3.4 que les rapports d’axes varient de quelques pour-cents au maximum (4.72%). L’ellipticité du système semble généralement surévaluée lorsque le nombre de particules baisse. Toutefois les rapports d’axes varient moins sensiblement pour $N \geq 3 \cdot 10^4$. Sur cette même figure, graphique de droite, nous pouvons également constater que les résultats des ajustements des couples potentiel-densité par les modèles polytropique et de sphère isotherme sont dépendants de N . L’indice polytropique du modèle H_{80} augmente notamment de plus de 20% lorsque N passe de 5571 à 73614 (voir le tableau C.2). Mais les indices polytropiques des deux autres modèles ne varient que très peu : 3.25% pour H_{10} et 2.36% pour H_{50} . En ce qui concerne la dispersion de vitesse de la sphère isotherme, elle varie de près de 21% pour H_{10} , de plus de 35% pour H_{80} et de seulement 2.46% pour H_{50} .

La figure 3.5 nous montre quant à elle très clairement que les rayons caractéristiques du système H_{10} sont très fortement dépendants de N . La valeur de R_{50}/R_d varie de 5.5 à 2.4 lorsque N varie de 5571 à 50947. Les autres rapports de rayons mesurés varient de façon moins spectaculaires, mais nous pouvons tout de même noter que R_{50}/R_d baisse de plus de 15% entre 7246 et 73636 particules pour le système H_{50} . Enfin, mis à part le rapport R_{50}/R_d pour le système H_{10} , nos mesures de rapports de rayons varient de moins de 3% lorsque le nombre de particules est supérieur à $3 \cdot 10^4$.

L’influence du nombre de particules peut être interprétée par la transition entre des systèmes avec

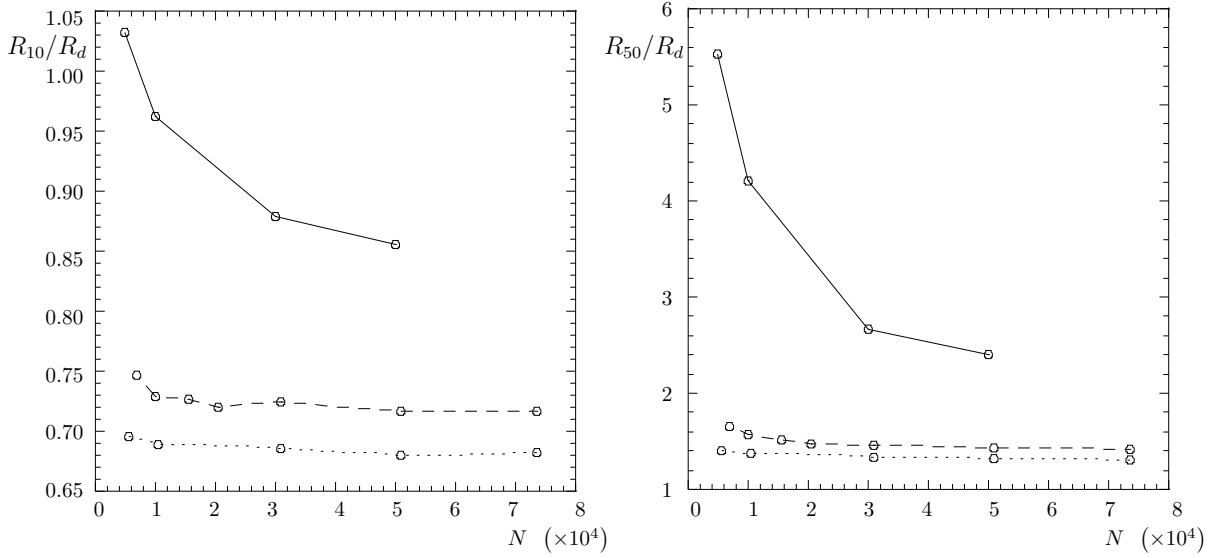


FIG. 3.5 – Évolution de R_{10}/R_d et de R_{50}/R_d en fonction du nombre de particules pour trois conditions initiales : H_{10} , H_{50} et H_{80} . Les lignes continues correspondent aux systèmes H_{10} , les lignes discontinues aux systèmes H_{50} et les lignes pointillées aux systèmes H_{80} .

peu de particules où les phénomènes individuels dominent et des systèmes avec un grand nombre de particules où les phénomènes collectifs deviennent prépondérants. Les inhomogénéités prennent également plus d'importance lorsque le système contient peu de particules et les observables statistiques, reposant sur des moyennes, telles que les résultats des ajustements, perdent de leur pertinence.

Nous avons donc choisi d'utiliser des systèmes composés d'environ $3 \cdot 10^4$ particules pour notre étude. Ce choix est à la fois la conséquence de cette étude préparatoire et de notre volonté d'effectuer un grand nombre de simulations. Un plus grand nombre de particules n'aurait vraisemblablement pas présenté beaucoup d'intérêt, puisque les valeurs des observables qui se sont avérées les plus intéressantes lors de notre travail sur la formation des systèmes auto-gravitants varient très peu lorsque l'on choisi $N > 3 \cdot 10^4$ (les rayons caractéristiques notamment).

3 Influence des conditions initiales sur la géométrie du système à l'équilibre

Plusieurs travaux ont déjà examiné l'influence des conditions initiales sur la géométrie des systèmes à l'équilibre issus de simulations d'effondrements gravitationnels. L.Aguilar et D.Merritt [3] ont étudié le cas de systèmes initiaux sphériques de densité $\rho \propto r^{-1}$ et déduit de leur analyse que l'ellipticité du système final dépend du rapport du viriel initial. Ils ont obtenu des ellipticités allant jusqu'à $e = 0.5$ pour des systèmes initialement sphériques et $e = 0.6$ pour des systèmes initialement aplatis. L'instabilité d'orbite radiale est invoquée par les deux auteurs pour expliquer l'aplatissement des systèmes. Nous avons étendu cette étude à une plus large gamme de conditions initiales afin de tester l'influence de paramètres autres que le rapport du viriel sur les rapports d'axes du système final. Les résultats de cette étude sont présentés dans les tableaux C.3–C.7, situés en annexe. Ces mêmes résultats sont également représentés sur la figure 3.6. Nous pouvons tout d'abord noter que l'ellipticité des systèmes homogènes (H_η , G_η^σ et M_η^k) ne varie pas en fonction du rapport du viriel initial. Tous les systèmes initialement homogènes engendrent des systèmes très proches de la sphéricité. L'ellipticité maximale

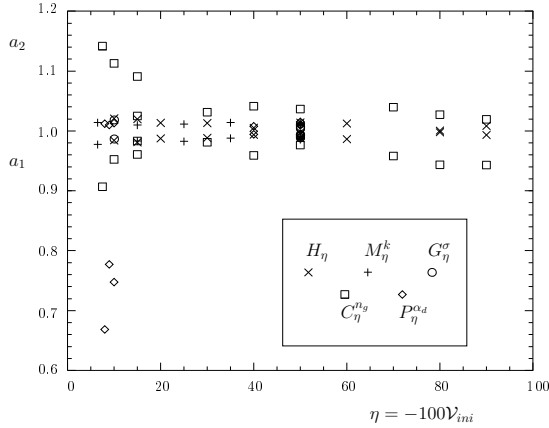


FIG. 3.6 – Petit (a_1) et grand (a_2) rapports d'axes représentés en fonction du rapport du viriel initial et du type de condition initiale.

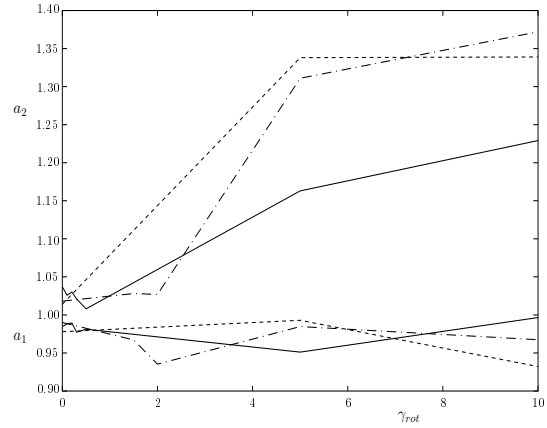


FIG. 3.7 – Petit (a_1) et grand (a_2) rapports d'axes représentés en fonction du paramètre de rotation pour trois conditions initiales différentes : $R_{10}^{\gamma_{rot}}$ (lignes pleines), $R_{30}^{\gamma_{rot}}$ (lignes mixtes) et $R_{50}^{\gamma_{rot}}$ (lignes interrompues).

mesurée pour un tel système est égale à 0.05. Ce résultat contredit celui obtenu par K.W.Min et C.S.Choi [40], puisque les simulations qu'ils ont effectuées ont conduit à la formation de systèmes triaxiaux à partir de conditions initiales homogènes lorsque le rapport du viriel initial était faible, ce qui correspond à $\eta \lesssim 20$. Cette différence peut être attribuée au nombre de particules utilisé par ces deux auteurs ($N = 8000$). Notre étude concernant l'influence du nombre de particules ne s'applique pas ici, puisque K.W.Min et C.S.Choi ont utilisé un code particule-grille, mais nous pouvons raisonnablement penser que ce paramètre a également une influence sur certaines observables lorsque le code utilisé est de type PM.

L'ellipticité des systèmes inhomogènes varie quant à elle en fonction de η . Ce résultat confirme ceux obtenus par L.Aguilar et D.Merritt, puisque leurs systèmes étaient également initialement inhomogènes. Nous pouvons également observer que η doit être faible ($\lesssim 15$) pour que l'aplatissement du système soit significatif. Au-delà de cette valeur, les systèmes à l'équilibre sont sphériques.

L'inhomogénéité induite par les grumeaux et celle résultant du profil de densité initial ne produisent pas le même résultat du point de vue de l'ellipticité. Ainsi, la présence de vingt grumeaux dans le système initial C_{07}^{20} donne naissance à un système dont l'ellipticité est égale à 0.19, alors que le système initial $P_{08}^{1.5}$ produit un système d'ellipticité 0.30. Il apparaît qu'un profil de densité favorise plus l'aplatissement du système lors de son effondrement.

Si nous comparons l'ellipticité obtenue avec vingt grumeaux ($e = 0.12$ pour le système C_{10}^{20}) avec celle obtenue avec trois grumeaux ($e = 0.07$ pour le système C_{10}^3), ou même sans grumeaux ($e = 0.04$ pour le système H_{10}), nous voyons que le nombre de grumeaux influe sur l'ellipticité finale. L'ellipticité augmente sensiblement avec le nombre de grumeaux.

De même, si nous comparons les résultats obtenus pour les différents profils de densité, c'est-à-dire $e = 0.20$ pour le système $P_{10}^{1.0}$, $e = 0.30$ pour le système $P_{08}^{1.5}$ et $e = 0.23$ pour le système $P_{09}^{2.0}$ entre autres (les autres rapports d'axes sont rassemblés dans le tableau C.5), il apparaît que l'ellipticité varie avec le paramètre α_d . Cette ellipticité augmente également fortement lorsque le paramètre du viriel initial η baisse.

Nous pensons que l'ellipticité du système final augmente lorsque α_d augmente, et que l'écart entre l'ellipticité de $P_{09}^{2.0}$ (0.23) et celle de $P_{08}^{1.5}$ (0.30) est dû au fait que le système $P_{09}^{2.0}$ fait partie des quelques systèmes dont le temps d'évolution, $1500 u_t$, est inférieur à $3000 u_t$. La violence de cet effondrement

($V_{ini} = -0.09$) nous a obligé à utiliser un pas de temps très faible, ce qui a entraîné un temps de calcul très important (environ une semaine). Nous avons donc décidé d'arrêter l'intégration en temps à $t = 1500 u_t$. Nous avons en effet fait le choix d'une étude la plus complète possible, ce qui nous a conduit à faire des concessions sur certains points, comme celui-ci, le temps dont nous disposions n'étant pas indéfiniment extensible. Il est possible que son évolution au-delà de $1500 u_t$ l'aurait amené à une plus grande ellipticité.

Ce type d'étude a également été mené par J.K.Cannizzo et T.C.Hollister [14], qui ont étudié l'influence du paramètre α_d sur la géométrie finale de leurs systèmes. Ils ont utilisé 10^4 particules, et ont fixé initialement le rapport $|K/\Phi|$ pour chaque particule² à la valeur 0.01. Ils ont observé que le rapport entre le grand et le petit axe varie de environ 1.38 pour $\alpha_d = 0.5$ à environ 1.79 pour $\alpha_d = 2.5$. Ces chiffres sont supérieurs aux nôtres, mais il sont compatibles si l'on considère que leur rapport du viriel initial est inférieur d'un facteur 8, au minimum, à ceux que nous avons utilisés.

Nous avons également testé l'influence de la rotation sur l'ellipticité finale. Ces résultats sont représentés sur la figure 3.7, ainsi que dans le tableau C.8. Plusieurs caractéristiques se dégagent de ces résultats. Tout d'abord, le système final s'écarte peu de la sphéricité si le facteur de rotation γ_{rot} est inférieur à 4. Lorsque $\gamma_{rot} > 4$, les rapports d'axes ne dépendent plus que faiblement de la quantité de rotation initiale. Il semble donc qu'il existe un seuil de rotation au delà duquel le système se déforme, mais que la déformation ne s'amplifie pas lorsque la rotation augmente.

Ensuite, le système final est allongé, et non aplati. Le grand rapport d'axe est en effet supérieur à 1, alors que le petit rapport d'axe reste proche de l'unité.

Enfin, la quantité de rotation présente initialement dans le système ne disparaît pas totalement au cours de l'effondrement. Le rapport μ_K entre l'énergie cinétique de rotation et l'énergie cinétique totale conserve typiquement 65% de sa valeur initiale. Or les observations montrent que les amas globulaires et les galaxies elliptiques ne sont pas en rotation rapide autour d'un de leurs axes. Il ne nous est pas donc possible de faire appel à la rotation afin d'expliquer les fortes ellipticités observées parmi les galaxies elliptiques.

Parmi les solutions que nous avons testées, l'effondrement de nuages inhomogènes froids est donc seul en mesure de reproduire, même partiellement, l'éventail des ellipticités observées. De nombreuses particules constituant ces systèmes se déplacent sur des orbites très radiales en raison de la violence de l'effondrement. La présence de ces orbites radiales provoque vraisemblablement une instabilité d'orbites radiales qui déforme le système (voir le paragraphe 3.5 du chapitre 1). Cependant, puisque les systèmes initialement homogènes ne présentent pas d'aplatissement significatif, il semble que la présence d'inhomogénéités soit indispensable au déclenchement de l'instabilité d'orbites radiales.

4 Taille du cœur et ségrégation

Nous avons observé au cours de notre étude un phénomène jusqu'ici inconnu et plus fin que la ségrégation entre les systèmes sphériques et les systèmes triaxiaux.

Nous avons comparé les densités $\rho(r)$ des états d'équilibre issus de nos simulations. Ces densités sont représentées, en fonction de r/R_{50} sur les figures 3.8, 3.9, 3.10, 3.11 et 3.12. Elles n'évoluent que très peu après l'effondrement, mis à part pour le système M_{07}^I , dont nous reparlerons par la suite.

Certaines de ces densités sont légèrement croissante en fonction du rayon au centre du système. Ceci est une erreur numérique provenant probablement de la méthode que nous avons employée pour calculer les densités (voir paragraphe 3.6 du chapitre 2). Lorsque nous calculons la densité à l'aide d'enveloppes sphériques contenant un nombre constant de particules, nous obtenons des densités constantes ou décroissantes au centre du système. Mais cette méthode implique quelques problèmes : si le nombre de particules contenues dans chaque enveloppe est petit, le résultat est très bruité pour les

²Ceci est différent de fixer le rapport du viriel initial à 0.01. En effet, le rapport du viriel concerne les énergies cinétiques et potentielles totales, et non individuelles.

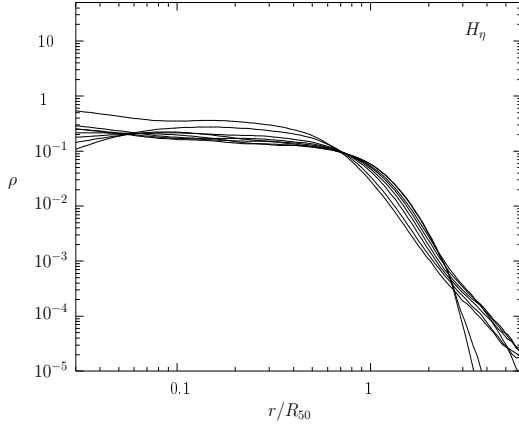


FIG. 3.8 – Densité finale $\rho(r/R_{50})$ pour tous les systèmes H_η .

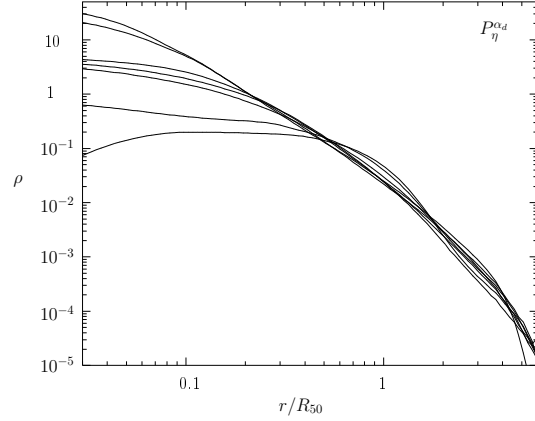


FIG. 3.9 – Densité finale $\rho(r/R_{50})$ pour tous les systèmes P_η^α .

rayons supérieurs à R_{50} . Et lorsque ce nombre est grand, le bruit disparaît mais la résolution spatiale diminue et la densité au cœur est mal déterminée.

Sur ces différentes figures, nous pouvons clairement distinguer deux types d'états d'équilibre :

- (i) des systèmes possédant une structure cœur/halo, c'est-à-dire une région centrale de densité quasi constante entourée d'une enveloppe de densité décroissante avec le rayon ;
- (ii) des systèmes dont la densité centrale est supérieure par un à deux ordre de grandeur à celle des systèmes de type (i), et dont la densité décroît régulièrement sans présenter de plateau dans la région centrale.

Il est très intéressant de remarquer que les systèmes de type (i) satisfont tous les critères

$$\ln(R_{50}/R_d) < 0.5 \quad \text{et} \quad \ln(R_{10}/R_d) < -0.05$$

alors que les systèmes de type (ii) satisfont

$$\ln(R_{50}/R_d) > 0.7 \quad \text{et} \quad \ln(R_{10}/R_d) > 0.1 .$$

Il apparaît sur la figure 3.13, sur laquelle nous avons représenté pour chaque système $\ln(R_{10}/R_d)$ en fonction de $\ln(R_{50}/R_d)$, que la majorité des systèmes que nous avons obtenus entre très clairement dans l'une ou l'autre de nos deux catégories. D'un côté, tous les systèmes de type H_η , G_η^σ et M_η^k (à l'exception de M_{07}^I), ainsi que les systèmes $P_{50}^{0.5}$ et P_{50}^1 , se trouvent dans la zone inférieure gauche du diagramme, correspondant aux systèmes de type (i). Ces systèmes sont homogènes, ou "légèrement" inhomogènes dans les cas des deux systèmes $P_{50}^{\alpha_d}$. De l'autre côté, tous les systèmes C_η^{20} et les systèmes P_{50}^2 et P_{09}^2 se trouvent dans le coin supérieur droit du diagramme, correspondant aux systèmes de type (ii). Ces systèmes possèdent une densité initiale fortement inhomogène.

Les cinq systèmes positionnés entre ces deux zones sont les systèmes C_{10}^3 , P_{10}^1 , $P_{40}^{1.5}$, $P_{08}^{1.5}$ et M_{07}^I . Les quatre premiers sont inhomogènes, et leur degré d'inhomogénéité se situe a priori entre celui des systèmes de type (i) et celui de systèmes de type (ii). Le système M_{07}^I présente quant à lui un comportement singulier qui mérite que l'on s'y attarde.

Sur les figures 3.14 et 3.15, nous avons représenté la densité des systèmes C_{10}^{20} et M_{07}^I à divers stades de leur évolution. Le profil de densité du système C_{10}^{20} varie très peu entre le temps $10 t_d$ et le temps $100 t_d$. Celui de M_{07}^I présente quant à lui les caractéristiques d'un système de type (i) au temps $10 t_d$ alors qu'il s'approche des caractéristiques d'un système de type (ii) au temps $100 t_d$. Il s'agit du

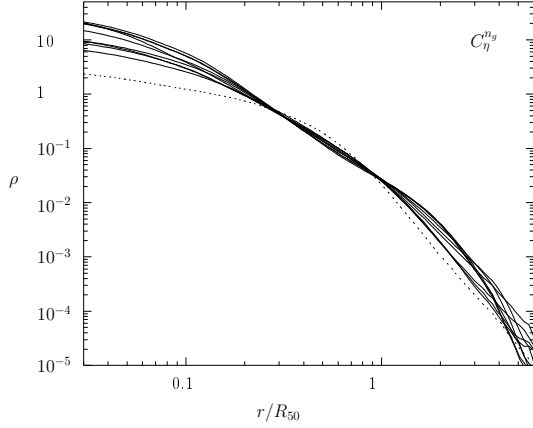


FIG. 3.10 – Densité finale $\rho(r/R_{50})$ pour tous les systèmes C_{η}^n . La ligne pointillé correspond au système C_{10}^3 .

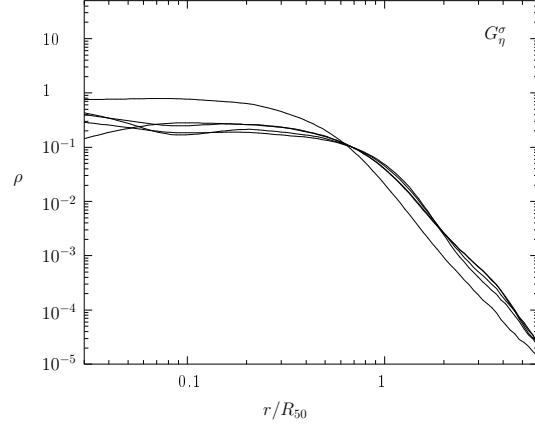


FIG. 3.11 – Densité finale $\rho(r/R_{50})$ pour tous les systèmes G_{η}^{σ} .

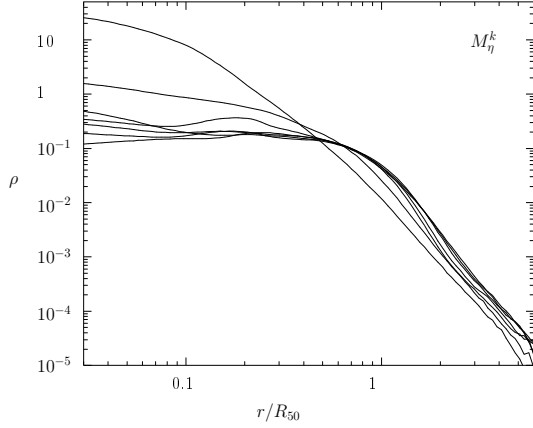
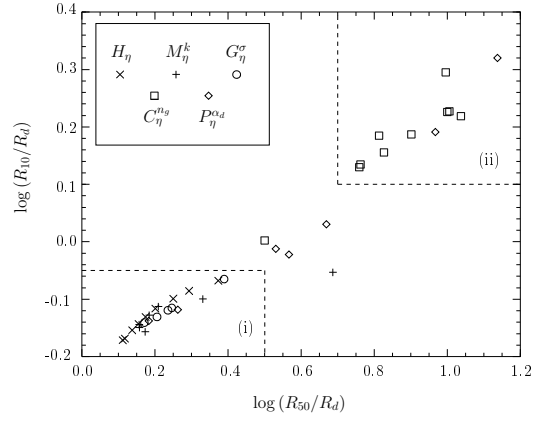
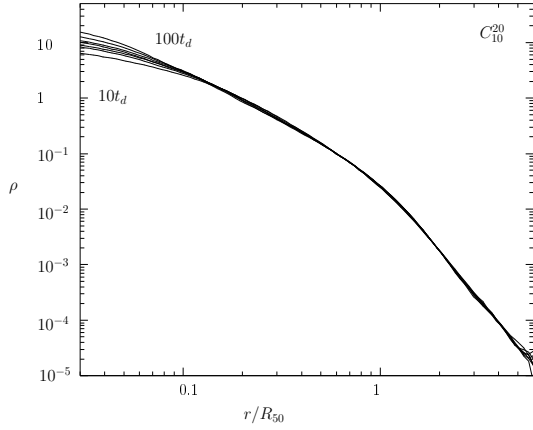
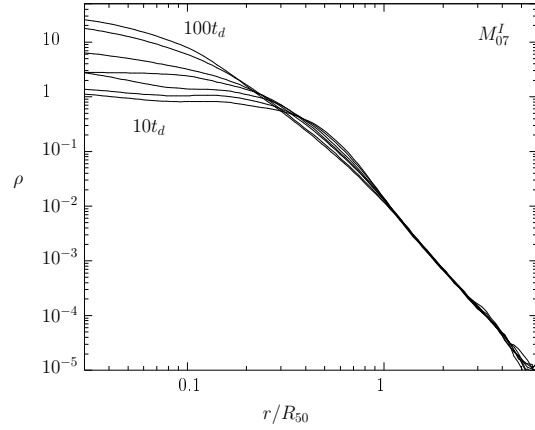
seul objet évoluant significativement en cours d'intégration temporelle sur le diagramme $\ln(R_{50}/R_d) - \ln(R_{10}/R_d)$. Initialement, il se trouve dans la zone (i) et migre petit à petit vers la zone (ii).

Si nous faisons exception des cinq cas particuliers que nous venons d'évoquer, nous nous trouvons donc en présence de deux types de systèmes présentant des profils de densité et des rapports de rayons caractéristiques tout à fait dissemblables. Nous interprétons cette ségrégation par l'instabilité d'Antonov.

Les systèmes ne présentant pas de cœur sont soit de type P_{η}^2 , soit de type C_{η}^{20} . Les systèmes P_{η}^2 possèdent par définition une surdensité centrale, qui s'accroît lors de la première phase de l'effondrement. Les systèmes de types C_{η}^{20} voient quant à eux leur surdensités locales (grumeaux) s'effondrer rapidement³ sur elles-mêmes puis rejoindre le centre du système lorsque l'ensemble du système s'effondre. Dans les deux cas, une forte surdensité de matière est présente au cœur du système très peu de temps après le début de l'effondrement. Le rapport $\rho_c/\rho(R)$ étant le facteur déclenchant de l'instabilité d'Antonov (voir le paragraphe 3.4 du chapitre 1), une surdensité centrale implique un risque accru d'instabilité. Si la surdensité est suffisante, l'instabilité se déclenche, et le cœur continue à se contracter, la densité centrale augmentant d'autant plus. C'est ce que nous avons observé. Si par contre cette surdensité est insuffisante pour déclencher l'instabilité, ce qui est le cas pour nos autres simulations, l'effondrement se poursuit, et le cœur grandit petit à petit jusqu'à former un plateau. Ceci explique également l'écart entre les densités centrales des deux types de systèmes.

Nous pouvons également proposer une explication de l'évolution du système M_{07}^I . Un cœur se forme dans ce système peu après l'effondrement, puisque ce système ne présente initialement aucune source importante d'inhomogénéités. Cependant, à la suite de l'effondrement particulièrement violent subi par ce système, une part importante des particules le constituant s'échappe. Ceci conduit le cœur à se contracter, et la densité centrale à augmenter, ce qui explique le déplacement de M_{07}^I sur le diagramme $\ln(R_{50}/R_d) - \ln(R_{10}/R_d)$. Si le cœur se contracte suffisamment sous l'effet de cette évaporation, l'instabilité d'Antonov pourra alors se déclencher, et le système continuera sa migration vers la zone (ii) du diagramme $\ln(R_{50}/R_d) - \ln(R_{10}/R_d)$. Ce mécanisme est souvent invoqué pour expliquer la présence d'amas globulaires au cœur effondré dans la Galaxie. À la fin de notre simulation, le système M_{07}^I

³ Les grumeaux s'effondrent beaucoup plus vite que le système dans son ensemble. En effet, le temps caractéristique de l'effondrement est du même ordre de grandeur que le temps de croisement, soit $\approx (G\rho_m)^{-1/2}$. La densité de nos grumeaux étant environ cent fois supérieure à celle du système dans son ensemble, le temps caractéristique de l'effondrement d'un grumeau est dix fois inférieur à celui du système.

FIG. 3.12 – Densité finale $\rho(r/R_{50})$ pour tous les systèmes M_η^k .FIG. 3.13 – $\ln(R_{10}/R_d)$ en fonction de $\ln(R_{50}/R_d)$ pour chaque état d'équilibre.FIG. 3.14 – Évolution de la densité $\rho(r/R_{50})$ pour le système C_{10}^{20} au cours du temps.FIG. 3.15 – Évolution de la densité $\rho(r/R_{50})$ pour le système M_{07}^I au cours du temps.

ne présente pas encore les caractéristiques des systèmes de type (ii), mais son évolution nous montre qu'il s'en rapprochait, puisque sa densité centrale ne cessait d'augmenter et le rayon de son cœur de baisser. Nous avons choisi de ne pas poursuivre l'évolution de ce système pour des raisons de coût numérique, que nous avons déjà évoquées dans la section 3 au sujet du système $P_{09}^{2,0}$.

5 Fonction de distribution à l'équilibre

Afin de comparer les résultats de nos simulations d'effondrements avec certaines solutions analytiques du système Boltzmann sans collisions-Poisson (1.36), nous avons ajusté les couples potentiel-densité de nos systèmes à l'équilibre par les modèles de la sphère isotherme et des polytropes (voir le paragraphe 3.9 du chapitre 2). La figure 3.18 représente l'ajustement par ces deux modèles du système final $P_{50}^{0.5}$. Les figures 3.16 et 3.17 représentent quant à elles respectivement l'indice polytropique γ et la dispersion de vitesse s^2 obtenus pour chacun des systèmes à l'équilibre. La valeur moyenne de la dispersion de vitesse s^2 obtenue par l'ajustement par une sphère isotherme est égale à $3.32 \cdot 10^{-2}$,

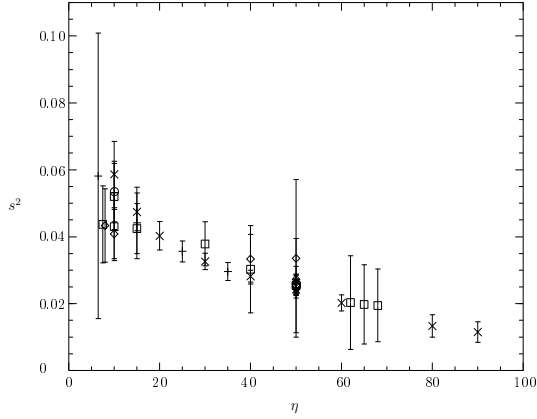


FIG. 3.16 – Dispersions de vitesse s^2 obtenues par ajustement des états d'équilibre post-effondrement par une sphère isotherme. Les barres d'erreur correspondent à l'erreur définie dans l'annexe A3.

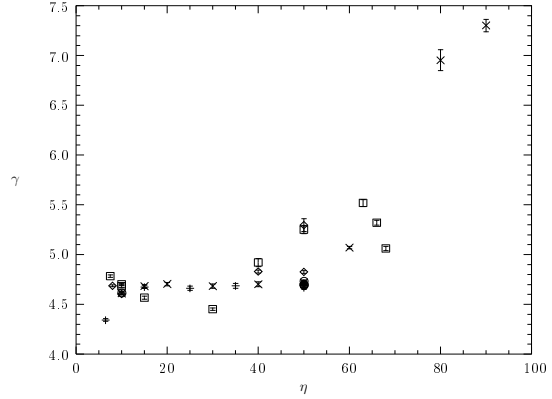


FIG. 3.17 – Indices polytropiques γ obtenus par ajustement des états d'équilibre post-effondrement par un polytrope. Les barres d'erreur correspondent à l'erreur définie dans l'annexe A3.

avec une écart à cette valeur égal à $1.18 \cdot 10^{-2}$ (36%). Les erreurs commises sur cet ajustement sont comprises entre $1.7 \cdot 10^{-3}$ et $5.7 \cdot 10^{-2}$, représentant souvent plus de 10% de la valeur de s^2 .

Cet ajustement est satisfaisant, mais de moins bonne qualité que celui que nous avons obtenu grâce aux polytropes. La valeur moyenne de l'indice polytropique γ obtenu est égale à 4.86, avec un écart à cette valeur égal à $5.9 \cdot 10^{-1}$ (12%) et une erreur pour chaque ajustement allant de $5.0 \cdot 10^{-3}$ à $6.5 \cdot 10^{-2}$, toujours inférieure à 3% de la valeur de γ . L'observation de la figure 3.17 nous permet de remarquer que deux points se détachent très largement des autres, correspondant aux deux systèmes présentant un paramètre du viriel initial η supérieur à 70, à savoir H_{88} et H_{79} . Bien que l'erreur commise lors de l'ajustement de ces deux systèmes soit très faible (moins de 1%), la valeur élevée de l'indice polytropique perturbe les statistiques précédentes. Si nous omettons donc ces deux systèmes, la valeur moyenne et l'écart à cette valeur deviennent respectivement 4.74 et $2.40 \cdot 10^{-1}$. La valeur de l'indice polytropique obtenue correspond globalement à celle du modèle de Plummer ($\gamma = 5$), solution analytique particulière du système BSC-Poisson (voir section 3.2 du chapitre 1). Seuls les deux systèmes précédemment cités, H_{88} et H_{79} , ne peuvent correspondre à un modèle de Plummer en raison de leur indice polytropique très largement supérieur à 5, égal respectivement à 7.37 et 6.86.

Ces résultats nous indiquent que les systèmes produits par effondrement gravitationnel lors de nos simulations ressemblent beaucoup à des polytropes. Cependant, certains diffèrent également assez peu d'une sphère isotherme. Il s'agit des systèmes de type H_η et M_η^k initialement chauds. De plus, il est intéressant de remarquer que les systèmes s'ajustant le moins bien par une sphère isotherme sont les systèmes de type (ii), c'est-à-dire les systèmes ne présentant pas de structure cœur-halo, alors que ceux présentant cette structure donnent de bons résultats. La valeur moyenne de s^2 pour les systèmes (ii) est égale à $3.38 \cdot 10^{-2}$, avec un écart à cette valeur égal à $1.11 \cdot 10^{-2}$ et une erreur moyenne de $1.2 \cdot 10^{-2}$. Pour les systèmes de type (i), nous trouvons une valeur moyenne de s^2 égale à $3.03 \cdot 10^{-2}$, un écart à cette valeur égal à $1.16 \cdot 10^{-2}$ et une erreur moyenne de $4.0 \cdot 10^{-3}$. Les résultats sont par contre comparable pour les systèmes de type (i) et les systèmes de type (ii) pour les ajustements par un polytrope. En effet, les valeurs moyenne de γ sont égales à 4.90 (ii) et 4.91 (i), les écarts à ces valeurs égaux à $3.4 \cdot 10^{-1}$ (ii) et $1.5 \cdot 10^{-1}$ (i) et les erreurs moyennes valent $2.0 \cdot 10^{-2}$ (ii) et $1.7 \cdot 10^{-2}$ (i). La qualité de l'ajustement par le modèle de Plummer provient donc probablement de la capacité de ce modèle à représenter des systèmes possédant des tailles de cœur très différentes, notamment très

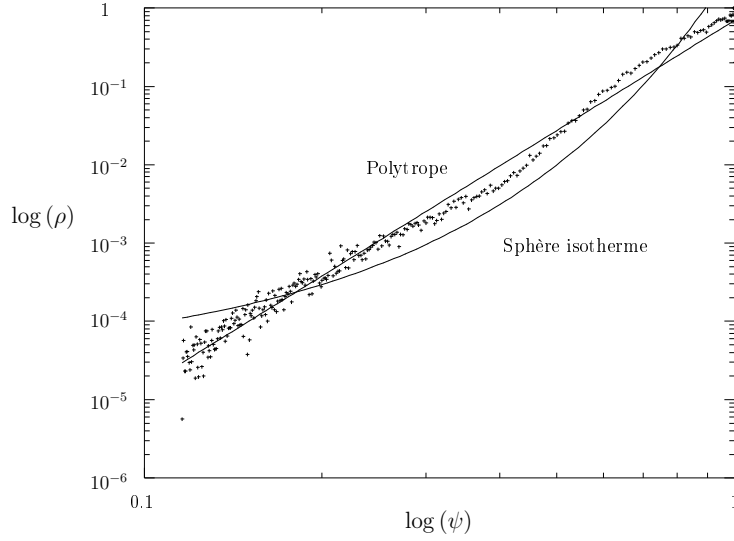


FIG. 3.18 – $\log(\rho)$ en fonction de $\log(\psi)$ pour le système $P_{50}^{0.5}$. Le polytrophe et la sphère isotherme obtenus par ajustement sont également représentés.

petites, grâce au paramètre b (voir l'équation (1.46)). Nous n'avons pas essayé d'utiliser d'autres types de couples densité–potentiel correspondant à des solutions analytiques de BSC–Poisson pour nos ajustement. Mais nous pensons que le modèle de King [30] donnerait des résultats équivalents à ceux que nous avons obtenus.

Comme les différents résultats analytiques connus le montrent, il existe donc plusieurs solutions analytiques différentes que nous pouvons comparer aux résultats de nos simulations. Il est probable que tout modèle pouvant décrire un système constitué d'un cœur et d'un halo de tailles variables s'ajusterait aux systèmes que nous avons produits.

6 Analyse thermodynamique

Après avoir étudié les propriétés morphologiques de nos états d'équilibre, nous avons analysé une propriété thermodynamique élémentaire de ces systèmes, leur température, que nous avons définie dans le paragraphe 3.8 du chapitre 2. La température que nous étudions dans ce paragraphe est celle de l'ensemble des particules présentes dans le nuage initial. Nous n'avons pas exclu du calcul de cette observable les particules éjectées au cours de l'effondrement.

Cette température évolue au cours de la simulation, mais atteint très rapidement une valeur d'équilibre, qui n'oscille que très peu. Nous avons représenté sur la figure 3.19 la température de cinq types de systèmes différents en fonction du temps. Nous pouvons observer qu'une fois la phase d'effondrement terminée, la température n'évolue plus que dans de faibles proportions.

De plus, comme nous étudions des systèmes isolés, l'énergie totale de ces systèmes est conservée au cours du temps, aux erreurs numériques près. En pratique, sur des évolutions durant $3000 u_t$, où u_t est notre unité de temps (voir le paragraphe 4 du chapitre 2), les variations d'énergie que nous avons observées n'ont pas excédé 1% de la valeur initiale de $\mathcal{H}$ dans la plupart des cas.

Il nous a donc paru intéressant de représenter la température en fonction de l'énergie totale pour chaque système sans rotation. Plus précisément, nous avons représenté la valeur moyenne des tempé-

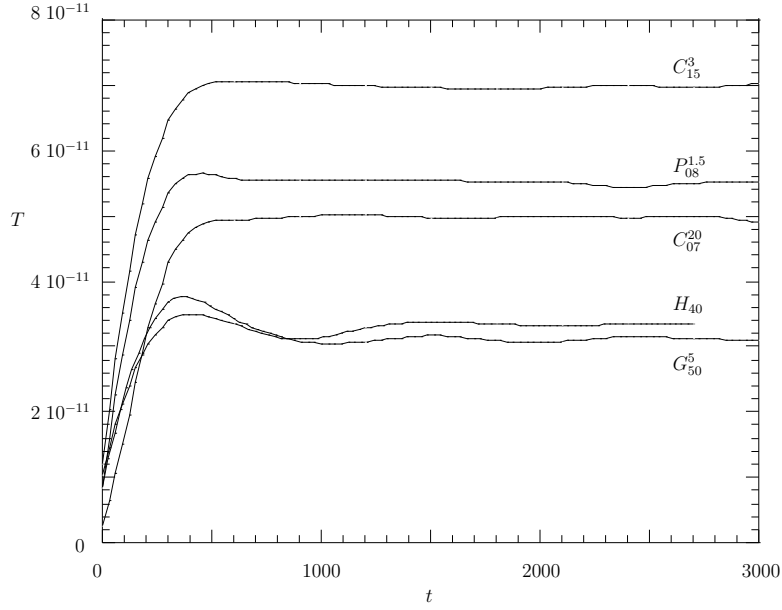


FIG. 3.19 – Évolution de la température en fonction du temps pour cinq systèmes différents.

ratures calculées toutes les 100 u_t entre le moment où la température atteint sa valeur à l'équilibre et le temps final. Le diagramme résultant est présenté sur la figure 3.20.

Une conclusion s'impose immédiatement à la vue de ce diagramme : en deçà d'une certaine valeur de l'énergie (environ -0.054 dans nos unités), tous les systèmes sont alignés dans le plan $\mathcal{H}-T$ et constituent ce que nous appellerons la Branche Basse 1 (*BB1*), alors que, au-delà de cette même valeur, nous observons deux comportements bien distincts. Une partie des systèmes avec $\mathcal{H} < -0.054$ demeure alignée avec les systèmes tels que $\mathcal{H} > -0.054$. Ces systèmes alignés avec la *BB1* forment la Branche Basse 2 ou *BB2*. Une autre partie des systèmes avec $\mathcal{H} < -0.054$ forme une famille séparée d'éléments alignés suivant une pente bien supérieure (Branche Haute ou *BH*).

Tous les systèmes de types H_η , G_η^σ et M_η^k avec $\eta > 25$ et tous les systèmes de types $P_\eta^{\alpha_d}$ et $C_\eta^{n_g}$ avec $\eta > 10$ se trouvent sur la *BB1* ou la *BB2*. Le système C_{10}^3 et tous les systèmes de types H_η , G_η^σ et M_η^k avec $\eta < 25$ se trouvent quant à eux sur la *BH*. La ségrégation n'a donc lieu que pour les systèmes issus d'effondrements violents, c'est-à-dire de conditions initiales froides. Lorsque les conditions initiales sont chaudes, tous les systèmes créés se situent dans la *BB1* du diagramme $\mathcal{H}-T$. Lorsque les conditions initiales sont froides et inhomogènes, les systèmes créés se trouvent sur la *BB2* alors que lorsqu'elles sont froides et homogènes ou quasi-homogènes, les systèmes créés se trouvent sur la *BH*.

La définition de notre température et la compréhension que nous avons de l'effondrement gravitationnel induisent l'analyse suivante de ces résultats.

Rappelons tout d'abord qu'un système initialement froid se réchauffe très fortement lors de l'effondrement gravitationnel, alors qu'un système initialement chaud se réchauffe beaucoup moins. Lorsqu'un système initialement froid et inhomogène s'effondre, la surdensité qui se forme rapidement au cœur (voir paragraphe 4) freine les particules traversant ce cœur de manière très efficace. Ce freinage est vraisemblablement le produit de la friction dynamique. Lorsqu'un système initialement froid et homogène s'effondre, le cœur se forme progressivement, et le freinage par friction dynamique est beaucoup moins efficace. Ceci implique qu'un nombre plus élevé de particules rapides peuvent échapper au puits de potentiel formé par le système et emporter une part plus importante de l'énergie cinétique totale

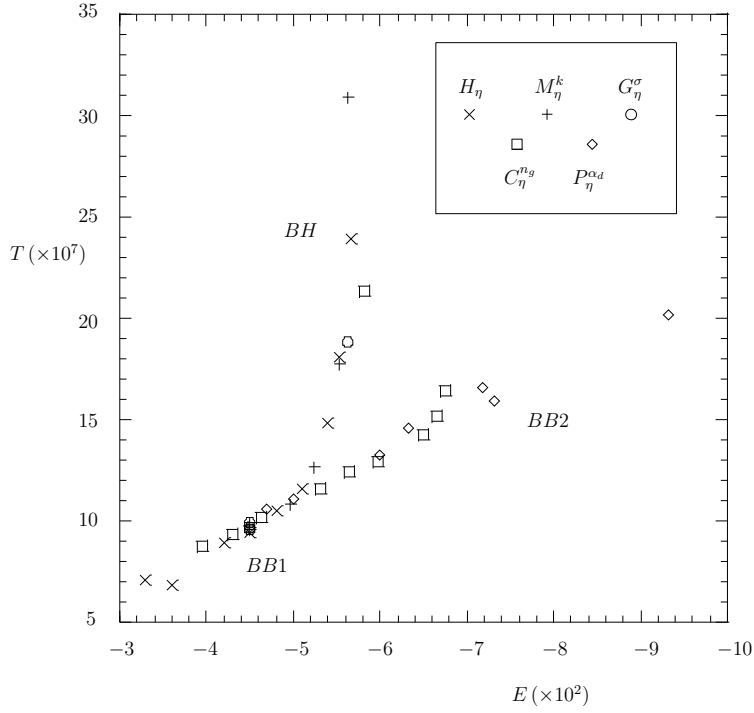


FIG. 3.20 – Diagramme température-énergie totale.

hors du système lié. Ceci est confirmé par le nombre de particules formant les systèmes liés finaux. Le système final C_{14}^{20} est par exemple formé de 82% du nombre initial de particules alors que le système final H_{15} est formé de 69% du nombre initial de particules. À rapport du viriel initial égal, et si ce rapport du viriel est supérieur à -0.5 , les systèmes homogènes perdent environ 10% de particules de plus que les systèmes inhomogènes. Pour des rapports du viriel inférieurs, il n'y a pas de différence significative.

Si nous considérons les systèmes liés seuls, afin de pouvoir envisager une comparaison entre nos résultats numériques et d'éventuels résultats observationnels, la ségrégation entre les systèmes initialement homogènes et ceux initialement inhomogènes est moins nette. La figure 3.21 représente le diagramme $\mathcal{H}_l - T_l$, où $\mathcal{H}_l$ est l'énergie du système lié et T_l sa température. Nous pouvons seulement constater qu'à énergie comparable, les systèmes de types H_η , G_η^σ ou M_η^k possèdent presque systématiquement une température supérieure aux systèmes de types C_η^{ng} ou P_η^{ad} lorsque $\mathcal{H}_l < -0.05$. Par exemple, le système C_{50}^{20} possède une énergie égale à -0.05299 et une température égale à $1.1701 \cdot 10^{-6}$ alors que l'énergie du système H_{30} vaut -0.05169 et sa température $1.3575 \cdot 10^{-6}$.

Si nous considérons que les systèmes liés sont virialisés, nous savons que leur énergie totale $\mathcal{H}_l$ est égale à leur énergie cinétique K_l au signe près. De plus, si nous appliquons notre définition de la température (voir le paragraphe 3.8 du chapitre 2) au système lié, nous savons encore que la température est proportionnelle à K_l/N_l , où N_l est le nombre de particules liées. Revenons à notre exemple. Sachant que le système lié C_{50}^{20} est constitué de 99.0% du nombre initial de particules alors que le système H_{30} ne contient plus que 81.3% des particules du nuage initial, nous pouvons de nouveau interpréter cette différence de température en terme de nombre de particules constituant le système lié.

Ainsi, pour un rapport du viriel initial donné, les nuages homogènes forment des systèmes contenant moins de particules que les nuages inhomogènes. Il en résulte que les systèmes liés formés à partir de conditions initiales homogènes sont plus chauds que ceux formés à partir de nuages inhomogènes à

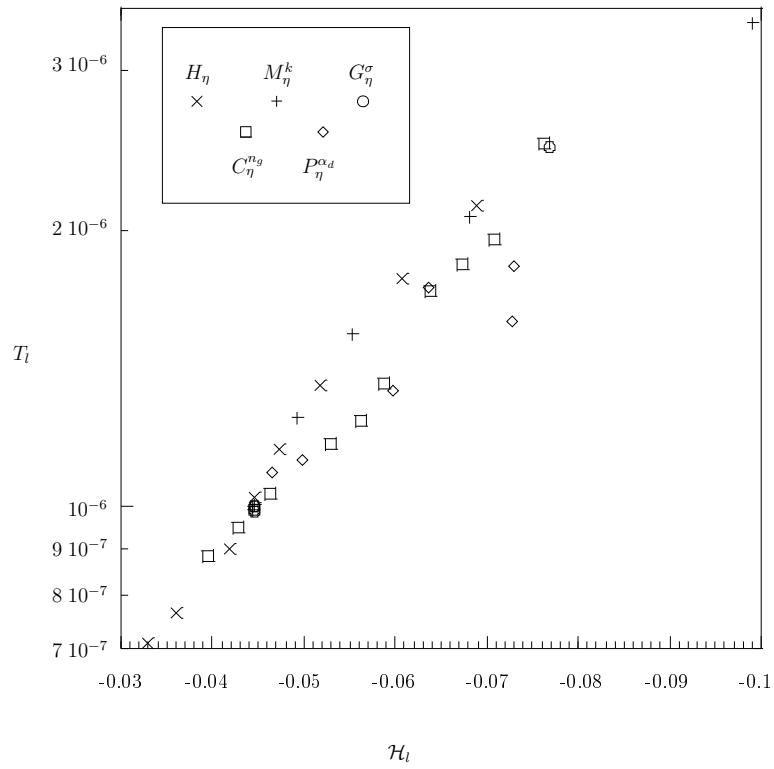


FIG. 3.21 – Diagramme température-énergie pour les systèmes liés uniquement.

énergie égale.

CONCLUSION

L'étude que nous avons réalisée confirme et complète tout à fait les résultats obtenus précédemment à l'aide de simulations numériques.

Ainsi, nous avons tout d'abord justifié l'utilisation de $N = 3 \cdot 10^4$ particules pour modéliser un système auto-gravitant en étudiant l'influence du nombre de particules sur les observables que nous avons l'intention de mesurer. Une autre étude préparatoire nous a permis de choisir un paramètre d'adoucissement du potentiel représentant un compromis intéressant entre la résolution spatiale, la conservation de l'énergie et les performances.

Nous avons de plus observé l'aplatissement des systèmes formés par effondrement gravitationnel d'un nuage froid et inhomogène de particules déjà mis en évidence entre autres par L.Aguilar et D.Merritt [3]. Nous avons de plus constaté que les systèmes initialement homogènes, à moins qu'ils ne subissent une forte rotation, forment après effondrement des systèmes sphériques, et ce quel que soit le rapport du viriel initial. L'instabilité d'orbites radiales, que nous invoquons pour expliquer la déformation de ces systèmes initialement froids, semble donc ne se déclencher qu'en présence d'inhomogénéités. Nous pensons plus précisément que l'instabilité d'orbites radiales ne se déclenche que si le système concentre très rapidement de la matière en son cœur, comportement que nous avons observé dans les systèmes inhomogènes mais pas dans les systèmes homogènes.

Nous avons ensuite mis en évidence un phénomène inobservé précédemment concernant la présence d'un large cœur au centre des systèmes formés. Nous avons montré que les nuages de particules initialement homogènes forment des systèmes possédant un large cœur de densité quasiment constante entouré d'un halo de densité décroissante. Les systèmes issus de nuages initialement inhomogènes possèdent un cœur de très petite taille (inférieure au rayon contenant 10% de la masse), et une densité uniformément décroissante avec le rayon. Leur densité centrale est de plus supérieure de un à deux ordres de grandeur à celle des systèmes initialement homogènes. Nous expliquons cette différence par l'instabilité d'Antonov. Cette instabilité résulte d'un fort contraste de densité entre le centre et l'extérieur du système et conduit à l'effondrement du cœur et à l'accroissement de sa densité. Les nuages inhomogènes que nous avons considérés possèdent la particularité de développer rapidement une forte surdensité de matière en leur centre, ils sont donc susceptibles de former un fort contraste de densité et de développer l'instabilité d'Antonov.

Nous avons enfin constaté que lorsqu'un nuage homogène s'effondre, le système lié formé comporte moins de particules que le système formé de l'effondrement d'un nuage inhomogène de même nombre de particules et de même rapport du viriel. Un système initialement homogène est également plus chaud qu'un système initialement inhomogène de même énergie. Ceci pourrait éventuellement être confirmé par des observations pouvant mesurer à la fois l'énergie totale et l'énergie cinétique de systèmes auto-gravitants.

Selon le scénario de formation "bottom-up" de formation des structures, les galaxies seraient issues de

milieux inhomogènes, alors que les amas globulaires, en raison de leur échelle de taille, auraient pour origine des milieux quasi homogènes.

Nos résultats se révèlent compatibles avec ces hypothèse. En effet, la très grande majorité des amas globulaires sont sphériques, de même que nos systèmes initialement homogènes. Leur fonction de distribution est correctement ajustée par un modèle de Plummer (structure cœur–halo), résultat que nous avons également mis en évidence numériquement, et ne possèdent pas de trou noir central d’après les observations menées jusqu’à présent.

Les galaxies elliptiques présentent quant à elles une ellipticité pouvant aller jusqu’à $e = 0.7$. Nous n’avons pas reproduit de si grandes ellipticités puisque notre système le plus elliptique est de type E3. Mais les résultats obtenus par L.Aguilar et D.Merritt montrent qu’il est possible d’atteindre de plus grandes ellipticités.

Nos simulations suggèrent de plus une ségrégation entre ces deux types d’objets dans le plan $\ln(R_{50}/R_d) - \ln(R_{10}/R_d)$. Les galaxies elliptiques correspondraient aux objets de type (ii) et les amas globulaires à ceux de type (i). Les rapports entre les rayons de densité et les rayons de 10% et de 50% de masse que nous avons calculés peuvent être confrontés aux données observationnelles. En effet, il existe une relation entre le rayon de densité et le rayon du cœur tel qu’il est défini pour les observations, à savoir le rayon pour lequel la densité de surface est égale à la moitié de sa valeur au centre. Des mesures de R_{10} et R_{50} permettraient la réalisation d’un diagramme $\ln(R_{50}/R_d) - \ln(R_{10}/R_d)$ et la vérification des résultats que nous avons obtenus.

S’il s’avérait exact que les galaxies elliptiques appartiennent au groupe (ii), il serait intéressant de mesurer précisément la densité centrale des systèmes obtenus par effondrement d’un nuage inhomogène. Puisque l’instabilité d’Antonov que nous avons identifiée provoque une forte augmentation de la densité centrale, elle pourrait être en mesure d’expliquer la présence de trous noirs au centre des galaxies. La résolution spatiale de nos simulations est probablement insuffisante pour mener une étude de ce genre, mais une nouvelle étude plus ciblée apporterait sans aucun doute des données permettant la confrontation entre les simulations et les observations.

Enfin, l’étude de l’évolution d’un amas globulaires au sein d’une galaxie hôte permettrait d’approfondir l’étude de l’instabilité d’Antonov et les effets de marées. En effet, les perturbations subies par l’amas seraient de nature à provoquer son évaporation partielle et éventuellement à déclencher l’instabilité d’Antonov selon le mécanisme exposé par J.Katz [28] et rappelé dans le paragraphe 3.4 concernant cette instabilité.

Il semble en tout état de cause que les simulations à N corps gravitationnels ne nous ont pas encore livré tous leurs secrets et qu’elles demeurent un sujet d’étude remarquablement intéressant.

ÉLÉMENTS DE CALCUL SCIENTIFIQUE

1 Génération de nombres pseudo-aléatoires : méthode de transformation inverse

La méthode de transformation inverse est utilisée pour générer des suites de nombres pseudo-aléatoires distribués selon une fonction f . Considérons la loi de densité continue et positive $f(x)$, et soit

$$F(x) = \int_{-\infty}^x f(t) dt$$

et u une variable aléatoire uniforme sur $[0, 1]$, de densité $\mathcal{U}$.

La probabilité pour que la variable aléatoire x de densité $f(x)$ soit comprise entre a et b est égale à

$$P_1 = P\{a < x < b\} = \int_a^b f(x) dx = F(b) - F(a) .$$

La probabilité pour que la variable aléatoire uniforme u soit comprise entre A et B est égale à

$$P_2 = P\{A < u < B\} = \int_A^B \mathcal{U}(u) du = B - A .$$

En prenant $A = F(a)$ et $B = F(b)$, nous obtenons

$$P_1 = P_2 = P\{A < u < B\} = P\{a < F^{-1}(u) < b\} ,$$

ce qui signifie que, si u est une variable aléatoire uniforme, alors $F^{-1}(u)$ est une variable aléatoire de loi de densité f .

Cette méthode nous permet d'obtenir des nombres aléatoires distribués selon des fonctions f telles que $f(x) \propto x^\alpha$ ou $f(x) \propto \exp(-x^2)$.

Examinons un exemple afin de noter quelques points pratiques à ne pas oublier. Prenons pour exemple une variable aléatoire comprise entre 0 et 1 de densité $f(r) = \alpha r^{-1}$. Cette densité doit être normalisée, c'est-à-dire

$$\int_0^1 f(r) dr = 1 .$$

Ceci nous pose un problème puisque cette intégrale diverge en 0. Comme nous résolvons le problème numériquement, nous allons contourner cette difficulté en utilisant un petit paramètre ϵ , et nous allons

générer des variables r comprises entre ϵ et 1. Nous pouvons maintenant normaliser f ainsi

$$f(r) = \frac{1}{\mathcal{S}_f} r^{-1} \text{ avec } \mathcal{S}_f = \int_{\epsilon}^1 r^{-1} dr = -\ln(\epsilon) .$$

Nous pouvons ensuite générer les variables r suivant f grâce à la méthode décrite ci-dessus, avec

$$F(r) = \frac{1}{\mathcal{S}_f} \int_{\epsilon}^1 r^{-1} dr = \frac{1}{\mathcal{S}_f} (\ln(r) - \ln(\epsilon)) , \quad (\text{A.1})$$

c'est-à-dire

$$r = F^{-1}(u) = \exp(\mathcal{S}_f u + \ln \epsilon) , \quad (\text{A.2})$$

où u est une variable aléatoire uniforme tirée entre 0 et 1.

Les nombres aléatoires uniformes sont quant à eux obtenus grâce à des algorithmes de type congruentiel, par exemple. Ces algorithmes fournissent des suites de nombres pseudo-aléatoires u_k définis par des relations du type

$$u_{k+1} = (a u_k + b) \bmod (m) ,$$

où a , b et m sont des paramètres du générateur. Nous avons utilisé un générateur de ce type, écrit par J.Barnes sur la base de la fonction RAN2 (voir [49]).

2 Schéma saute-mouton (ou leap-frog)

Le problème que nous simulons peut être résumé par le système suivant

$$\begin{cases} \frac{d\mathbf{r}}{dt} = \mathbf{v} , \\ \frac{d\mathbf{v}}{dt} = \mathbf{f}(\mathbf{r}) . \end{cases} \quad (\text{A.3})$$

Pour intégrer ce système dans le temps, nous utilisons une méthode de différences finies appelée le saute-mouton (ou leap-frog).

Le système (A.3) est approché grâce aux différences finies :

$$\frac{\mathbf{r}_{new} - \mathbf{r}_{old}}{\Delta t} = \mathbf{v}_{new} , \quad (\text{A.4})$$

$$\frac{\mathbf{v}_{new} - \mathbf{v}_{old}}{\Delta t} = \mathbf{f}(\mathbf{r}_{old}) , \quad (\text{A.5})$$

où Δt est appelé *pas de temps*. La particularité de ce schéma réside dans le décalage entre l'évaluation des vitesses et celle des positions. En effet, les positions sont évaluées à chaque pas de temps entiers, alors que les vitesses sont évaluées tous les demi-pas de temps (voir figure A.1).

Nous pouvons étudier la consistance¹ de ce schéma en considérant les positions et vitesses exactes $\bar{\mathbf{r}}$

¹Voir [12] pour plus de détails concernant les études de consistance et de stabilité.

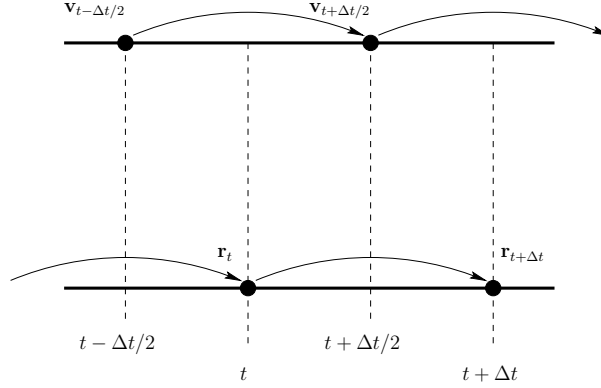


FIG. A.1 – Schéma saute-mouton.

et $\bar{\mathbf{v}}$. Nous pouvons écrire les développements de Taylor suivant :

$$\begin{aligned}\bar{\mathbf{r}}_{n+1} &= \bar{\mathbf{r}}_{n+1/2} + \frac{\Delta t}{2} \frac{d\bar{\mathbf{r}}_{n+1/2}}{dt} + \left(\frac{\Delta t}{2}\right)^2 \frac{d^2\bar{\mathbf{r}}_{n+1/2}}{dt^2} + O(\Delta t^3), \\ \bar{\mathbf{r}}_n &= \bar{\mathbf{r}}_{n+1/2} - \frac{\Delta t}{2} \frac{d\bar{\mathbf{r}}_{n+1/2}}{dt} + \left(\frac{\Delta t}{2}\right)^2 \frac{d^2\bar{\mathbf{r}}_{n+1/2}}{dt^2} + O(\Delta t^3), \\ \bar{\mathbf{v}}_{n+1/2} &= \bar{\mathbf{v}}_n + \frac{\Delta t}{2} \frac{d\bar{\mathbf{v}}_n}{dt} + \left(\frac{\Delta t}{2}\right)^2 \frac{d^2\bar{\mathbf{v}}_n}{dt^2} + O(\Delta t^3), \\ \bar{\mathbf{v}}_{n-1/2} &= \bar{\mathbf{v}}_n - \frac{\Delta t}{2} \frac{d\bar{\mathbf{v}}_n}{dt} + \left(\frac{\Delta t}{2}\right)^2 \frac{d^2\bar{\mathbf{v}}_n}{dt^2} + O(\Delta t^3),\end{aligned}$$

où $\bar{\mathbf{r}}_k$ et $\bar{\mathbf{v}}_k$ désignent respectivement la position et la vitesse exacte au temps $k\Delta t$. Les erreurs commises lors de l'évaluation des nouvelles positions et vitesses par (A.4)–(A.5), données par

$$\begin{aligned}\mathcal{R}_{n+1}^{\mathbf{r}} &= \frac{\bar{\mathbf{r}}_{n+1} - \bar{\mathbf{r}}_n}{\Delta t} - \bar{\mathbf{v}}_{n+1/2}, \\ \mathcal{R}_{n+1/2}^{\mathbf{v}} &= \frac{\bar{\mathbf{v}}_{n+1/2} - \bar{\mathbf{v}}_{n-1/2}}{\Delta t} - \mathbf{f}(\mathbf{r}_n),\end{aligned}$$

sont égales à

$$\begin{aligned}\mathcal{R}_{n+1}^{\mathbf{r}} &= \frac{1}{2} \left(2 \frac{d\bar{\mathbf{r}}_{n+1/2}}{dt} \right) - \bar{\mathbf{v}}_{n+1/2} + O(\Delta t^3) = O(\Delta t^3), \\ \mathcal{R}_{n+1/2}^{\mathbf{v}} &= \frac{1}{2} \left(2 \frac{d\bar{\mathbf{v}}_n}{dt} \right) - \mathbf{f}(\mathbf{r}_n) + O(\Delta t^3) = O(\Delta t^3).\end{aligned}$$

Le schéma saute-mouton est donc d'ordre 2.

Il est important de noter également que le schéma que nous employons est un schéma que nous pourrions qualifier d'*à limiteur d'accélération*. Ainsi, comme nous l'avons déjà mentionné dans la section 1.2 du chapitre 2 lors du calcul de l'accélération $\dot{\mathbf{v}}_i$ subie par la particule i , si une particule j se trouve à une distance inférieure à ε de i , alors sa contribution à l'accélération $\dot{\mathbf{v}}_i$ ne sera pas prise en compte. Ceci permet de prévenir une trop forte augmentation de l'accélération de la particule i , bien entendu au détriment de l'exactitude de l'intégration de sa trajectoire. Cette limitation de l'accélération suggère que si la résolution en espace, donc ε , diminue, l'accélération maximale subie

par une particule augmentera. L'intégration en temps ne pourra alors pas s'effectuer avec une aussi grande précision à moins que le pas de temps ne diminue, puisque l'erreur commise sur le calcul de la position $\mathbf{r}^n$ au temps n dépend du produit $\mathbf{v}^n \Delta t$. Il existe de ce fait une condition de type CFL, que nous avons expérimentalement observée (voir la figure 3.3). Cette condition impose l'utilisation d'un pas de temps tel que $\Delta t/\varepsilon$ soit inférieur à une constante. Si cette condition n'est pas remplie, l'erreur commise augmente très rapidement.

Cette limitation de l'accélération provoque donc une erreur lors de l'estimation des nouvelles vitesses, mais permet de limiter l'erreur due au schéma "saute-mouton".

3 Méthode des moindres carrés

La méthode des moindres carrés permet d'ajuster un ensemble de points (X_i, Y_i) par une fonction f en minimisant la somme Σ^2 des carrés des distances entre les Y_i et les $f(X_i)$

$$\Sigma^2 = \sum_{i=1}^n (Y_i - f(X_i))^2 .$$

Dans le cas d'un ajustement par une fonction affine $f(X_i) = a X_i + b$, Σ^2 est une fonction de a et b . Σ^2 est minimale à la condition que

$$\frac{\partial \Sigma^2}{\partial a} = 0 \quad \text{et} \quad \frac{\partial \Sigma^2}{\partial b} = 0 ,$$

c'est-à-dire

$$-2 \sum_{i=1}^n (Y_i - (a X_i + b)) X_i = 0 \quad \text{et} \quad -2 \sum_{i=1}^n (Y_i - (a X_i + b)) = 0 . \quad (\text{A.6})$$

En manipulant les équations (A.6), il vient

$$a = \frac{\text{cov}(X, Y)}{\text{var}(X)} \quad \text{et} \quad b = E(Y) - a E(X) .$$

Rappelons que l'espérance est définie, dans notre cas (probabilité uniforme) par

$$E(X) = \frac{1}{N} \sum_{i=1}^N X_i ,$$

la variance par

$$\text{var}(X) = \frac{1}{N} \sum_{i=1}^N (X_i - E(X))^2$$

et la covariance par

$$\text{cov}(X, Y) = E[(X - E(X))(Y - E(Y))] .$$

L'erreur effectuée peut être estimée par la distance Σ^2 entre la variable Y et l'estimateur $f(X)$ égale à

$$\Sigma^2 = \text{var}(Y) - \frac{\text{cov}^2(X, Y)}{\text{var}(X)} .$$

NOTIONS DE CALCUL PARALLÈLE ET DISTRIBUÉ – INTRODUCTION À LA BIBLIOTHÈQUE MPI

1 Calcul parallèle

1.1 Présentation du calcul parallèle

De nombreuses applications numériques scientifiques nécessitent des ressources informatiques de plus en plus performantes. Cette croissance des besoins a pour origine la volonté des scientifiques de résoudre des problèmes de plus en plus complexes, ou bien de résoudre le même problème de plus en plus rapidement. L'évolution technologique des ordinateurs – matérielle, logicielle – n'étant pas suffisante à elle seule pour répondre à ces besoins de puissance de calcul, il est nécessaire de trouver une solution algorithmique à ce problème.

Une solution largement employée consiste à diviser le problème étudié en plusieurs parties, et à résoudre chacune de ces fractions de problème en même temps. Si n tâches, demandant chacune un temps d'exécution $t^{(1)}$, sont exécutées séquentiellement sur un seul processeur, le temps total de calcul est égal à $nt^{(1)}$. Si maintenant ces n tâches sont exécutées en même temps sur n processeurs différents, le temps total de calcul reste égal à $t^{(1)}$.

Le calcul parallèle repose sur cette idée : plusieurs processeurs effectuant de manière coopérative un calcul ou un travail bien précis. Nous pouvons illustrer cette "définition" à l'aide d'un exemple simple, comme le calcul d'un produit matrice vecteur $v' = Mv$, avec $M \in \mathbb{R}^{n \times n}$ et $v \in \mathbb{R}^n$. Ce calcul nécessite n^2 opérations lorsqu'il est effectué de façon séquentielle. Nous disposons de n processeurs, c'est-à-dire autant que la matrice a de lignes. Il est donc logique de demander à chaque processeur de calculer le produit scalaire d'une ligne de matrice par le vecteur v . Pour obtenir v' , il suffit alors de rassembler les résultats obtenus par chacun des processeurs. Ainsi, si tous les processeurs effectuent les n opérations nécessaires au produit scalaire en même temps, nous aurons obtenu notre résultat en un temps proportionnel à n et non n^2 .

Le calcul parallèle est mis en place sur des machines composées de plusieurs processeurs, appelées machines parallèles. Ces machines peuvent être classées en plusieurs catégories suivant, entre autres, leur type de mémoire ou le flot d'instructions traitées par les processeurs. Concernant le flot d'instructions, la plupart des machines étant maintenant de classe MIMD (Multiple Instruction, Multiple Data – les différents processeurs effectuent des tâches différentes sur des données différentes), ce sont les seules que nous considérerons.

Parmi ces machines MIMD, nous pouvons distinguer les machines à mémoire partagée et les machines à mémoire distribuée. Dans une machine à mémoire partagée, tous les processeurs partagent la même mémoire (voir figure (B.1)). Lorsque l'un des processeurs a terminé un accès à une zone de la mémoire, les autres processeurs peuvent avoir à leur tour accès à cette même zone. Par exemple, lorsqu'un processeur a terminé un calcul et en a écrit le résultat dans la mémoire, les autres processeurs peuvent avoir accès à ce résultat. A l'inverse, les processeurs d'une machine à mémoire distribuée possèdent

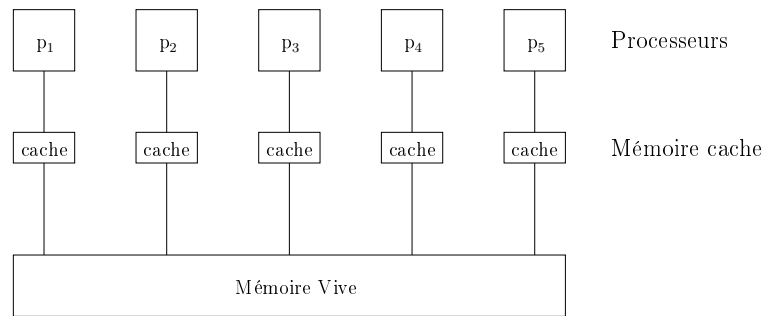


FIG. B.1 – Machine à mémoire partagée.

chacun leur propre mémoire (voir figure (B.2)). Les données présentes sur la mémoire attachée à un processeur ne sont donc pas directement accessibles aux autres processeurs. Lorsqu'un processeur a a terminé un calcul et en a écrit le résultat sur sa mémoire, si un autre processeur b a besoin de ce résultat pour effectuer son propre calcul, a doit lui communiquer ce résultat au travers du réseau reliant les différents processeurs entre eux.

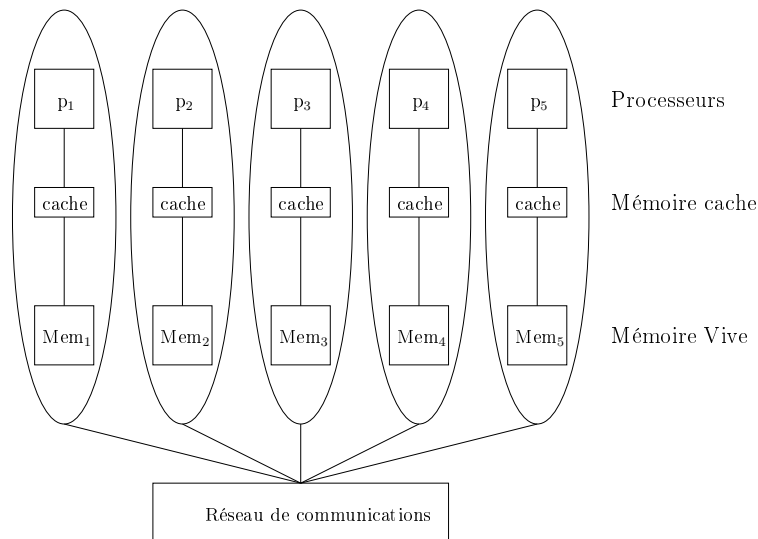


FIG. B.2 – Machine à mémoire distribuée.

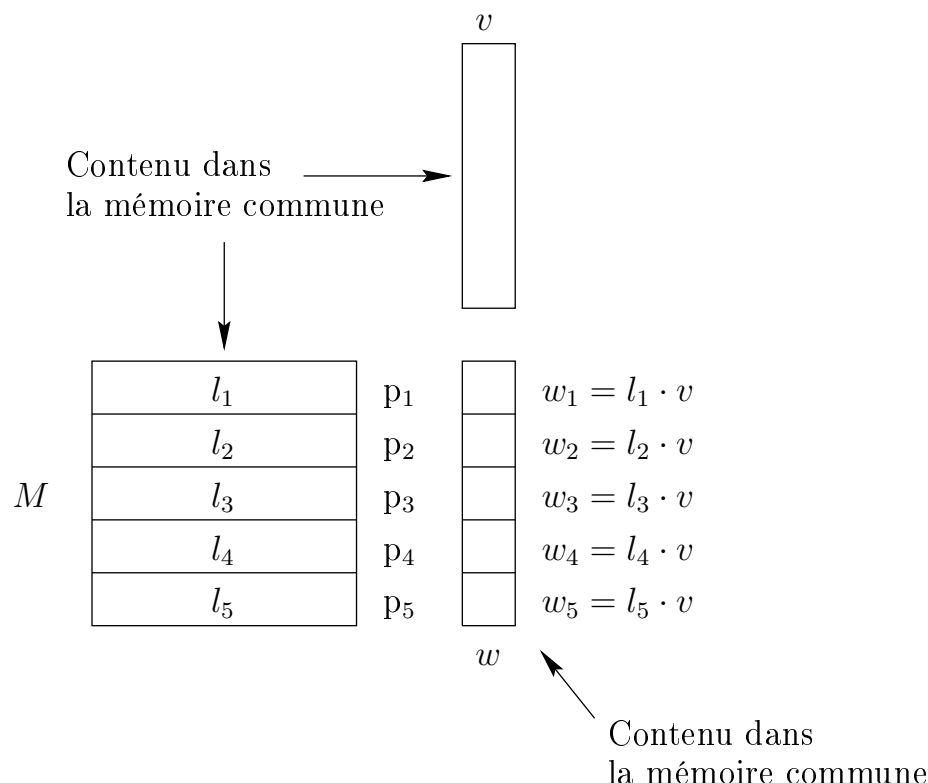


FIG. B.3 – Exemple du calcul d'un produit matrice vecteur sur une machine à mémoire partagée.

1.2 Machines à mémoire partagée, machines à mémoire distribuée : les différences pratiques

Comme nous l'avons mentionné plus haut, on peut distinguer deux grands types de machines parallèles : les machines à mémoire partagée et les machines à mémoire distribuée. Les premières sont des calculateurs dans lesquels plusieurs processeurs partagent une même mémoire physique. Dans les secondes, les processeurs travaillent chacun avec une mémoire qui leur est propre.

Bien évidemment, la parallélisation ne peut s'effectuer de la même manière sur ces deux types de calculateurs. Dans le cas des machines à mémoire partagée, chaque processeur a à sa disposition les mêmes informations que les autres, alors que les processeurs d'une machine à mémoire partagée doivent échanger les informations nécessaires à la réalisation de leur travail. Revenons à notre exemple précédent, celui du produit matrice vecteur, afin de noter les différences entre nos deux types de machines. La matrice et le vecteur sont stockés sur le disque dur de la machine à mémoire partagée, et sur le disque dur d'un seul ordinateur composant la machine à mémoire distribuée. Nous souhaitons que le résultat soit écrit, à la fin du calcul, sur un seul emplacement. Dans la machine à mémoire partagée, chaque processeur peut avoir accès à chaque élément de la matrice et du vecteur, puisque les données contenues dans l'unique disque dur sont chargées dans l'unique mémoire. Il suffit donc au programmeur de désigner, pour chaque processeur, le numéro l_n de la ligne de matrice qu'il doit multiplier avec le vecteur. Le résultat est ensuite rangé à la l_n -ième position du vecteur résultat dans la mémoire commune, puis est écrit sur le disque dur (voir la figure (B.3)). Le calcul est achevé.

Dans la machine à mémoire distribuée, la matrice et le vecteur ne sont connus à l'instant initial que du processeur dont le disque dur contient les fichiers de données (en supposant qu'une seule machine

dispose des données). Le premier travail de ce processeur consiste donc à distribuer à chacun des processeurs la ligne de matrice qui lui est destinée. Le calcul s'effectue ensuite sur chaque processeur. À ce moment, le résultat est éparpillé entre les mémoires associées à chaque processeur. Un des processeurs doit donc centraliser les résultats, en demandant aux autres de lui faire parvenir leur élément de v' . Une fois que ce processeur a tous les éléments de v' en mémoire, il peut écrire le résultat complet sur le disque dur qui lui est associé (voir la figure (B.4)).

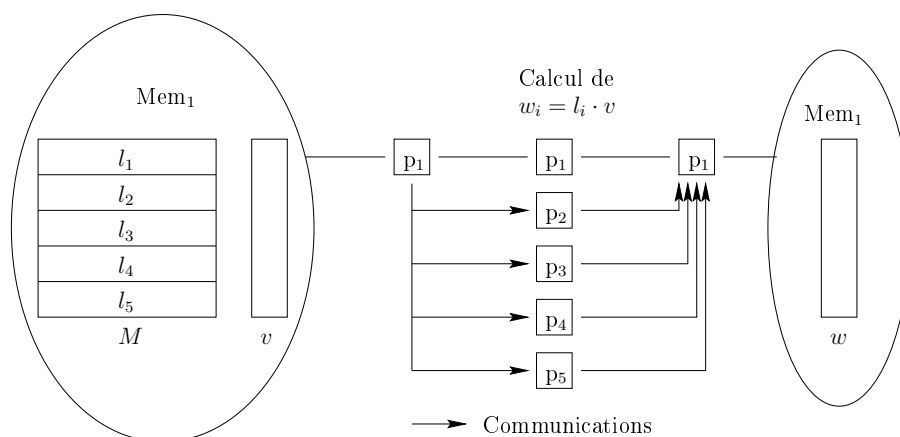


FIG. B.4 – Exemple du calcul d'un produit matrice vecteur sur une machine à mémoire distribuée.

1.3 La solution logicielle et matérielle utilisée

Le parallélisme n'est pas adapté à tous les types d'algorithme, mais, par chance, les méthodes que nous utilisons se prêtent à une parallélisation simple et efficace.

La parallélisation du treecode a été réalisée par Daniel Pfenniger (Observatoire de Genève), alors que nous avons parallélisé les autres codes nécessaires à nos études.

Tous les travaux numériques présentés dans ce document ont été réalisés sur une grappe d'ordinateurs de type Beowulf dont le Laboratoire de Mathématiques Appliquées de l'ENSTA s'est doté il y a 4 ans. Une grappe Beowulf est une machine composée de plusieurs ordinateurs reliés entre eux par un réseau de type Ethernet, Myrinet ou autre. Ces ordinateurs sont autonomes et disposent d'un ou plusieurs processeurs, de mémoire, d'un disque dur, et d'une interface réseau. Une grappe Beowulf est donc une machine à mémoire distribuée et non à mémoire partagée.

Ce type de machines présente de nombreux avantages. Leur rapport performance/coût est remarquable. Il suffit de quelques ordinateurs de bureau standards et d'un réseau Ethernet pour constituer une grappe. De plus, il est possible de faire évoluer une grappe sans aucune difficulté, en ajoutant quelques ordinateurs ou en changeant le réseau de communication (pour passer d'un réseau Ethernet 10Mbps à un réseau Gigabit Ethernet par exemple). Ces deux caractéristiques différencient une grappe d'un super-calculateur classique, dont le coût est très supérieur et qui ne sera pas forcément évolutif. La grappe PCP est composée, dans sa configuration actuelle, de 10 ordinateurs bi-processeurs Intel Celeron 400MHz, 10 ordinateurs bi-processeurs Intel Pentium III 1GHz et 10 ordinateurs bi-processeurs AMD Athlon 1900+ et un ordinateur bi-processeur AMD Athlon 2000+ – ce dernier est la machine maître, c'est-à-dire la machine qui commandera l'exécution globale des applications parallèles. A ces 31 ordinateurs, qui sont les nœuds de calcul (des ordinateurs ont pour seule fonction de servir au calcul scientifique), il faut ajouter un ordinateur servant de passerelle entre la grappe PCP et le réseau informatique du laboratoire, et un ordinateur servant de machine d'administration. Le réseau

de communication était de type Ethernet 100Mbps lorsque nous avons effectué nos calculs, mais il a depuis évolué vers un réseau de type Ethernet Gigabyte.

La parallélisation des codes repose sur la bibliothèque d'échange de message MPI LAM¹. L'utilisation de ce type d'architecture machine et du calcul distribué s'explique par diverses raisons, parmi lesquelles la disponibilité de logiciels libres performants et l'adéquation entre les méthodes numériques employées et le calcul distribué.

2 MPI – Message Passing Interface

2.1 Le passage de message et le standard MPI

L'utilisation d'une machine et du parallélisme à mémoire distribuée requiert un mécanisme de communication entre les différents processeurs composant la machine parallèle, le passage de message. Cela consiste à demander aux processeurs d'émettre et de recevoir des messages vers ou en provenance des autres processeurs de la machine. Ces messages transitent par le réseau de connexions, dans notre cas un réseau Ethernet. La figure (B.5) illustre ces passages de messages. Une application parallélisée

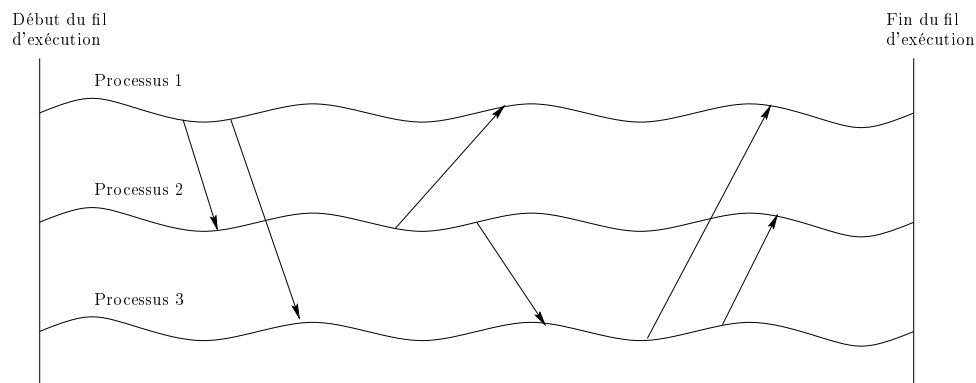


FIG. B.5 – Exemple de passage de message entre trois processus.

avec MPI est donc composée de processus tournant chacun sur un processeur différent, et travaillant en collaboration, en échangeant des informations, de manière à aboutir plus rapidement au résultat recherché.

Si nous reprenons notre exemple de calcul de produit matrice vecteur (figure B.4), les premières communications sont alors des messages provenant du processeur p_1 et destinés aux processeurs p_2 à p_5 , le message destiné à p_i contient la ligne l_i de la matrice. Les communications suivantes sont des messages provenant des processeurs p_2 à p_5 et destinés à p_1 : chacun de ces messages contenant un élément de résultat.

Les échanges de messages sont coûteux en matière de temps d'exécution. Il est donc préférable de réduire leur nombre en regroupant les envois. L'envoi d'un seul message volumineux est en effet plus rapide qu'une succession d'envois de taille plus réduites. Le temps de transfert t_t d'un message est divisé en un temps fixe et incompressible t_i , correspondant à l'initialisation de l'envoi, et un temps t_e dépendant de la taille du message et de la bande passante du réseau², soit

$$t_t = t_i + t_e (\text{taille, bande passante}) .$$

¹Voir le site Internet <http://www.lam-mpi.org/>

²La bande passante est la quantité maximale d'information pouvant transiter par le réseau par unité de temps. Elle s'exprime par exemple en Mo par seconde.

L'initialisation de l'envoi interrompt l'exécution du code, alors que l'envoi proprement dit, c'est-à-dire le temps pendant lequel le message transite sur le réseau, peut être, au moins en partie, superposé aux calculs. L'envoi de messages multiples plutôt que d'un seul message regroupant toutes les données multiplie inutilement le temps t_i et est donc plus coûteux que l'envoi d'un message unique regroupant l'ensemble de l'information.

MPI est une interface de programmation d'applications (API) permettant de programmer très simplement des échanges de messages entre différents processus³. MPI peut être utilisée dans des programmes écrits en langage Fortran, C ou C++, sur des environnements hétérogènes. MPI LAM est une distribution gratuite de MPI, performante, et compatible avec un certain nombre d'outils intéressants tels que XMPI, outil graphique de visualisation et de diagnostics de l'exécution d'une application parallèle.

2.2 Nœuds et processus

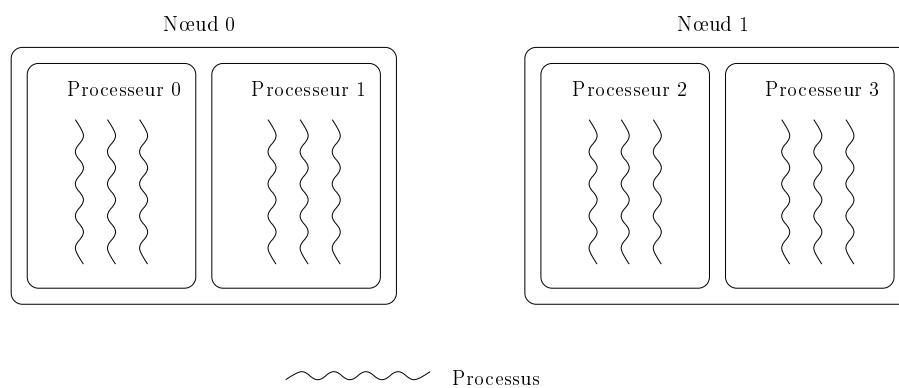


FIG. B.6 – Distinction entre nœuds, processeurs et processus.

Avant d'exécuter un programme parallélisé à l'aide de MPI LAM, il est nécessaire d'établir une liste de processeurs sur lesquels les processus seront lancés. Prenons l'exemple de la machine PCP du Laboratoire de Mathématiques Appliquées de l'ENSTA. Les nœuds de calcul de PCP sont numérotés de pcp1 à pcp31, pcp31 étant la machine maître. Les applications sont lancées à partir de la machine maître : cette machine doit être la première sur notre liste. Ensuite, il suffit d'ajouter le nombre de machines désiré à la liste, sachant qu'il est possible d'exécuter plusieurs processus sur un même processeur, mais que ceci va au détriment de l'efficacité. Par exemple, la liste de nœuds suivante conviendra à l'exécution de 4 processus :

```
pcp31
pcp1
pcp2
pcp3
```

Les ordinateurs de PCP contenant deux processeurs, il est possible d'inscrire deux fois le nom de chaque machine dans la liste de nœuds. Ceci signifie simplement que les deux processeurs des machines inscrites deux fois seront utilisés. Ainsi, la liste suivante convient également à l'exécution de 4 processus :

```
pcp31
pcp31
pcp1
pcp1
```

Une commande `lamboot fichier-noeuds` permet de préparer les nœuds contenus dans la liste à l'exécution d'une application MPI (le démon lamd est démarré sur ces nœuds). Une fois cette commande exécutée, les nœuds sont numérotés de n0 à np, où p+1 est le nombre de machines différentes

³Un processus est un programme en cours d'exécution.

présentes dans le fichiers des nœuds. Ainsi, dans notre première liste, `pcp31` correspond à `n0`, `pcp1` à `n1`, `pcp2` à `n2` et `pcp3` à `n3`. Dans notre deuxième liste, `pcp31` correspond à `n0`, `pcp1` à `n1`. Une fois le fichier contenant la liste de nœuds édité, et la commande `lamboot fichier-noeuds` exécutée, il faut éditer un fichier contenant la liste des processus que nous souhaitons voir exécutés sur les nœuds. Cette liste associe un processus à un nœud de calcul. Elle se présente sous la forme suivante, qui correspond à notre premier exemple de liste de nœuds :

```
n0 -s h executable1
n1 -s h executable2
n2 -s h executable3
n3 -s h executable4
```

(`-s h` représente une option que nous ne détaillerons pas ici). Les exécutables numérotés de 1 à 4 peuvent être tous identiques, ou tous différents, mais ils doivent avoir été conçus pour fonctionner ensemble, c'est-à-dire que les appels aux fonctions MPI contenus dans ces exécutables doivent être compatibles entre eux (si l'exécutable 1 envoie un message à l'exécutable 2, celui-ci doit le recevoir). Si une machine est présente deux fois dans la liste de nœud, elle peut l'être également deux fois dans la liste d'exécutables. Par exemple, le fichier d'applications suivant pourrait correspondre à notre deuxième exemple de liste de nœuds :

```
n0 -s h executable1
n0 -s h executable2
n1 -s h executable3
n1 -s h executable4
```

L'application consistant en l'exécution parallèle des exécutables contenus dans la liste des exécutables est prête à être lancée. Il suffit pour cela de taper une simple commande :

```
mpirun -0 fichier-applications.
```

Il est donc extrêmement simple d'utiliser une application parallélisée grâce à MPI LAM. Mais avant cela, il est nécessaire d'inclure dans les exécutables les instructions MPI nécessaires à la parallélisation de l'application.

Tous les détails syntaxiques présentés par la suite correspondent à l'utilisation de procédures MPI par un programme écrit en Fortran77. Comme il est d'usage avec l'utilisation de bibliothèques de programmes, la compilation des exécutables doit contenir le chemin vers les bibliothèques utilisées.

Pour plus de détails sur la programmation parallèle à l'aide de MPI, vous pouvez vous référer à [20], ou bien aux documents suivants, disponibles sur Internet : [43] ou [44].

2.3 Communicateurs et communications point à point

Un exécutable MPI doit obligatoirement contenir les trois éléments suivants :

- inclusion du fichier `mpif.h`,
- appel à la procédure `Mpi_Init`, avant tout autre appel à une procédure MPI,
- appel à la procédure `Mpi_Finalize`, après tout autre appel à une procédure MPI.

Ces trois éléments n'ont aucune action particulière sur la parallélisation effective de l'application, mais ils sont nécessaires.

Les processus composant notre application parallèle sont regroupés dans des communicateurs. Il existe par défaut un communicateur, appelé `Mpi_Comm_World`, contenant l'ensemble des processus. L'utilisateur a ensuite la possibilité de définir d'autres communicateurs. Nous verrons l'utilité des communicateurs par la suite, mais nous ne considérerons pour l'instant que le communicateur par défaut.

Il est nécessaire d'identifier chaque processus, ou plus précisément que chaque processus puissent s'identifier, afin que les messages puissent être adressés à un processus particulier plutôt qu'à n'importe lequel d'entre eux. Cette identification est possible grâce au rang du processus. Le rang est un entier associé à un communicateur compris entre 0 et p , où $p + 1$ est le nombre total de processus contenu dans le communicateur. Il est également nécessaire, dans une très grande majorité des cas, que chaque processus sache combien de processus sont contenus dans son communicateur. Ces deux informations sont obtenues respectivement à l'aide des procédures `Mpi_Comm_Rank` et `Mpi_Comm_Size`. Ces quelques procédures nous permettent déjà de réaliser l'exemple très simple d'application parallèle suivant, l'équivalent du traditionnel "Bonjour le monde".

```

Program Bonjour

Include 'mpif.h'
Integer rang
Integer ierr

Call Mpi_Init(ierr)
Call Mpi_Comm_Rank(Mpi_Comm_World,rang,ierr)
Print *, 'Bonjour, ici le processus #',rang
Call Mpi_Finalize(ierr)

End

```

L'exécution de 4 processus comme celui-ci sur 4 processeurs produit le résultat suivant⁴ :

```

Bonjour, ici le processus # 0
Bonjour, ici le processus # 1
Bonjour, ici le processus # 2
Bonjour, ici le processus # 3

```

Une application parallèle dont les processus ne communiquent pas entre eux ne présente pour nous aucun intérêt. Nous allons donc maintenant introduire les commandes standards d'envoi et de réception de message. Ces commandes se nomment `Mpi_Send` et `Mpi_Recv`. La commande d'envoi s'utilise de la manière suivante :

`Mpi_Send(contenu,taille,type,destinataire,identifiant,communicateur,erreur)`

ce qui signifie qu'un message composé du contenu de la variable `contenu`, consistant en `taille` éléments de type `type`, portant l'identifiant `identifiant`, est adressé au processus portant le numéro `destinataire` dans le communicateur `communicateur`. La variable `erreur` contient un code d'erreur. Et celle de réception s'écrit :

`Mpi_Recv(contenu,taille,type,emetteur,identifiant,communicateur,status,erreur)`

ce qui signifie qu'un message consistant en `taille` éléments de type `type`, portant l'identifiant `identifiant`, provenant du processus portant le numéro `emetteur` dans le communicateur `communicateur`, sera stocké dans la variable `contenu`. La variable `erreur` contient un code d'erreur, et la variable `status` indique l'état de la réception du message.

Voici un exemple d'utilisation d'envoi et de réception par deux processus :

```

Program Bonjour

Include 'mpif.h'
Integer ierr
Integer message
Integer rang
Integer reçu
Integer status(Mpi_Status_Size)
Integer taille

Call Mpi_Init(ierr)
Call Mpi_Comm_Rank(Mpi_Comm_World,rang,ierr)
Call Mpi_Comm_Size(Mpi_Comm_World,taille,ierr)

If(rang.Eq.0) Then
  Print *, 'Nombre de processus : ',taille
  message = 2
  Call Mpi_Send(message,1,Mpi_Integer,1,0,Mpi_Comm_World,ierr)
  Print *, 'Bonjour, ici le processus #',rang,
&      ', message envoye = ',message

```

⁴Les affichages provenant des différents processus n'ont aucune raison de s'effectuer dans l'ordre de la numérotation : le processus 1 peut afficher son résultat avant le processus 0 et après le 3. L'ordre des affichages provenant de processus différents est donc hasardeux et anecdotique.

```

Endif
If(rang.Eq.1) Then
    Call Mpi_Recv(recu,1,Mpi_Integer,0,0,Mpi_Comm_World,
&        status,ierr)
    Print *,'Bonjour, ici le processus #',rang,
&        ', message reçu = ',recu
Endif
Call Mpi_Finalize(ierr)

End

```

L'exécution de ce programme sur deux processeurs produit le résultat suivant :

```

Nombre de processus : 2
Bonjour, ici le processus # 0, message envoyé = 2
Bonjour, ici le processus # 1, message reçu = 2

```

Ces commandes d'envoi et de réception sont bloquantes, ce qui signifie que l'exécution du code est bloquée jusqu'à ce que l'envoi ou respectivement la réception soit terminée. Lors d'un envoi standard, MPI choisit de faire ou non une copie temporaire du contenu du message. Si une copie est effectuée dans un tampon, l'envoi est considéré comme achevé lorsque la recopie est effectuée, sinon, il ne se termine que lorsque la commande de réception correspondante est exécutée par le processus destinataire. La commande de réception standard, quant à elle, n'est considérée comme achevée que lorsque le message a effectivement été réceptionné.

Il est donc possible qu'un processus perde un temps précieux à attendre qu'un envoi ou une réception se termine. Imaginons le cas suivant : un processus p_1 effectue un envoi standard vers un processus p_2 à un instant t_1 . Le contenu du message n'est pas recopié dans un tampon. Le processus p_2 exécute la commande de réception au temps t_2 . Le processus p_1 n'a rien fait entre les temps t_1 et t_2 . Si la commande de réception est exécutée par p_2 au temps t_1 et que la commande d'envoi est exécutée par p_1 au temps t_2 , le processus p_2 ne fait rien entre t_1 et t_2 (voir la figure B.7).

Il existe des commandes d'envoi et de réception non bloquantes, qui permettent au processus de poursuivre son exécution sans attendre que l'envoi ou la réception soit complétée. Cependant, ces commandes non bloquantes sont à utiliser avec précautions. Il est par exemple dangereux de réutiliser la mémoire tampon contenant les messages trop rapidement après un envoi non bloquant avant de s'assurer que le message a bien été envoyé. Il est également fortement déconseillé de tenter d'utiliser des données reçues avec une réception non bloquante sans s'être au préalable renseigné sur l'état de la réception.

Il existe enfin d'autres types d'envoi et de réception point à point que nous n'exposerons pas ici, mais qui sont détaillés dans les ouvrages cités précédemment.

2.4 Communications collectives et communicateurs (compléments)

Lorsque le nombre de processus utilisés est grand, il est important d'optimiser les communications afin que les processus ne perdent pas trop de temps à dialoguer entre eux. Il existe déjà des routines de communications collectives dans MPI. Une de ces fonctions est particulièrement utile, puisqu'elle permet à un processus d'envoyer le même message à tous les autres processus. Il s'agit de la routine `Mpi_Bcast` (*broadcast*, ou diffusion), qui s'emploie de la manière suivante :

```
Mpi_Bcast(contenu,taille,type,emetteur,communicateur,erreur)
```

ce qui signifie que le processus `emetteur` envoie le message `contenu`, consistant en `taille` éléments de type `type`, vers tous les autres processus du communicateur `communicateur`. La variable `erreur` est encore une fois destinée à recevoir un code d'erreur.

Les communications se déroulent de la façon représentée sur les figures B.8 et B.9.

Il est possible d'apprécier le gain de temps réalisé par des considérations très simples. Si l'on admet qu'un envoi de message est effectué en une unité de temps, l'envoi successif d'un message de p_1 vers les 7 autres processus prend, dans le meilleur des cas, 7 unités de temps. Par contre, en nous plaçant une fois encore dans un cas optimal (les envois de la diffusion s'enchaînent sans que rien ne les interrompent), la diffusion est effectuée en 3 unités de temps. Plus généralement, l'envoi d'un message

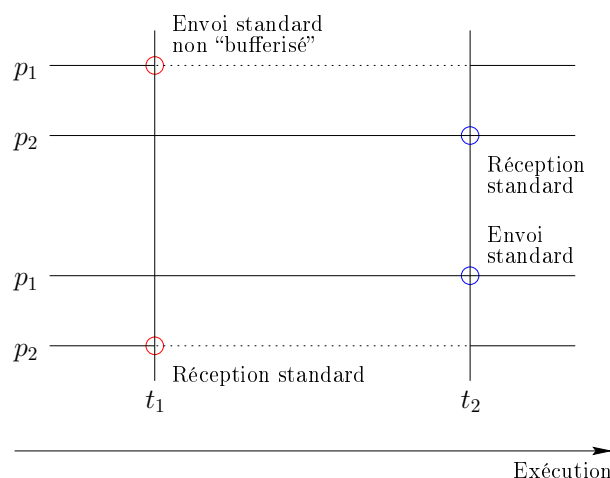


FIG. B.7 – Exemple de processus en attente à cause des communications bloquantes. Les lignes pointillées représentent les temps d’attente, les lignes pleines les périodes d’activité.

d’un processus vers $p - 1$ autres processus par une diffusion s’effectue en $\lceil \log_2(p) \rceil$ unités de temps. Le gain de temps est donc conséquent. Il existe trois autres types de communications collectives très utiles, et très employées. Ces commandes sont les suivantes : **Mpi_Scatter** (dispersion), **Mpi_Gather** (assemblage) et **Mpi_Reduce** (réduction). Leur fonction est illustrée par la figure B.10.

La commande de dispersion permet de partager un ensemble de données en parties égales, et d’envoyer à chaque processus du communicateur une de ces parties. Le processus de rang 0 (p_1 sur notre schéma) reçoit la première partie, le processus de rang 1 (p_2) la deuxième, et ainsi de suite. Si le processus émetteur est le processus de rang p , il conserve pour lui la $p + 1^{\text{ème}}$ partie des données.

La commande d’assemblage a l’effet inverse de la commande de dispersion. Elle permet à un processus de collecter et d’assembler des données provenant de tous les processus du communicateur. Ainsi, le processus qui collecte les informations recevra la première partie des données du processus de rang 0, la deuxième du processus de rang 1, et ainsi de suite. Si le processus collecteur est le processus de rang p , sa propre partie des données est la $p + 1^{\text{ème}}$ partie de l’ensemble. Comme pour l’assemblage, toutes les parts doivent avoir la même dimension.

La commande de réduction permet quant à elle d’effectuer une opération sur des données provenant de tous les processus du communicateur. Il existe plusieurs opération définies par défaut, comme l’addition ou le maximum. Voici un exemple simple : chaque processus envoie un entier et le processus collecteur est chargé d’additionner ces entiers. Cette opération peut être effectuée par une réduction avec l’opération addition.

Les communications collectives définies dans MPI ont la particularité d’être appelées de la même manière par les processus émetteurs et récepteurs. Ainsi, l’appel à la fonction **Mpi_Bcast** doit être le même pour le processus qui diffuse l’information et pour ceux qui la reçoivent. Le code suivant donne un exemple de ceci.

```
Program Madiffusion

Include 'mpif.h'
Integer rang
Integer message
Integer ierr

Call Mpi_Init(ierr)
Call Mpi_Comm_Rank(Mpi_Comm_World, rang, ierr)
If(rang.Eq.0) Then
    message = 1
Else
    message = 0
Endif
```

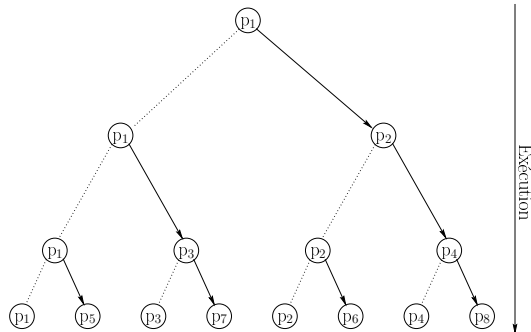


FIG. B.8 – Structure en arbre d'une diffusion. Les flèches représentent les envois de message.

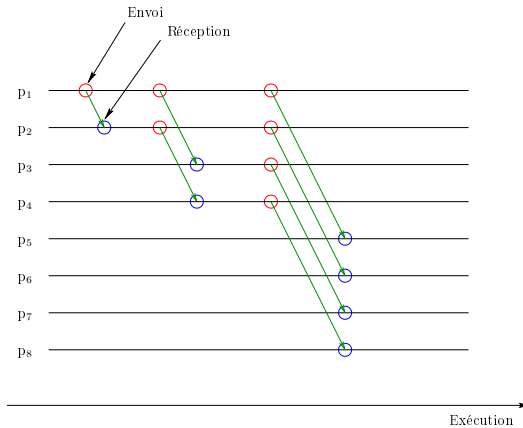


FIG. B.9 – Exemple d'une diffusion avec 8 processus.

```

Call Mpi_Bcast(message,1,Mpi_Integer,0,Mpi_Comm_World,ierr)
Print *, 'Bonjour, ici le processus #',rang,
&      ', message reçu : ',message
Call Mpi_Finalize(ierr)

End

```

Le résultat produit par 4 processus est le suivant :

```

Bonjour, ici le processus # 0, message reçu : 1
Bonjour, ici le processus # 1, message reçu : 1
Bonjour, ici le processus # 2, message reçu : 1
Bonjour, ici le processus # 3, message reçu : 1

```

Les communications collectives telles que `Mpi_Bcast` nous permettent d'entrevoir l'intérêt que représente les communicateurs. Si le programmeur désire diffuser un message à un ensemble de processus, mais pas à tous les processus, il lui est possible de définir un communicateur, nommons le `Comm1`, contenant les processus concernés par ce message. Il lui suffit ensuite d'utiliser la commande `Mpi_Bcast` avec le communicateur `Comm1` en argument.

Il existe quelques commandes permettant de définir de nouveaux communicateurs, telles que `Mpi_Comm_Split` (se référer, par exemple, à [20] pour des précisions sur son utilisation).

2.5 Types composés et autres fonctions

Nous avons vu précédemment qu'il est nécessaire de préciser le type de données qui sont transmises dans un message. Il existe des types définis par MPI, qui pour le Fortran sont `Mpi_Character`, `Mpi_Complex`, `Mpi_Double_Precision`, `Mpi_Integer`, `Mpi_Logical` et `Mpi_Real`, auxquels nous pouvons également ajouter `Mpi_Byte` et `Mpi_Packed`, qui n'ont pas leur équivalent en Fortran.

Il est possible de définir de nouveaux types de données à partir de ces types de bases. Il est par exemple très aisé de définir un type colonne qui correspondrait à une colonne de la matrice A de notre exemple. La commande `Mpi_Type_Vector` a pour fonction de nous aider à définir de telles types.

Le type `Mpi_Packed` évoqué précédemment permet à l'utilisateur d'empaqueter des données dans le tampon d'envoi, puis d'envoyer l'ensemble de ces données en une seule fois. L'empaquetage se réalise avec la commande `Mpi_Pack` et le dépaquetage avec la commande `Mpi_Unpack`. Ce type permet de réunir dans un même envoi différentes données, et donc de minimiser le nombre de messages échangés.

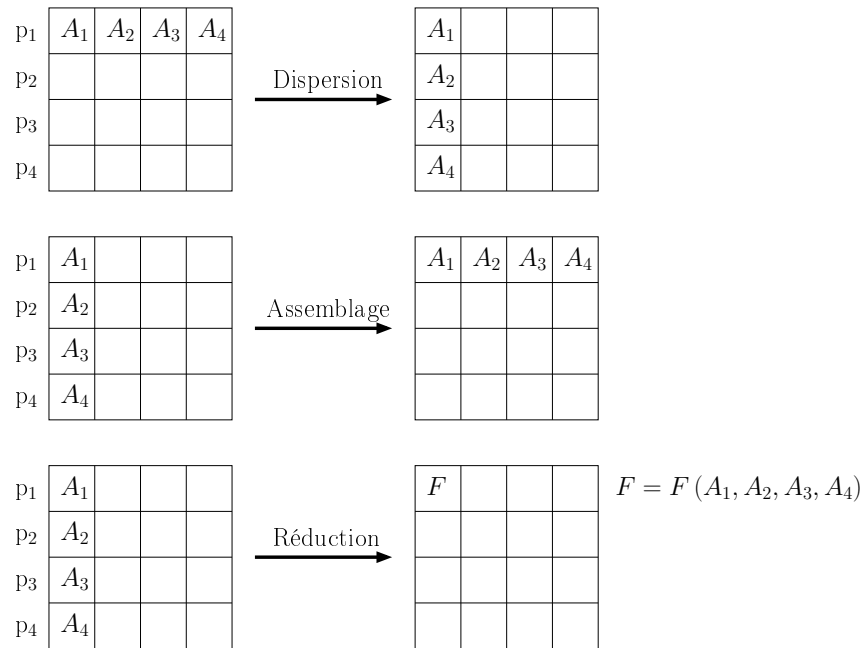


FIG. B.10 – Fonction des commandes de dispersion, d'assemblage et de réduction.

Parmi les autres fonctionnalités proposées par MPI, nous pouvons ajouter la possibilité de créer des communicateurs suivant la topologie de la machine parallèle. Ceci peut se révéler utile si la machine parallèle utilisée possède une topologie particulière (de type grille par exemple). Il existe de nombreuses fonctions permettant ensuite d'exploiter cette topologie du communicateur.

Enfin, il existe une fonction permettant de mesurer le temps d'exécution d'un processus, `Mpi_Wtime`. Cette fonction est l'équivalent de la fonction `time` de Fortran77.

TABLEAUX

Les tableaux suivants regroupent l'ensemble des résultats numériques utilisés dans les analyses que nous avons effectuées. Les paramètres et observables sont définis dans le chapitre 2, mis à part les erreurs commises sur les ajustements Σ_γ^2 et $\Sigma_{s^2}^2$, qui sont définies dans la section 3 de l'annexe A. Un facteur multiplicateur a été appliqué aux valeurs de s^2 ($\times 10^2$), de Σ_γ^2 ($\times 10^2$), de $\Sigma_{s^2}^2$ ($\times 10^3$), de T ($\times 10^7$) et de $\mathcal{H}$ ($\times 10^2$).

ε	Δt	$\Delta \mathcal{H}$	$\mathcal{V}_{fin}$	a_1	a_2	R_{50}	γ
0.001	0.030	45.18	-1.0470	0.986	1.023	0.697	4.930
0.005	0.030	12.13	-1.0015	0.979	1.023	0.474	4.835
0.010	0.030	4.71	-0.9973	0.994	1.018	0.440	4.792
0.050	0.030	0.30	-0.9831	0.994	1.027	0.427	4.812
0.001	0.015	28.74	-0.9991	0.971	1.008	0.581	4.890
0.005	0.015	4.42	-0.9916	0.961	1.026	0.441	4.809
0.010	0.015	1.04	-0.9925	0.993	1.032	0.424	4.781
0.050	0.015	0.09	-0.9934	0.976	1.013	0.419	4.764
0.001	0.008	16.18	-1.0006	0.973	1.012	0.496	4.921
0.005	0.008	1.02	-1.0047	0.961	1.010	0.419	4.807
0.010	0.008	0.22	-1.0009	0.977	1.018	0.419	4.792
0.050	0.008	0.09	-1.0013	0.983	1.024	0.416	4.787
0.001	0.003	5.29	-1.0013	0.981	1.025	0.440	4.861
0.005	0.003	0.14	-1.0006	0.996	1.021	0.419	4.776
0.010	0.003	0.03	-1.0063	0.962	1.012	0.415	4.845
0.050	0.003	0.08	-0.9868	0.979	1.031	0.425	4.767

TAB. C.1 – Test de l'influence de la valeur du paramètre de lissage du potentiel ε et du pas de temps Δt .

Modèle	N	a_1	a_2	R_{10}/R_d	R_{50}/R_d	γ	s^2
H_{10}	5571	1.053	0.972	1.032	5.523	4.491	-0.0489
H_{10}	10344	1.013	0.966	0.962	4.204	4.637	-0.0500
H_{10}	30998	1.018	0.978	0.879	2.674	4.586	-0.0576
H_{10}	50947	1.010	0.983	0.855	2.410	4.578	-0.0591
H_{50}	7246	1.023	0.951	0.747	1.648	4.817	-0.0248
H_{50}	10344	1.014	0.979	0.729	1.566	4.763	-0.0250
H_{50}	15569	1.015	0.979	0.727	1.518	4.756	-0.0247
H_{50}	20497	1.021	0.987	0.720	1.469	4.792	-0.0244
H_{50}	30998	1.014	0.978	0.724	1.453	4.729	-0.0249
H_{50}	50947	1.003	0.980	0.717	1.426	4.742	-0.0245
H_{50}	73636	1.016	0.984	0.716	1.422	4.706	-0.0247
H_{80}	5571	1.029	0.975	0.695	1.408	5.948	-0.0152
H_{80}	10344	1.014	0.992	0.689	1.376	6.243	-0.0144
H_{80}	30998	1.010	0.991	0.686	1.328	6.773	-0.0125
H_{80}	50947	1.006	0.994	0.680	1.322	6.957	-0.0116
H_{80}	73614	1.003	0.989	0.682	1.308	7.162	-0.0112

TAB. C.2 – Test de l'influence du nombre de particules pour les modèles H_{10} , H_{50} et H_{80} .

η	$\Delta\mathcal{H}$	$\mathcal{V}_{fin}$	a_1	a_2	R_{10}	R_{50}	R_{90}	R_d	γ	Σ_γ^2	s^2	$\Sigma_{s^2}^2$	T	$-\mathcal{H}$
88	0.0	-1.00	1.02	0.99	3.39	6.53	10.6	5.02	7.37	6.1	1.12	3.2	7.12	3.30
79	0.0	-0.99	1.00	1.00	3.11	6.01	9.8	4.59	6.86	4.3	1.37	3.1	6.82	3.60
60	0.0	-0.99	1.01	0.98	2.41	4.73	11.5	3.41	5.05	1.0	2.04	2.2	8.95	4.20
50	0.0	-1.01	1.01	0.98	2.04	4.09	15.1	2.81	4.73	1.8	2.49	1.7	9.44	4.50
40	0.1	-1.00	1.01	0.99	1.72	3.51	25.7	2.30	4.72	2.1	2.79	1.8	10.51	4.80
30	0.0	-1.00	1.02	0.99	1.36	2.88	253.2	1.75	4.66	1.8	3.27	2.5	11.59	5.10
20	0.0	-1.00	1.01	0.99	0.95	2.22	874.1	1.17	4.68	1.4	4.03	5.2	14.81	5.39
15	0.0	-1.00	1.01	0.99	0.74	1.89	1143.0	0.88	4.66	1.4	4.73	6.3	18.07	5.53
10	1.4	-1.00	1.02	0.98	0.52	1.59	1448.0	0.60	4.59	1.2	5.76	11.1	23.89	5.66

TAB. C.3 – Conditions initiales de type homogène : H_η

η	n_g	$\Delta\mathcal{H}$	$\mathcal{V}_{fin}$	a_1	a_2	R_{10}	R_{50}	R_{90}	R_d	γ	Σ_γ^2	s^2	$\Sigma_{s^2}^2$	T	$-\mathcal{H}$
10	03	0.1	-1.00	1.03	0.96	0.55	1.85	1241.0	0.58	4.61	1.2	5.21	10.7	21.36	5.82
67	20	0.8	-1.00	1.01	0.94	0.93	6.14	16.1	0.66	5.00	1.2	1.98	10.1	8.72	3.96
65	20	0.9	-1.00	1.02	0.96	0.93	5.63	13.5	0.63	5.27	1.8	2.00	11.4	9.38	4.29
61	20	0.6	-1.00	1.05	0.98	0.86	5.15	12.3	0.58	5.43	2.8	2.09	13.3	10.14	4.64
48	20	0.6	-1.00	1.03	0.99	0.81	4.24	11.8	0.59	5.15	2.6	2.61	13.0	11.58	5.31
39	20	0.5	-1.00	1.04	0.99	0.76	3.84	12.8	0.40	4.72	1.5	3.20	10.2	12.40	5.65
29	20	1.2	-1.00	1.04	0.99	0.72	3.42	15.4	0.56	4.42	1.1	3.81	6.3	12.95	5.97
14	20	0.3	-1.00	1.09	0.98	0.64	2.72	39.5	0.48	4.56	1.3	4.27	7.5	14.23	6.50
10	20	1.4	-0.99	1.13	0.99	0.61	2.54	345.5	0.54	4.70	0.9	4.25	10.2	15.21	6.66
07	20	0.3	-1.00	1.14	0.92	0.57	2.44	224.6	0.45	4.74	1.1	4.41	11.2	16.39	6.76

TAB. C.4 – Conditions initiales de type grumeaux : $C_\eta^{n_g}$

η	α_d	$\Delta\mathcal{H}$	$\mathcal{V}_{fin}$	a_1	a_2	R_{10}	R_{50}	R_{90}	R_d	γ	Σ_γ^2	s^2	$\Sigma_{s^2}^2$	T	$-\mathcal{H}$
50	0.5	0.0	-1.01	1.01	0.99	1.84	3.92	13.5	2.53	4.66	1.2	2.65	2.3	10.61	4.69
50	1	0.0	-1.00	1.01	0.99	1.56	3.77	12.1	2.01	4.77	0.6	2.78	3.4	11.06	5.00
10	1	0.1	-0.99	1.00	0.80	0.69	2.71	382.2	0.70	4.61	0.8	4.05	8.1	14.56	6.32
8	1.5	0.1	-1.00	1.01	0.71	0.62	2.63	25.1	0.61	4.63	0.7	4.42	10.4	16.56	7.18
50	2	1.7	-0.99	1.02	0.99	0.53	3.20	9.2	0.34	5.30	6.5	3.35	23.6	15.91	7.32
40	1.5	0.1	-1.00	1.00	0.99	0.97	3.31	11.0	1.03	4.71	0.8	3.44	6.0	13.27	5.99
9	2	1.6	-0.99	1.01	0.78	0.38	2.51	10.6	0.18	4.64	1.4	5.25	19.6	20.18	9.31

TAB. C.5 – Conditions initiales avec une densité en loi de puissance : $P_\eta^{\alpha_d}$

η	k	$\Delta\mathcal{H}$	$\mathcal{V}_{fin}$	a_1	a_2	R_{10}	R_{50}	R_{90}	R_d	γ	Σ_γ^2	s^2	$\Sigma_{s^2}^2$	T	$-\mathcal{H}$
7	I	5.0	-0.98	1.02	0.99	0.25	2.04	1721.0	0.37	4.26	1.0	5.16	57.5	30.89	5.62
15	III	0.6	-1.00	1.01	0.98	0.68	2.04	366.8	0.97	4.65	1.3	4.35	11.6	17.73	5.53
25	II	0.4	-1.01	1.01	0.98	1.18	2.51	225.5	1.54	4.66	1.9	3.59	3.3	12.64	5.23
35	I	0.2	-1.01	1.01	0.99	1.55	3.18	39.8	2.09	4.66	2.2	2.97	2.7	10.88	4.95
51	III	0.2	-1.00	1.01	0.98	1.79	4.19	14.2	2.83	4.67	0.9	2.39	3.6	9.77	4.48
50	I	0.1	-1.00	1.02	0.98	1.93	4.09	15.1	2.87	4.60	1.3	2.45	3.2	9.90	4.50
50	II	0.1	-1.00	1.01	0.98	2.03	4.13	15.1	2.87	4.70	1.7	2.45	2.1	9.56	4.49

TAB. C.6 – Conditions initiales avec un spectre de masse : M_η^k

η	σ	$\Delta\mathcal{H}$	$\mathcal{V}_{fin}$	a_1	a_2	R_{10}	R_{50}	R_{90}	R_d	γ	Σ_γ^2	s^2	$\Sigma_{s^2}^2$	T	$-\mathcal{H}$
48	1	0.0	-1.00	1.00	1.00	1.72	3.98	14.6	2.23	4.66	0.5	2.64	2.6	9.69	4.50
49	2	0.0	-1.00	1.02	0.99	1.74	4.02	14.4	2.28	4.65	0.6	2.61	2.9	9.69	4.50
50	3	0.0	-1.00	1.00	1.00	1.90	4.08	14.1	2.56	4.72	0.8	2.52	2.3	9.59	4.50
12	4	0.4	-0.99	1.03	0.98	0.56	1.77	1312.0	0.64	4.56	1.1	5.33	10.1	18.85	5.63
50	5	0.0	-1.00	1.01	0.99	1.98	4.11	14.8	2.72	4.71	1.4	2.50	1.8	9.89	4.50

TAB. C.7 – Conditions initiales avec une dispersion de vitesse gaussienne : G_η^σ

Nom	η	γ_{rot}	$\Delta\mathcal{H}$	$\mathcal{V}_{fin}$	$\mu_{K,ini}$	$\mu_{K,fin}$	a_1	a_2	R_{10}	R_{50}	R_d
$R_{10}^{0.1}$	10	0.1	0.77	-0.99	$4.29 \cdot 10^{-4}$	$2.87 \cdot 10^{-4}$	1.03	0.99	0.51	1.61	0.59
$R_{10}^{0.2}$	10	0.2	0.96	-1.00	$9.75 \cdot 10^{-4}$	$9.39 \cdot 10^{-4}$	1.03	0.99	0.52	1.60	0.60
$R_{10}^{0.3}$	10	0.3	0.72	-1.00	$1.63 \cdot 10^{-3}$	$2.12 \cdot 10^{-3}$	1.02	0.98	0.51	1.62	0.58
$R_{10}^{0.5}$	10	0.5	0.48	-1.00	$3.09 \cdot 10^{-3}$	$2.15 \cdot 10^{-3}$	1.01	0.98	0.52	1.64	0.59
$R_{10}^{5.0}$	10	5.0	1.32	-0.99	$4.66 \cdot 10^{-2}$	0.0	1.16	0.95	0.50	2.81	0.50
$R_{10}^{10.0}$	10	10.0	3.83	-1.00	$6.29 \cdot 10^{-2}$	$1.88 \cdot 10^{-2}$	1.23	1.00	0.50	2.94	0.72
$R_{12}^{0.0}$	12	0.0	0.79	-1.00	0.0	$1.82 \cdot 10^{-3}$	1.02	0.98	0.60	1.70	0.70
$R_{12}^{0.3}$	12	0.3	0.79	-0.99	$1.79 \cdot 10^{-3}$	$1.36 \cdot 10^{-3}$	1.03	0.99	0.61	1.74	0.71
$R_{12}^{0.5}$	12	0.5	0.79	-1.00	$3.40 \cdot 10^{-3}$	$2.09 \cdot 10^{-3}$	1.01	0.98	0.60	1.76	0.70
$R_{12}^{1.0}$	12	1.0	0.76	-1.00	$8.29 \cdot 10^{-3}$	$5.61 \cdot 10^{-3}$	1.05	0.97	0.78	1.73	0.96
$R_{12}^{0.3}$	15	0.3	0.83	-1.00	$2.00 \cdot 10^{-3}$	$1.68 \cdot 10^{-3}$	1.02	0.98	0.74	1.93	0.88
$R_{15}^{1.0}$	15	1.0	0.35	-0.99	$9.27 \cdot 10^{-3}$	$5.11 \cdot 10^{-3}$	1.02	0.97	0.72	2.06	0.84
$R_{15}^{1.5}$	15	1.5	0.32	-1.00	$1.58 \cdot 10^{-2}$	$1.03 \cdot 10^{-2}$	1.09	1.00	0.70	2.23	0.79
$R_{20}^{0.3}$	20	0.3	0.28	-1.00	$2.30 \cdot 10^{-3}$	$1.84 \cdot 10^{-3}$	1.02	0.98	0.95	2.22	1.16
$R_{20}^{0.5}$	20	0.5	0.24	-1.00	$4.38 \cdot 10^{-3}$	$3.55 \cdot 10^{-3}$	1.01	0.99	0.95	2.24	1.16
$R_{20}^{1.0}$	20	1.0	0.29	-1.00	$1.07 \cdot 10^{-2}$	$7.91 \cdot 10^{-3}$	1.03	0.99	0.91	2.30	1.09
$R_{20}^{1.5}$	20	1.5	0.17	-1.00	$1.82 \cdot 10^{-2}$	$1.26 \cdot 10^{-2}$	1.03	0.97	0.88	2.40	1.05
$R_{20}^{2.0}$	20	2.0	0.22	-1.00	$2.64 \cdot 10^{-2}$	$1.78 \cdot 10^{-2}$	1.07	0.96	0.85	2.56	0.99
$R_{20}^{5.0}$	20	5.0	0.25	-1.00	$6.66 \cdot 10^{-2}$	$2.82 \cdot 10^{-2}$	1.28	0.99	0.68	3.22	0.66
$R_{20}^{10.0}$	20	10.0	1.61	-1.00	$8.99 \cdot 10^{-2}$	$3.51 \cdot 10^{-2}$	1.30	0.96	0.66	3.36	0.59
$R_{30}^{1.5}$	30	1.5	0.26	-1.00	$2.25 \cdot 10^{-2}$	$1.76 \cdot 10^{-2}$	1.03	0.97	1.23	2.90	1.54
$R_{30}^{2.0}$	30	2.0	0.22	-1.00	$3.26 \cdot 10^{-2}$	$2.24 \cdot 10^{-2}$	1.03	0.94	1.17	2.95	1.43
$R_{30}^{5.0}$	30	5.0	0.51	-1.00	$8.25 \cdot 10^{-2}$	$4.35 \cdot 10^{-2}$	1.31	0.98	0.89	3.56	0.92
$R_{30}^{10.0}$	30	10.0	0.80	-1.00	$1.11 \cdot 10^{-1}$	$4.72 \cdot 10^{-2}$	1.37	0.97	0.76	3.73	0.70
$R_{50}^{5.0}$	50	5.0	0.06	-1.00	$1.07 \cdot 10^{-1}$	$6.41 \cdot 10^{-2}$	1.34	0.99	1.24	4.10	1.40
$R_{50}^{10.0}$	50	10.0	0.53	-1.00	$1.46 \cdot 10^{-1}$	$6.35 \cdot 10^{-2}$	1.34	0.93	0.98	4.32	0.98
$R_{50}^{20.0}$	50	20.0	0.85	-1.00	$1.60 \cdot 10^{-2}$	$6.61 \cdot 10^{-2}$	1.34	0.94	0.90	4.30	0.80

TAB. C.8 – Conditions initiales en rotation $R_{\eta}^{\gamma_{rot}}$.

BIBLIOGRAPHIE

- [1] S.J.AARSETH, *Dynamical evolution of clusters of galaxies, I*, MNRAS, **126**, 223, 1963.
- [2] S.J.AARSETH, *NBODY2 : A direct N-body integration code*, New Astronomy, **6**, 277, 2001.
- [3] L.AGUILAR ET D.MERRITT, *The structure and dynamics of galaxies formed by cold dissipationless collapse*, ApJ, **354**, 33, 1990.
- [4] V.A.ANTONOV, *Most probable phase distribution in spherical star systems and condition for its existence*, IAU Symposium 113, Dynamics of Globular Clusters, édité par J.Goodman et P.Hut, Dordrecht, D. Reidel Publishing Co., 1985. Traduction d'un article précédemment publié dans Vest. Leningrad Univ., **7**, 135, 1962.
- [5] V.I.ARNOLD, *Mathematical methods of classical mechanics*, Springer-Verlag New York, 1978.
- [6] J.E.BARNES ET P.HUT, *Force-calculation algorithm*, Nature, **324**, 446, 1986.
- [7] J.E.BARNES ET P.HUT, *Error analysis of a tree code*, ApJSS, **70**, 389, 1988.
- [8] J.BINNEY ET M.MERRIFIELD, *Galactic Astronomy*, Princeton University Press, 1998.
- [9] J.BINNEY ET S.TREMAINE, *Galactic Dynamics*, Princeton University Press, 1987.
- [10] A.M.BLOCH, P.S.KRISCHNAPRASAD, J.E.MARSDEN ET T.S.RATIU, *Dissipation induced instabilities*, Ann. de l'Inst. Henri Poincaré : Analyse non linéaire, **11**, 37, 1994.
- [11] C.M.BOILY, C.J.CLARKE AND S.D.MURRAY, *Collapse and evolution of flattened star clusters*, MNRAS, **302**, 399, 1999.
- [12] A.-S.BONNET-BENDHIA ET E.LUNÉVILLE, *Résolution numérique des équations aux dérivées partielles*, Cours MA201 de l'École Nationale Supérieure de Techniques Avancées, édition 2002-2003.
- [13] P.BRÉMAUD, *Introduction aux probabilités*, Springer-Verlag Berlin Heidelberg, 1984.
- [14] J.K.CANNIZZO ET T.C.HOLLISTER, *Cold dissipationless collapse of spherical systems : sensitivity to the initial density law*, ApJ, **400**, 58, 1992.
- [15] S.CASERTANO ET P.HUT, *Core radius and density measurements in N-Body experiments : Connections with theoretical and observational definitions*, ApJ, **298**, 80, 1985.
- [16] S.CHANDRASEKHAR, *An introduction to the study of stellar structure*, Dover Publications, 1958.
- [17] J.-L.DELCROIX ET A.BERS, *Physique des plasmas*, InterÉditions/CNRS Éditions, 1994.
- [18] R.A.GERBER, *Global Effects of Softening n-Body Galaxies*, ApJ, **466**, 724, 1996.
- [19] B.GIDAS, W.-M.NI ET L.NIRENBERG, *Symmetry of positive solutions of nonlinear elliptic equations in $\mathbb{R}^n$* , Math. Anal. and Applications, Part A, Advances in Math. Suppl. Studies 7A, édité par L.Nachbin, Academic Press, 369, 1981.

- [20] W.GROPP, E.LUSK ET A.SKJELLUM, *Using MPI : Portable Parallel Programming with the Message-Passing Interface*, seconde édition, The MIT Press, 1999.
- [21] D.C.HEGGIE ET R.D.MATHIEU, *Standardised Units and Time Scales* dans *The Use of Supercomputers in Stellar Dynamics*, édité par P.Hut & S.L.W.McMillan, Springer Verlag, 1986.
- [22] L.HERNQUIST ET J.E.BARNES, *Are some N-body algorithms intrinsically less collisional than others ?*, ApJ, **349**, 562, 1990.
- [23] J.D.JACKSON, *Classical electrodynamics*, 2^{de} édition, John Wiley & Sons, 1975.
- [24] J.H.JEANS, *On the theory of star-streaming and the structure of the universe*, MNRAS, **76**, 70, 1915.
- [25] H.KANDRUP, *The secular instability of axisymmetric collisionless star clusters*, ApJ, **380**, 511, 1991.
- [26] H.KANDRUP, *Theoretical techniques in modern galactic dynamics*, Notes de cours de l'Université de Floride.
- [27] H.KANDRUP ET J.-F.SYGNET, *A simple proof of dynamical stability for a class of spherical clusters*, ApJ, **298**, 27, 1985.
- [28] J.KATZ, *La thermodynamique des systèmes autogravitants*, Notes de quatre exposés donnés au Laboratoire de Physique Théorique de l'Université Paris XI, Orsay, du 15 novembre au 15 décembre 2001.
- [29] S.KAZANTZIDIS, J.MAGORRIAN ET B.MOORE, *Generating equilibrium dark matter halos : inadequacies of the local maxwellian approximation*, ApJ, **601**, 37, 2004.
- [30] I.R.KING, *The Structure of Star Clusters. I. An Empirical Density Law*, AJ, **67**, 471, 1962.
- [31] R.KRIKORIAN, *Linéarisation et stabilité des équations différentielles*, Cours AO102 de l'École Nationale Supérieure de Techniques Avancées, édition 2002-2003.
- [32] P.KROUPA, *On the variation of the Initial Mass Function*, MNRAS, **322**, 231, 2001.
- [33] R.M.KULSRUD ET J.W.-K.MARK, *Collective instabilities and waves for inhomogeneous stellar systems. I. The necessary and sufficient energy principle*, ApJ, **160**, 471, 1970.
- [34] L.LANDAU ET E.LIFCHITZ, *Mécanique - Physique théorique, Tome I*, deuxième édition, Éditions Mir, 1966.
- [35] G.LAVAL, C.MERCIER ET R.PELLAT, *Necessity of the energy principles for magnetostatic stability*, Nucl. Fusion, **5**, 156, 1965.
- [36] D.LYNDEN-BELL ET R.WOOD, *The gravo-thermal catastrophe in isothermal spheres and the onset of red-giant structure for stellar systems.*, MNRAS, **138**, 495, 1968.
- [37] G.MEYLAN, *The Internal Dynamics of Globular Clusters*, preprint, astro-ph/9912495, 1999.
- [38] E.A.MILNE, *The analysis of stellar structure*, MNRAS, **91**, 4, 1930.
- [39] E.A.MILNE, *The analysis of stellar structure II*, MNRAS, **92**, 610, 1932.
- [40] K.W.MIN ET C.S.CHOI, *Cold collapse of a spherical stellar system*, MNRAS, **238**, 253, 1989.
- [41] P.MORRISON, *The Maxwell-Vlasov equations as a continuous hamiltonian system*, Phys.Letter A, **80**, 383, 1980.
- [42] J.E.NAVARRO ET S.D.M.WHITE, *Simulations of dissipative galaxy formation in hierarchically clustering universes*, MNRAS, **265**, 271, 1993.
- [43] OHIO SUPERCOMPUTER CENTER, *MPI Primer / Developing with LAM*, 1996.
- [44] P.S.PACHECO, *A User's Guide to MPI*, 1998.
- [45] T.PADMANABHAN, *Antonov instability and gravothermal catastrophe – Revisited*, ApJSS, **71**, 651, 1989.

- [46] J.PEREZ, *Stabilité des systèmes de particules en interaction gravitationnelle*, Thèse de doctorat de l'université Paris VII, 1995.
- [47] J.PEREZ ET J.-J.ALY, *Stability of spherical stellar systems - I. Analytical results*, MNRAS, **280**, 689, 1996.
- [48] J.PEREZ ET M.LACHIEZE-REY, *A symplectic approach to gravitational instability*, ApJ, **465**, 54, 1996.
- [49] W.H.PRESS, S.A.TEUKOLSKY, W.T.VETTERLING ET B.P.FLANNERY, *Numerical Recipes*, Cambridge University Press, 1989.
- [50] E.E.SALPETER, *The Luminosity Function and Stellar Evolution*, ApJ, **121**, 161, 1955.
- [51] L.SCHWARTZ, *Méthodes mathématiques pour les sciences physiques*, Éditions Hermann, 1965.
- [52] CH.THEIS ET R.SPURZEM, *On the evolution of shape in N-body simulations*, A&A, **341**, 361, 1999.
- [53] S.C.TRAGER, I.R.KING ET S.DJORGOVSKI, *Catalogue of galactic globular-cluster surface-brightness profiles*, AJ, **109**, 218, 1995.
- [54] T.S.VAN ALBADA, *Dissipationless galaxy formation and the $r^{1/4}$ law*, MNRAS, **201**, 939, 1982.
- [55] T.S.VAN ALBADA ET J.H.VAN GORKOM, *Experimental stellar dynamics for systems with axial symmetry*, A&A, **54**, 121, 1977.
- [56] R.E.WHITE ET S.J.SHAWL, *Axial Ratios And Orientations For 100 Galactic Globular Star Clusters*, ApJ, **317**, 246, 1987.

ÉTUDE DU SYSTÈME COUPLÉ BOLTZMANN SANS COLLISIONS–POISSON POUR LA GRAVITATION.
SIMULATIONS NUMÉRIQUES DE LA FORMATION DES SYSTÈMES AUTO-GRAVITANTS.

Résumé : Nous étudions la formation et les propriétés des systèmes auto-gravitants à l’aide de simulations numériques à N corps d’effondrements gravitationnels.

Nous effectuons dans un premier temps une synthèse des principaux résultats analytiques concernant les équations de Boltzmann sans collisions et de Poisson, qui modélisent les systèmes gravitationnels non collisionnels ainsi que certaines solutions analytiques de ce système couplé d’équations.

Nous présentons ensuite les codes de calcul utilisés pour les simulations. Nous avons parallélisé certains de ces codes, nous introduisons donc le calcul parallèle et la bibliothèque d’échange de message MPI. Nous exposons enfin les résultats de nos simulations, et leurs analyses. Nous déduisons de ces analyses divers résultats pouvant expliquer différentes caractéristiques des systèmes auto-gravitants ainsi que les conditions initiales nécessaires au déclenchement des instabilités d’Antonov et d’orbites radiales.

Mots clés : dynamique des systèmes auto-gravitants, amas globulaires, galaxies elliptiques, équation de Boltzmann sans collisions, équation de Poisson, méthodes analytiques, simulations numériques à N corps, calcul parallèle.

ANALYSIS OF THE GRAVITATIONAL COUPLED COLLISIONLESS BOLTZMANN–POISSON EQUATIONS
AND NUMERICAL SIMULATIONS OF THE FORMATION OF SELF-GRAVITATING SYSTEMS.

Abstract: We study the formation of self-gravitating systems and their properties by means of N -body simulations of gravitational collapse.

First, we summarize the major analytical results concerning the collisionless Boltzmann equation and the Poisson’s equation which describe the dynamics of collisionless gravitational systems. We present a study of some analytical solutions of this coupled system of equations.

We then present the software used to perform the simulations. Some of this has been parallelized and implemented with the aid of MPI. For this reason we give a brief overview of it.

Finally, we present the results of the numerical simulations. Analysis of these results allows us to explain some features of self-gravitating systems and the initial conditions needed to trigger the Antonov instability and the radial orbit instability.

Keywords: dynamics of self-gravitating systems, globular clusters, elliptical galaxies, collisionless Boltzmann equation, Poisson’s equation, analytical methods, numerical N -body simulations, parallel computing.