

THÈSE / UNIVERSITÉ DE BRETAGNE OCCIDENTALE
sous le sceau de l'Université Bretagne Loire
pour obtenir le titre de
DOCTEUR DE L'UNIVERSITÉ DE BRETAGNE OCCIDENTALE
Mention : Informatique
École Doctorale SICMA

présentée par

Sébastien Le Yaouanq

préparée au Centre Européen de Réalité Virtuelle,
Laboratoire CNRS Lab-STICC

Co-simulation redondante d'échelles de
modélisation hétérogènes pour une
approche phénoménologique.

Thèse soutenue le 17 juin 2016
devant le jury composé de :

Samir Otmane (Rapporteur)
Professeur, Université d'Evry Val d'Essonne

Éric Ramat (Rapporteur)
Professeur, Université du Littoral Côte d'Opale

Vincent Chevrier (Examinateur)
Professeur, Université de Lorraine

Jacques Tisseau (Directeur de thèse)
Professeur, ENI Brest

Pascal Redou (Encadrant)
Maître de Conférence, HDR, ENI Brest

Christophe Le Gal (Encadrant)
Directeur général, Docteur, CERVVAL

UNIVERSITÉ BRETAGNE LOIRE

— Mémoire de thèse —

Spécialité : Informatique

Co-simulation redondante d'échelles de modélisation hétérogènes
pour une approche phénoménologique

SÉBASTIEN LE YAOUANQ

Soutenue le 17 juin 2016 devant la commission d'examen :

Rapporteurs

Samir	OTMANE	Professeur, Université d'Evry Val d'Essone
Eric	RAMAT	Professeur, Université du Littoral Côte d'Opale

Examineurs

Vincent	CHEVRIER	Professeur, Université de Lorraine
Pascal	REDOU	Maître de conférences, HDR, ENI Brest
Christophe	LE GAL	Directeur général, Docteur, CERVVAL

Directeur

Jacques	TISSEAU	Professeur, ENI Brest
---------	---------	-----------------------

Co-simulation redondante d'échelles de modélisation hétérogènes pour une approche phénoménologique

Mémoire de thèse de doctorat

SÉBASTIEN LE YAOUANQ

— juin 2016 —



UBL

Université Bretagne Loire

Cité Internationale

1 place Paul Ricoeur, CS 54417, 35044 RENNES CEDEX

Lab-STICC

Laboratoire en sciences et techniques de l'information,
de la communication et de la connaissance

Université de Bretagne Occidentale, 20 Avenue Le Gorgeu, 29200 Brest

CERV

Centre Européen de Réalité Virtuelle

Ecole Nationale d'Ingénieurs de Brest, 25 rue Claude Chappe, 29280 Plouzané

CERVVAL

140 Avenue Graham Bell, 29280 Plouzané

Contact

SÉBASTIEN LE YAOUANQ - leyaouanq@cervval.com

<http://www.cervval.com>



Thèse financée par Cervval

Remerciements

Il est coutume de dire qu'une thèse ne se réalise pas seul. Je ne peux que le confirmer tant ce travail de recherche et moi-même avons bénéficié de soutiens durant toutes ces années. Je tiens donc à remercier chaleureusement toutes les personnes qui ont contribué, et qui contribueront encore, à ma vie de chercheur.

Je remercie sincèrement Éric Ramat et Samir Otmane d'avoir accepté de rapporter cette thèse, ainsi que Vincent Chevrier qui a présidé mon jury de thèse. Leurs remarques et leurs questions pertinentes m'ont prouvé qu'un sujet de recherche n'est jamais réellement clos. Je tiens également à souligner la bienveillance dont ils ont fait preuve vis à vis de cette thèse « industrielle ». C'est un comportement qu'il serait agréable de rencontrer plus souvent dans le milieu de la recherche.

Vient ensuite le tour de celui qui a suscité bon nombre de vocations de chercheurs. Jacques Tisseau, c'est un peu le guru du CERV, celui grâce à qui tout a commencé. Je le remercie d'avoir été ce directeur de thèse qui vous pousse en permanence à expliciter l'implicite.

Pour encadrer cette thèse « magnifique », il ne fallait rien de moins qu'un tandem de champions. Christophe Le Gal et Pascal Redou m'ont montré l'importance de toujours rester critique et que la première qualité d'un chercheur est d'apprendre à réfléchir par soi-même. Je les remercie de tous leurs efforts pour répandre cet esprit scientifique au plus grand nombre.

Je suis certes l'auteur de ces quelques pages, mais cette thèse est en réalité le fruit du travail de toute l'équipe de CERVVAL. Sa compétence n'a d'égal que la bonne humeur qui y règne et c'est ce qui me donne chaque jour l'envie d'en faire partie. Je souhaite remercier tout particulièrement Cykril et Pierre-Antoine qui ont complété mon équipe encadrante. Cette thèse, c'est aussi un peu la leur, au propre comme au figuré.

J'exprime toute ma gratitude à la direction de CERVVAL pour la confiance qu'elle m'a accordée et pour la latitude dont j'ai bénéficié pour mener à bien ces travaux.

Un grand merci à toutes les personnes du CERV qui m'ont fait découvrir les différentes facettes du monde de la recherche. L'équipe a beaucoup évolué depuis mon arrivée dans ces

contrées, tant au niveau de ses membres que des thèmes abordés, de sorte que j'ai presque l'impression d'être un vestige. Il restera de mon passage quelques belles collaborations *in virtuo* avec Fabrice, Gobi, Marc, Manu et j'en oublie. Par ailleurs, je remercie Richard Garfield et les « Sorciers de la Côte » d'avoir un peu provoqué le début de cette expérience magique.

Mes plus profonds remerciements vont à ma famille qui m'a soutenu, encouragé et aidé tout au long de ce parcours. En particulier, je dédie ce mémoire à ma maman qui a toujours œuvré pour m'offrir les meilleures chances de réussite. Qu'elle voit dans ce travail l'aboutissement de ses nombreux efforts.

Enfin, je ne saurais clore ces remerciements sans citer les deux êtres qui me sont les plus précieux, ceux qui donnent un sens à tout le reste. Aurore et Théo, il me tarde de voir notre famille s'agrandir à nouveau et de profiter de chaque instant à vos côtés. Puisse notre chance insolente perdurer.

Avant-propos

Cette thèse est née du lien étroit entre le Centre Européen de Réalité Virtuelle, centre de recherche pluri-disciplinaire dépendant de l'École Nationale d'Ingénieurs de Brest et rattaché au Laboratoire en Sciences et Techniques de l'Information, de la Communication et de la Connaissance (Lab-STICC), et la société CERVVAL, fondée en 2003 par des acteurs de l'industrie et de la communauté scientifique, qui développe et met en œuvre des outils de simulation de la complexité.

Les méthodes développées par CERVVAL trouvent leur origine dans les recherches sur la réalité virtuelle menées au sein du CERV. Elles permettent de construire de façon incrémentale, et en relation étroite avec les experts du domaine, des outils sur mesure d'aide à la décision, intégralement dédiés aux besoins du client et à sa problématique du moment.

Les travaux présentés dans ce document s'inscrivent dans cet objectif industriel. Ils visent à définir un cadre méthodologique pour la construction de simulateurs interactifs qui allient vitesse d'exécution et précision grâce à des simulations redondantes d'échelles de modélisation hétérogènes.

Ces réalisations ont donné lieu à plusieurs collaborations fructueuses dont les applications sont décrites dans les chapitres 2 et 5. Tout d'abord, nous avons contribué au développement du modèle MACMA (*MultiAgent Convective Mantle*), en coopération avec le Laboratoire Domaines Océaniques de l'Institut Universitaire Européen de la Mer (IUEM, UMR CNRS 6538). Un second partenariat a été noué avec les sociétés TECHNIP et BUREAU VERITAS dans le domaine de l'aide au design de structures offshore pour les milieux polaires.

Table des matières

Remerciements	vii
Avant-propos	ix
Table des matières	xi
Table des figures	xv
Introduction	1
1 Systèmes et expérimentation	5
1.1 Systèmes et complexité	6
1.2 Modéliser	9
1.2.1 Qu'est-ce qu'un modèle ?	9
1.2.2 Pourquoi modéliser ?	9
1.2.3 Comment construire un modèle ?	9
1.2.4 Comment valider un modèle ?	10
1.3 Expérimenter	11
1.3.1 <i>In vivo</i>	12
1.3.2 <i>In vitro</i>	13
1.3.3 <i>In silico</i>	13
1.3.4 <i>In virtuo</i>	13
1.4 Systèmes multi-agents	15
1.4.1 Définitions	15
1.4.2 Autonomies	17
1.4.3 Approche phénoménologique	18
1.5 Simulation	22
1.5.1 De temps en temps	22
1.5.2 Politiques d'ordonnancement	23
1.5.3 Méthodes numériques pour les systèmes différentiels	24
1.5.4 Méthodes numériques et systèmes multi-agents	28
1.6 Synthèse	28

2	Approche phénoménologique et Systèmes multi-agents	31
2.1	Éléments de géophysique	32
2.1.1	Mécanique interne du système Terre	33
2.1.2	Dynamique des plaques	34
2.1.3	Histoire thermique de la Terre	36
2.2	Le Modèle MACMA : <i>MultiAgent Convective Mantle</i>	37
2.2.1	Vue d'ensemble	38
2.2.2	Agents-interaction pour le bilan des forces	41
2.2.3	Agents-entité pour la cinématique	42
2.2.3.1	Modifications structurales	44
2.2.4	Agents-énaction pour les phénomènes thermiques	46
2.2.5	Bilan	50
2.3	Expérimentations <i>in virtuo</i>	51
2.3.1	Influence des lois empiriques	51
2.3.2	Influence des paramètres	54
2.4	Synthèse	54
3	Échelles de modélisation	59
3.1	Niveaux de description	60
3.1.1	Modèles macroscopiques	60
3.1.2	Modèles microscopiques	61
3.1.3	Hierarchie	61
3.2	Méthodes numériques multi-échelles	62
3.2.1	Phénomènes critiques ou séparation des échelles	62
3.2.2	Super-paramétrisation	64
3.2.3	Méthode quasicontinuum	66
3.2.4	Approche « sans-équation »	67
3.2.4.1	Définitions	67
3.2.4.2	Opérateurs	68
3.2.4.3	Time-stepper macroscopique	69
3.2.4.4	Méthodes d'intégration projectives	70
3.2.4.5	Patch dynamics	71
3.2.5	Méthodes multi-échelles hétérogènes	73
3.2.6	Bilan	74
3.3	Approches de multi-modélisation	74
3.3.1	Paradigmes de modélisation	74
3.3.2	Intégration formelle	76
3.3.3	Couplage de simulateurs et co-simulation	77
3.3.4	Synchronisation de co-simulation	78
3.3.4.1	Modèle maître-esclave	78
3.3.4.2	Modèle parallèle ou distribué	79
3.3.5	Bus de co-simulation	80
3.4	Synthèse	82
4	Simulations redondantes d'échelles de modélisation hétérogènes	85
4.1	Exemple : le phénomène de diffusion	86
4.1.1	Modèle macroscopique	86
4.1.1.1	Éléments théoriques	86

4.1.1.2	Implémentation	87
4.1.2	Modèle microscopique	88
4.1.2.1	Éléments théoriques	88
4.1.2.2	Implémentation	89
4.1.3	Détermination du coefficient de diffusion	89
4.1.4	Détermination de paramètres en cours de simulation	91
4.2	Architecture pour la co-simulation redondante d'échelles de modélisation hétérogènes	93
4.2.1	Vue générale	94
4.2.2	Mécanisme de contrôle	94
4.2.3	Requêtes de paramétrage	96
4.2.4	Bus de co-simulation et modèles auxiliaires	97
4.2.5	Discussions	100
4.2.6	Bilan	102
4.3	Détermination implicite d'un jeu de paramètres	104
4.3.1	Stratégie de co-simulation	104
4.3.2	Implémentation	108
4.3.3	Exemple d'application : diffusion de deux molécules antagonistes	108
4.3.4	Discussions	111
4.4	Synthèse	115
5	Application à la conception de structures offshore	117
5.1	Structures offshore et milieux polaires	118
5.1.1	Contexte	118
5.1.2	Interactions glace-structure	119
5.1.2.1	Modes de rupture	119
5.1.2.2	Phénomène d'empilement	121
5.1.3	Synthèse	123
5.2	ICE-MAS : Ice Multi-Agent Simulator	124
5.2.1	Approche multi-modèles	124
5.2.1.1	Interactions solide/solide	125
5.2.1.2	Hydrodynamique	127
5.2.1.3	Rupture de la plaque de glace	129
5.2.2	Bilan	132
5.3	Co-simulation explicite de la dynamique du courant	134
5.3.1	Rôle de l'agent- <i>watchdog</i>	135
5.3.2	Traitement de la requête	137
5.3.3	Intégration de la requête	139
5.3.4	Résultats et discussion	139
5.3.5	Bilan	141
5.4	Co-simulation implicite du phénomène d'empilement	141
5.4.1	Abstraction de l'empilement des blocs	142
5.4.2	Rôle de l'agent- <i>watchdog</i>	143
5.4.3	Traitement de la requête	145
5.4.4	Résultats et discussion	146
5.4.5	Bilan	150
5.5	Synthèse	150

Conclusion	153
Publications scientifiques	157
Références bibliographiques	159
A Méthodes numériques multi-agents	171
Résumé de thèse	183

Table des figures

1.1	Exemple de système : la bicyclette	6
1.2	Boucle de rétro-action et marché financiers	8
1.3	Méthodes d'expérimentation des systèmes complexes	12
1.4	Modélisation du houlogénérateur « Bilboquet »	14
1.5	Couplage de simulateurs de structures flottantes et d'états de mer	21
1.6	Agents « entité » versus agents « interaction »	21
1.7	Ordonnancements synchrone et asynchrone	24
1.8	Erreur d'approximation de fonction par un schéma numérique	26
2.1	Structure interne de la Terre	33
2.2	Schéma simplifié du couplage entre plaques tectoniques et manteau convectif	34
2.3	Schéma des principales forces exercées sur une plaque	35
2.4	Vue d'ensemble de l'environnement de simulation de MACMA	39
2.5	Schéma logique de l'algorithme d'évolution du système	40
2.6	Comportement des fosses de subduction basée sur un critère structurel	43
2.7	Suture de plaques due à la subduction d'une dorsale	45
2.8	Arrivée d'un continent en butée de plaque	46
2.9	Formation d'un supercontinent suivi d'un processus d'océanisation	47
2.10	Paramétrisation du schéma d'advection sous un continent isolant	50
2.11	Comportement des fosses de subduction basé sur un critère cinématique . . .	52
2.12	Comparaison de l'impact des lois de migration des fosses de subduction . . .	53
2.13	Étude de l'impact du seuil de rupture continentale	55
2.14	Comparaison de configurations continentales à la surface du manteau	55
2.15	Synthèse de la composition d'un modèle multi-agents mixte	56
3.1	Échelles spatio-temporelles caractéristiques des systèmes vivants	60
3.2	Maillage adaptatif pour la simulation d'allées de Von Karman	63
3.3	Super-paramétrisation	65
3.4	Méthode quasicontinuum	67
3.5	Méthode d'intégration projective	71
3.6	Représentation d'un pas de temps du schéma Gap-tooth	72
3.7	Représentation schématique de l'Heterogeneous Multiscale Method	74

3.8	Paradigmes de modélisation et leurs relation au temps, à l'espace et aux éléments modélisés	75
3.9	Encapsulation de formalismes	76
3.10	Formalisme DEVS et couplage	77
3.11	Co-simulation de modèles hétérogènes	77
3.12	Diagrammes de communication des différents modèles de co-simulation	79
3.13	Illustration du problème de partage de données en mode « maître-esclave » .	80
3.14	Modèle générique de co-simulation distribuée	82
4.1	Simulation de deux niveaux de description du phénomène de diffusion	90
4.2	Mesure du coefficient de diffusion de la caféine	91
4.3	Comparaison des distributions de concentration	92
4.4	Intégration d'un agent-watchdog	97
4.5	Diagramme de séquence du processus par co-simulation	99
4.6	Architecture de co-simulation	101
4.7	Temps de résolution d'une requête	103
4.8	Sélection d'un jeu de paramètre	106
4.9	Construction d'une simulation interactive par co-simulation	107
4.10	Architecture de co-simulation exploratoire	109
4.11	Erreur de comparaison temporelle	112
4.12	Erreur d'estimation de trajectoire	113
4.13	Interaction en cours de co-simulation	114
5.1	Structures offshore en milieux polaires	118
5.2	Phénomène de rupture par écrasement d'une plaque de glace	120
5.3	Phénomène de rupture par flexion d'une plaque de glace	121
5.4	Évolution de la charge associée à la rupture par flexion	122
5.5	Rupture par empilement	123
5.6	Diagramme d'interactions de ICE-MAS	125
5.7	Interaction solide/solide	126
5.8	Flottabilité d'un bloc	127
5.9	Calcul de la traînée de la structure	128
5.10	Dynamique du courant autour de la structure	129
5.11	Détermination du rayon de fracture circonférentielle	130
5.12	Distribution du poids et de la flottabilité des blocs sur la plaque	131
5.13	Simulation de ruptures par flexion	131
5.14	Détachement hiérarchique d'éclats de glace par écrasement	133
5.15	Combinaison des ruptures par flexion et par écrasement	133
5.16	Fréquence de co-simulation et volume	136
5.17	Maillage et simulation OPENFOAM	138
5.18	Influence de la co-simulation du courant sur la charge	140
5.19	Abstraction du phénomène d'empilement des blocs	143
5.20	Simulation du modèle de Marchenko	144
5.21	Traduction du modèle de Marchenko à ICE-MAS	146
5.22	Itérations de co-simulation du modèle de Marchenko	147
5.23	Comparaison de ICE-MAS et du modèle de Marchenko piloté	148
5.24	Paramétrage du processus de co-simulation	149

A.1	Méthode d'Euler explicite	173
A.2	Méthode d'Euler implicite	175
A.3	Méthode d'Euler explicite multi-agents à pas de temps égaux	176
A.4	Méthode d'Euler explicite multi-agents à pas de temps différents	176
A.5	Méthode d'Euler implicite multi-agents à pas de temps égaux	178
A.6	Cas de stabilité de la méthode d'Euler implicite multi-agents	179
A.7	Cas d'instabilité de la méthode d'Euler implicite multi-agents	179
A.8	Choix des différents pas de temps pour la méthode d'Euler implicite multi-agents	180
A.9	Raideur et méthode d'Euler implicite multi-agents	180

Introduction



Philippe Geluck - LE TOUR DU CHAT EN 365 JOURS

Pour bon nombre de secteurs industriels, la modélisation et la simulation des systèmes complexes sont devenues des étapes incontournables du processus de développement de nouveaux produits. Particulièrement, la conception de systèmes critiques soumis à de multiples phénomènes couplés nécessite de prévoir quelle sera leur réponse à certaines contraintes, dans l'optique d'optimiser les coûts de fabrication tout en garantissant la plus grande sécurité. Par exemple, le design de structures offshore pour les milieux polaires requiert entre autres données, des estimations des efforts maximaux causés par les morceaux de glace dérivants [Lau, 2001] [Timco et Croasdale, 2006] [Fransson et Bergdahl, 2009]. De cette manière, les structures peuvent être dimensionnées au plus juste pour supporter pérennément ces conditions extrêmes.

Pour traiter ce type de problèmes, deux approches sont souvent opposées. D'un côté, nous disposons aujourd'hui de quantités de données considérables et de modèles de plus en plus fins qui nous permettent de décrire un système dans les moindres détails. Cette approche microscopique cherche généralement à reproduire le comportement individuel des entités qui composent le système, grâce aux systèmes multi-agents par exemple [Seghrouchni et al., 2010]. Cette méthode présente pourtant des inconvénients majeurs. Ainsi, la précision des modèles

s'accompagne souvent de temps de calculs bien trop importants et rendent ces modèles inexploitable pour simuler la totalité du système. La démarche classique pour contourner ce problème consiste à réduire drastiquement le domaine de simulation en temps et en espace, ce qui ne favorise pas une étude globale.

Une autre piste, que nous appuyons, consiste à changer de point de vue sur le système et à le décomposer, non plus en entités autonomes, mais en ses phénomènes constitutifs pour étudier l'impact de chacun d'entre eux. Cette approche, dite phénoménologique, nous permet de décrire la dynamique du système telle que nous la percevons, ce qui se révèle particulièrement profitable lorsque nous sommes confrontés à des phénomènes qui ne sont que partiellement connus. Ce positionnement macroscopique rend possible la simulation efficace du système. Car si les résultats obtenus par une simulation macroscopique sont quantitativement moins bons que ceux obtenus par une simulation microscopique, ils sont néanmoins mieux adaptés à l'étude du système dans sa globalité. De plus, les temps de calculs plus raisonnables des modèles macroscopiques autorisent la construction de simulateurs interactifs. Pour autant, ce changement de niveau de modélisation n'est pas sans coût. Les lois descriptives que nous souhaitons utiliser impliquent généralement l'introduction de nombreux paramètres qui peuvent être difficilement identifiables dans le contexte.

Pour répondre à ce problème, nous proposons de combiner les différents points de vue de modélisation. Les lois descriptives étant le résultat d'abstractions qui résument un ou plusieurs phénomènes sous-jacents, nous faisons l'hypothèse qu'une simulation d'un modèle microscopique est équivalente à une simulation d'un modèle macroscopique correctement paramétré, à temps simulé égal et pour un phénomène donné. Aussi, notre idée est de simuler conjointement ces niveaux de description hétérogènes et d'utiliser des simulations microscopiques pour paramétrer un modèle macroscopique incomplet. La redondance de ces simulations permet d'envisager des simulateurs interactifs qui concilieraient la précision des simulations microscopiques et l'efficacité des simulations macroscopiques.

Pour mettre en œuvre notre proposition, nous nous appuyerons notamment sur un concept habituellement associé à la conception de systèmes électroniques : la co-simulation [Comete et Abib, 1998]. L'idée générale de la co-simulation est de coupler des simulateurs d'origines diverses et de formalismes différents, développés de manière totalement séparée. Ces aspects nous seront particulièrement précieux au moment de manipuler des échelles de modélisation hétérogènes.

Enfin, nous profitons de cette introduction générale pour souligner que ces travaux ont été menés dans un contexte industriel. Cette thèse a été financée par la société CERVVAL qui développe pour ses clients des outils sur mesure d'aide à la conception et à l'optimisation de systèmes complexes. Aussi, nos propositions s'inscrivent dans une démarche applicative et visent à apporter des solutions concrètes aux problèmes liés à la simulation interactive de ces systèmes.

Ces travaux ont fait l'objet de plusieurs publications à des conférences internationales dans les domaines de :

- ▷ la géophysique ([Combes, 2011], [Grigne et al., 2012], [Grigné et al., 2011], [Combes et al., 2012]),
- ▷ la modélisation de systèmes complexes ([Le Yaouanq et al., 2011a], [Le Yaouanq et al.,

- 2011b], [Le Yaouanq et al., 2015]),
- ▷ de l'ingénierie offshore pour les conditions arctiques ([Septseault et al., 2014], [Septseault et al., 2015], [Dudal et al., 2015], [Béal et al., 2016]).

Organisation du mémoire

La présentation de nos travaux s'étale sur cinq chapitres. Nous débuterons par une description des différents concepts liés à la modélisation et à la simulation des systèmes complexes au moyen des systèmes multi-agents. Nous poursuivrons avec une preuve par l'exemple de la pertinence de l'approche phénoménologique et de son implémentation par une approche multi-agents mixte, mais aussi de ses limitations. Nous présenterons ensuite un état de l'art des méthodes de modélisation redondantes et multi-formalismes. Puis, en nous appuyant sur ces éléments, nous proposerons une architecture logicielle pour la simulation conjointe d'échelles de modélisation hétérogènes ainsi que deux stratégies de co-simulation qui répondent aux problèmes de paramétrage des modèles descriptifs. Enfin, nous éprouverons l'ensemble de nos propositions dans le cadre du développement d'un outil d'aide à la conception de structures offshore pour les milieux polaires.

Nous donnons ci-dessous un résumé plus détaillé de ces cinq chapitres.

Le chapitre 1, intitulé « Systèmes et expérimentation », présente le cadre de notre étude et détaille tour à tour les différents aspects liés à la modélisation et l'expérimentation des systèmes complexes *in virtuo*. Nous décrivons ensuite les différentes approches pour construire des maquettes numériques à l'aide des systèmes multi-agents et de leurs déclinaisons. Pour cela, nous donnons une vue générale des problématiques liées à la simulation des systèmes multi-agents et en particulier la mise en œuvre de méthodes numériques adaptées. Nous concluons ce chapitre en affichant les atouts de l'approche phénoménologique pour la simulation des systèmes complexes et de son implémentation par des systèmes multi-agents mixtes.

Le chapitre 2, intitulé « Approche phénoménologique et systèmes multi-agents », présente la construction du modèle MACMA à laquelle nous avons contribué dans la première partie de cette thèse. Nous détaillons l'implémentation de l'approche phénoménologique pour décrire les mécanismes de refroidissement du manteau terrestre grâce à une approche multi-agents mixte. Nous terminons ce chapitre par une analyse paramétrique du modèle qui illustre les difficultés des modèles descriptifs.

Le chapitre 3, intitulé « Échelles de modélisation », propose une revue des méthodes de modélisation redondante d'échelles de modélisation hétérogènes. Nous débutons ce chapitre par la mise en avant d'une forme de hiérarchie des niveaux de description des modèles et leurs particularités. Ensuite, nous présentons plusieurs méthodes multi-échelles qui visent à exploiter cette hiérarchie des modèles pour nourrir les modèles de haut niveau grâce à des simulations redondantes de modèles plus fins. Nous présentons enfin la technique de la co-simulation que nous souhaitons adapter pour concevoir des simulations redondantes.

Le chapitre 4, intitulé « Simulation redondante d'échelles de modélisation hétérogènes », présente l'architecture logicielle que nous proposons pour corriger les failles des modèles

phénoménologiques fortement paramétrés. Tout au long de ce chapitre, nous nous appuyons sur le cas d'école du phénomène de diffusion pour lequel nous commençons par décrire un modèle macroscopique et un modèle microscopique. Suite à cela, nous détaillons une architecture de co-simulation qui offre la possibilité de simuler conjointement plusieurs niveaux de description pour un même système, sans impact sur le temps de simulation. Cette première implémentation aboutit à l'expression d'une stratégie de co-simulation pour la détermination explicite de paramètre à la demande. Nous proposons enfin une seconde stratégie qui consiste à utiliser des simulations microscopiques ponctuelles comme des outils de validation et de sélection implicite de jeux de paramètres pour les cas où la première stratégie échoue.

Le chapitre 5 intitulé « Application à la conception de structures offshore », illustre la mise en œuvre de nos différentes propositions à travers le développement d'un outil d'aide à la conception de structures offshore. Tout d'abord, nous situons ce contexte applicatif et particulièrement nous décrivons les interactions glace-structure qui motivent la création d'un simulateur multi-modèles. Par la suite, nous donnons des détails sur la construction du simulateur ICE-MAS et sur la manière dont les différents phénomènes y sont représentés et superposés. Nous expliquons alors comment optimiser, grâce à notre architecture de co-simulation, d'une part la précision des modèles hydrodynamiques et d'autre part le temps de calcul pour le rendre compatible avec un prototypage rapide.

Enfin, nous terminons ce mémoire par une synthèse de nos contributions et nous énonçons plusieurs pistes d'études complémentaires.

Chapitre 1

Systèmes et expérimentation



Philippe Geluck - LE TOUR DU CHAT EN 365 JOURS

La raison d'être de ce travail de thèse est de proposer des outils pour la compréhension de ce que nous appelons les systèmes complexes. Ce premier chapitre a pour but de présenter le cadre de notre étude et de détailler tour à tour un ensemble de concepts habituellement liés à ce type de démarche. Ainsi, dans un premier temps nous discuterons des différents aspects de la modélisation des systèmes complexes et de leur expérimentation, notamment *in virtuo*. Nous décrirons ensuite les différentes approches, dont celles développées au CERV, pour construire des maquettes numériques à l'aide des systèmes multi-agents. Enfin, nous évoquerons les problématiques liées à la simulation des systèmes multi-agents et en particulier la mise en œuvre de méthodes numériques adaptées.

1.1 Systèmes et complexité

Le terme « système » fait partie de ceux qui ont autant de définitions que d'auteurs qui écrivent à son sujet. Le dénominateur commun, qui fait consensus entre ces définitions, nous dit qu'un système désigne un ensemble d'éléments qui interagissent selon certains principes ou règles. En grec ancien, « *sustema* » signifie « organisation, ensemble » et provient du verbe « *sunistemi* », qui peut se traduire par « mettre en rapport, instituer, établir ». La description d'un système donné est donc constituée de la nature des éléments qui le composent, des interactions entre ces derniers et des limites du système. La figure 1.1 illustre un exemple simple de système, la bicyclette. Cette vue en éclaté nous montre que le système principal peut être décomposé en un ensemble de sous-systèmes. La fonction principale de l'objet, ici la locomotion, résulte alors de l'assemblage de ces sous-systèmes et des relations entre ces derniers.

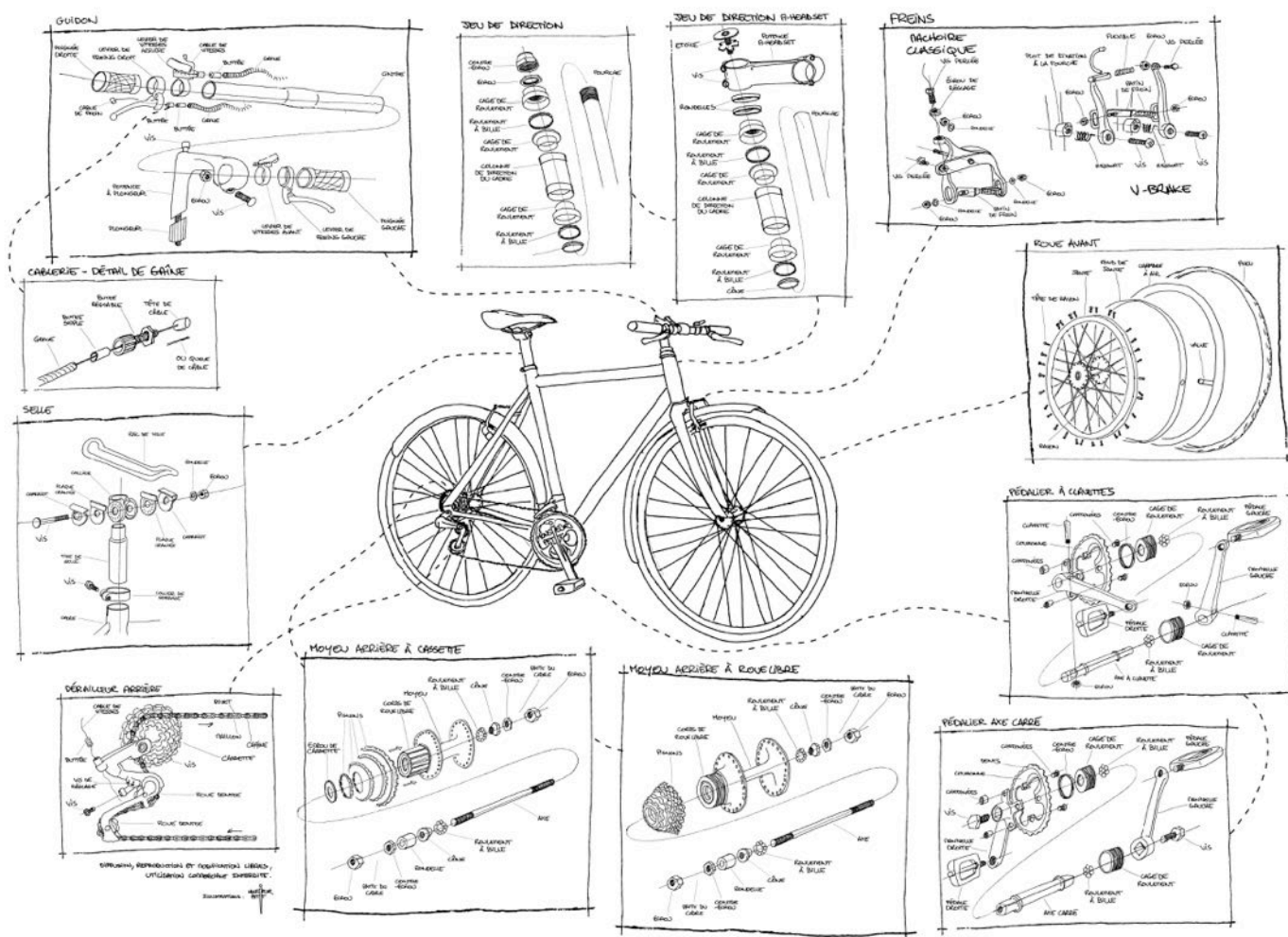


FIGURE 1.1 – Exemple de système : la bicyclette – Une bicyclette est un assemblage de sous-systèmes qui interagissent pour assurer une fonction, ici la locomotion. Source : <http://www.ohcyclo.org>

Dans cette vue simplifiée, les différents sous-systèmes sont compartimentés. Il n'y a en effet pas de rétroaction entre eux. Parfois, c'est un peu plus complexe.

La complexité est une propriété que l'on retrouve dans beaucoup de systèmes. Cette notion est largement utilisée dans la littérature. On considère souvent à tort que quelque chose est « complexe » quand il est difficile à comprendre ou à appréhender. Pour une définition plus rigoureuse, il faut recourir à l'étymologie. Le terme « complexe » vient du latin *complexus* qui signifie **entremêlé**. Ainsi, un système est dit complexe quand les interactions qui relient ses entités sont interdépendantes.

Par exemple, une réaction chimique, comme la dissolution d'un comprimé effervescent dans de l'eau, est simple car quelques équations permettent de décrire exactement l'évolution du système. Au contraire, sa diffusion dans un organisme est complexe car le métabolisme, la diversité des propriétés biochimiques et des réactions mises en œuvre, font que les équations ne suffisent plus à expliquer la totalité du système. Et lorsque l'homme et son libre arbitre entrent dans le système, sa complexité augmente encore [Tisseau, 2001]. La compréhension et donc la prédictibilité de ces systèmes deviennent difficiles dès lors qu'on ne connaît pas *a priori* les effets indirects que les interactions peuvent avoir sur le comportement global, du fait de leurs interdépendances.

Une approche réductionniste voudrait que l'on décompose un phénomène complexe en ses sous-phénomènes enchaînés par des liens de causalité. La théorie des systèmes s'oppose à cette démarche et étudie au contraire la finalité du système. Le phénomène complexe ne doit pas être décomposé, mais plutôt perçu comme une boîte noire. On considère non plus le fonctionnement du phénomène mais son comportement et ses interactions avec les autres phénomènes. En effet, l'enchevêtrement des interactions conduit à une boucle de rétroaction, comme décrite par Norbert Wiener pour ce qu'il a appelé la cybernétique [Wiener, 1950]. Platon déjà utilisait le terme « *kubernetike* » pour désigner le pilote, ou gouvernail, d'un bateau. Sous la plume de Wiener, la cybernétique est la science qui s'intéresse aux relations entre les éléments d'un système, par l'étude de l'information et des principes d'interaction. L'idée générale de la cybernétique est que ce sont au final les interactions qui dirigent la dynamique des systèmes et qu'elles maintiennent une pseudo-stabilité de ces systèmes. Cette stabilité est le résultat de processus de rétroaction ou *feed-back*, c'est-à-dire l'action en retour d'un effet sur sa propre cause. En général, la rétroaction a un effet variable au cours du temps, en fonction des conditions, des délais de transmission et de l'inertie du système. La rétroaction sera tantôt positive, jouant le rôle d'amplificateur d'un phénomène, tantôt négative, provoquant un amortissement qui permet une régulation. Le comportement global du système émerge alors des différentes boucles de rétroaction. Ces boucles sont illustrées par la figure 1.2 sur le thème des marchés financiers, mais on les retrouve dans les domaines de la biologie, de la psychologie, ou des systèmes mécaniques comme par exemple le thermostat.

Etudier un système complexe revient donc à considérer un système dans lequel des entités interagissent. Ces relations, en altérant les comportements de ces entités, vont modifier la dynamique du système. Une fois la dynamique globale du système acquise, on pourra alors envisager d'affiner nos connaissances sur le système [Le Moigne, 1993] et ainsi essayer de prédire quel sera son comportement dans certaines situations critiques.

Aujourd'hui, il reste difficile d'exprimer une théorie générale capable de formaliser ce genre de système. La compréhension des systèmes complexes passe alors par leur modélisation.



FIGURE 1.2 – **Boucle de rétro-action et marché financiers** – Différentes études sur les cours de la bourse ont révélé que les cours ne sont pas indéterministes, ce qui signifie que des causes (notamment des informations financières, économiques, technologiques, géopolitiques et psychologiques) impactent les cours de manière déterministe, entraînant une dynamique du processus de formation des prix dépendante de conditions initiales. Certains organismes se sont ainsi intéressés à la finance comportementale, et notamment aux sur-réactions des agents à des informations sur la santé de l'entreprise à court terme. En sur-évaluant ou sous-évaluant les prix, les prix s'écarteraient de leur « vraie » valeur d'équilibre à court terme, amorçant par la suite un repli plus gradué vers leur « vraie » valeur. Source : Image de KEVIN KALLAUGHER, <http://www.kaltoons.com>

1.2 Modéliser

1.2.1 Qu'est-ce qu'un modèle ?

La démarche de modélisation est une approche scientifique qui a pour but de produire de la connaissance sur des phénomènes naturels ou artificiels. Cette démarche est utilisée pour comprendre, prédire voire contrôler des objets ou des phénomènes, qu'ils soient déjà existants ou encore au stade de la conception. Le principe est de créer une représentation simplifiée, appelée « modèle », d'un phénomène pour pouvoir l'étudier. Quand on aborde le sujet de la modélisation, en informatique du moins, il est d'usage de citer Minsky [Minsky, 1965], qui définit un modèle comme suit :

Pour un observateur B , un objet A^ est un modèle d'un objet A (ou d'un phénomène) si B peut utiliser A^* pour répondre à des questions sur A*

Généralement, on distingue deux grandes familles de modèles :

- ▷ les modèles **analogiques** sont une représentation physique du système réel, par exemple un modèle réduit, un prototype ou un échantillon.
- ▷ les modèles **numériques** décrivent le système sous la forme d'équations mathématiques, et éventuellement d'une simulation informatique.

1.2.2 Pourquoi modéliser ?

La modélisation trouve son intérêt lorsqu'il s'agit d'étudier des phénomènes ou des objets qui sont difficilement manipulables ou observables, de par leur taille, leur dangerosité, ou tout simplement parce qu'ils n'existent pas.

En fonction du contexte, la modélisation peut revêtir des importances différentes. Par exemple, pour un chercheur, l'objectif est en général de comprendre le système réel et d'obtenir de nouvelles connaissances théoriques sur celui-ci.

Dans un cadre industriel, plus que de comprendre, il s'agit le plus souvent d'optimiser le système étudié. La modélisation représente plus une contrainte qu'un véritable choix, principalement à cause du coût d'expérimentations sur le système réel. La démarche de modélisation consiste alors en l'élaboration de prototypes ou maquettes sur lesquels sont effectuées des études paramétriques à des fins de dimensionnement et d'optimisation.

1.2.3 Comment construire un modèle ?

Proposer un modèle et en tirer des résultats exploitables n'est pas une activité anodine. Pour que les résultats obtenus soient utilisables, il est nécessaire de respecter une certaine démarche. Depuis plusieurs dizaines d'années, les recherches dans ce domaine ont fait

de l'activité de modélisation une science à part entière en proposant théories, outils et vocabulaires spécifiques.

Pour construire son modèle, l'observateur doit définir un ensemble de questions qu'il exprimera selon un certain paradigme, c'est-à-dire un modèle cohérent de pensée, un système de représentation du monde. Ceci le conduira à choisir les phénomènes et les objets à représenter, ainsi que les observations et mesures à effectuer sur le modèle pour répondre aux questions qu'il se pose. C'est ce qu'on appelle le **cadre expérimental** [Zeigler et al., 1976].

Un modèle est donc associé à un observateur, à un paradigme et à un cadre expérimental. De ce fait, il est aisé de voir qu'il peut exister autant de modèles possibles d'un système que de combinaisons de modélisateurs, de paradigmes et de questions sur ce système.

D'autre part, l'observateur construira son modèle selon un certain point de vue, en fonction de l'objectif de l'étude. Ces points de vue sont multiples pour un même phénomène. On parle alors de **niveaux d'abstraction** ou d'**échelles de modélisation**. Par exemple, le phénomène de diffusion d'un médicament peut être décrit au niveau de la cellule, de l'organe, de l'organisme, etc. Notons que les niveaux d'abstraction n'ont de sens que si on les confronte les uns avec les autres. L'échelle de l'organe est dite *microscopique* par rapport à celle de l'organisme et *macroscopique* vis-à-vis de celle de la cellule. Nous verrons dans la suite de ce manuscrit, et c'est tout l'enjeu de cette thèse entre autres travaux, que l'on peut tirer avantage de cette multiplicité des échelles de modélisation.

1.2.4 Comment valider un modèle ?

Toujours d'après Minsky,

Un modèle A^ est un bon modèle du phénomène A , du point de vue de l'observateur B si les réponses fournies par A^* correspondent à celles que B aurait obtenues en travaillant sur A*

Ainsi, pour valider un modèle, il est nécessaire de comparer les résultats fournis par ce modèle et d'évaluer l'erreur induite par la modélisation par rapport au fonctionnement du système réel. Dans le cas d'un système existant, on pourra effectuer des mesures en conditions réelles ou sur une maquette le cas échéant. Si le système est inexistant ou inaccessible, il faudra alors uniquement compter sur l'expertise du modélisateur. Si les sorties du modèle sont jugées suffisamment proches des données mesurées sur le système qu'il décrit, il est considéré comme valide.

Il convient cependant d'être vigilant par rapport à la définition de la validité d'un modèle. En effet, il est important de noter qu'un modèle, quel qu'il soit, demeure une représentation abstraite du système étudié. Quelle que soit la corrélation des résultats du modèle avec les données du système réel, ils n'ont pour autant de valeur que dans le domaine de validité du modèle défini par le cadre expérimental. Ainsi, le modélisateur doit toujours garder à l'esprit que les résultats obtenus par un modèle dépendent uniquement de la façon dont ce dernier a été construit. Cette remarque vaut autant pour les moyens d'évaluation du modèle que

pour les choix de conception du modèle. Un modèle ne peut être qu'une vision subjective du système étudié. Par conséquent, on préférera parler d'**évaluation** d'un modèle plutôt que de validation [Bommel, 2009].

Par ailleurs, la pertinence d'un modèle ne doit pas forcément être décidée au seul regard de la précision des résultats. S'il imite convenablement le comportement du système et permet au modélisateur d'affiner ses connaissances, il a un intérêt. Par exemple, un modèle très précis mais lent à simuler peut être moins utile qu'une simulation interactive en fonction du contexte de l'étude. De plus, le simple fait de construire le modèle implique le modélisateur dans une démarche de réflexion qui est aussi importante que l'analyse des résultats. De ce fait, l'évaluation d'un modèle par rapport à un autre devra prendre en compte l'objectif de ce modèle [Grimm, 1999] :

There are no true or false models because all models are, necessarily and deliberately, false to some degree. The only useful general classification of models is 'useful' and 'not useful' with respect to the purpose for which the model was designed.

Nous avons jusqu'ici éludé la question de l'acquisition des résultats d'un modèle. Nous avons vu précédemment qu'on considère qu'un système n'est pas complexe quand il est possible de préjuger de son état par la « simple » résolution d'un système d'équations. On parle alors de modèles *analytiques*. Par opposition, l'étude des systèmes complexes implique l'usage de modèles *numériques*, composés de règles qui décrivent l'évolution du système qu'on ne peut connaître *a priori*. Il s'agit alors d'expérimenter ces modèles de la même façon qu'on éprouverait un phénomène réel, dans le but d'obtenir des données exploitables.

1.3 Expérimenter

L'expérience est une observation provoquée dans le but de faire naître une idée.

CLAUDE BERNARD [Bernard, 1865]

Les principales qualités d'un modèle reposent sur ses capacités à décrire, suggérer, expliquer, prédire et simuler. La simulation, ou l'expérimentation, consiste à soumettre le modèle à un ensemble de contraintes et d'actions afin de tester son comportement. C'est interagir avec le modèle afin de vérifier qu'il est cohérent au regard des questions de modélisation, des données réelles et du comportement théorique qu'il devrait avoir. Au delà de l'aspect scientifique, c'est également un moyen pour l'ingénieur d'apprendre à utiliser le système expérimenté et de prévoir son fonctionnement dans certaines situations. L'expérimentation autorise l'observation du comportement d'un système, dont les règles de fonctionnement ne permettent pas de déduire le comportement global.

Les résultats d'une simulation permettent alors d'affiner nos connaissances sur le système réel et de formuler de nouvelles hypothèses que l'on pourra intégrer au modèle. On peut donc considérer que la simulation fait partie intégrante du processus de modélisation, comme le montre la figure 1.3.

Comme nous l'avons vu, un modèle peut prendre différentes formes. Il s'agit d'adapter les modes d'expérimentation en fonction du type de modèle mis à l'épreuve.

1.3.1 *In vivo*

Il arrive dans certaines situations que l'expérimentation doive se faire directement sur le système réel. Par exemple, pour étudier la tectonique des plaques, il est possible de construire des modèles pour étudier les phénomènes qui conduisent aux déplacements des plaques mais le seul moyen de savoir ce qui se passe réellement reste de faire des mesures de terrain. Dans l'industrie, il peut également être trop coûteux de construire une maquette d'un outil en production pour son calibrage suite à des changements de conditions aux limites ou des tests de nouvelles procédures qui n'avaient pas été imaginées à l'époque de sa construction.

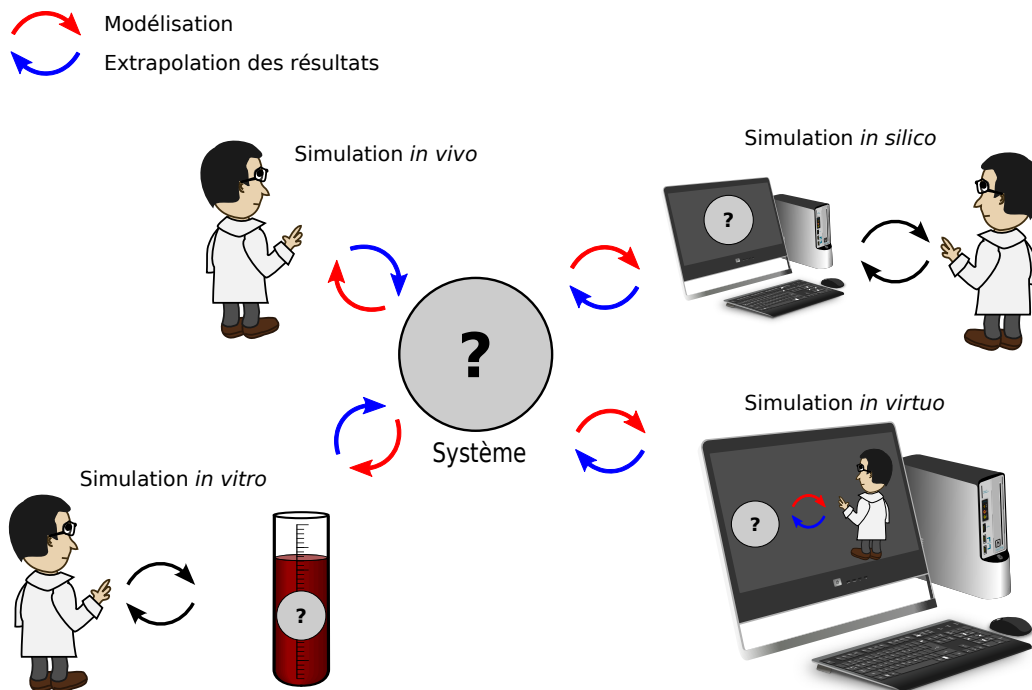


FIGURE 1.3 – **Méthodes d'expérimentation des systèmes complexes** – La démarche expérimentale consiste à modéliser le système étudié, puis à le simuler pour enfin extrapoler les résultats obtenus au système réel. Ce schéma illustre tout d'abord les trois types classiques de simulation : *in vivo*, c'est-à-dire sur le système réel ou une copie proche, *in vitro* sur une maquette simplifiée ou un sous-ensemble représentés ici par un tube à essais, et enfin *in silico*, avec un modèle numérique sur ordinateur. L'expérimentation *in virtuo* ajoute une notion d'interaction entre le modèle et son concepteur qui devient partie prenante de la simulation. Illustration inspirée de [Crépin, 2013].

1.3.2 *In vitro*

La simulation d'un modèle analogique passe par une expérimentation *in vitro* sur un échantillon ou sur maquette réduite construite par analogie avec le système réel. Les similitudes entre la maquette et le système améliorent alors la compréhension du système étudié. Par exemple, les essais en bassin sur des prototypes de houlogénérateurs permettent de mieux évaluer les capacités de production du système réel (figure 1.4). Cependant, ce type d'expérimentation a un coût non négligeable, d'autant que le nombre de paramètres à tester est souvent très grand. Qui plus est, il est souvent difficile d'avoir un modèle *in vitro* fiable. Du sang dans un tube à essai n'est pas du sang pulsé dans les artères, du fait que le changement d'échelle n'est pas naturel pour certaines propriétés du fluide comme la viscosité ou le nombre de Reynolds.

1.3.3 *In silico*

Les progrès technologiques dans tous les domaines scientifiques ont conduit à une explosion de données que seule l'informatique est à même de traiter. La notion de calcul *in silico* [Sieburg, 1990] a alors été introduite. Celui-ci entreprend l'élaboration *a priori* d'un modèle mathématique, qui est ensuite simulé numériquement sur ordinateur, pour aboutir à l'extrapolation *a posteriori* des différents résultats.

La simulation d'un modèle numérique passe le plus souvent par la résolution numérique de systèmes d'équations mathématiques (linéaires, non linéaires, différentielles, etc.). En effet, la détermination analytique de solutions se heurte souvent à des difficultés qui tiennent aussi bien aux caractéristiques des équations à résoudre (non-linéarité, couplage) qu'à la complexité des conditions aux limites et à la nécessité de prendre en compte des échelles spatio-temporelles très différentes. L'étude de la cinétique de réactions chimiques, le calcul des déformations d'un solide sous l'effet de contraintes thermo-mécaniques, ou la caractérisation du rayonnement électromagnétique d'une antenne, sont des exemples classiques d'implémentation sur ordinateur de systèmes d'équations différentielles. Ainsi le modèle numérique obtenu par discrétisation du modèle théorique est devenu aujourd'hui un outil indispensable pour dépasser les limitations théoriques.

Toutefois, l'expérimentation *in silico*, en se détournant des contraintes qu'impose la manipulation de l'objet réel, s'est également débarrassée de l'interaction entre l'expérimentateur et son expérience. L'expérimentation *in virtuo* vise à restaurer cette interaction.

1.3.4 *In virtuo*

Plus récemment, la possibilité d'interagir avec un programme en cours d'exécution, grâce notamment à la réalité virtuelle [Fuchs et al., 2003], a ouvert la voie à une véritable expérimentation *in virtuo* des modèles numériques. Il est désormais possible de perturber un modèle en cours d'exécution, de modifier dynamiquement les conditions aux limites, de supprimer ou d'ajouter des éléments en cours de simulation. Ceci confère aux modèles

numériques un statut de maquette virtuelle, infiniment plus malléable que les modèles classiques ou même que le système réel. Dépasseant la simple observation de l'activité du modèle numérique en cours d'exécution, l'utilisateur peut tester la réactivité et l'adaptabilité du modèle en fonctionnement.

Une expérimentation *in virtuo* est donc une expérimentation conduite dans un univers de modèles numériques en interaction et à laquelle l'homme participe. A l'instar du biologiste qui réalise des expérimentations *in vitro*, l'expérimentation *in virtuo* permet d'observer le phénomène comme si l'on disposait d'un microscope virtuel déplaçable et orientable à volonté, et capable de mises au point variées. L'utilisateur spectateur - acteur - créateur peut ainsi se focaliser sur l'observation d'un type de comportement particulier, observer l'activité d'un sous-système ou bien l'activité globale du système. A tout moment, l'utilisateur peut interrompre le phénomène et faire un point précis sur les corps en présence et sur les interactions en cours ; puis il peut relancer la simulation là où il l'avait arrêtée. A tout moment, l'utilisateur peut perturber le système en modifiant une propriété d'un élément (état, comportement), en retirant des éléments ou en ajoutant de nouveaux éléments. Il peut ainsi tester un comportement particulier, et plus généralement une idée, et immédiatement en observer les conséquences sur le système en fonctionnement.

Nous avons vu que la construction même du modèle permettait l'acquisition de connaissances sur le modèle et sur le système étudié, et qu'au delà du modèle, c'était sa construction qui était féconde. Dans une approche *in virtuo* qui augmente les interactions entre le système et le modélisateur, on peut penser que cette activité de modélisation est favorisée, parce que l'expérimentation *in virtuo* d'un modèle numérique lui assure une véritable présence et implique un véritable vécu de l'expérimentation que ne suggère pas la simple analyse de résultats numériques [Tisseau, 2001].

Nous allons à présent nous intéresser aux systèmes multi-agents, qui affichent des caractéristiques qui nous permettent d'envisager la mise en œuvre des expérimentations *in virtuo* telles que nous venons de les décrire.



FIGURE 1.4 – **Modélisation du houlogénérateur « Bilboquet »** – L'expérimentation *in virtuo* d'une maquette virtuelle dans le cadre de la conception d'un système, comme illustré à gauche, peut permettre d'identifier les paramètres sensibles du système. La simulation ne vise pas à remplacer les essais réels, ici sur un modèle réduit en bassin (à droite), mais intervient en complément de ceux-ci dans le but de limiter les plages de paramètres à tester. Source : Images CERVVAL, <http://www.cervval.com>

1.4 Systèmes multi-agents

Devant la complexité des systèmes que nous souhaitons modéliser et expérimenter, il est nécessaire de recourir à des outils adaptés, tels que les systèmes multi-agents (SMA).

Si les systèmes multi-agents sont aujourd'hui une branche de l'intelligence artificielle, le concept d'intelligence distribuée tire son inspiration de l'étude d'éco-systèmes d'insectes sociaux comme les essaims de fourmis. Le terme « *swarm intelligence* » a été énoncé dans le contexte de systèmes artificiels pour désigner la propriété de systèmes de robots non-intelligents qui montrent collectivement un comportement intelligent [Beni, 2005]. Cette analogie est également souvent reprise lorsque l'on évoque les systèmes multi-agents.

Ce paradigme de modélisation met en avant le fait que pour résoudre un problème informatique complexe, il est parfois plus simple de concevoir des programmes relativement petits (les agents) en interaction qu'un seul gros programme que nous qualifierons par la suite de monolithique. Cette idée se décline tout naturellement pour l'étude de systèmes complexes.

Aussi, l'approche multi-agents consiste à décrire individuellement, ou par groupe d'individus, les entités qui composent un système complexe sous la forme d'agents autonomes, potentiellement hétérogènes, qui interagissent au sein d'un environnement dynamique et partagé. La dynamique globale du système émerge alors de la superposition des interactions entre les différents agents, pour former une sorte d'intelligence collective ou intelligence en essaim [Bonabeau et al., 1999] [Chevrier et St Dizier, 2005].

La spécificité du paradigme multi-agents est qu'il fait, au contraire des paradigmes de modélisation par objets ou par composants, la distinction entre le niveau de résolution d'un problème et le niveau de contrôle de la résolution de ce problème [Seghrouchni et al., 2010]. De ce fait, les agents n'ont besoin que d'une connaissance limitée de leur environnement et leur autonomie leur permet de s'adapter dynamiquement aux changements imprévus qui interviennent dans l'environnement.

Les champs d'application des systèmes multi-agents sont aussi nombreux que variés. Dans le monde de la recherche, ils sont le plus souvent utilisés pour la simulation de phénomènes complexes [Navarrete Gutierrez, 2012], la robotique [Williams et al., 2014] ou la résolution et l'optimisation de problèmes [Eddy et Gooi, 2011]. Nous les trouvons également pour la conception de programmes, tels que les *Softbots* (software robots) de Etzioni [Etzioni et al., 1993], qui constituent une sorte d'interface logicielle de haut niveau entre l'utilisateur et le système d'exploitation d'un ordinateur.

Dans l'industrie, les systèmes multi-agents sont présents dans les jeux vidéo [Asselin, 2013] et l'animation cinématographique [Lind, 2001], mais aussi dans le domaine de la finance pour du *trading* automatique [Mariano et al., 2001].

1.4.1 Définitions

Pour présenter les principales caractéristiques des SMA, nous pouvons prendre appui sur l'approche *Voyelles* [Demazeau, 1995], complétée par J. Tisseau.

A-gents

S'il peut être de différentes natures, comme nous le verrons par la suite, deux définitions du terme « agent » se dégagent [Wooldridge, 2001].

La première tend à faire consensus et définit un agent comme étant un objet informatique autonome qui évolue et coexiste avec d'autres agents au sein d'un environnement. Il dispose en général de trois primitives comportementales qui forment un cycle. L'agent sait **percevoir** partiellement son environnement, sans en avoir une représentation totale. Cela lui permet de **décider** quel comportement il va adopter. Enfin, l'agent est capable d'**agir** sur son environnement et potentiellement sur les autres agents.

La deuxième définition est sujette à plus de contestations entre chercheurs et vient étendre la première. Pour certains chercheurs en intelligence artificielle, un agent peut également intégrer des notions anthropomorphiques comme le savoir, l'expérience, des intentions [Shoham, 1993] voire des émotions [Bates et al., 1992]. On parle alors d'agent « cognitif », capable de planifier sa décision [Wooldridge et Jennings, 1994] et de changer son propre plan d'action. La distinction se fait d'autant plus que les agents de ce type ont au minimum un but personnel, voire même la connaissance d'un objectif partagé avec la communauté qu'ils forment avec les autres agents.

Dans le cadre de nos travaux, nous nous limiterons aux agents dits « réactifs » dont le comportement est plus de type stimulus/réponse.

E-nvironnement

L'environnement est l'espace dans lequel évoluent les agents. Il peut être composé d'objets divers et d'autres agents. Au niveau local, il représente le médium d'interaction entre les agents et introduit les contraintes sur les relations inter-agents. Au niveau global, il est une sorte de mémoire collective et met à disposition des agents des informations sur son état ou sur les lois physiques du système simulé.

I-nteractions

La diversité des entités génère une diversité d'interactions entre elles. On rencontre des phénomènes physiques comme les collisions, des réactions chimiques, des échanges d'informations, etc. Ce sont ces interactions qui sont à la base des différents comportements tels que la coopération, la compétition ou la coordination. Ce concept est un point-clé des systèmes auto-organisés.

O-rganisation

L'organisation est une propriété essentielle des systèmes multi-agents. Elle est définie comme étant un ensemble de règles sociales, tel un code de la route, qui structure l'ensemble des entités en définissant des rôles et des contraintes entre ces rôles. Ces rôles peuvent être imposés aux agents, ou négociés entre eux. Ils peuvent être connus à l'avance ou émerger de situations conflictuelles ou collaboratives, résultats des interactions entre agents. Ainsi, les organisations, qu'elles soient prédéfinies ou émergentes, structurent les systèmes multi-agents en niveaux d'organisation, un niveau donné étant à la fois un agrégat d'entités au niveau inférieur, et une entité du niveau supérieur.

U-tilisateur

Dans le cadre d'une expérimentation *in virtuo*, l'utilisateur peut intervenir à tout moment dans la simulation. L'approche *Voyelles* a été étendue par Jacques Tisseau

en 2001 [Tisseau, 2001] en ajoutant la voyelle *U* pour désigner cet utilisateur. L'interaction de l'utilisateur avec le modèle peut se faire de plusieurs façons. Il peut par exemple intégrer, sous la forme d'un avatar, un agent absent en temps normal dans le système. Par l'intermédiaire de périphériques multisensoriels, il est immergé dans la simulation et a la possibilité de communiquer avec les autres agents, de les modifier, d'en créer de nouveaux ou d'en détruire. Dans un autre contexte, l'utilisateur peut se substituer à un agent et ainsi remplacer le processus de décision de l'agent. Dans les deux cas, il pourra par cette expérience évaluer l'impact de ses actions dans la simulation de son modèle. De ce fait, les systèmes multi-agents concilient les avantages des simulations numériques (reproductibilité, résolutions spatiale et temporelle) et ceux des sciences expérimentales.

1.4.2 Autonomies

Les modules de perception, de décision et d'action des agents constituent la base de leur autonomie. Cette notion d'autonomie est un point essentiel de l'expérimentation *in virtuo* et de l'étude des systèmes complexes d'une manière générale.

En effet, il n'existe pas de modèle capable d'expliquer *a priori* un système complexe. Ce manque de modèle de comportement global conduit à répartir le contrôle au niveau des entités qui composent ces systèmes et ainsi à autonomiser les modèles de ses composants. L'observation de l'évolution simultanée de ces composants permet alors de mieux appréhender le comportement d'ensemble du système global.

L'approche multi-agents nous permet de mettre en œuvre cette autonomisation des modèles en introduisant les concepts d'autonomie de conception et d'exécution.

Autonomie de conception

L'autonomie de conception d'un système multi-agents présente principalement deux avantages sur les simulations *in silico*.

Tout d'abord, l'étude d'un système complexe mêle généralement plusieurs champs disciplinaires. Par exemple, l'élaboration d'une éolienne fait intervenir conjointement des aérodynamiciens, des électroniciens et des mécaniciens entre autres experts de domaines différents. Même si une vue d'ensemble est indispensable, la réalisation de chaque élément reste relativement cloisonnée au départ. Les systèmes multi-agents et l'expérimentation *in virtuo* suivent une démarche similaire. Dans l'absolu, l'autonomie des agents permet de garantir cette incrémentalité de conception d'un modèle, puisqu'il « suffit » d'insérer le nouvel agent dans le système en définissant ses interactions possibles avec les autres agents. En pratique cependant, il est parfois difficile de découpler certaines équations en raison des intrications des phénomènes et des paradigmes de modélisation fixés.

Après la phase de conception, les différents modules sont assemblés pour former un prototype dont on teste le comportement général. Le second avantage qu'offre la relative autonomie de conception des systèmes multi-agents réside dans le fait que les paramètres de chaque brique du modèle sont facilement identifiables. Ainsi, il est possible d'évaluer exhaustivement l'impact de chacun d'entre eux.

Autonomie d'exécution

La dynamique des systèmes étudiés impose aux agents de prendre en compte les changements dans l'environnement, ainsi que d'adapter leurs réactions aux stimuli tant externes qu'internes. Les conditions aux limites, qui définissent un ensemble de contraintes sur le système, sont susceptibles de changer sans arrêt, que les causes soient connues ou non (interactions, perturbations, modification de l'environnement). Ceci est d'autant plus vrai quand l'homme intervient dans la simulation car il peut provoquer des modifications tout à fait imprévisibles initialement. Le modèle doit donc être capable de percevoir ces changements pour adapter son comportement en cours de simulation. L'autonomie de chaque agent est concrétisée grâce à un mécanisme d'ordonnancement qui exécute les cycles des agents les uns après les autres. Nous détaillerons par la suite les différents modes d'ordonnancement envisageables ainsi que leur influence sur l'expérimentation.

1.4.3 Approche phénoménologique

Les systèmes multi-agents ressemblent à première vue à la modélisation et à l'expérimentation réductionnistes de systèmes complexes, de par leurs définitions relativement proches. La notion d'autonomie des agents vient ajouter des possibilités de modélisation et d'exécution très appréciables quant à l'incrémentalité des modèles et à l'expérimentation. Ces avantages font en règle générale consensus dans le domaine des simulations multi-agents.

Nous allons à présent introduire un aspect moins traditionnel de l'approche multi-agents qui offre des perspectives avantageuses, notamment dans un contexte industriel.

Phénoménologie et expertise métier

L'approche phénoménologique consiste à modéliser chaque phénomène sur la base des observations ou des résultats expérimentaux à notre disposition. D'un certain point de vue, le résultat obtenu ne constitue pas véritablement un bon modèle, pas plus qu'un historique de température n'est un bon modèle météorologique, car sa validité et son aptitude à prédire le comportement du système sont restreintes aux données qui le composent. Toutefois, cette méthode présente des avantages tangibles pour la modélisation de systèmes complexes.

Tout d'abord, elle est particulièrement pertinente dans la mesure où nous ne maîtrisons souvent que partiellement la mécanique interne des phénomènes que nous étudions. Il sera alors plus aisé de modéliser leurs effets, tels qu'ils sont observés, sur les composants du système. Dans ce cas, il peut être utile de faire appel à des experts du métier cible qui, de par leur expérience, permettront de définir des lois empiriques et de calibrer les différents paramètres introduits.

Ensuite, certains phénomènes sont difficilement exprimables sous la forme d'équations. C'est notamment le cas des comportements humains qui nécessitent d'être modélisés à l'aide d'un autre formalisme, comme par exemple des machines à état, des réseaux de neurones, etc. Les systèmes multi-agents ont d'ailleurs fait leurs preuves dans différents domaines qui

s'attachent à la simulation de tels comportements, comme par exemple le trafic routier [Herviou, 2006] et les procédures d'évacuation de foules [Néron et al., 2012].

Nous verrons d'ailleurs que l'approche multi-agents, et plus particulièrement l'approche phénoménologique, permet de mixer les formalismes au sein d'une même simulation. Ces aspects font l'objet d'une démonstration sur un cas d'application au chapitre 2 de ce manuscrit.

Limites des systèmes multi-agents « classiques »

Dans la plupart des travaux, les systèmes multi-agents sont considérés comme une méthode de modélisation dite **individu-centrée**. En effet, un agent représentera souvent un individu d'une colonie (humain, fourmi, robots, etc.) ou une colonie entière avec un niveau de description supérieur, un programme informatique autonome (softbots), une molécule, etc. Si cette approche a montré un intérêt certain dans beaucoup de domaines, elle est néanmoins sujette à discussions.

Un reproche fréquemment formulé à l'encontre des SMA concerne la difficulté de passer à l'échelle les simulations de modèles individu-centrés. En effet, les systèmes qui nous intéressent sont souvent composés d'un très grand nombre d'entités.

Pour illustrer nos propos, prenons l'exemple de la simulation de la prolifération inter-cellulaire d'un virus dans un corps humain. Le nombre de cellules propres à une personne adulte est de l'ordre de 10^{14} . Dans l'hypothèse où nous serions à même de modéliser correctement le comportement d'une cellule, il resterait impossible de simuler la prolifération du virus à l'échelle réelle, au vu de nos capacités de calcul actuelles. Bien sûr, pour pallier cet écueil, nous avons abordé la possibilité d'user de niveaux de descriptions supérieurs. Mais de cette façon, le modèle n'est plus à même de remplir son objectif initial, puisque la dynamique inter-individuelle n'est plus observable. Il devient alors difficile d'étudier et d'expliquer l'impact de variations individuelles sur le comportement global.

Ensuite, il est parfois difficile de décrire un comportement de manière individuelle. Une communication directe suppose l'intervention de deux entités, de même qu'une collision, une réaction chimique, etc. Or avec l'approche multi-agents classique, la modélisation de ces événements entraîne généralement une redondance des calculs ou la mise en place de protocoles d'interaction complexes.

Il serait alors intéressant de pouvoir prendre en compte cet aspect de la modélisation des systèmes complexes tout en gardant les avantages de l'approche multi-agents. Pour ce faire, nous allons procéder à un changement de point de vue.

Déclinaisons de l'approche agents

Revenons un instant à la proposition de la cybernétique. Elle met en exergue que l'étude d'un système complexe doit se focaliser davantage sur les interactions entre les entités qui composent le système que sur les entités elles-mêmes, car ce sont ces interactions qui forment la dynamique du système et animent les constituants du système. Il convient alors de procéder

à une **décomposition phénoménologique** du système et de modéliser chaque phénomène en tant qu’entité autonome à part entière [Varela, 1979].

Plusieurs travaux ont abordé ce changement de point de vue dans l’optique de la simulation multi-agents. Ils reposent tous sur la même idée d’autonomisation des phénomènes, particulièrement traitée au CERV ¹.

En premier lieu, M. Parenthoën a proposé d’adapter la pensée énaïve de Varela aux systèmes multi-agents [Parenthoën, 2004] en autonomisant les modèles. Ces modèles autonomes deviennent alors les agents-**énaction** qui décomposent le monde en phénomènes. Ces agents vont agir sur un milieu qui sera créé et structuré par les agents eux-mêmes grâce à la notion de balise ou médiateur d’interactions. On voit là trois différences majeures avec les systèmes multi-agents classiques. La première est que le découpage du système n’est pas « topologique » mais « sémantique », chaque agent n’est pas représentatif d’une partie du système mais d’une partie de ce qui se passe au sein du système. La seconde est que les données représentatives du monde ne sont pas portées dans la structure des agents mais par une autre structure (le milieu) qui possède sa propre dynamique. Enfin les échanges d’informations entre les agents ne se font pas de façon directe, elles sont médiées par le milieu. Ces travaux ont notamment été appliqués à l’animation phénoménologique d’une mer virtuelle [Le Gal et al., 2007], qui a ensuite été couplée à différents simulateurs de flexibles, de bateaux et de structures offshore. La figure 1.5 illustre ces réalisations.

Kerdélo a ensuite introduit les agents-**réaction** [Kerdélo, 2006], qui ont fait l’objet d’une validation mathématique [Redou et al., 2005], dans le contexte de la simulation des réactions chimiques de la coagulation sanguine. Desmeulles a généralisé cette démarche et proposé de la nommer la **réification des interactions** [Desmeulles, 2006] dans le méta-modèle RÉISCOP ². Ces types d’agents ont la particularité de ne pas décomposer le monde en petits atomes actifs (par exemple 1 agent = 1 particule d’eau, 1 molécule) mais en petits phénomènes élémentaires (1 agent = 1 poussée d’archimède, 1 réaction chimique, etc.). Nous avons donc affaire à des agents-**interaction** qui réduisent les constituants du système, anciens agents-individu, à de simples états de la mémoire. Ils diffèrent des agents-énaction dans le sens où une interaction n’a de raison d’être qu’entre deux éléments spécifiques. La figure 1.6 illustre le passage de l’autonomie des agents à l’autonomie des interactions. Ce procédé de modélisation instaure une délocalisation de la dynamique du système, des individus vers les interactions. Notons que les agents-interaction constituent une méthode numérique particulièrement adaptée pour l’expérimentation de systèmes différentiels [Redou et al., 2007], car chaque agent calcule une partie considérée comme indépendante du système et garantit ainsi la flexibilité de modélisation et d’exécution dont nous avons parlé auparavant. Nous détaillerons ce fonctionnement plus en détail dans la section suivante.

Enfin, les travaux de Kubera [Kubera, 2010] s’inscrivent également dans cette logique et ont donné naissance au modèle formalisé IODA ³. Il se différencie de ses prédécesseurs en se concentrant sur l’aspect cognitiviste des interactions entre individus.

1. Centre Européen de Réalité Virtuelle

2. Réification des Interactions, Structures, Constituants, Organisations, Phénomènes

3. Interaction Oriented Design of Agent simulations

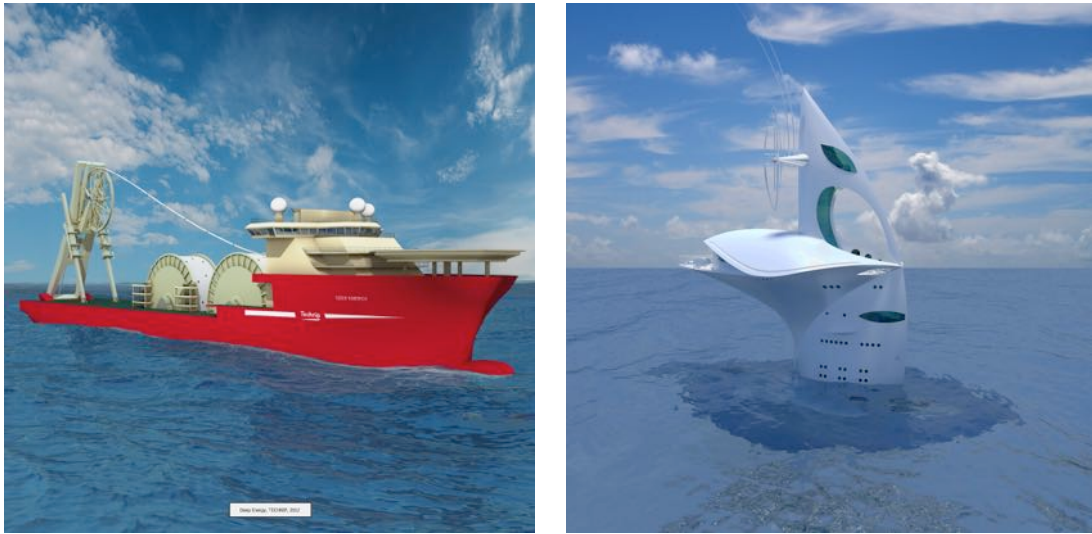


FIGURE 1.5 – **Couplage de simulateurs de structures flottantes et d'états de mer** – Un des avantages des simulations multi-agents réside dans le fait que différents modules développés séparément peuvent être facilement couplés. Sur ces images, nous pouvons voir un simulateur d'états de mer auquel nous avons associé des simulateurs de structures flottantes. À gauche, un simulateur du Deep Energy, navire de pose d'équipements en eaux profondes et à droite une maquette numérique du vaisseau d'exploration sous-marin SeaOrbiter. Source : Images CERVVAL, <http://www.cervval.com>, <http://www.seaorbiter.com>

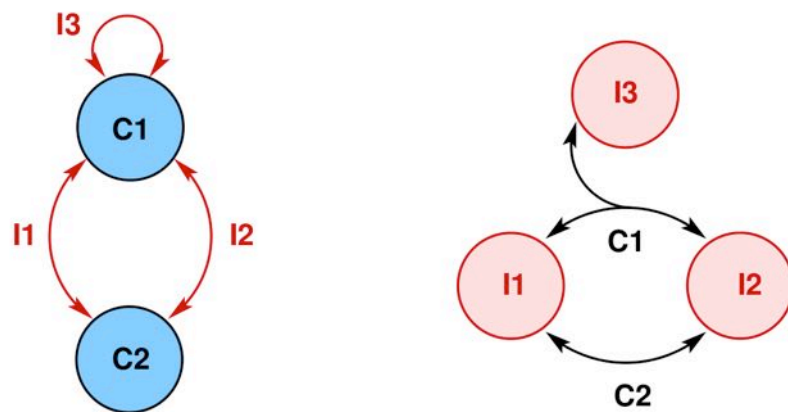


FIGURE 1.6 – **Agents « entité » versus agents « interaction »** – Ces deux graphes illustrent deux points de vue des systèmes multi-agents. À gauche, le graphe représente une méthode de modélisation individus-centrée avec des nœuds « composants » et des arcs « interactions ». Le graphe de droite représente l'approche interactions-centrée dans laquelle les arcs sont transformés en nœuds et inversement [Desmeulles, 2006].

1.5 Simulation

Notre objectif est d'expérimenter *in virtuo* des modèles de systèmes dynamiques, construits sur le paradigme multi-agents. Pour ce faire, nous devons recourir à la simulation numérique qui historiquement est apparue en même temps que l'informatique pour les besoins du projet Manhattan pendant la seconde guerre mondiale, afin de modéliser le processus de détonation nucléaire.

De même que la modélisation par agents peut être vue comme une discrétisation spatiale du système, le principe de la simulation implique de définir une discrétisation temporelle au moyen d'une variable appelée le temps de simulation. La simulation d'un modèle consiste à faire évoluer le temps dans le modèle et ainsi calculer, pour chaque valeur du temps de simulation, un état temporaire qui constitue un résultat partiel. Le changement d'état du modèle dépend de l'état précédent, de paramètres d'entrée et du temps de simulation.

Nous avons vu qu'une simulation multi-agents est composée d'une multitude d'objets actifs autonomes qui exécutent chacun leur comportement « perception-décision-action », dans le cas des agents réactifs auxquels nous nous limitons, à chaque itération de la simulation. Nous allons maintenant détailler l'organisation du temps de calcul des agents au sein d'une simulation. Notons que ce qui suit vaut autant pour les agents-individu que pour les agents-interaction.

1.5.1 De temps en temps

Dans le contexte de la simulation numérique, la définition du temps et sa gestion sont souvent une source de confusion. Nous commençons par donner ci-dessous quelques notions qui seront utilisées tout au long de ce mémoire.

Temps réel

Le temps réel ou temps physique correspond à celui que nous subissons, celui de notre montre. Le mode d'exécution temps réel signifie que l'on maîtrise la durée d'exécution des cycles de simulation. Pour les simulateurs dits *interactifs*, c'est-à-dire pour lesquels l'homme est dans la boucle de simulation, il est nécessaire de respecter la contrainte temps réel. Cela signifie que la latence entre la réalisation d'une action par l'utilisateur et sa prise en compte dans le monde virtuel par les modèles autonomes qui le composent soit « acceptable ».

Temps virtuel

Le temps simulé ou temps virtuel correspond au temps de la simulation informatique, celui que « perçoivent » les objets actifs, et qui s'écoule dans l'univers virtuel. La quantité de calculs à réaliser par unité de temps virtuel peut varier au cours de l'évolution du système simulé. Ainsi chaque cycle de simulation est effectué le plus rapidement possible mais potentiellement trop lentement pour maintenir la contrainte temps réel. Dans la mesure où il n'y a plus de corrélation entre les horloges réelles et virtuelles, la durée virtuelle d'un cycle, ou pas de temps, devra être une entrée de la simulation.

Temps de simulation

Le temps de simulation correspond au temps d'exécution de la simulation. En d'autres termes, il s'agit du temps nécessaire pour faire correspondre le temps virtuel avec le temps réel.

Nous allons à présent détailler l'organisation des calculs au cours d'un pas de simulation.

1.5.2 Politiques d'ordonnancement

Si l'autonomie des agents est une fonctionnalité séduisante dans le contexte de l'expérimentation *in virtuo*, son implémentation n'est pas triviale pour autant et pose un certain nombre de questions. Dans la majorité des articles présentant les applications orientées agents, la question de l'ordonnancement des calculs est le plus souvent éludée car le modèle numérique est considéré comme étant strictement équivalent au modèle théorique. L'ordonnancement est alors perçu comme un détail d'implémentation. Or le choix de la méthode d'ordonnancement constitue une hypothèse forte de la modélisation. En effet, nous étudions des systèmes dynamiques dont la simulation est indissociable de la gestion du temps.

L'ordonnanceur est le programme chargé de faire progresser la simulation de l'instant t à l'instant $t + 1$. Il doit également faire évoluer les agents en lançant leur cycle comportemental appelé **activité**, successivement. Au niveau du système, le cycle correspond à l'appel unique de l'ensemble des activités. La question de la synchronicité devient cruciale dès lors qu'on manipule des objets actifs. Pratiquement, nous distinguons deux modes d'ordonnancement séquentiel des activités lors d'un cycle : le mode synchrone et le mode asynchrone [Harrouet, 2000].

Synchrone Un cycle d'ordonnancement synchrone se déroule de la manière suivante :

- ▷ perception : tous les agents perçoivent l'état de l'environnement de l'instant t ;
- ▷ décision : les agents décident à partir des perceptions de l'instant t ;
- ▷ action : les agents font des modifications sur l'environnement perceptibles uniquement à partir de l'instant $t + 1$;

Les perceptions d'un agent sont les mêmes, qu'importe son ordre d'exécution. Il n'y a donc pas de causalité entre les activités durant un cycle.

Asynchrone Lors d'un ordonnancement asynchrone, les cycles de chaque agent sont exécutés successivement, de manière indivisible. Ainsi, il existe un lien de causalité entre les différentes activités au sein d'un même cycle de simulation. Cette méthode présente au moins deux avantages.

Tout d'abord, l'intérêt des simulations multi-agents est de pouvoir mettre en œuvre des mécanismes comme la compétition entre les phénomènes. Avoir un lien de causalité entre les activités au cours d'un même cycle permet d'assurer la réalisation de ces mécanismes. Ensuite, l'asynchronisme permet de s'affranchir des problèmes des accès concurrents aux données de la simulation. Les agents perçoivent et agissent chacun leur tour sur l'environnement.

La causalité inhérente à l'asynchronisme peut amener une activité à être favorisée par rapport à une autre si la séquence de leur appel est invariable, ce qui risque

d'introduire un biais dans la simulation. Une solution acceptable pour pallier ce problème est de brasser les activités à chaque cycle. Pour être efficace, ce brassage doit assurer un tirage aléatoire équiprobable pour chaque activité. Bien entendu, il s'agit d'un tirage sans remise puisque les activités ne doivent être appelées qu'une et une seule fois par cycle. Ce type d'ordonnement est qualifié de « chaotique ». Une composante stochastique est donc introduite dans la simulation. De cette manière, elle n'est plus déterministe. Le biais éventuel ne s'annule qu'à partir d'un grand nombre de cycles, puisqu'au premier cycle, le choix d'une activité plutôt qu'une autre n'est pas du tout justifié et n'a même aucun sens.

La comparaison des modes d'ordonnement synchrone et asynchrone est représentée par la figure 1.7.

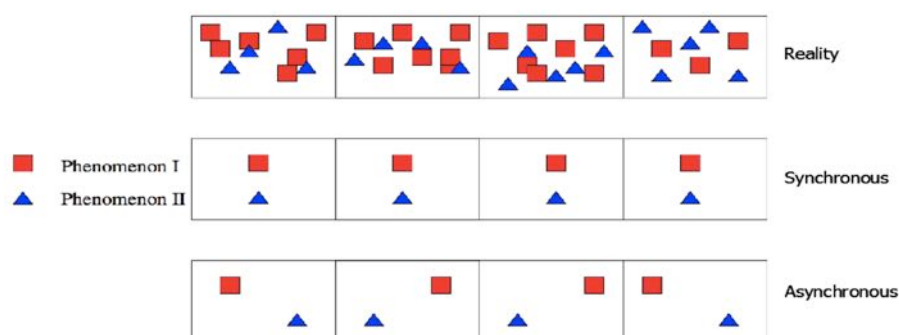


FIGURE 1.7 – **Ordonnements synchrone et asynchrone** – Les « chronogrammes » ci-dessus illustrent les deux types d'ordonnement synchrone et asynchrone des agents. Dans l'hypothèse synchrone, chaque phénomène est activé une fois par pas de temps. Puisque tous les événements se produisent simultanément, il n'y a pas de lien de causalité entre les phénomènes. Au contraire, dans l'hypothèse asynchrone, si chaque phénomène n'est également activé qu'une seule fois, ils ne le sont pas de manière simultanée. Ainsi, une causalité existe entre les phénomènes car comme nous pouvons le voir au premier pas de temps, le phénomène représenté par un triangle bleu perçoit un monde dans lequel le phénomène en rouge a déjà agi.

L'ordonnement « asynchrone chaotique » paraît particulièrement adapté dans le cadre de l'expérimentation *in virtuo*. Il permet non seulement de mettre en œuvre facilement l'autonomie d'exécution des systèmes multi-agents, mais également de mieux rendre compte de l'aspect parfois non déterministe des systèmes complexes. Cependant, nous allons voir dans la section suivante que son utilisation ne peut pas être systématique et que le mode d'ordonnement est en réalité imposé par les méthodes numériques que nous utilisons pour simuler des systèmes différentiels.

1.5.3 Méthodes numériques pour les systèmes différentiels

Les phénomènes que nous souhaitons expérimenter sont classiquement modélisés par des équations différentielles, organisées en système. Celui-ci ne peut cependant être que très

rarement résolu analytiquement dans le contexte. Dans de tels cas, il est possible d'utiliser une méthode numérique pour tenter de déterminer une approximation de la solution. Une telle méthode démarre depuis un état initial connu ou estimé et trouve des approximations successives qui devraient converger vers la solution sous certaines conditions. Notons que, même quand une méthode directe existe, une méthode numérique peut être préférable car elle est souvent plus efficace en terme de temps computationnel et d'interactivité de la simulation.

Dans le cas général de la résolution de systèmes d'équations différentielles ordinaires, les *problèmes de Cauchy* s'expriment classiquement sous la forme [Hairer et al., 1993] :

$$\begin{cases} y'(t)=f(t, y(t)) \\ y(0)=y_0 \end{cases} \quad (1.1)$$

Nous supposons que le problème de Cauchy (1.1) admet une unique solution sur un intervalle $[0, T]$ (T fini). Pour résoudre numériquement (1.1), la méthode naturelle est de découper l'intervalle $[0, T]$ en N intervalles, de longueurs non nécessairement identiques.

$$0 = t_0 < t_1 < \dots < t_n < \dots < t_{N-1} < t_N = t_0 + T \quad 0 \leq n \leq N$$

Le principe de la résolution numérique d'un système d'équation différentielles consiste à approcher les valeurs des solutions exactes $y(t_n)$ par des valeurs numériques y_n , au moyen d'itérations successives de pas d'intégration. Chaque pas consiste à passer de l'état y_n à l'état y_{n+1} en utilisant une méthode numérique. Pour illustrer leur fonctionnement, nous prendrons comme exemple les méthodes d'Euler qui n'ont qu'un intérêt pédagogique de par leur faible précision.

Méthodes explicites

Les méthodes explicites sont les méthodes de résolution numérique de (1.1) qui peuvent s'écrire sous la forme

$$y_{n+1} = y_n + h_n \Phi_f(t_n, y_n, h_n) \quad (1.2)$$

où la longueur h_n du n ième pas est

$$h_n = t_{n+1} - t_n \quad (1.3)$$

et $\Phi_f : [t_0, t_0 + T] \times \mathbb{R} \times \mathbb{R} \rightarrow \mathbb{R}$ est une fonction que l'on supposera continue.

Dans le cas de la méthode d'Euler explicite nous avons :

$$y_{n+1} = y_n + h_n f(t_n, y_n) \quad (1.4)$$

Comme indiqué précédemment, la résolution numérique d'un système d'équations différentielles engendre forcément une erreur numérique représentée par la figure 1.8. Pour valider les calculs, celle-ci doit être estimée. A l'instant t_{n+1} , cette erreur peut-être décomposée en deux parties :

- **l'erreur de consistance** ϵ_n , correspondant à l'erreur qui vient d'être commise sur le pas de temps $[t_n, t_{n+1}]$

$$\epsilon_n = y(t_{n+1}) - [y(t_n) + h_n f(t_n, y(t_n))] \quad (1.5)$$

- **l'erreur globale** e_n qui provient de tous les pas de temps antérieurs

$$e_n = y(t_n) - y_n \quad (1.6)$$

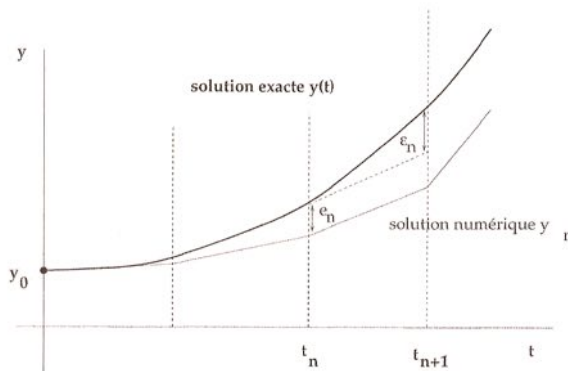


FIGURE 1.8 – **Erreur d'approximation de fonction par un schéma numérique** – La résolution numérique d'un système d'équations différentielles ne peut qu'approximer sa solution analytique. À chaque pas de temps, la solution numérique commet une erreur de consistance ϵ_n . La somme de ces erreurs constitue l'erreur globale e_n .

Malgré ces erreurs, on peut montrer que le schéma est :

- **consistant** si $\sum_{n=0}^N |e_n|$ tend vers 0 quand h_n tend vers 0
- **stable** si e_n reste borné pour tout ϵ_n

Si une méthode est consistante et stable, on peut en déduire qu'elle est **convergente** [Ascher et Petzold, 1998], c'est à dire que

$$\lim_{h \rightarrow 0} \max_{0 \leq n \leq N} |e_n| = 0 \quad (1.7)$$

Le coût du schéma d'Euler explicite est simplement déterminé par les évaluations de la fonction f à chaque pas de temps. Economiser du temps de calcul revient donc à réduire N , à augmenter h_n . L'idée motrice pour accélérer la simulation est de se dire que quand on remarque que l'erreur qu'on commet est tolérable, on s'autorise pour les quelques pas de temps qui suivent un pas plus large, et quand au contraire, on décèle une erreur devenant dangereusement grande on réduit le pas de temps. On parle alors de stratégies de contrôle de pas. Le choix du pas est cependant soumis à certaines contraintes introduites ci-après.

Méthodes implicites

Nous avons vu que les méthodes explicites approximent la valeur de y_{n+1} à partir de la valeur de y_n à t_n et de l'intégration sur le pas de temps h_n de la fonction Φ_f , elle-même

étant la dérivée de la fonction $f(t_n, y_n, h_n)$. Les méthodes implicites ont pour particularité d'utiliser en plus une approximation de y_{n+1} à l'instant t_{n+1} pour le calcul de Φ_f .

Nous avons donc :

$$y_{n+1} = y_n + h_n \Phi_f(t_n, h_n, t_{n+1}, y_{n+1}, h_{n+1}) \quad 0 \leq n \leq N-1 \quad (1.8)$$

Voyons maintenant pourquoi ce type de méthode est utilisé.

La A -stabilité est une propriété imposée lorsque sur un certain intervalle on désire choisir un grand pas d'intégration h pour « gagner du temps ». On dit qu'une méthode est A -stable si

$$\forall h > 0, \lim_{n \rightarrow \infty} y_n = 0 \quad (1.9)$$

pour toute équation (dite "test")

$$y' = -\lambda y(t) \quad \text{pour } \operatorname{Re}(\lambda) > 0. \quad (1.10)$$

Si l'on introduit notre exemple suivant, le schéma d'Euler implicite :

$$y_{n+1} = y_n + h_n f(t_{n+1}, y_{n+1}) \quad 0 \leq n \leq N-1, \quad (1.11)$$

donne l'équation suivante :

$$y_{n+1} = \prod_{k=0}^n \frac{1}{1 + \lambda h_k} y_0 \quad (1.12)$$

Il est clair que $y_n \rightarrow 0$ quand $n \rightarrow +\infty$ pour tout n indépendamment de h_n [Le Bris, 2005].

La méthode d'Euler implicite est donc A -stable, ce qui n'est pas le cas d'Euler explicite. On pourra ainsi prendre un grand pas de temps pour accélérer la résolution. En contrepartie, le calcul de y_{n+1} à partir de y_n revient à résoudre un système, non linéaire en général, à chaque itération. Dans l'évaluation du coût global de la méthode implicite, il faudra tenir compte de celui de chaque pas de temps en plus du nombre de pas.

Raideur d'un système d'équations différentielles

Lors de la résolution numérique d'un système d'équations différentielles, la valeur du pas d'intégration h_n est dictée à la fois par la précision numérique souhaitée mais aussi par la stabilité. Un exemple classique illustrant cette problématique est celui du système masse-ressort [Béal et al., 2008]. La raideur, physique cette fois, du ressort nous impose un pas de temps très faible pour conserver la stabilité de la simulation. On dira que le système est **raide** si la stabilité de la méthode numérique employée induit une contrainte sur le pas de temps plus forte que l'exigence de précision.

1.5.4 Méthodes numériques et systèmes multi-agents

Dans le cas d’une simulation multi-agents, le problème global peut être décomposé en sous-systèmes. Les agents résolvent chacun une partie du calcul, à des vitesses adaptables et hétérogènes. Les agents sont localement itérés puis leurs résultats sont superposés.

Par exemple, pour un système à deux agents (représentés par les fonctions f_1 et f_2) résolvant le schéma d’Euler explicite suivant

$$y_{n+1} = y_n + h_n \varphi_{f_1+f_2}(t_n, y_n, h_n) \quad (1.13)$$

on obtient équiprobablement [Redou et al., 2007] :

$$\begin{aligned} y_* &= y_n + h_n \varphi_{f_1}(t_n, y_n, h_n) \\ y_{n+1} &= y_{n*} + h_n \varphi_{f_2}(t_n, y_*, h_n) \end{aligned} \quad (1.14)$$

et

$$\begin{aligned} y_* &= y_n + h_n \varphi_{f_2}(t_n, y_n, h_n) \\ y_{n+1} &= y_{n*} + h_n \varphi_{f_1}(t_n, y_*, h_n) \end{aligned} \quad (1.15)$$

Le lecteur intéressé trouvera en annexe un exemple détaillé de l’utilisation des méthodes d’intégration numérique multi-agents à travers la simulation d’un système masse / ressort.

Nous pouvons voir ici que l’approche asynchrone a des limites, notamment dans le cadre de simulations numériques complexes pour lesquelles des schémas de calcul implicites sont indispensables [Redou et al., 2010]. En effet, la résolution du système implicite introduit par le calcul de y_{n+1} à partir de y_n par un agent suppose que chaque autre agent a déjà réalisé sa partie du calcul global.

1.6 Synthèse

Cette première partie nous a permis de présenter la complexité des systèmes qui vont nous intéresser par la suite. Nous avons mis en évidence que ces systèmes doivent être expérimentés afin de mieux les appréhender et ainsi gagner en prédictibilité. C’est ce que propose l’expérimentation *in virtuo* au moyen des systèmes multi-agents.

On retiendra que la modélisation multi-agents autorise la description d’un système à l’aide d’informations, d’agents et d’interactions plutôt que par des variables et des équations décrivant le système dans sa globalité. Elle permet donc au modèle d’être au plus proche de la réalité. Tout comme un système complexe, un système multi-agents est un ensemble d’éléments organisés par des interactions, en constante évolution vis-à-vis de son environnement et qui ne dispose pas d’un contrôleur global. Dans ce système émerge un

comportement qui résulte de la somme des individualités. De plus, l'autonomie des agents autorise une incrémentalité dans la modélisation d'un système complexe et dans l'exécution de son modèle.

L'approche multi-agents a été déclinée de nombreuses fois au cours des dernières années et propose des méta-modèles qui permettent une description plus phénoménologique des systèmes complexes. Les interactions sont au cœur de ces nouvelles approches et viennent remplacer les entités en tant qu'agents.

Chacun de ces modèles a ses applications de prédilection. Les agents-réaction, comme le laisse entendre leur nom, conviennent ainsi bien au calcul de systèmes chimiques. Les agents-interaction ont fait quant à eux leurs preuves dans le cadre d'applications en biologie ou en mécanique. Enfin, les agents-énaction, du fait de leur mode de fonctionnement, conviennent bien à la simulation de milieux homogènes peuplés de différents phénomènes connus de façon descriptive.

Si les méta-modèles présentés ont le défaut d'être auto-exclusifs, ils apportent tous des outils intéressants que nous souhaiterions combiner au sein d'une même simulation. En effet, certaines applications ne peuvent choisir aisément entre un seul de ces modèles. La simulation de comportement de structures offshore en est un exemple :

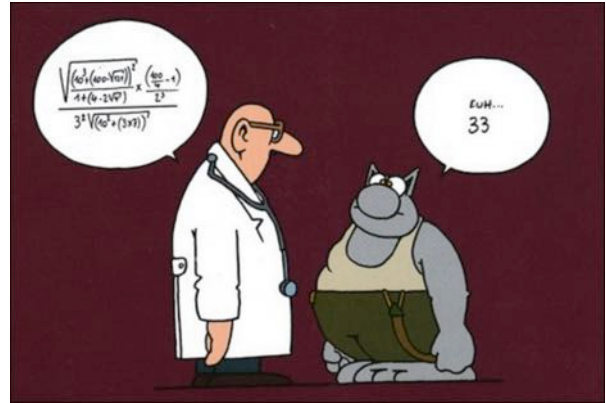
- Des phénomènes de natures très diverses interviennent dans ce type de simulation (mécanique, hydrodynamique, thermodynamique, ...). Le fait de pouvoir programmer ces phénomènes indépendamment les uns des autres est souhaitable du point de vue « génie logiciel » ;
- Les structures mécaniques, nous l'avons vu, se modélisent bien par agents-interaction. La preuve en est que les règles de mécanique élémentaire (celles apprises au lycée) sont quasiment toujours la description de l'interaction entre deux éléments. *A contrario*, l'état de la mer est bien plus efficacement représenté par un ensemble d'agents-énaction.

Le premier objectif de cette thèse est de démontrer l'utilité d'une approche multi-agents mixte. Plutôt que de tenter de trouver un modèle d'agent universel, capable de représenter tous types de phénomènes ou d'individus, on s'attachera ici à faire cohabiter les modèles existants. C'est l'idée directrice des travaux présentés dans le chapitre 2.

Par ailleurs, la modélisation phénoménologique que nous soutenons s'appuie sur des modèles dits descriptifs. Ces derniers vont nous permettre de garantir à la fois l'interactivité de nos simulations mais également d'exprimer les phénomènes mal connus. Cependant, ils imposent un paramétrage fort qui peut être difficile à identifier, non seulement concernant les conditions initiales, mais également au cours de l'évolution du système, de par sa dynamique et les interventions de l'expérimentateur. Nous étudierons dans le chapitre 3 différentes pistes pour résoudre ce problème en nous appuyant sur des échelles de modélisation hétérogènes.

Chapitre 2

Approche phénoménologique et Systèmes multi-agents



Philippe Geluck - LE TOUR DU CHAT EN 365 JOURS

Les systèmes multi-agents ont montré leur efficacité pour la description et l'expérimentation de systèmes complexes. Le chapitre précédent a mis en avant leurs avantages par rapport aux méthodes numériques classiques, et notamment leur faculté à adopter différents formalismes pour exprimer les comportements des agents. Plus particulièrement, nous avons vu qu'une approche phénoménologique permettait à la fois de s'affranchir des limites de passage à l'échelle inhérentes à un point de vue individus-centré, mais également de combler certaines lacunes du modèle par des connaissances expérimentales.

Le but de ce nouveau chapitre est d'illustrer la mise en œuvre de cette approche dans le cadre de la construction d'un simulateur des phénomènes de réchauffement du manteau terrestre. Pour cela, nous allons nous appuyer sur les travaux de thèse en physique de M. Combes [Combes et al., 2010] auxquels nous avons participé en implémentant le modèle MACMA (*MultiAgent Convective Mantle*). Ces travaux ont été réalisés en collaboration avec le Laboratoire Domaines Océaniques de l'Institut Universitaire Européen de la Mer (IUEM, UMR CNRS 6538). Nous nous intéresserons ici aux particularités de l'architecture logicielle

retenue.

Ces travaux utilisent la modélisation multi-agents comme un outil puissant pour déterminer dans quelle mesure un simple changement local de comportement géophysique peut affecter la dynamique terrestre à grande échelle. L’objectif est l’étude par une méthode nouvelle des relations implicites qui lient le refroidissement de la Terre à long terme avec les processus rapides à l’échelle des temps géologiques tels que la tectonique des plaques. Cette approche permettra aussi de tester l’importance relative des phénomènes de couplage thermomécanique qui animent notre planète. Nous utiliserons donc les systèmes multi-agents pour expérimenter des modèles, et construire un laboratoire virtuel de géophysique qui puisse rendre compte de tels processus.

La géophysique n’est pas notre domaine d’expertise et comme pour bon nombre de travaux inter-disciplinaires, il est important de définir le contexte dans lequel vont s’inscrire nos propos. Ainsi, nous commencerons par donner quelques éléments de géophysique qui permettront une compréhension intuitive chez le lecteur qui ne serait pas familier avec les concepts de base de la géodynamique terrestre. Nous limiterons donc au maximum l’écriture d’équations lourdes et nous nous concentrerons plutôt sur leur sens au sein des différents modèles présentés. Nous invitons le lecteur intéressé par le détail des équations et des hypothèses géophysiques à consulter le manuscrit à l’origine de MACMA [Combes, 2011].

Ensuite, nous exposerons la mise en œuvre de l’approche phénoménologique pour la modélisation du système Terre et de sa mécanique. En particulier, nous verrons que différents types d’agents peuvent être couplés pour aborder des problèmes que les approches classiques peinent à résoudre.

Enfin, nous montrerons les bénéfices que peuvent apporter les systèmes multi-agents pour l’analyse et la compréhension de systèmes complexes, de par leur flexibilité et leur expressivité.

2.1 Éléments de géophysique

La planète Terre est un objet relativement mal connu. Si la dynamique de sa surface est de mieux en mieux contrainte depuis que nous la scrutons par voie satellitaire, sa dynamique interne reste globalement une énigme. En effet, la connaissance indirecte que nous avons de la structure interne est issue principalement des mesures de champ de pesanteur et de vitesses d’ondes sismiques, qui sont autant d’observables qui n’ont pas de causes uniques. Nous sommes donc tributaires d’un certain nombre d’hypothèses faites sur la composition du manteau et sur les couplages physico-chimiques qui régissent son fonctionnement.

Cette section a pour but d’exposer ces hypothèses et d’esquisser un portrait dynamique de la Terre. Nous tenterons d’expliquer le fonctionnement du système Terre en tant que machine thermique, qui transforme une énergie d’origine principalement radioactive en énergie cinétique de convection, pour finalement la dissiper par conduction à sa surface, au contact des océans et de l’atmosphère.

2.1.1 Mécanique interne du système Terre

Un premier niveau de description de la structure thermique et mécanique de la Terre est présenté sur la figure 2.1. Le centre de la Terre se compose d'une graine solide incluse dans un noyau externe liquide. Ce noyau est entouré du manteau, qui est un milieu de fortes viscosités, siège des mouvements de convection thermique. Une fine lithosphère rigide recouvre le manteau. Cette lithosphère est divisée en plaques supportant les continents, qui sont moins denses.

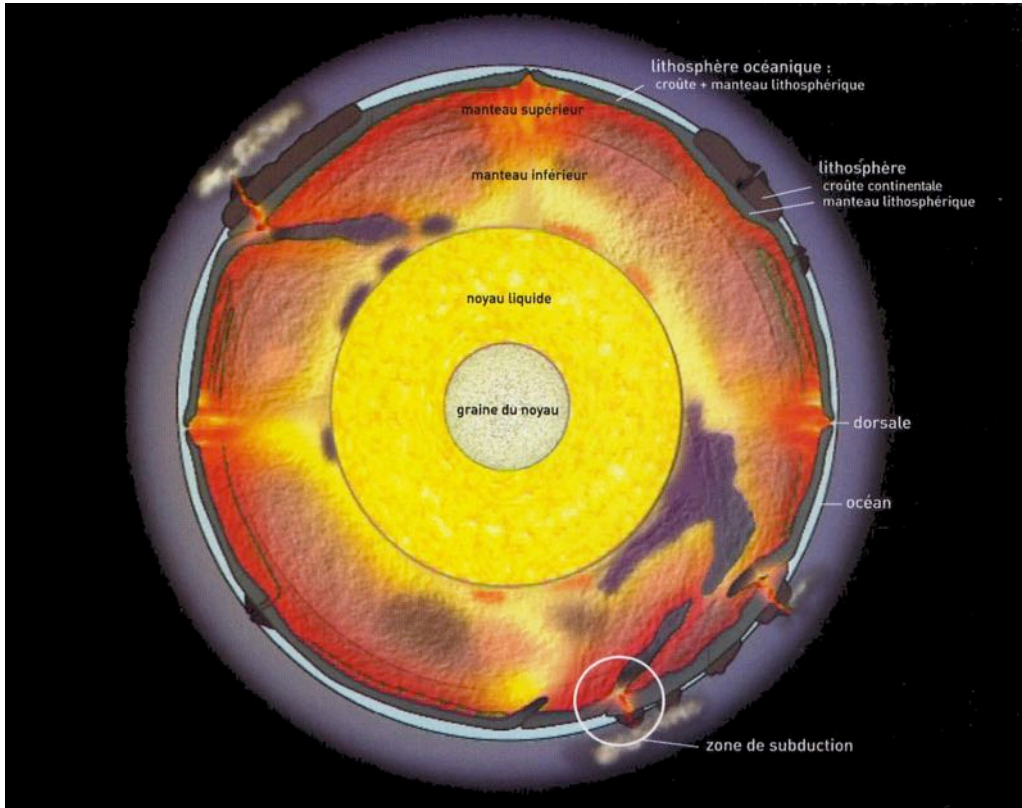


FIGURE 2.1 – **Structure interne de la Terre** – Les quatres couches superposées de la Terre sont la graine (solide), le noyau externe (liquide), le manteau (visqueux) et la lithosphère (rigide). Les continents, moins denses, sont portés par la lithosphère. Source : A. Meunier, *Université de Poitiers*

La température de l'interface noyau/manteau est $T_n = 1900K$, et on fixe la température à la surface de la Terre à $T_0 = 300K$. Ce saut de température entraîne des mouvements de convection dans le manteau, car la densité d'une roche (fluide de très forte viscosité) décroît avec sa température. La poussée d'archimède entraîne donc vers la surface des panaches chauds, qui sont freinés par le frottement visqueux du milieu plus froid dans lequel ils se déplacent. La convection thermique se met en place dans le manteau sous la forme de cellules de convection, dites de Rayleigh-Bénard [Grigné, 2003], qui entraînent dans leur mouvement les plaques lithosphériques. Par conduction, une fraction du flux de chaleur des cellules est alors dissipée au contact de ces plaques.

L'entraînement de la lithosphère par la convection mantellique est à l'origine de la

tectonique des plaques. La remontée des panaches chauds en surface se traduit par des épanchements de matériau mantellique au niveau des dorsales, qui sont des zones de production de plancher océanique (croissance de la lithosphère rigide). Pour conserver une surface constante de la planète, il est logique de concevoir des zones de disparition du plancher océanique, appelées zones de subduction. Dans ces zones, une plaque de forte densité replonge dans le manteau en étant chevauchée par une plaque de densité moindre. La figure 2.2 illustre l'ensemble de ces mécanismes.

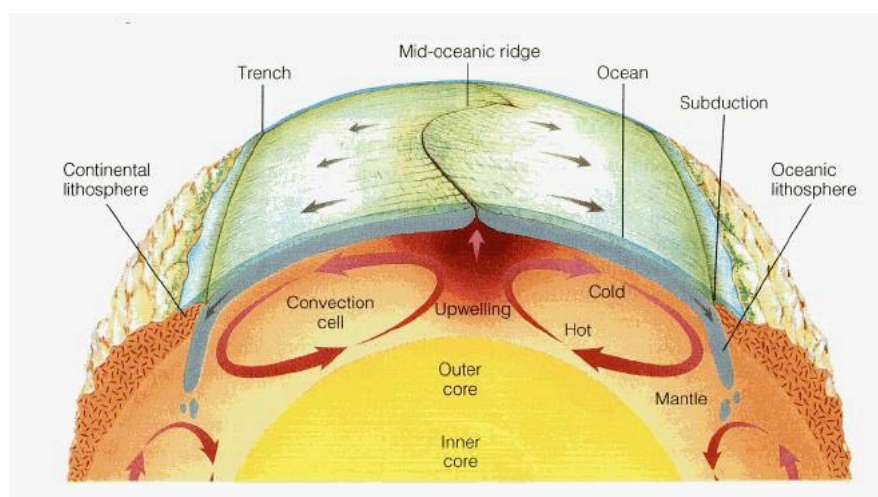


FIGURE 2.2 – **Schéma simplifié du couplage entre plaques tectoniques et manteau convectif** – Les cellules de convection du manteau sont bornées par des courants chauds ascendants à l'aplomb des dorsales, et des courants froids descendants au niveau des zones de subduction. Source : <http://www.planet-terre.ens-lyon.fr>

Nous allons à présent nous intéresser à la nature des forces qui mettent en mouvement les plaques lithosphériques.

2.1.2 Dynamique des plaques

De nombreuses études se sont attachées à décrire les forces exercées sur une plaque et à évaluer leur intensité relative. Nous en retiendrons les principales, qui sont présentées sur la figure 2.3. Ces forces peuvent être réparties en deux catégories. D'un côté nous avons des forces motrices qui tendent à déplacer une plaque lithosphérique : la traction gravitaire de la plaque plongeante (**Slab pull**), la force de succion (**Suction**) et la résultante des forces de pression au niveau des dorsales (**Ridge push**). De l'autre côté, les forces résistives suivantes vont s'opposer au mouvement de la plaque : les frottements visqueux avec le manteau (**Drag force**, **Slab force**) et la résistance visqueuse intra-plaque (**Bending**), non représentée sur la figure.

Toutes les études s'accordent à dire que la force motrice la plus importante est la traction gravitaire de la plaque plongeante immergée dans le manteau. Le poids de la plaque plongeante (*slab*) est parfois décomposé en deux contributions. On retrouve correctement les vitesses des plaques terrestres actuelles si l'on considère que le poids de la partie plongeante située dans le

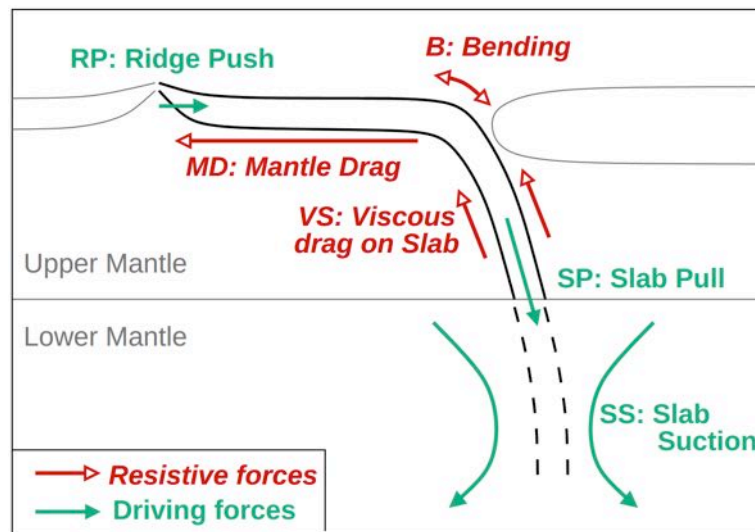


FIGURE 2.3 – **Schéma des principales forces exercées sur une plaque** – Les plaques tectoniques sont soumises à un ensemble de forces motrices et résistives [Grigné et al., 2012]

manteau supérieur est supporté par la plaque horizontale en surface, et que le poids du reste de la plaque est supporté par le manteau inférieur, ce qui a pour conséquence d'entraîner la plaque horizontale *via* la circulation résultante du manteau [Conrad et Lithgow-Bertelloni, 2002]. Cette seconde contribution sera appelée force de succion (*suction force*) dans la suite de cette étude. Notons que cette force entraîne la circulation du manteau visqueux dans les deux cellules adjacentes, ce qui peut contribuer significativement au déplacement d'une plaque qui ne dispose pas de traction gravitaire de plaque plongeante.

La troisième force motrice considérée est la résultante des forces de pression au niveau d'une dorsale. Elle résulte d'un déséquilibre de pression existant à des profondeurs de moins de 100 km, du fait de l'élévation de l'axe de la dorsale jusqu'à 2000 m au dessus des fonds océaniques. En effet, la répartition de masse dans les enveloppes superficielles de la Terre montre que les bosses et les creux de la surface topographique (montagnes, nappes phréatiques, etc.) sont pour l'essentiel compensés par des déficits ou des excès de masse, conduisant à un équilibrage des forces de pression sous une ligne imaginaire appelée « niveau de compensation isostatique ». La lithosphère se situant au-dessus de ce niveau, la pression produite à une profondeur donnée par l'excès d'élévation de la ride océanique est plus forte que la pression exercée par la partie plus vieille de la lithosphère. L'amplitude de cette force est toutefois évaluée à un ordre de grandeur en dessous de la traction gravitaire exercée par une plaque plongeante.

Le frottement visqueux exercé par le manteau sur la plaque horizontale et sur la plaque plongeante est en première approximation considéré comme newtonien, c'est-à-dire que le taux de déformation du manteau est proportionnel à la contrainte appliquée. Cette force étant aussi proportionnelle à la surface de friction, une plaque sera d'autant plus freinée qu'elle sera étendue. Enfin, une dernière force résistive est associée à la déformation visqueuse interne à la plaque, au niveau des zones de subduction. Dans ce contexte, la plaque plie sous son propre poids, et la puissance développée par la traction gravitaire de la plaque plongeante est partiellement consommée par la dissipation visqueuse due à la déformation

de la lithosphère. On considère, pour exprimer cette force, que la lithosphère est un fluide extrêmement visqueux. La résultante de cette force est connue sous le nom anglais de *bending* et son intensité varie selon les auteurs, de celle du *ridge push* à celle de la force de succion (jusqu'à 40% de la traction gravitaire).

2.1.3 Histoire thermique de la Terre

Un scénario généralement admis consiste à dire qu'au cours de son histoire thermique, la Terre a d'abord vu son manteau se solidifier sous l'effet de la pression, et donc se réchauffer, laissant un océan de magma en surface. La cristallisation progressive de cet océan a conduit à la création de plaques solides qui se déstabilisaient régulièrement pour participer à la convection du système.

L'énergie thermique produite par la radioactivité du manteau et la chaleur du noyau est libérée au niveau des dorsales et à travers la lithosphère, ce qui a pour effet de faire chuter la température du manteau. On estime que le manteau tel qu'il est décrit ici (fluide visqueux recouvert d'une lithosphère rigide) a atteint sa température maximale il y a 3 milliards d'années, et que celle-ci était plus élevée de seulement $200K$ [Jaupart et al., 2007].

Pour étudier la capacité de la Terre à évacuer son énergie radioactive, il est pratique de définir le nombre d'Urey comme le rapport de l'énergie radioactive produite dans le manteau sur l'énergie diffusée en surface. Ce rapport vaut actuellement $0,33 \pm 0,15$, ce qui illustre un régime de refroidissement du manteau de l'ordre de $120K/Ga$ ¹. Dans ces conditions, il semble que le taux de refroidissement est aujourd'hui plus fort que sa valeur moyenne dans le passé car sinon la température du manteau aurait été plus élevée dans le passé. Les modèles classiques ne parviennent pas à expliquer cette différence et conduisent à une « catastrophe thermique » [Labrosse et Jaupart, 2007] lorsqu'on remonte le temps à partir de la valeur actuelle des pertes de chaleur en surface. Il semble donc que d'autres phénomènes doivent être pris en compte pour compléter les études de refroidissement du manteau.

En premier lieu, il a été démontré que la présence des continents est une donnée importante du problème. Les mesures de terrain indiquent que le flux géothermique à travers les continents est 5 à 10 fois plus faible que le flux à travers le plancher océanique. Les continents jouent donc un rôle d'isolants thermiques [Grigné et al., 2007].

Ensuite, toujours à partir de mesures, nous pouvons affirmer que la distribution du flux de chaleur n'est pas uniforme selon l'âge de la lithosphère océanique. Le flux surfacique est beaucoup plus élevé pour une lithosphère très jeune, puis décroît rapidement jusqu'à $30 Ma$ ², avant de stagner après $80 Ma$ [Erickson et al., 1975].

Il apparaît donc que l'étude des mécanismes de refroidissement du manteau doit tenir compte des processus à court terme de la tectonique des plaques. Cependant, nous avons vu que la lithosphère est composée d'un nombre variable de plaques, et que si leur dynamique est relativement bien décrite, la mobilité de leurs frontières reste mal comprise. En particulier, nous ne pouvons que formuler des hypothèses à propos de leurs formation et disparition, à

1. $1 Ga = 1$ milliard d'années

2. $1 Ma = 1$ million d'années

défaut d'avoir pu les observer dans l'histoire récente. Cette remarque s'applique également aux continents. Si nous savons qu'il existe des cycles continentaux dans l'histoire de la Terre, conduisant à des ouvertures et des fermetures régulières des océans [Wilson, 1966], nous ne pouvons que présumer des mécanismes de ces événements.

Cependant, les outils de modélisation utilisés classiquement en géophysique ne nous permettent pas d'intégrer facilement de tels événements structuraux. C'est à ce problème que veut répondre le modèle MACMA en s'appuyant sur les atouts des systèmes multi-agents et particulièrement de l'approche phénoménologique.

2.2 Le Modèle MACMA : *MultiAgent Convective Mantle*

Les simulations numériques classiques de la géodynamique terrestre sont confrontées à un obstacle qui caractérise la plupart des systèmes complexes : bien que les processus mécaniques et thermiques impliqués dans le refroidissement de la planète soient relativement bien décrits à l'échelle du laboratoire, plusieurs processus majeurs de la tectonique des plaques restent mal compris à l'heure actuelle.

À ce constat s'ajoute le fait que nous ne sommes pas toujours en mesure d'exprimer nos hypothèses sous la forme d'une équation mathématique. Ceci constitue un obstacle notamment pour l'étude de l'histoire thermique de la Terre par des approches classiques.

Nous devons donc changer de point de vue concernant la simulation des systèmes complexes qui est classiquement réalisée à partir d'équations différentielles de plus en plus élaborées, résolues sur des maillages de plus en plus fins. Nous proposons une alternative qui consiste à simuler le système en superposant des modèles analytiques et phénoménologiques qui ajoutent aux équations classiques de conservation la connaissance empirique que nous avons du système.

D'un point de vue informatique, la mise en œuvre de cette démarche passe par l'emploi de systèmes multi-agents en adoptant une approche mixte. En effet, l'utilisation conjointe d'agents-entité, dotés de comportements individuels, et d'agents-interaction, plus à même d'exprimer la dynamique des phénomènes en jeu, nous offre la liberté de mêler différents paradigmes de modélisation.

Ainsi, les objectifs du modèle MACMA (*MultiAgent Convective Mantle*) sont au nombre de trois :

- étudier le couplage entre tectonique des plaques et convection mantellique, à partir de **mécanismes explicites qui peuvent être examinés indépendamment**,
- simuler une tectonique à la surface du manteau avec des frontières de plaques **mobiles**, et par conséquent, un nombre de plaques qui n'est **pas fixé a priori**,
- **superposer des lois analytiques, empiriques et phénoménologiques** pour décrire la dynamique terrestre en tenant compte des **phénomènes observés** qui restent pour l'heure mal compris.

Nous commencerons par donner une vue d'ensemble de l'environnement de simulation

et de ses entrées et sorties. Puis, nous détaillerons la réalisation de chacun des agents et la manière dont ils s'intègrent au reste de la simulation. Enfin, nous décrirons l'implémentation des lois empiriques qui permettent de combler les manques du modèle.

2.2.1 Vue d'ensemble

L'environnement de simulation de notre modèle est une vue en coupe d'une planète tellurique de mêmes dimensions que la Terre. Nous proposons une simulation en deux dimensions afin d'étudier au mieux les interactions constitutives de la dynamique des plaques et de leurs frontières, en s'affranchissant des effets géométriques touchant à la forme des plaques en surface. Le modèle repose sur une description géophysique et algorithmique à trois niveaux de complexité.

Le premier niveau décrit la structure interne du système. Le noyau de la planète est considéré comme un ensemble de conditions aux limites thermiques et mécaniques pour le manteau convectif.

Le second niveau vient paver la surface du manteau par des plaques lithosphériques rigides. Ces plaques sont soumises à des forces connues, que nous détaillerons par la suite, qui leur donnent une vitesse. Dans ce modèle, les plaques sont considérées comme les couches limites thermiques des cellules de convection du manteau, comme le montre la vue en coupe de la figure 2.4. Une plaque peut être composée de plusieurs sections qui sont parfois séparées par une suture symbolisant une discontinuité d'âge du plancher océanique ou un lien nouvellement créé entre une lithosphère océanique et une lithosphère continentale. Chaque plaque est délimitée par deux interfaces. Ces interfaces peuvent être de deux natures : soit une dorsale (création de matière au niveau des remontées chaudes des cellules), soit une subduction (disparition de matière au niveau des courants froids descendants). La géométrie cylindrique et les conditions aux limites périodiques pour des plaques inélastiques assurent la conservation de la matière dans le système.

Le troisième niveau de complexité superpose à cet ensemble des continents légers et rigides. Les continents sont fixés sur les plaques, et cette superposition permet de décrire des plaques comportant à la fois des sections océaniques et des sections continentales, à l'instar des plaques africaine, sud-américaine, ou encore australienne.

L'état initial de la simulation peut être choisi en détail par l'utilisateur. Il est en effet possible de choisir la valeur de tous les paramètres physiques du modèle, ainsi que la position et l'âge de chaque frontière de plaque et de chaque extrémité de continent. La distribution des âges dans l'état initial est ensuite calculée selon une progression linéaire des âges entre ces extrémités [Combes, 2011].

Le principe de fonctionnement de ce modèle bidimensionnel est le suivant. La température du manteau est mise à jour à chaque pas de temps en opposant les productions de chaleur et le flux d'énergie surfacique. Cette température sert à déterminer les nouvelles viscosités du manteau supérieur et du manteau inférieur, qui sont des paramètres influents dans le bilan des forces appliquées à chaque plaque. Chaque plaque est un agent qui calcule sa vitesse à partir du bilan des forces qui lui sont appliquées. Les vitesses de plaques sont ensuite utilisées

pour déterminer le mouvement des frontières, afin de calculer leur nouvelle position à chaque pas de temps. Cette nouvelle géométrie permet de mettre à jour la distribution des âges du plancher océanique, qui sera utilisée pour calculer le flux de chaleur en surface.

L'évolution quasi-statique du système tient aussi compte de modifications structurelles telles que l'initiation de subductions ou la création de dorsales. Ces transformations de la structure de la planète sont régies par des lois empiriques, et interviennent toujours au début de chaque pas de temps. La simulation du système est donc basée sur une alternance structure/dynamique : la structure est mise à jour tant qu'il y a des transformations à appliquer (par exemple, en cas de coïncidence de deux frontières de plaque sur la même position), puis les vitesses des agents mobiles sont calculées, avant de boucler le cycle par un bilan thermique du système. Cette succession d'étapes est illustrée par la figure 2.4.

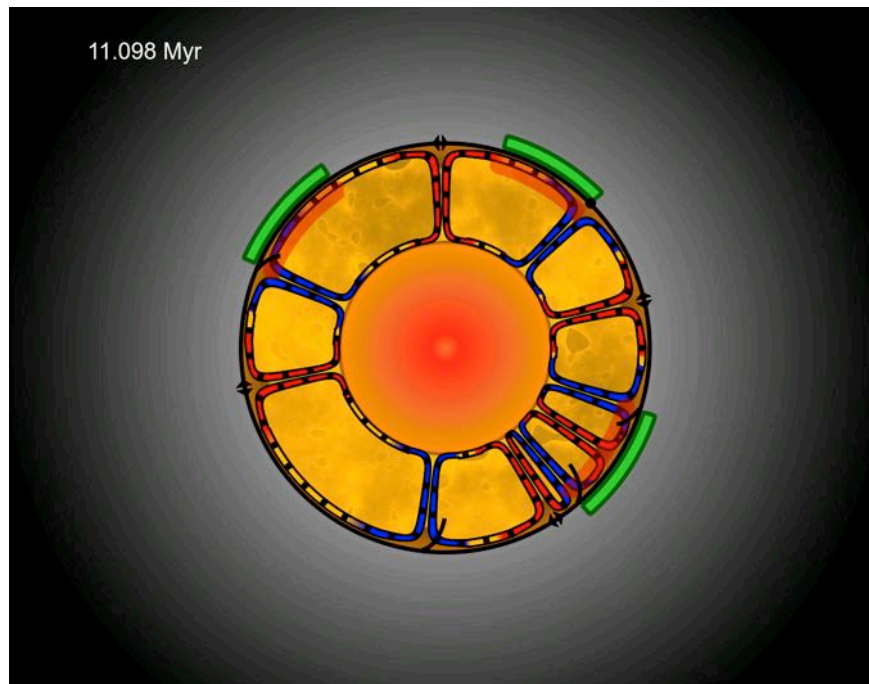


FIGURE 2.4 – Vue d'ensemble de l'environnement de simulation de MACMA – Ce modèle 2D présente un manteau pavé des cellules de convection (courants chauds ascendants en rouge, courants froids descendants en bleu). La surface du manteau est recouverte de plaques tectoniques (en noir), dont les sections continentales sont parfaitement isolantes et indéformables (en vert). Les frontières de plaques sont des dorsales océaniques (triangles noirs) ou des zones de subduction (plaques plongeantes noires). Les sutures de plaques sont représentées par des agrafes (pastilles noires), dont un exemple est visible à l'extrémité d'un continent situé en haut à droite de l'image. Des zones sous-continentales (en orange) indiquent une élévation superficielle de la température. Le temps écoulé depuis le départ de la simulation est indiqué en haut à gauche de l'écran (1 *Myr* = 1 million d'années).

Cette description générale du modèle met en évidence la diversité des modèles que nous souhaitons associer, ce qui constitue généralement un frein pour la modélisation des systèmes complexes. Dans le chapitre 1, nous avons vu que l'approche multi-agents peut être déclinée pour s'adapter au point de vue de modélisation. Nous allons à présent montrer que nous pouvons tirer avantage de l'autonomie de conception que nous offre l'approche multi-agents

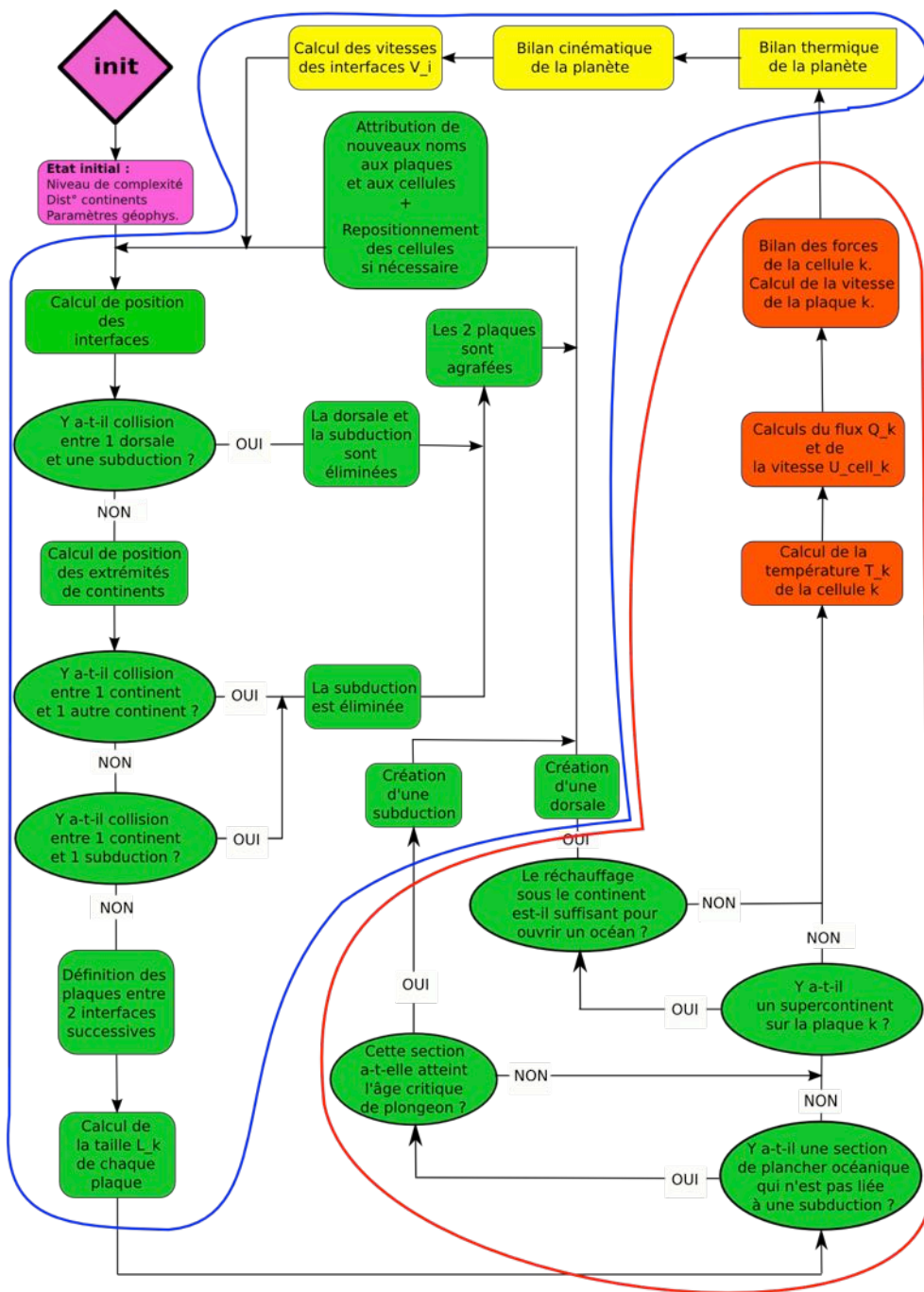


FIGURE 2.5 – **Schéma logique de l'algorithme d'évolution du système** – La simulation démarre sur un état initial choisi par l'utilisateur. Ensuite, chaque pas de temps est divisé en trois étapes. La première étape (en vert) consiste à mettre à jour la structure globale du système (création/disparition des interfaces, définition des plaques). La seconde (en orange) permet de calculer la nouvelle dynamique des plaques (bilan des forces). La dernière (en jaune) dresse les bilans thermique et cinématique du système pour calculer le déplacement des interfaces.

et combiner ces déclinaisons au sein d'une approche mixte.

2.2.2 Agents-interaction pour le bilan des forces

Le modèle MACMA propose de déterminer la vitesse d'une plaque en tenant compte des forces présentées dans la section 2.1. Les forces motrices sont : la résultante des forces de pression au niveau d'une dorsale (*ridge push RP*), la traction gravitaire de la partie plongeante d'une plaque (*slab pull SP*, dans le manteau supérieur uniquement), et la force de succion dérivant du poids de la plaque plongeante dans le manteau inférieur (*slab suction SS*). Les forces résistives sont : le frottement fluide dérivant du couplage mécanique entre le manteau et la plaque horizontale (*mantle drag MD*), le couplage entre le manteau et la plaque plongeante (*viscous drag on slab VS*, pour sa partie supérieure uniquement), et la force illustrant la dissipation visqueuse à l'intérieur de la plaque, lors de son pliage au niveau des zones de subduction (*bending force B*). En négligeant l'inertie du système, à l'échelle d'une plaque lithosphérique, on peut écrire que la somme des forces exercées sur une plaque est nulle à chaque instant. L'équilibre de ces forces projeté horizontalement pour chaque plaque s'écrit donc :

$$RP + SP + B + MD + VS + SS = 0 \quad (2.1)$$

Nous remarquons qu'une force désigne, en physique, l'interaction entre deux corps qui induit une modification du vecteur vitesse de l'un d'eux. Dans ce contexte, l'utilisation d'agents-interaction pour la modélisation des forces appliquées à une plaque nous semble clairement justifiée. De cette manière chaque phénomène est considéré indépendamment des autres, ce qui confère au moins deux avantages au modèle MACMA.

Tout d'abord, en phase de conception, il est possible de modifier l'expression d'une force, ou d'en prendre en compte une nouvelle, sans remettre en cause l'implémentation du reste du modèle.

Ensuite, en phase de simulation, la décomposition phénoménologique permet d'adapter le calcul de la vitesse aux conditions aux limites de la plaque que forment ses interfaces. En effet, une plaque bordée de deux subductions n'est pas affectée par le *RP*, de même qu'une plaque qui n'a pas encore plongé n'est pas tractée par le *SP*. L'approche multi-agents permet tout particulièrement de prendre en compte des changements de conditions aux limites dynamiquement, c'est-à-dire en cours de simulation, ce qui représente généralement un obstacle pour les méthodes classiques de modélisation en géophysique.

À chaque pas de temps, chaque agent-interaction calcule sa contribution au bilan des forces pour une plaque donnée. Sachant que les trois forces résistives dépendent de la vitesse de la plaque en régime permanent, l'application du principe fondamental de la dynamique nous permet de calculer la vitesse U de la plaque à chaque pas de temps, comme le montre l'équation (2.2).

$$U = \frac{RP + SP + SS}{B_U + MD_U + VS_U} \quad (2.2)$$

Cette vitesse est ensuite utilisée pour définir les déplacements des différents éléments du système.

2.2.3 Agents-entité pour la cinématique

Si l'approche interactions-centrée offre des avantages indéniables pour la modélisation de la dynamique d'un système, certains comportements individuels restent plus simples à exprimer au moyen d'agents-entité. En particulier, le point de vue individus-centré est souvent le plus pertinent lorsqu'il s'agit de régir des contraintes spatiales.

Pour ce qui est de MACMA, plusieurs éléments structurels peuvent être modélisés par des entités autonomes : les plaques, les interfaces et les continents. Ces éléments disposent de comportements individuels qui déterminent leur cinématique, mais également leurs potentiels changements d'état qui découlent de leurs déplacements.

Nous nous intéressons dans un premier temps aux lois de mouvement de chacun des éléments. Nous avons vu que les vitesses de plaque résultent du bilan des forces calculé par les agents-interaction. Ces vitesses servent de base pour calculer les déplacements de chacun des éléments du système. Par exemple, du fait de leur faible densité, les continents sont astreints à se déplacer à la vitesse de la plaque sur laquelle ils sont fixés. En revanche, en ce qui concerne les interfaces, la seule connaissance des vitesses de plaques n'est pas suffisante. En effet, ces mouvements demeurent actuellement mal compris, et l'absence d'expression analytique pour les décrire empêche de les inclure adéquatement dans les approches numériques classiques. Ainsi, pour rendre compte de la diversité des comportements observés sur Terre, nous devons recourir à un ensemble de lois empiriques que nous implémentons au sein de chaque agent.

Migration de la fosse de subduction

Le déplacement des zones de subduction est un problème qui résiste depuis longtemps aux tentatives de modélisation de la tectonique des plaques [Heuret et Lallemand, 2005]. La relation entre le mouvement de la fosse et le régime compressif ou extensif de la plaque supérieure n'est pas explicite à l'heure actuelle. C'est pourquoi le modèle MACMA utilise un mécanisme simple qui rend compte de la plupart des phénomènes observés, en se basant sur un critère structurel pour déterminer le mouvement de la fosse de subduction.

Ainsi, à chaque pas de temps, un agent-entité « subduction » détermine sa vitesse en analysant son voisinage immédiat, à savoir les deux plaques qu'il délimite.

À l'initiation d'une subduction, la fosse se déplace à une vitesse V_t égale à la vitesse de la plaque supérieure V_{up} , et cela reste vrai de manière générale pour la suite. Un tel régime est dit compressif ou neutre [Lallemand et al., 2008] car la plaque plongeante reste au contact de la plaque supérieure.

Ce comportement n'est modifié qu'en cas d'apparition d'un régime extensif, c'est-à-dire quand la plaque supérieure présente une vitesse qui l'éloigne rapidement de la zone de subduction concernée. Nous avons choisi pour cela de nous appuyer sur un critère structurel

présenté sur la figure 2.6. Ce critère conduit à l'apparition d'une dorsale océanique devant la fosse de subduction quand le plongeon d'une plaque implique que deux fosses se succèdent. Si cette règle n'est pas dénuée de logique scientifique, elle semble moins correspondre à la réalité de terrain que les lois du régime compressif. Elle intervient surtout ici en tant que garde-fou pour garantir l'alternance dorsale/subduction nécessaire au reste du modèle.

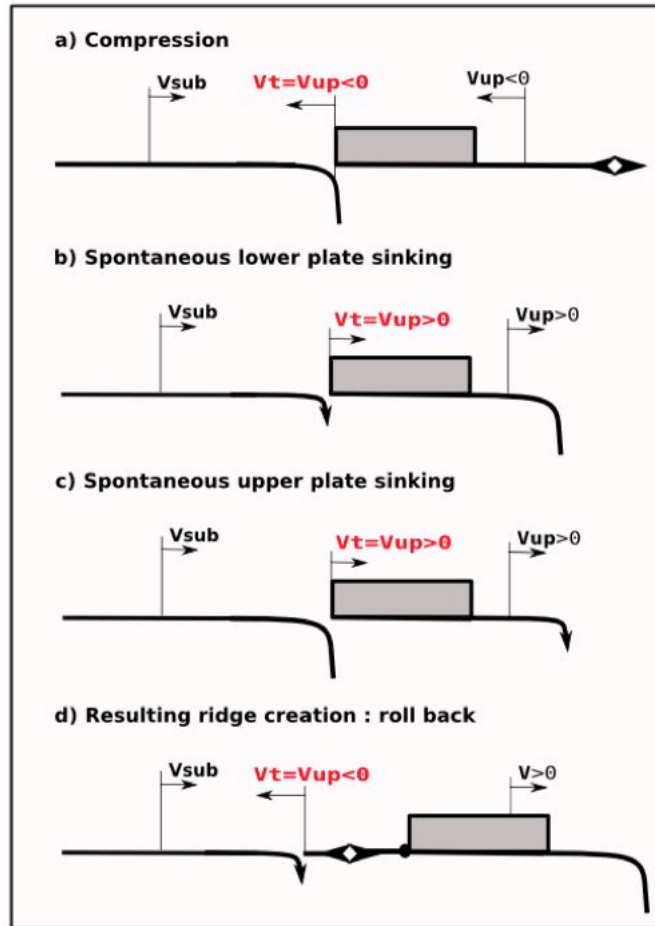


FIGURE 2.6 – **Comportement des fosses de subduction basée sur un critère structurel** – Les vitesses sont comptées positivement dans le sens de déplacement de la plaque plongeante. a) Si la plaque supérieure ne possède pas de *slab*, le régime est considéré comme compressif, et la migration de la fosse est prescrite par le mouvement de la plaque supérieure : $V_t = V_{up}$. Si au contraire la plaque supérieure est entraînée par un *slab pull*, elle s'éloigne rapidement de la fosse et le régime est considéré comme extensif. Ceci a lieu en cas d'initiation d'une subduction (b) pour la plaque inférieure ou (c) pour la plaque supérieure. Dans les deux cas, une dorsale océanique est créée devant la fosse qui sépare les deux plaques, et le plancher nouvellement formé est suturé sur l'ancienne plaque supérieure, tandis que de l'autre côté, la plaque plongeante part en roll-back suivant $V_t = V_{up}$ (d)

Production de plancher océanique

Pour ce qui est des dorsales, le modèle de déplacement est plus simple. Nous faisons l'hypothèse que la production de plancher océanique au niveau des dorsales est symétrique.

Si l'on considère la dorsale océanique séparant les plaques A et B de vitesses respectives U_A et U_B dans un référentiel absolu, alors le taux de production de plancher océanique total s'écrit simplement pour cette dorsale $U_B - U_A$, en comptant positivement les mouvements dans le sens de déplacement de B . L'hypothèse d'accrétion asymétrique implique que le taux de production pour chaque plaque s'écrit en fonction de la vitesse V_r de l'axe de la dorsale :

$$\frac{U_B - U_A}{2} = V_r - U_A = U_B - V_r \quad (2.3)$$

ce qui conduit à une expression simple de la vitesse de migration des agents « dorsale » :

$$V_r = \frac{U_A + U_B}{2} \quad (2.4)$$

2.2.3.1 Modifications structurelles

Les agents-entité, que sont les frontières de plaques et les continents, se déplacent à une vitesse qui varie selon leur type. Ces différences de vitesses conduisent inévitablement les agents à se rencontrer, étant donné le contexte géométriquement borné imposé par MACMA. Il est donc nécessaire de définir les comportements à adopter dans ces situations, afin de garantir la viabilité de la structure spatiale du système.

Encore une fois, il s'agit ici de proposer un ensemble de lois empiriques qui permettent de reproduire les observations de terrain, même si dans le cas présent, les échelles de temps des phénomènes considérés font qu'aucun d'entre eux n'a pu être observé en conditions réelles dans l'histoire moderne. Les lois proposées dans MACMA se composent donc plus d'hypothèses que de certitudes scientifiques. Néanmoins, nous verrons que la méthode multi-agents combinée à cette approche phénoménologique permet, si ce n'est de valider ces hypothèses, au moins de les étayer par des résultats comparables aux observables disponibles.

Plusieurs cas de figures peuvent se présenter et nous poussent à envisager des modifications structurelles :

- si deux frontières de plaques (une dorsale et une zone de subduction) coïncident à la même position,
- si une agrafe se retrouve à la même position qu'une zone de subduction,
- si le bord d'un continent se retrouve à la même position qu'une zone de subduction,
- si le bord d'un continent se retrouve à la même position qu'un autre bord de continent.

Dans chaque cas, l'algorithme applique une ou plusieurs lois incluses dans le modèle géophysique, afin que le système retrouve un état viable avec un pavage en surface bien délimité. Après l'application d'une règle, l'algorithme recommence une nouvelle passe en partant de l'origine, et ainsi de suite tant qu'il y a des règles à appliquer. Les règles exercées dans chacun des cas cités plus haut sont exposées ci-après.

Collisions d'interfaces

Tout d'abord, il faut noter que seules une dorsale et une subduction peuvent être amenées à entrer en collision. En effet, deux dorsales ne peuvent que se repousser mutuellement par la production de plancher océanique et deux subductions ne peuvent se succéder, hormis dans le cas d'une plaque entièrement continentale. De cette manière, nous pouvons garantir une cohérence des températures pour deux cellules de convection juxtaposées.

Pour le premier cas, la dorsale est subduite, et la plaque plongeante se détache : ces deux agents sont éliminés et remplacés par un agent « agrafe ». Ce dernier marque la discontinuité d'âge du plancher océanique et devient une zone passive qui pourra entrer en subduction dans la suite de la simulation, sous certaines conditions que nous détaillerons plus loin. Ce comportement est illustré par la figure 2.7.

Une agrafe adopte la même loi de vitesse qu'un continent, à savoir qu'elle se place à la vitesse de la plaque sur laquelle elle est fixée. De ce fait, poussée par une dorsale, elle peut être subduite avant d'entrer en subduction et ainsi faire disparaître la discontinuité d'âge du plancher avec elle.

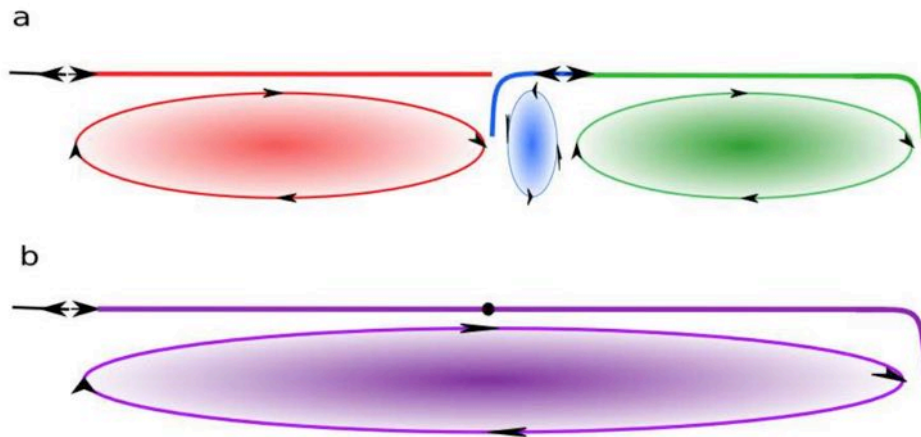


FIGURE 2.7 – **Suture de plaques due à la subduction d'une dorsale** – (a) L'axe de la dorsale migre vers la fosse de subduction à la vitesse V_r égale à la demi-somme des vitesses des plaques verte et bleue. (b) Si aucune ne peut plonger, les deux plaques deviennent solidaires, et une agrafe vient marquer la discontinuité d'âge du plancher océanique.

Continents « insubmersibles »

Du fait de leur faible densité, les continents sont astreints à se déplacer à la vitesse de la plaque sur laquelle ils sont fixés. Pour autant, il ne leur est pas possible de plonger dans le manteau. Leur migration les mène toujours à l'extrémité de la plaque et, par conséquent, à une subduction. Quand le continent arrive en butée de plaque, le *slab* se détache et le bord du continent coïncide alors avec la frontière de la plaque, comme l'illustre la figure 2.8. Les deux plaques sont alors suturées par une agrafe, sans modifier la distribution des âges du plancher.

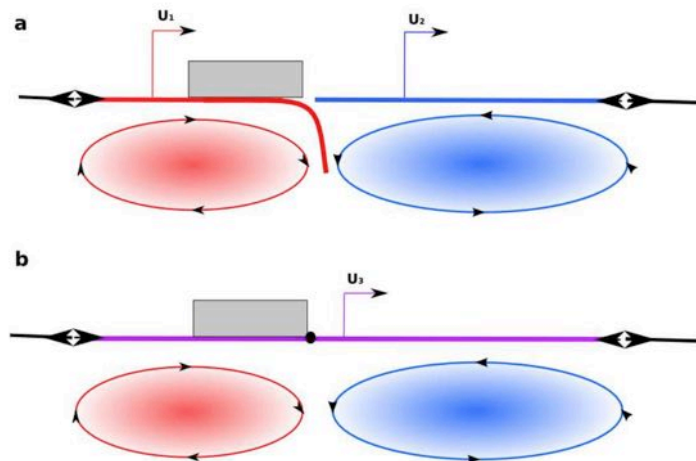


FIGURE 2.8 – **Arrivée d'un continent en butée de plaque** – (a) Quand la totalité de la section océanique précédant un continent a été subduite, le *slab* se détache sous l'action de son propre poids, et la traction gravitaire associée disparaît. (b) Ensuite, soit la plaque adjacente est suffisamment épaisse pour plonger (non représenté), soit les plaques sont suturées (cas illustré).

Formation de supercontinents

Deux continents sur des plaques adjacentes peuvent être conduits l'un vers l'autre si celles-ci sont délimitées par une subduction. Cette situation est présentée sur la figure 2.9. Nous assistons alors à une collision continentale qui donne naissance à un supercontinent en détruisant les deux précédents agents « continent » et s'accompagne de la fusion des deux plaques.

2.2.4 Agents-énaction pour les phénomènes thermiques

Nous avons décrit jusqu'ici les aspects dynamique et cinématique de MACMA qui permettent de simuler une tectonique des plaques réaliste au vu des données de terrain. Ainsi, nous sommes en mesure d'obtenir une distribution des âges du plancher océanique comparable à celle de la Terre, ce qui est absolument fondamental pour l'étude des phénomènes thermiques qui prennent place dans le manteau. En effet, nous avons vu dans la section 2.1 que le flux de chaleur en surface est étroitement lié à la distribution des âges du plancher océanique.

Nous allons à présent détailler le modèle thermique mis en œuvre pour le calcul de la température du manteau au cours du temps. Ces variations de température induisent des phénomènes qui ont un impact rétroactif sur la tectonique des plaques. Pour décrire ces phénomènes, qui mêlent lois analytiques et empiriques, nous choisissons d'utiliser des agents-énaction tels que nous les avons dépeints dans le chapitre 1. Ceux-ci sont particulièrement adaptés dans le cadre de MACMA au regard de l'approche phénoménologique que nous avons retenue.

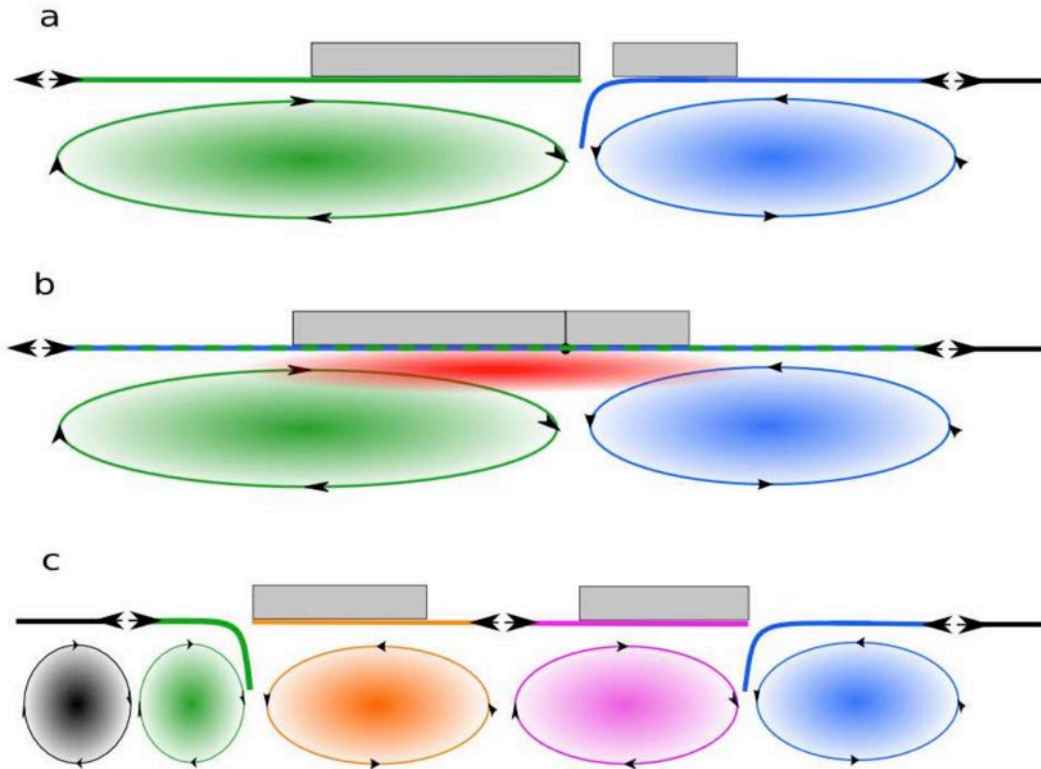


FIGURE 2.9 – **Formation d'un supercontinent suivi d'un processus d'océanisation**
 – (a) La traction gravitaire exercée par une plaque plongeante fait converger les deux continents qui l'entourent. (b) La collision continentale s'accompagne de la suture des deux plaques, et si le nouveau continent formé est de taille suffisante, une zone d'advection divergente est créée sous ce dernier. (c) La contrainte cisailante créée sous la plaque peut dépasser la valeur critique de rupture de la lithosphère continentale, menant à l'ouverture d'un nouvel océan.

Calcul du flux de chaleur

Nous avons vu précédemment que la variation de la température moyenne du manteau s'explique par un déséquilibre entre les pertes de chaleur en surface et la production de chaleur d'origine radioactive à l'intérieur dans le volume du manteau [Jaupart et Mareschal, 2011]. Pour autant, les modèles classiques peinent à remonter l'histoire thermique de la Terre car il leur est très difficile de prendre en compte la dynamique du système. Notamment, leurs problèmes résident dans la gestion des conditions aux limites (nombre d'éléments variable et non fixé *a priori*) et dans l'impossibilité d'exprimer de façon formelle certains processus.

L'ensemble des lois que nous avons décrit jusqu'à présent vise à prendre en compte les évolutions structurales de la lithosphère soumise à une tectonique des plaques, elle-même fortement dépendante de la viscosité du manteau et donc de sa température.

La production de chaleur radioactive étant supposée connue, il reste à évaluer la perte de chaleur en surface. Pour ce faire, nous introduisons un agent-énaction « fuite thermique » qui est en charge de calculer le flux surfacique moyen de chaque tronçon de lithosphère océanique. En effet, les continents sont ici considérés comme parfaitement isolants.

Des balises sont placées au niveau de chaque limite de tronçon océanique pour mesurer l'âge de la lithosphère en ces points, car nous avons vu que son épaisseur augmente au fur et à mesure qu'elle vieillit. Les âges sont mis à jour à chaque pas de temps en fonction des migrations de chaque interface. L'agent calcule ensuite la nouvelle température du manteau en opposant la somme des flux surfaciques à la production de chaleur radioactive.

Initiation d'une subduction

Le plongeon spontané d'une plaque océanique n'a pas été observé dans les reconstructions récentes de la tectonique des plaques [Becker et Faccenna, 2009]. Ce processus reste mal compris, mais il semble néanmoins qu'une contrainte verticale suffisamment forte puisse déséquilibrer la flottabilité d'une plaque sur le manteau et initier une subduction [Nikolaeva et al., 2010]. Cette contrainte croît au cours du temps en raison de l'épaississement de la lithosphère produite par les dorsales en vieillissant. Par ailleurs, toutes les études de paramétrisation du flux de chaleur se basent sur l'âge maximal du plancher océanique [Labrosse et Jaupart, 2007], ce qui implique qu'au delà de cet âge le plancher doit disparaître. Il est donc nécessaire d'ajouter à notre modèle un processus d'initiation de la subduction pour donner corps à ces hypothèses. Nous proposons ici un mécanisme simple qui dépend à la fois de la température du manteau et de l'âge actuel supposé du plongeon spontané de la lithosphère océanique.

Nous considérons ici qu'au niveau des marges passives, la lithosphère océanique peut continuer de s'épaissir, en supposant la présence d'un continent ou d'une suture de plaques marquant une discontinuité d'âge du plancher océanique. Sous cette hypothèse forte, la lithosphère océanique peut atteindre sa valeur critique de déstabilisation et commencer à plonger.

À l'heure actuelle, le plancher océanique le plus vieux observé sur Terre, au niveau des fosses de subduction, atteint 180 Ma. Ce constat nous permet de supposer que l'âge critique

de plongeon spontané est actuellement proche de cette valeur. En réalité, en considérant que la contrainte seuil est constante au cours de la phase tectonique de l'histoire thermique de la Terre, nous obtenons une expression de l'âge critique de plongeon dépendant de la température du manteau.

Les marges passives susceptibles de plonger sont identifiées : ce sont les bords de continents et les agraes. Nous ajoutons donc un agent « plongeon spontané » en charge de placer des balises en ces points pour mesurer l'âge de la lithosphère et calculer la contrainte verticale qui y est exercée. Si cette contrainte dépasse la valeur seuil, l'agent met à jour la structure et crée une nouvelle subduction en lieu et place de la balise.

Ouverture d'un océan

Le modèle MACMA combine jusqu'ici des processus de subduction des dorsales, de suppression des zones de convergence (suture de plaques), et de création des fosses de subduction. Pour que le système présente une activité tectonique sur le long terme, il est nécessaire d'ajouter à cette liste un processus de création de dorsales pour garantir le renouvellement du plancher océanique.

Un tel processus est déjà présent sous certaines conditions dans le mécanisme proposé pour la migration des zones de subduction. Cependant, ce mécanisme correspond sans doute moins à la dynamique terrestre que le reste du modèle et vient surtout combler les lacunes du modèle de migration. Ainsi, afin de rendre compte des cycles continentaux [Wilson, 1966], nous allons maintenant proposer un mécanisme semi-empirique d'océanisation.

Le principe en est le suivant : les continents étant considérés comme des isolants thermiques du fait notamment de leur forte production d'énergie radioactive [Jaupart et al., 2007], une couche superficielle sous-continentale, de largeur a et d'épaisseur h , est réchauffée par la production d'énergie radioactive *in situ* qui ne peut être évacuée par le haut. Il s'ensuit une advection latérale de part et d'autre du continent isolant, créant ainsi un cisaillement à la base de la lithosphère qui augmente avec l'élévation progressive de la température. Ce phénomène tend à créer une nouvelle marge active, une dorsale, au milieu du continent. Nous disposons d'ordres de grandeurs sur le temps de vie des supercontinents qui ont existé dans l'histoire de la tectonique des plaques (entre 100 et 200 Ma selon [Li et al., 2008]), ce qui nous permet de contraindre la valeur des paramètres *a priori* libres du modèle : la force de cohésion F_{lim} de la lithosphère continentale et l'épaisseur h de la couche d'advection. La figure 2.10 illustre la paramétrisation du schéma d'advection sous le couvercle isolant.

Notre démarche vise à mettre en place un agent-énaction « brisure de continents » en charge d'évaluer le cisaillement sous-continentale à chaque pas de temps. Cet agent pose des balises sous les continents qui ont une taille a supérieure à celle de l'Océanie, car la zone d'advection ne peut se former sous un continent que s'il est assez grand. Nous faisons de plus l'hypothèse qu'une océanisation ne peut avoir lieu que si le supercontinent est bordé de zones de subduction de chaque côté car dans le cas contraire, l'action d'un *ridge push* supplémentaire (si le continent était bordé par au moins une dorsale) augmenterait le seuil de contrainte de telle façon que la rupture continentale serait inaccessible par ce mécanisme dans des temps acceptables géologiquement. Quand ce cisaillement atteint la valeur critique de rupture de la lithosphère continentale, l'agent-énaction brise le continent en deux sous-continentes pour

laisser une nouvelle dorsale se mettre en place. Cela se traduit par la suppression de l’agent-entité « super-continent » et par la création de deux nouveaux continents qui vont s’éloigner l’un de l’autre sur leurs nouvelles plaques respectives, en raison du nouveau bilan des forces auquel va s’ajouter le *RP* de la dorsale.

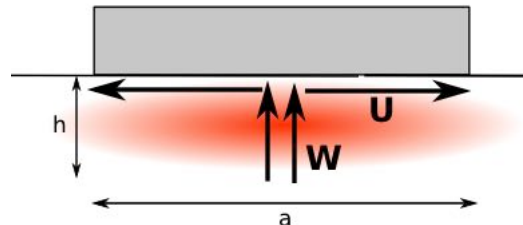


FIGURE 2.10 – **Paramétrisation du schéma d’advection sous un continent isolant**

– Le continent est ici totalement isolant, ne laissant pas s’évacuer par le haut la chaleur apportée par la désintégration radioactive. Le fluide voit donc sa température augmenter sur une épaisseur h et une largeur a , entraînant une remontée du fluide à la vitesse W et une évacuation latérale à la vitesse U .

La figure précédente 2.9 présente une collision continentale suivie d’une élévation superficielle de la température, et d’une océanisation par rupture de la lithosphère continentale.

2.2.5 Bilan

Nous avons présenté dans cette section un modèle géophysique de refroidissement du manteau terrestre, basé sur une approche multi-agents mixte. La superposition des comportements des agents-entité, interaction et éraction permet de décrire le couplage entre convection mantellique et tectonique des plaques, en tenant compte d’une grande partie de sa complexité. En effet, le modèle MACMA calcule la vitesse de chaque plaque en fonction de ses caractéristiques individuelles (âge, taille, histoire) mais aussi des caractéristiques de son environnement immédiat (viscosité du manteau, nature des frontières de plaques). Ces vitesses sont utilisées pour déterminer le mouvement des limites de plaques : la tectonique qui émerge des processus locaux de notre modèle présente donc un nombre de plaques qui n’est pas fixé *a priori*. D’autre part, notre approche permet de créer ou de supprimer des frontières de plaques, en s’appuyant sur des lois empiriques qui rendent compte d’un certain nombre de données de terrain, ce qui n’est pas réalisable pour des simulations numériques classiques. Les processus empiriques proposés permettent de décrire notamment l’initiation des subductions, la migration des fosses de subduction et l’ouverture de nouveaux océans par rupture de la lithosphère continentale. Si ces lois sont parfois basées sur des hypothèses fortes, et si elles ne décrivent pas la totalité des phénomènes observés sur Terre, elles ont néanmoins l’avantage d’être cohérentes avec les échelles caractéristiques des mécanismes observés à la surface de la planète. De plus, le découpage phénoménologique du modèle permet de modifier aisément les lois empiriques que nous venons de citer, ainsi que leurs paramètres. C’est cet aspect du modèle que nous allons maintenant illustrer en nous appuyant sur deux expérimentations *in vitro*.

2.3 Expérimentations *in virtuo*

Dans cette section, nous démontrons l'intérêt de l'approche *in virtuo* pour la simulation de systèmes complexes, en nous concentrant sur l'étude de l'influence de la modification locale de deux paramètres du modèle. Le lecteur intéressé trouvera d'autres exemples d'expérimentations *in virtuo* et d'analyse de l'impact des choix de modélisation dans les publications dédiées [Combes, 2011] [Combes et al., 2012].

2.3.1 Influence des lois empiriques

Dans la mesure où le choix d'une loi empirique par rapport à une autre n'est parfois pas évident dans le contexte, il convient de tirer avantage de la méthode multi-agents sur laquelle se fonde le modèle MACMA pour procéder à un ensemble de tests comparatifs.

Deux lois empiriques décrivant le mouvement des zones de convergence du plancher océanique ont été testées : la première repose sur un critère structurel dont l'implémentation a été décrite dans la section précédente, tandis que la seconde repose sur un critère cinématique.

Le critère cinématique vise à déterminer le régime extensif ou compressif d'une plaque à partir de la vitesse de la plaque qui fait face à la subduction. La figure 2.11 présente les différents cas considérés par l'agent-entité « subduction » pour déterminer sa vitesse de migration.

La distinction avec le critère structurel intervient lorsque la plaque supérieure est déjà bordée par une autre subduction. En effet, dans cette situation, la vitesse V_{up} est nulle et les deux subductions deviennent statiques. Notons que ce critère cinématique est physiquement plus légitime que le critère structurel : il tient notamment compte du taux de subduction pour déterminer le régime de déformation. Par ailleurs, il correspond davantage aux relations cinématiques proposées à partir d'études statistiques ou d'expériences analogues en laboratoire.

La figure 2.12 illustre l'impact des deux lois envisagées pour rendre compte du mouvement des subductions. Nous constatons que les deux lois fonctionnent à court terme, mais au bout de quelques centaines de millions d'années, le critère cinématique débouche sur une interruption de l'activité tectonique, due à la disparition de toutes les dorsales, sans que ces dernières soient remplacées. Ce résultat pose le problème du mouvement des frontières de plaques à l'échelle globale : pour que la tectonique reste active au cours de l'histoire thermique, il doit exister suffisamment de créations de zones d'accrétion pour compenser les disparitions de dorsales, sans quoi le système convergerait vers une tectonique intermittente [Silver et Behn, 2008].

Si cette expérimentation ne valide aucunement la loi empirique avec un critère structurel, elle a le mérite de révéler les défauts du critère cinématique au regard du reste du modèle et nous permet de l'écarter.

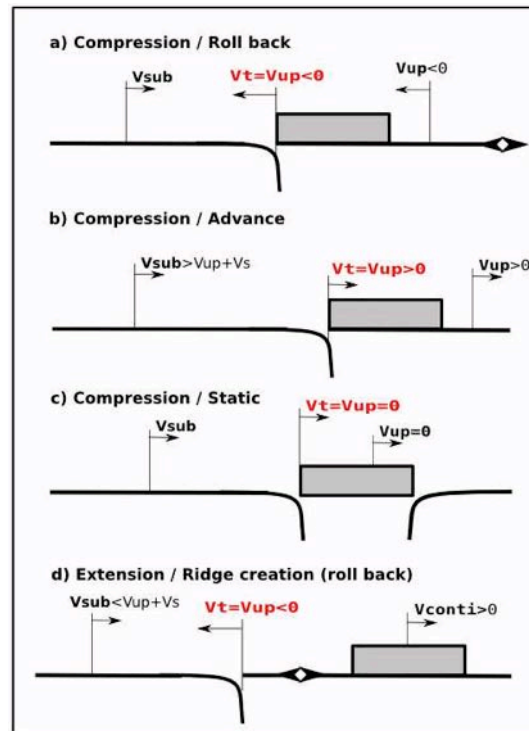


FIGURE 2.11 – **Comportement des fosses de subduction basé sur un critère cinématique** – Les vitesses sont comptées positivement dans le sens des subductions. Pour un taux de déformation négatif ($V_d < 0$), le régime est compressif et la fosse peut être soit en retrait (*roll-back*, *a*), soit en avance ($V_{sub} > V_{up} + V_s$, *b*), soit en position statique (continent bordé de zones de subduction, *c*). Pour $V_d > 0$, le régime est extensif et une dorsale est créée entre la fosse et la plaque supérieure (*d*). Cette dernière se déplace alors vers la fosse de subduction, qui recule en état de *roll-back* ($V_t = V_{up}$).

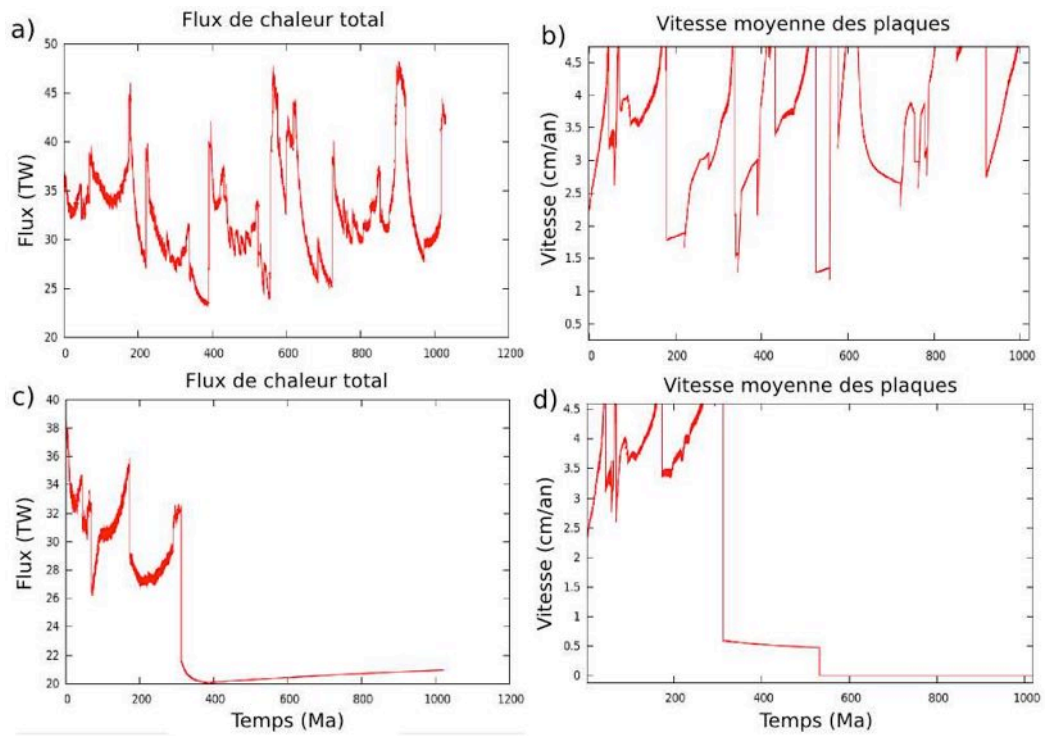


FIGURE 2.12 – **Comparaison de l'impact des lois de migration des fosses de subduction** – a) Évolution du flux de chaleur pour un critère structural. b) Évolution correspondante de la vitesse moyenne de plaques. c) Évolution du flux de chaleur pour un critère cinématique. d) Évolution correspondante de la vitesse moyenne des plaques, illustrant une interruption manifeste de l'activité tectonique.

2.3.2 Influence des paramètres

Le plus souvent, les modèles sont construits pour servir de support à des études paramétriques. Une étude paramétrique est une suite de tests du modèle pour lesquels on fait varier un ou plusieurs paramètres, comme par exemple des propriétés de matériaux, des températures, etc. Dans le cas d'un système complexe, du fait des nombreux liens directs ou indirects entre les phénomènes, nous savons que la variation, même minime, d'un paramètre peut entraîner des résultats de simulation complètement différents. C'est ce que nous avons voulu démontrer par cette deuxième expérimentation.

Entre autres paramètres, nous nous sommes intéressés à la valeur du seuil de rupture continentale F_{lim} utilisée par l'agent-énaction « brisure de continents ». En effet, cet agent permet de garantir une tectonique active à la surface du manteau par la création de dorsales et l'ouverture d'océans. Pour autant, ce mécanisme doit être paramétré avec précaution, car il contrôle par conséquent le taux de renouvellement de plancher océanique.

La figure 2.13 présente les résultats de trois tests avec des valeurs de F_{lim} différentes. Nous remarquons que le comportement du système semble similaire au début des simulations mais qu'à partir de 2 Ga années, les cycles continentaux se raccourcissent et que la taille moyenne des continents diminue. Ce seuil temporel s'explique par le fait que l'agent-énaction prend aussi en compte les propriétés du manteau, telles que sa viscosité et sa température, pour décider de briser un continent.

Cette expérience met en évidence la nature de l'impact de F_{lim} sur le comportement thermique du système : quand le seuil de rupture continentale diminue, l'ouverture des océans est facilitée, ce qui rend l'apparition des nouvelles dorsales plus facile et plus fréquente. Ces événements favorisent grandement la production de plancher océanique jeune, qui va de pair avec un flux de chaleur important et un refroidissement efficace.

Nous avons montré ici l'intérêt de l'expérimentation *in virtuo* de maquettes numériques, notamment à des fins d'études paramétriques qui révèlent les intrications et les liens de causalité entre les phénomènes. Il faut cependant garder à l'esprit que comme toutes les études paramétriques, celle que nous avons présentée ici n'est évidemment valable que dans le cadre de notre modèle, car elle dépend dans une large mesure des lois empiriques utilisées pour décrire les processus d'océanisation et de déstabilisation de la lithosphère océanique.

2.4 Synthèse

L'élaboration du modèle MACMA a été initiée à partir du constat suivant : nous ne savons pas tout écrire sous forme mathématique. En effet, la simulation des mécanismes de refroidissement terrestre se heurte actuellement à l'incapacité d'inclure, dans les équations de la convection, des termes qui rendraient compte des événements tectoniques tels que la création et la disparition des dorsales et des zones de subduction.

Pour tenter de contourner cet obstacle, nous avons contribué au développement d'un modèle original de la géophysique terrestre. Nous avons adopté une approche multi-agents

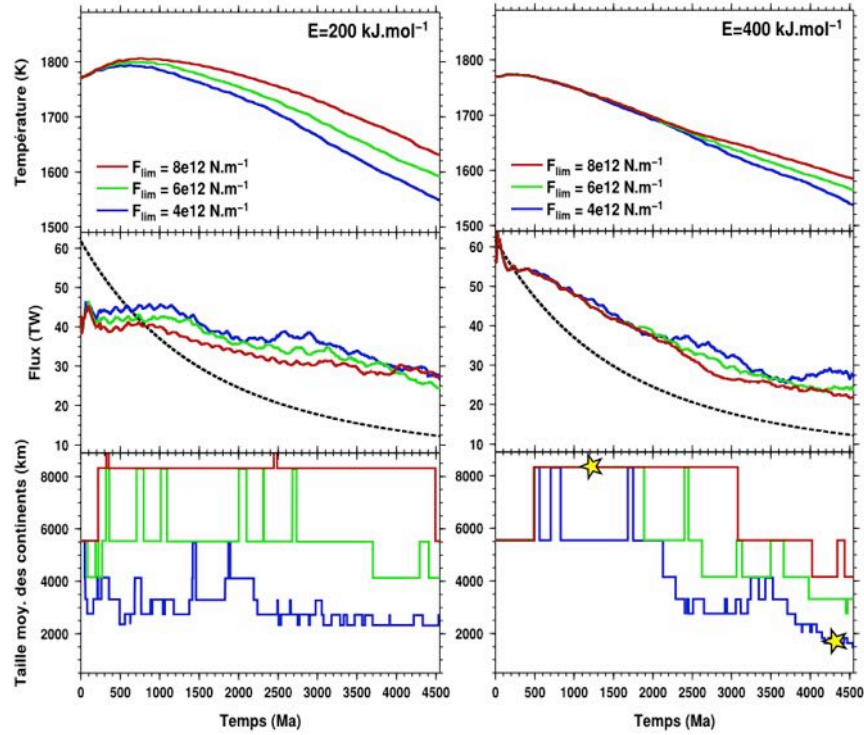


FIGURE 2.13 – **Étude de l'impact du seuil de rupture continentale** – Évolution de la taille moyenne des continents, du flux de chaleur et de la température du manteau en fonction du seuil de rupture continentale F_{lim} . Les étoiles placées sur les courbes de taille moyenne des continents correspondent aux configurations présentées sur la figure 2.14.

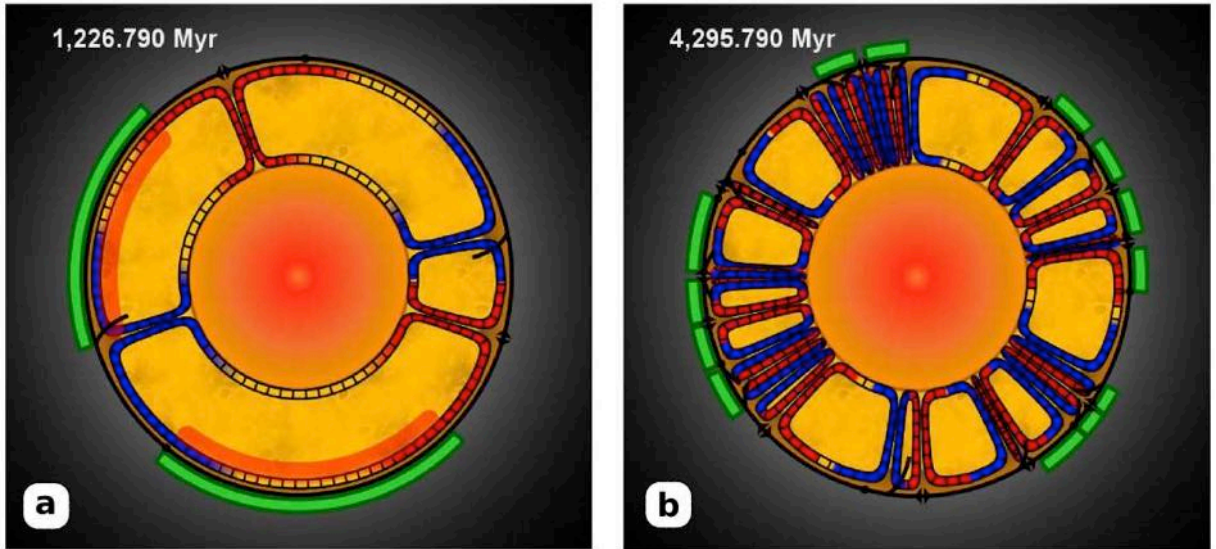


FIGURE 2.14 – **Comparaison de configurations continentales à la surface du manteau** – a) un régime de continents larges pour une viscosité faible et un seuil F_{lim} élevé, et b) un régime de continents étroits pour une viscosité élevée et un seuil F_{lim} faible.

mixte qui permet de superposer des modèles analytiques et empiriques pour ajouter aux équations classiques de conservation, des lois de comportement qui permettent de décrire un certain nombre de processus tectoniques qui restent aujourd'hui mal compris. Le système est alors décomposé en phénomènes répartis entre agents de différentes natures, d'ordinaire auto-exclusives. Nous utilisons ici les atouts de chaque type d'agents : les agents-interaction sont les plus à même de décrire la dynamique du système en calculant le bilan des forces, tandis que les comportement individuels, tels que les lois de migration, sont plus efficacement mis en œuvre par des agents-entités. Enfin, les agents-énaction procurent une certaine liberté pour la représentation et l'autonomisation de phénomènes qui résultent des interactions indirectes des constituants du système.

Notre méthode de modélisation peut être résumée par la figure 2.15. Plus qu'un nouveau méta-modèle, nous proposons donc un ensemble de briques élémentaires qu'il est possible de combiner pour construire des modèles multi-agents avec une approche mixte.

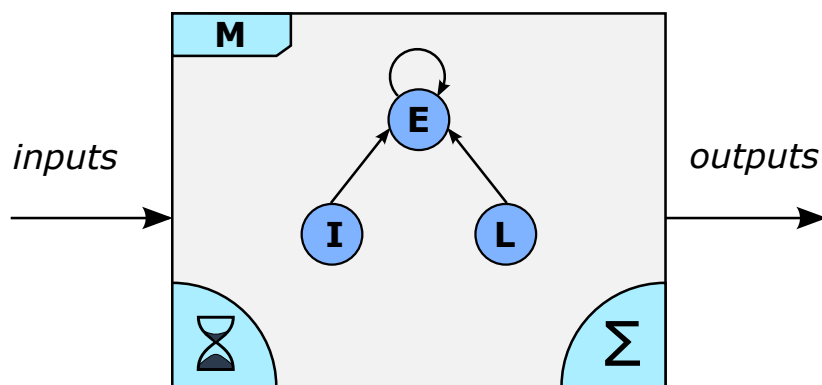


FIGURE 2.15 – **Synthèse de la composition d'un modèle multi-agents mixte** – Un modèle multi-agents (M) comprend plusieurs éléments de base. Des agents de tous types (entité, interaction, énaction, etc.) peuvent être combinés pour décrire un système complexe de la manière plus efficace. Un ordonnanceur, représenté par le sablier, est chargé de faire avancer le temps virtuel et de l'appel successif des activités des agents. Un schéma d'intégration numérique (Σ) est adjoint au modèle dès lors que les actions des agents sont décrites à l'aide d'équations différentielles.

Par ailleurs, nous tirons plusieurs enseignements de la méthode développée dans le cadre de cette étude.

Tout d'abord, le choix d'utiliser plusieurs types d'agents est ici plus une question de génie logiciel que de simulation. Comme nous l'avons vu, l'emploi d'un agent plutôt qu'un autre facilite la modélisation des différents phénomènes en fonction du contexte. En dépit de cet atout, nous aurions pu opter pour des solutions équivalentes en terme de résultats numériques, par exemple en utilisant uniquement des agents-entité. Seulement, une telle démarche aurait constitué un frein à la modélisation en imposant un formalisme pas forcément adapté à toutes les situations. C'est pourquoi nous insistons sur le fait que la modélisation par agents, quelle que soit leur nature, doit être un moyen et non une fin. Nous préconisons donc la flexibilité des outils de modélisation et même de mixer les atouts de chacun d'eux pour favoriser le

processus de réflexion sur le modèle.

Ensuite, nous souhaitons souligner la pertinence de l'approche phénoménologique dans le contexte de la modélisation de systèmes complexes. Là où d'autres méthodes, multi-agents ou non, s'attachent souvent à décrire le système le plus exactement possible, l'utilisation de lois descriptives permet ici de synthétiser un ensemble de mécanismes sous-jacents, certes de façon peut-être plus grossière mais également plus efficace du point de vue du temps de calcul. Nous voyons bien que le modèle ne constitue ici qu'un prototype dans lequel nous ne pouvons avoir qu'une confiance limitée au vu de l'incertitude introduite par les lois empiriques. Dans ces conditions, nous pensons qu'il est plus judicieux de prendre de la hauteur vis-à-vis de la modélisation du système et de considérer des lois continues pour décrire des phénomènes quand cela est possible, plutôt que de suivre une voie réductionniste. De cette manière, il est plus aisé d'expérimenter le modèle et d'en évaluer les capacités.

Les expérimentations *in virtuo* que nous avons réalisées sur le modèle MACMA viennent appuyer ce dernier point. La décomposition phénoménologique du système nous a permis de substituer la loi de migration des subductions par une autre, sans remettre en cause le reste du modèle. Nous avons également pu observer l'intérêt de l'approche multi-agents pour réaliser des études paramétriques. Il est tout à fait possible de réaliser ce genre de tests sur d'autres types de modèles, mais nous retiendrons que la décomposition du système en éléments autonomes permet de mieux évaluer l'impact d'une modification locale sur le comportement global.

Si la maquette numérique gagne en malléabilité, le fait de recourir à des lois descriptives peut pourtant rendre l'analyse des résultats paradoxalement plus difficile, au moment de justifier la valeur de paramètres qui, comme nous l'avons indiqué plus tôt, résument le comportement de sous-systèmes qui eux-mêmes peuvent présenter les caractéristiques d'un système complexe.

C'est pourquoi nous pensons qu'il est important de redonner du sens aux paramètres des lois descriptives en les situant comme le résultat de simulations plus raffinées de ces sous-systèmes. Le chapitre suivant s'inscrit dans cette logique et présente un ensemble de techniques de simulations analogues.

Chapitre 3

Échelles de modélisation



Philippe Geluck - LE TOUR DU CHAT EN 365 JOURS

Les chapitres précédents nous ont donné l'occasion de livrer notre définition des systèmes complexes, de proposer une approche phénoménologique pour modéliser et expérimenter *in virtuo* de tels systèmes au moyen des systèmes multi-agents. Ils ont également mis en évidence certaines limites de ces méthodes. En effet, les modèles descriptifs que nous mettons en avant sont souvent décrits à l'aide d'équations différentielles. De ce fait, ils introduisent des paramètres qui peuvent être difficiles à identifier selon le contexte.

Pour contourner cet obstacle, nous partons du constat suivant : les paramètres des modèles descriptifs de haut niveau reflètent la dynamique sous-jacente du système. De plus, un même phénomène peut être décrit à différentes échelles selon le point de vue adopté pour étudier le système. Ce chapitre vise à donner une vue d'ensemble des méthodes existantes pour tirer avantage de cette hétérogénéité des niveaux de modélisation. Ainsi, nous nous intéresserons particulièrement aux notions de niveaux de description, de transfert d'échelle et à leur utilisation au sein de simulations multi-échelles.

3.1 Niveaux de description

Nous avons vu que lorsqu'on modélise un système, les éléments et les phénomènes à représenter sont choisis en fonction de ce que l'on souhaite observer. Ainsi, à la manière d'un microscope, nous fixons le degré de grossissement adéquat dans le cadre de l'étude du système. Ce degré définit l'échelle de modélisation, c'est-à-dire le rapport de taille entre un objet et sa représentation. Ce rapport est nécessairement de deux natures. En effet, comme l'illustre la figure 3.1, il existe une relation forte entre l'échelle spatiale utilisée pour décrire l'objet et l'échelle temporelle à laquelle opèrent les interactions, aussi bien internes qu'externes, qui le sollicitent.

Pour étudier et représenter ces interactions, il est nécessaire d'adapter les méthodes de modélisation en fonction du niveau de description désiré. Nous aborderons dans cette section les différences entre les modèles dits **macroscopiques** et les modèles dits **microscopiques**. Nous verrons également que ces qualificatifs induisent une notion de hiérarchie des échelles de modélisation que nous souhaitons mettre à profit.

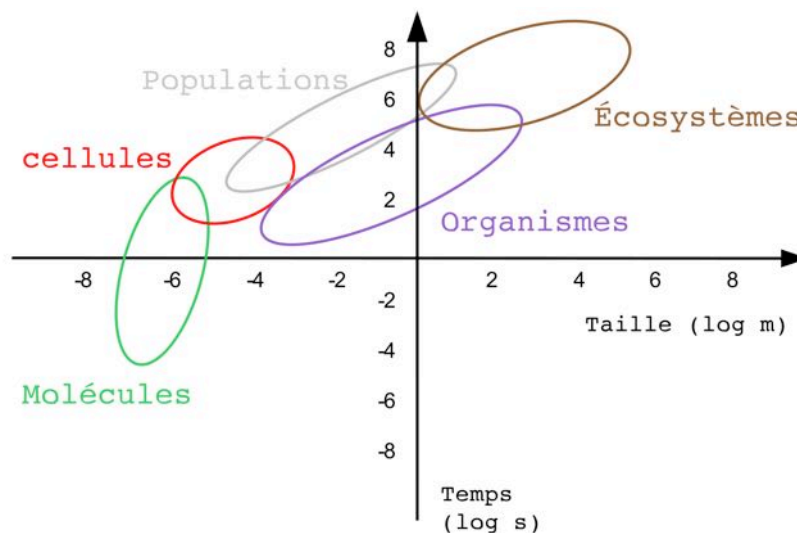


FIGURE 3.1 – Échelles spatio-temporelles caractéristiques des systèmes vivants – On peut noter que l'augmentation des dimensions selon une échelle va de pair avec l'augmentation des dimensions sur l'autre [Lales, 2007]. Ainsi, l'étude de la réplication d'une protéine peut être réalisée à l'échelle de la journée, alors que si on s'intéresse au devenir d'un organe, il sera nécessaire d'augmenter l'échelle de temps.

3.1.1 Modèles macroscopiques

Les phénomènes qui animent les systèmes complexes impliquent généralement un grand nombre d'éléments en interaction. Durant des siècles, les scientifiques sont parvenus à décrire ces phénomènes à l'aide de modèles macroscopiques en ignorant souvent tout ou partie de leur mécanique interne. Aujourd'hui encore, nous sommes fréquemment amenés à utiliser des lois empiriques pour décrire un comportement dont nous ne connaissons pas le fonctionnement

sous-jacent.

Ces modèles, souvent à bases d'équations différentielles, reposent principalement sur des lois d'évolution de quelques propriétés du système, comme l'énergie ou la vitesse. Parmi ces modèles macroscopiques, nous pouvons citer les équations de Navier-Stokes pour la dynamique des fluides, les équations fondamentales de la thermodynamique ou encore les lois de réaction en cinétique chimique. Bien qu'ils ne prennent *a priori* pas en compte les propriétés individuelles de chacun des éléments qui composent le système, ils sont capables de reproduire assez fidèlement les effets perceptibles des phénomènes.

3.1.2 Modèles microscopiques

En dépit de leur attractivité, les modèles macroscopiques ne permettent généralement pas de rendre compte de toute la complexité d'un phénomène. C'est alors qu'interviennent des modèles dits **microscopiques** qui décrivent plus finement les interactions entre les composants, voire composants de composants, du système. Ces modèles s'attachent à modéliser les comportements individuels des composants. L'exemple le plus couramment cité est celui de la dynamique moléculaire qui consiste à modéliser et simuler l'évolution de particules dans des domaines comme la dynamique des fluides, la chimie ou la physique des particules.

S'il paraît évident, en partant du principe que les interactions locales sont bien connues, que de tels modèles microscopiques garantissent une meilleure précision qu'un modèle macroscopique plus simple, ce changement de niveau de description a un coût. L'utilisation de modèles microscopiques implique d'adapter les échelles spatiales et/ou temporelles du domaine de simulation, ce qui influe fortement sur le pas de temps de simulation. Dans ces conditions, au vu des capacités de calcul actuelles, il devient impossible de simuler un système sur plusieurs heures de temps réel dans un temps de simulation raisonnable. Et quand bien même nous disposerions d'une puissance de calcul suffisante, il n'est pas dit que le gain de précision soit vraiment profitable pour l'étude du système. En effet, certains comportements n'apparaîtront à l'observateur que s'il parvient à s'extraire de son point de vue local.

3.1.3 Hiérarchie

Nous avons jusqu'ici distingué les niveaux de description macroscopique et microscopique, ainsi que les modèles qui leur sont généralement attachés. Notons qu'un troisième niveau est fréquemment mentionné dans la littérature, à savoir le niveau mésoscopique qui se veut être une échelle intermédiaire entre les deux niveaux cités précédemment. Cette distinction nous semble cependant un peu abrupte et nous souhaitons lui apporter quelques éclaircissements.

Tout d'abord, il nous semble important de relativiser cette classification et de souligner le fait que les termes macroscopique et microscopique n'ont de sens que s'ils sont associés. Prenons l'exemple d'une réaction chimique. Les éléments modélisés peuvent être des molécules ou des représentants d'espèces chimiques, selon que l'on s'intéresse à la transformation de molécules ou à l'évolution de concentrations. Pour autant, d'un certain point de vue, nous

pouvons considérer que la population n'est en fait qu'un individu à un niveau de description supérieur, de même qu'une molécule est le résultat de l'assemblage d'atomes à un niveau de description inférieur, qu'un atome est le théâtre d'interactions entre particules élémentaires à un niveau encore plus microscopique, etc. Nous nous attacherons donc dans ce manuscrit à situer les échelles de modélisation que nous utiliserons les unes par rapport aux autres.

Cette organisation nous conduit à constater qu'il existe une notion de hiérarchie entre les différents niveaux de description. Il s'agit ici, non pas de degrés d'autorité distincts, mais plutôt d'une sorte d'emboîtement à la manière de poupées russes, comme le décrit Rosen dans sa théorie de la hiérarchie [Rosen, 1969]. Il avance dans cette théorie qu'un même système peut être décrit à différentes échelles et que des interactions font le lien entre les niveaux de description. Plus qu'un simple lien, nous verrons par la suite que les niveaux inférieurs constituent en fait le corps des niveaux supérieurs.

La posture classique pour modéliser un système est de choisir un niveau de description et d'ignorer tous les autres. Si cette approche est souvent suffisante, face à la complexité des systèmes que nous souhaitons simuler nous sommes amenés à adopter des méthodes dites **multi-échelles** qui ont, comme leur nom l'indique, la particularité de mêler plusieurs échelles de modélisation.

3.2 Méthodes numériques multi-échelles

De nombreux problèmes relevant des systèmes complexes font apparaître des phénomènes qui opèrent à des échelles d'espace ou de temps différentes. De plus, nous avons vu que plusieurs niveaux de description peuvent être considérés pour un même phénomène. Le traitement numérique efficace de tels problèmes nécessite des méthodes spécifiques. Une des difficultés majeures, outre la puissance de calcul nécessaire pour la résolution numérique, est de parvenir à coupler ces différentes échelles de simulation.

Nous débuterons cette section par une classification des problèmes multi-échelles et des méthodes qui leur sont applicables. Nous présenterons ensuite plusieurs approches, issues de domaines applicatifs variés, qui permettent de prendre en compte cette diversité des niveaux de description.

3.2.1 Phénomènes critiques ou séparation des échelles

Une des étapes dans l'étude d'un système multi-échelles consiste à identifier certaines caractéristiques de ce système. Ces caractéristiques conditionnent le choix de la méthode de modélisation à utiliser. Nous trouvons généralement dans la littérature un consensus autour de la distinction des problèmes multi-échelles en deux classes.

La première classe rassemble les systèmes composés de **phénomènes critiques**. Ces phénomènes, telles la propagation de fissures dans un matériau ou la transition de phase, ont pour particularité d'imposer localement un niveau de description microscopique. En effet, dans ces situations, les variables d'état du système présentent des fluctuations de toutes tailles

et de toutes durées. Ainsi, dans un même modèle, il s'agira de faire cohabiter des échelles spatiales et temporelles différentes afin de prendre en compte ces phénomènes. Concrètement, des modèles microscopiques seront utilisés localement alors que des modèles macroscopiques couvriront l'ensemble du domaine de simulation. Pour ce faire, plusieurs méthodes sont classiquement utilisées parmi lesquelles nous pouvons citer la décomposition de domaine, les méthodes de splitting et les maillages adaptatifs, illustrés par la figure 3.2. Ces méthodes proposent peu ou prou la même approche de résolution séparée et séquentielle des modes lents et rapides du système. Par exemple, la partie raide du système peut être calculée à l'aide d'un schéma d'intégration plus lent mais plus précis et avec un petit pas de temps, tandis que la partie non raide utilisera un schéma plus rapide et un grand pas de temps. Si ces méthodes donnent de bons résultats, elles ne parviennent pas totalement à gommer l'inconvénient lié à l'utilisation de modèles microscopiques. En effet, leur temps de calcul reste souvent élevé car le pas de temps général de la simulation est généralement contraint par celui de la plus petite échelle considérée.

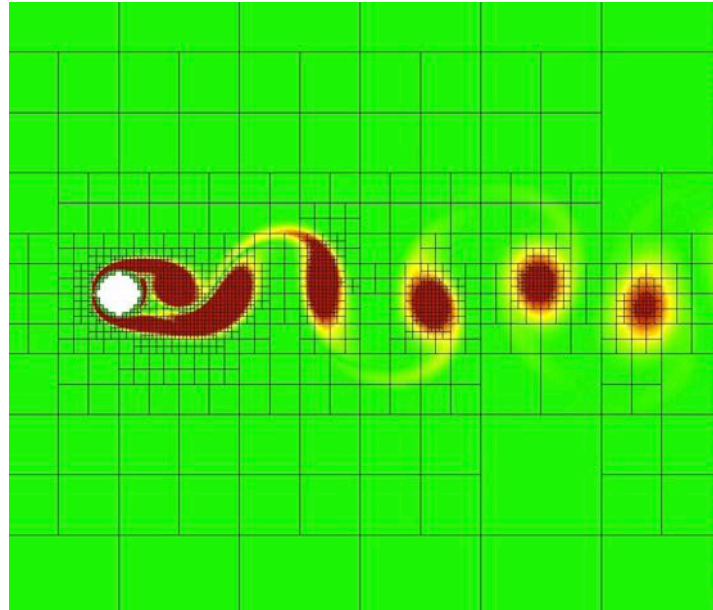


FIGURE 3.2 – **Maillage adaptatif pour la simulation d'allées de Von Karman** – La discrétisation spatiale devient plus fine uniquement là où les vortices provoquent des variations de vitesse de fluide plus importantes [Delaux et al., 2010]

La seconde classe regroupe les systèmes dans lesquels un niveau de description n'influence les niveaux supérieurs que par l'intermédiaire de **paramètres effectifs**. Ces paramètres effectifs résument un phénomène sous-jacent, au sens où ils suffisent à rendre compte de ses conséquences observables aux échelles supérieures. Cette abstraction est possible quand le système affiche une **séparation des échelles**, par opposition aux systèmes présentant des phénomènes critiques.

L'approche phénoménologique que nous proposons se concentre sur cette seconde classe de problèmes. Les phénomènes modélisés à l'échelle macroscopique par des lois physiques ou empiriques comprennent souvent des paramètres effectifs, dont les valeurs sont parfois difficiles à définir. Un exemple courant, que nous reprendrons dans le chapitre 4 de ce mémoire, est celui de la diffusion. Selon l'échelle d'observation, ce phénomène est décrit tour à tour par

une loi de diffusion, une distribution de probabilité de présence ou la loi de Fick. Dans chacune de ces descriptions, nous retrouvons un paramètre effectif, le coefficient de diffusion D , seule trace aux échelles supramoléculaires de l'agitation thermique moléculaire à l'origine du mouvement de diffusion. Pour évaluer ce paramètre, la méthode classique consiste à mesurer expérimentalement le déplacement des particules diffusantes dans un milieu. Il s'agit donc de s'appuyer sur l'échelle microscopique pour paramétrer des lois macroscopiques.

Les méthodes numériques que nous présentons ci-après visent à proposer des mécanismes similaires pour simuler le plus efficacement possible les systèmes multi-échelles.

3.2.2 Super-paramétrisation

De nos jours, la simulation trouve un intérêt dans de nombreux domaines applicatifs, mais s'il en est un qui l'utilise depuis ses origines, c'est celui de la météorologie. Cette discipline a pour objet l'étude des phénomènes atmosphériques tels que les nuages, les précipitations ou le vent dans le but de comprendre comment ils se forment et évoluent. Purement descriptive à l'origine, la météorologie moderne repose principalement sur la mécanique des fluides et la thermodynamique, mais aussi la chimie et les mathématiques, pour établir des prévisions de l'évolution du temps à court comme à long terme. Ces théories ont permis d'élaborer des modèles utilisés depuis des décennies [Randall et al., 2003] mais restent confrontées à des problèmes, notamment en termes de paramétrisation des nuages, qui ont récemment poussé les météorologues à développer de nouvelles méthodes de simulation multi-échelles.

En effet, plusieurs modèles existent selon l'échelle de modélisation considérée. Les GCMs¹ décrivent, à l'aide des équations de Navier-Stokes, l'évolution de l'atmosphère, sur un maillage, à l'échelle synoptique. Les phénomènes de l'échelle synoptique se caractérisent par des distances de plusieurs centaines à plusieurs milliers de kilomètres et des durées de plusieurs jours. La résolution de ces modèles nécessite donc de faire des hypothèses fortes sur les conditions aux limites et les effets des phénomènes aux échelles inférieures.

Au début des années 1980 sont apparus les CSRM² qui donnent la dynamique individuelle ou de groupes de nuages, à l'échelle du kilomètre et de la minute. Pour ce faire ils prennent en compte des lois d'organisation et d'interaction des structures nuageuses, mais aussi des phénomènes tels que les turbulences et la convection thermique. Plus précis, ils sont pourtant inutilisables en l'état pour une simulation à l'échelle planétaire en raison de leur coût computationnel. Par ailleurs, ils nécessitent toujours d'être paramétrés.

Aujourd'hui, ces modèles reviennent en force suite à la proposition de Grabowski [Grabowski et Smolarkiewicz, 1999] de combiner les forces des deux niveaux de description. Son idée est d'utiliser un CSRM dans chaque maille du GCM pour mettre en place ce qu'il appelle une super-paramétrisation, illustrée par la figure 3.3. Ainsi, les hypothèses formulées pour résoudre le GCM reposent dorénavant sur des données physiquement réalistes. De même, les résultats du GCM sont utilisés pour paramétrer le CSRM pour tenir compte de phénomènes macroscopiques et notamment du comportement des mailles adjacentes. Il

1. General Circulation Models

2. Cloud-System Resolving Models

est en effet important de noter que chaque maille est résolue indépendamment des autres et dispose de ses propres conditions aux limites. Cet aspect constitue à la fois un inconvénient et un avantage. D'une part, le paramétrage du CSRM implique d'être capable de traduire les données en provenance du GCM. De retour à l'échelle macroscopique, le GCM doit également gérer la discontinuité relative des résultats provenant de mailles adjacentes. D'autre part, cette décomposition du domaine de simulation en éléments indépendants ouvre la voie à une parallélisation des calculs.

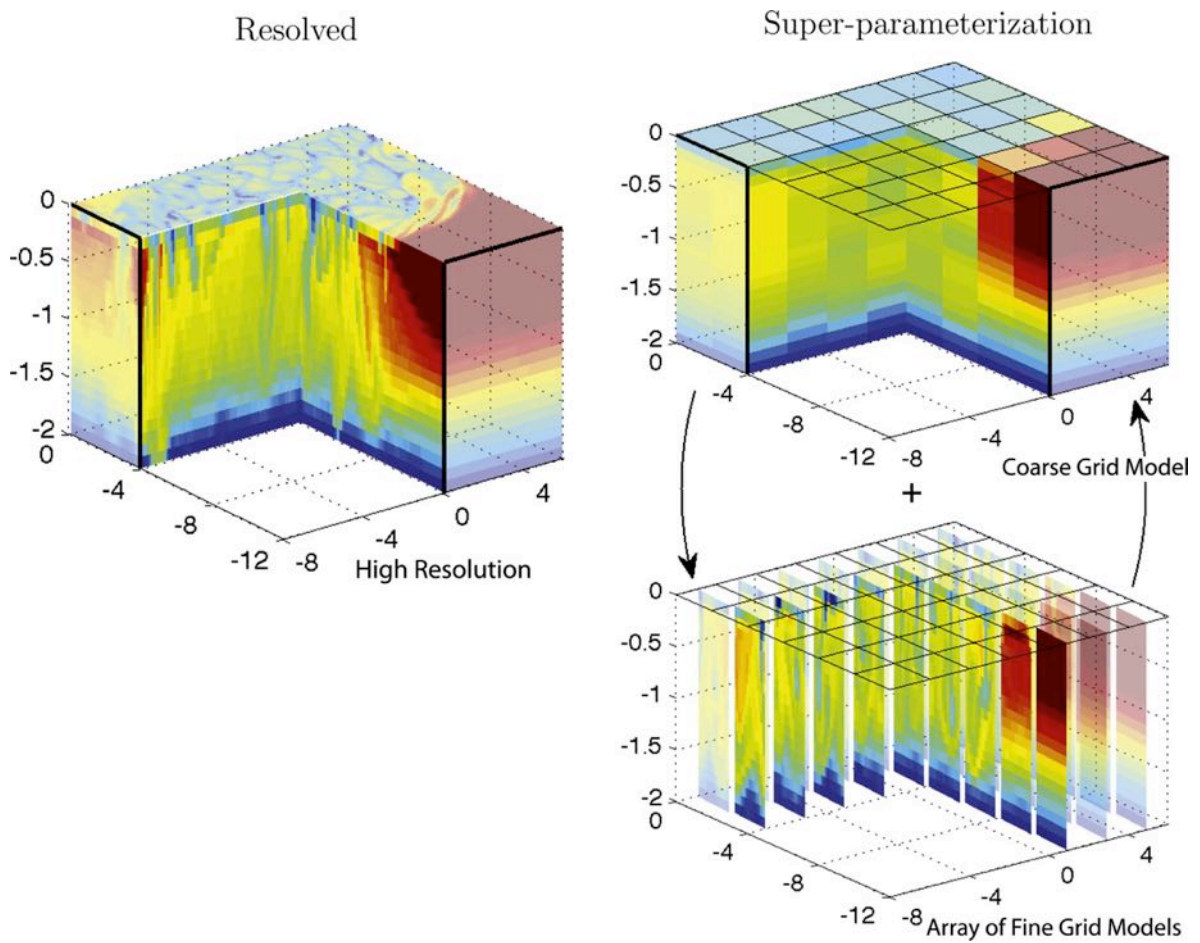


FIGURE 3.3 – **Super-paramétrisation** – Vue 3D du champ de température (rouge = chaud, bleu = froid) pour deux simulations d'une cheminée de convection. A gauche : Simulation haute résolution à l'échelle micro de panaches. A droite : Utilisation d'un modèle super-paramétrisé dans lequel un modèle à gros grains (CG) embarque dans chacune de ses mailles un modèle plus fin (FG). Les modèles CG et FG ont une influence l'un sur l'autre et échangent des informations [Campin et al., 2011].

3.2.3 Méthode quasicontinuum

La mécanique des milieux continus (continuum) est le domaine de la mécanique qui s'intéresse à la déformation des solides et à l'écoulement des fluides. L'hypothèse des milieux continus consiste à considérer des milieux dont les propriétés macroscopiques, comme la densité ou l'élasticité, sont continues. Une telle hypothèse permet d'avoir recours aux outils mathématiques reposant sur les fonctions continues et/ou dérivables. En pratique, cela revient à considérer que le volume élémentaire de matière du système modélisé est suffisamment grand pour que l'on puisse négliger à tout instant les fluctuations du nombre de particules (atomes ou molécules) qui le composent.

Cependant, le développement des nanotechnologies ces dernières années a conduit les mécaniciens à s'intéresser à des échelles physiques de plus en plus petites. Particulièrement, l'étude des phénomènes de fracture et de fatigue des systèmes microélectromécaniques a rapidement mis en exergue la nécessité de prendre en compte qu'un matériau est en réalité composé de particules discrètes. De nouveaux modèles se focalisant sur ce niveau de description sont apparus, mais ont rencontré le même écueil que la majorité des modèles microscopiques, à savoir l'impossibilité de simuler un nombre significatif de particules.

Il apparaît néanmoins que même dans le cas de nano-systèmes, il n'est pas forcément nécessaire de se contraindre à l'échelle atomique sur l'ensemble du domaine simulé. En effet, s'il est important d'adopter ce niveau de description pour certaines régions du problème, d'autres peuvent se résumer à un milieu continu.

La méthode quasicontinuum (QC) a été développée dans le but de coupler les deux types de modélisation [Tadmor et al., 1996]. L'idée est de considérer la description atomique comme le modèle « exact » du comportement du matériau, mais à la fois de reconnaître que le grand nombre d'atomes en jeu rend le problème intraitable à cette échelle.

Le principe de la QC est de construire un modèle continu à partir du modèle atomique en introduisant des contraintes cinématiques. Pour ce faire, l'idée est de sélectionner des *représentants* d'atomes afin de réduire les degrés de liberté de chaque atome et ainsi le nombre de calculs. Nous construisons ensuite un maillage triangulaire d'éléments finis, chaque représentant étant utilisé comme noeud [Ming et Yang, 2009]. À chaque pas de simulation, le déplacement des représentants est calculé ainsi que l'énergie qui leur est associée. Ces données servent de base à une stratégie de raffinement de maillage qui conduit à considérer chaque atome comme un représentant dans les zones où l'énergie calculée dépasse un seuil prédéfini. Ce processus est illustré par la figure 3.4. De cette manière, le nombre de calculs est diminué tout en gardant la précision du modèle atomique quand cela est nécessaire.

Les méthodes multi-échelles que nous avons présentées jusqu'ici reposent sur le fait qu'il existe un modèle pour chaque niveau de description. Leur objectif est de réduire le nombre de calculs en n'utilisant l'échelle microscopique que lorsque cela est nécessaire, par exemple dans des zones critiques ou pour le paramétrage du modèle macroscopique. Nous allons maintenant introduire une méthode qui traite les problèmes pour lesquels nous ne disposons pas de modèle macroscopique.

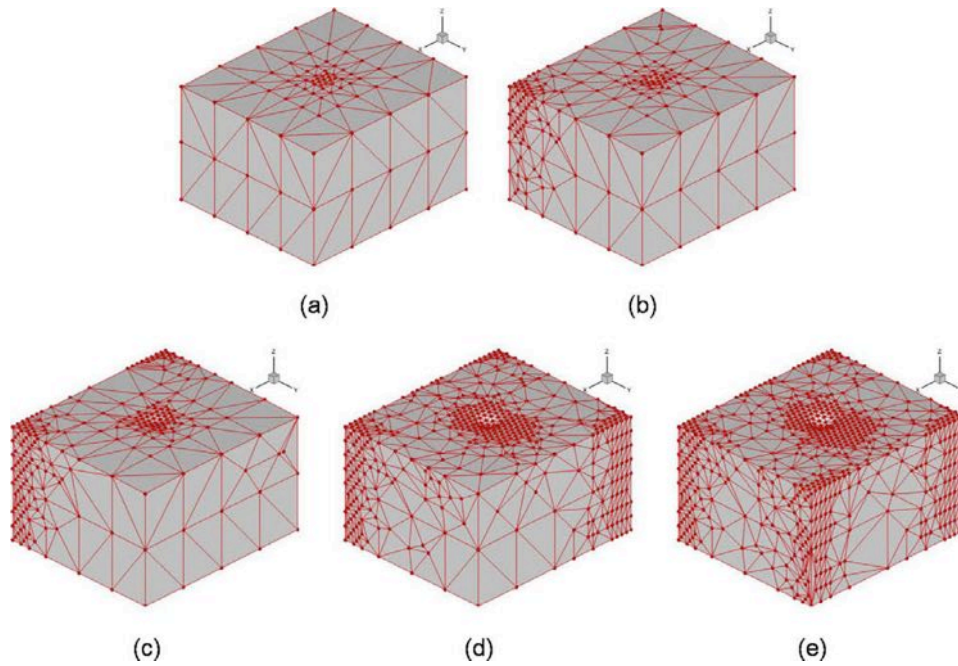


FIGURE 3.4 – **Méthode quasicontinuum** – Simulation d’une nano-indentation sur un tétraèdre. On exerce une indentation successivement sur chacun des sommets du tétraèdre. Les représentants d’atomes et le maillage qui leur est associé sont redéfinis à chaque pas de temps en fonction de l’ampleur de la déformation calculée [Kwon et al., 2009].

3.2.4 Approche « sans-équation »

L’étude d’un système complexe débute souvent par la construction d’un modèle microscopique qui se concentre sur le comportement individuel de ses composants. Pour autant, il ne faut pas perdre de vue que c’est le comportement du système dans son ensemble que nous cherchons à simuler. Nous nous appuyons alors généralement sur des modèles macroscopiques qui abstraient les détails des modèles microscopiques et définissent des lois d’évolution plus globales du système. Malheureusement, dans certains cas, le modèle d’évolution du système peut ne pas exister pour ce niveau de description.

Pour ce type de problèmes, nous aurons alors recours à des méthodes dites **sans équation**. Celles-ci reposent sur la création d’un time-stepper macroscopique se basant sur le time-stepper microscopique connu et sur des opérateurs de traduction entre les échelles de modélisation. Nous détaillons ici les mécanismes qui donnent corps à cette méthode et permettent d’obtenir un point de vue macroscopique sur le système simulé tout en s’épargnant une simulation coûteuse de la totalité du système à l’échelle microscopique.

3.2.4.1 Définitions

Considérons une loi d’évolution microscopique

$$\partial_t u(x, t) = f(u(x, t)) \quad (3.1)$$

pour laquelle $u(x, t)$ décrit l'état des variables microscopiques x à l'instant t . Nous supposons qu'il existe un modèle macroscopique équivalent, mais qu'il ne peut être exprimé de manière explicite. Nous notons ce modèle

$$\partial_t U(X, t) = F(U(X, t)) \quad (3.2)$$

dans lequel $U(X, t)$ représente l'état des variables macroscopiques X à l'instant t .

Nous introduisons un time-stepper pour le modèle microscopique (3.1)

$$u(x, t + \delta t) = s(u(x, t), \delta t) \quad (3.3)$$

où δt est le pas de temps microscopique. L'approche sans-équation de [Kevrekidis et al., 2003] propose d'effectuer la même opération pour le modèle macroscopique (3.2)

$$\bar{U}(X, t + \Delta t) = \bar{S}(\bar{U}(X, t), \Delta t) \quad (3.4)$$

où Δt dénote le pas de temps macroscopique, et les barres ont été ajoutées pour mettre en valeur le fait que le time-stepper macroscopique est seulement une approximation du time-stepper de (3.2) puisque ce modèle n'est pas connu explicitement. Pour définir ce time-stepper macroscopique, nous introduisons deux opérateurs chargés de traduire les variables entre les deux niveaux de description.

3.2.4.2 Opérateurs

Pour passer de l'échelle macroscopique à l'échelle microscopique, nous appliquons un opérateur de *lifting*

$$\mu : U(X, t) \mapsto u(x, t) = \mu(U(X, t)). \quad (3.5)$$

À l'inverse, pour le passage de l'échelle microscopique à l'échelle macroscopique, nous utilisons un opérateur de *restriction*

$$\mathcal{M} : u(x, t) \mapsto U(X, t) = \mathcal{M}(u(x, t)). \quad (3.6)$$

La plupart du temps, l'opérateur de restriction (3.6) est assez simple à définir dès lors que les variables macroscopiques sont identifiées. Par exemple, lorsque le modèle microscopique représente l'évolution d'un ensemble de particules, la restriction consiste à calculer les moments d'ordre inférieur du système comme la densité ou l'énergie, le plus souvent en

moyennant les moments d'ordre supérieur tels que la position ou la vitesse. Pour ce faire, plusieurs techniques mathématiques peuvent être employées, parmi lesquelles les méthodes de calcul des moyennes [Pavliotis et Stuart, 2008], d'homogénéisation [Kavenoky, 1978] ou de renormalisation [Lesne, 2003]. Nous ne détaillerons pas ici ces différentes méthodes car elles sont trop dépendantes de la forme du modèle microscopique.

La construction de l'opérateur de *lifting* (3.5) est généralement plus compliquée. En reprenant notre exemple d'un modèle particulaire, l'objectif est à présent de faire correspondre des moments d'ordre inférieur avec des conditions initiales pour chaque particule. La difficulté est que nous ne connaissons pas de fonction explicite pour lier ces moments dans ce sens. Au mieux, c'est parmi une infinité de fonctions que nous devons faire un choix. Nous verrons par la suite que ce choix peut être guidé dans la mesure où nous disposons de plus d'informations sur l'échelle macroscopique.

Le plus souvent, les moments d'ordre supérieur sont initialisés aléatoirement. Cette démarche introduit une **erreur de *lifting*** qui ne peut être maîtrisée que si la séparation des échelles temporelles est importante. À partir de cette hypothèse, on procède à une correction appelée le *healing* [Ross et Mohanty, 2008], qui consiste à itérer le modèle microscopique jusqu'à ce qu'il se relaxe, ou converge, vers un état d'équilibre.

Quand cela est possible, nous pouvons également nous appuyer sur notre connaissance du système pour procéder à un *lifting* efficace.

3.2.4.3 Time-stepper macroscopique

À partir d'un état macroscopique $\bar{U}(X, t^*)$ à un instant t^* , nous construisons le time-stepper (3.2) de la manière suivante :

- **Lifting**
Créer des conditions initiales $\bar{u}(x, t^*)$ cohérentes vis à vis de $\bar{U}(X, t^*)$ à l'aide de l'opérateur de *lifting* (3.5).
- **Simulation**
Utiliser le time-stepper microscopique (3.1) pour calculer l'état $\bar{u}(x, t)$ pour $t \in [t^*, t^* + \Delta t]$.
- **Restriction**
Obtenir l'état macroscopique $\bar{U}(X, t^* + \Delta t)$ à partir de l'état microscopique $\bar{u}(x, t^* + \Delta t)$ à l'aide de l'opérateur de restriction (3.6).

En considérant $\Delta t = k\delta t$, le time-stepper macroscopique peut donc être écrit de la sorte :

$$\bar{U}(X, t + \Delta t) = \bar{S}(\bar{U}(X, t), \Delta t) = \mathcal{M}(s^k(\mu(\bar{U}(X, t), \delta t))) \quad (3.7)$$

Avec cet algorithme, nous sommes en mesure d'observer l'évolution temporelle de variables macroscopiques. Nous ne pouvons pas vraiment parler ici de simulation du système à l'échelle macroscopique, mais plutôt d'une succession d'états, puisque il n'existe toujours

pas de modèle pour ce niveau de description. Par ailleurs, le gain n'est pour le moment qu'au niveau de la connaissance du système. En effet, le modèle microscopique doit encore être calculé entièrement à chaque pas de temps. Nous allons à présent voir comment la démarche du time-stepper peut être utilisée pour gagner en temps de calcul.

3.2.4.4 Méthodes d'intégration projectives

L'algorithme du time-stepper macroscopique peut être adapté pour accélérer la simulation du système en utilisant des méthodes d'intégration projectives. L'idée générale de ces méthodes est de ne plus simuler le modèle microscopique en permanence, mais uniquement sur de petites durées, que nous noterons τ . La trajectoire du modèle microscopique sert alors de référence pour le time-stepper macroscopique qui l'intègre sur un plus gros pas de temps, comme nous pouvons le voir sur la figure 3.5.

Nous ajustons donc l'algorithme précédent en considérant le pas de temps $\Delta t \gg \delta t$:

- **Lifting**

Créer des conditions initiales $\bar{u}(x, t^*)$ cohérentes vis à vis de $\bar{U}(X, t^*)$ à l'aide de l'opérateur de *lifting* (3.5)

- **Simulation**

Utiliser le time-stepper microscopique (3.1) pour calculer l'état $\bar{u}(x, t)$ pour $t \in [t^*, t^* + \tau]$

- **Restriction**

Obtenir l'état macroscopique $\bar{U}(X, t^* + \tau)$ à partir de l'état microscopique $\bar{u}(x, t^* + \tau)$ à l'aide de l'opérateur de restriction (3.6)

- **Projection**

Calculer $\bar{U}(X, t^* + \Delta t) = \bar{U}(X, t^*) + \Delta t \frac{\bar{U}(X, t^* + \tau) - \bar{U}(X, t^*)}{\tau}$

La dernière étape de projection repose ici sur une différence finie pour évaluer la dérivée du modèle microscopique, après relaxation. Cette dérivée est intégrée sur le pas de temps du time-stepper macroscopique Δt .

Si l'intégration projective permet d'accélérer la simulation au niveau macroscopique, deux points doivent cependant être soulignés.

Tout d'abord, le choix des pas de temps Δt et δt , ainsi que la durée de simulation microscopique τ , ne sont pas anodin. En effet, tout comme pour les schémas d'intégration explicites classiques, si Δt est choisi trop grand au regard de la dynamique du système au niveau macroscopique, la simulation aura tendance à diverger. Si τ est choisi trop petit par rapport au temps de relaxation des variables macroscopiques, la dérivée retournée au modèle macroscopique sera faussée et causera également la divergence de la simulation.

Par ailleurs, τ devra être adapté en fonction de la qualité de l'opérateur de *lifting*. Si l'erreur introduite par ce dernier est grande, le *heal* demandera plus de temps. Or il est évident que plus τ est réduit, plus le gain de temps simulé est significatif.

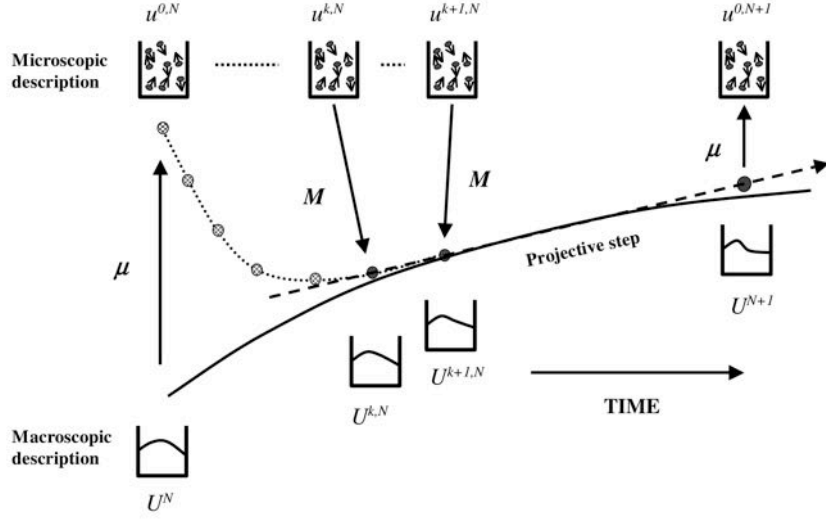


FIGURE 3.5 – **Méthode d'intégration projective** – On applique l'opérateur de lifting μ sur un état macroscopique U^N pour obtenir un état microscopique initial $u^{0,N}$. Le modèle microscopique est simulé sur plusieurs pas de relaxation pour estimer la variation des variables macroscopiques, évaluées à l'aide de l'opérateur de restriction $\mathcal{M}$. Une fois cette variation stabilisée, elle sert de trajectoire au time-stepper macroscopique qui est extrapolé sur un pas de temps plus grand [Kevrekidis et al., 2003].

3.2.4.5 Patch dynamics

Les méthodes numériques multi-échelles peuvent généralement être classées dans deux catégories. La première regroupe les méthodes qui ont pour objectif d'optimiser la discrétisation spatiale du domaine simulé. C'est par exemple le cas de la méthode quasi-continuum qui alterne les niveaux de description continus et atomiques, en fonction du besoin de précision. La deuxième classe cherche à accélérer la résolution temporelle des modèles. Comme nous venons de le voir, les méthodes d'intégration projectives suivent cette voie en ne revenant au modèle microscopique, et donc à un petit pas de temps, que pour évaluer la trajectoire de modèle macroscopique sur un pas de temps plus grand.

La méthode *Patch dynamics* est une combinaison de ces deux stratégies [Samaey et al., 2007]. Elle propose d'associer une méthode d'intégration projective temporelle classique à une discrétisation spatiale similaire appelée le schéma *gap-tooth* [Gear et al., 2003].

Le principe du schéma *gap-tooth* est de considérer des petits intervalles (les *teeth*) séparés par des fossés (les *gaps*). La figure 3.6 donne une représentation de cette discrétisation pour un modèle à une dimension. Pour un modèle en trois dimensions, les dents deviennent des boîtes englobantes autour d'un point de référence.

A première vue, le schéma *gap-tooth* n'est qu'une version spatiale des méthodes d'intégration projectives. Cet aspect spatial implique cependant de prendre en compte de nouvelles contraintes. En effet, pour la version temporelle, c'est l'intégralité du modèle microscopique qui est reconstruite et simulée avec un petit pas de temps. Ici, seule une sous-partie du système correspondant à la taille de la dent est simulée, ce qui suppose d'être capable d'exprimer

le comportement du modèle aux frontières de chacune de ces dents.

Ainsi, l'algorithme suivant est appliqué pour chaque dent :

- **Lifting**

Créer des conditions initiales $\bar{u}(x, t^*)$, cohérentes vis à vis de $\bar{U}(X, t^*)$ à l'aide de l'opérateur de *lifting* (3.5) et du profil spatial macroscopique (c'est à dire des boîtes avoisinantes).

- **Simulation**

Utiliser le time-stepper microscopique (3.1) pour calculer l'état $\bar{u}(x, t)$ pour $t \in [t^*, t^* + \tau]$ avec des conditions aux limites appropriées.

- **Restriction**

Obtenir l'état macroscopique $\bar{U}(X, t^* + \tau)$ à partir de l'état microscopique $\bar{u}(x, t^* + \tau)$ à l'aide de l'opérateur de restriction (3.6).

- **Projection**

Calculer $\bar{U}(X, t^* + \Delta t) = \bar{U}(X, t^*) + \Delta t \frac{\bar{U}(X, t^* + \tau) - \bar{U}(X, t^*)}{\tau}$.

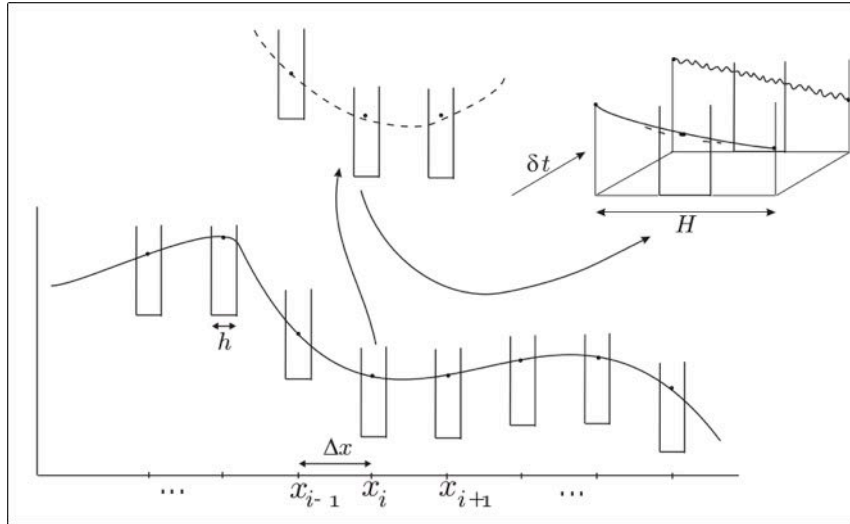


FIGURE 3.6 – **Représentation d'un pas de temps du schéma Gap-tooth** – Des boîtes de taille h sont placées autour de chaque point x_i du maillage macroscopique. Elles contiennent un modèle microscopique partiel reconstruit à chaque pas de temps à partir de leurs données du pas de temps précédent, de celles des boîtes avoisinantes et de conditions aux limites. Le modèle microscopique est ensuite calculé dans chaque boîte et les résultats sont moyennés pour reformer le modèle macroscopique [Samaey et al., 2007].

Dans le cadre de la méthode *Patch dynamics*, les dents sont appelées les *patches*. Pour chaque *patch*, le modèle microscopique est simulé pendant une durée τ dans l'idée d'obtenir un état stable et une trajectoire qui sera intégrée grâce à une méthode projective.

Nous avons décrit l'ensemble des éléments qui composent l'approche « sans-équation ». Comme son nom l'indique, cette approche consiste à obtenir un niveau d'observation macroscopique pour lequel nous ne disposons pas de modèle. De plus, dans la grande majorité

des cas, le système ne peut être simulé entièrement à l'échelle microscopique. Les techniques de discrétisation spatiales et temporelles que nous avons présentées permettent de diminuer fortement le nombre de calculs.

Pour autant, il apparaît assez clairement que cette approche ne peut fonctionner que pour un nombre limité de systèmes. En effet, dès lors que le time-stepper macroscopique ne consiste qu'en une extrapolation, si le système a tendance à évoluer rapidement à l'échelle macroscopique, il sera difficile d'obtenir un comportement crédible pour ce niveau de description. C'est cet écueil que les méthodes multi-échelles hétérogènes proposent de contourner.

3.2.5 Méthodes multi-échelles hétérogènes

Nous avons pu remarquer que les méthodes multi-échelles décrites jusqu'ici ont certains points communs. L'approche retenue est souvent d'abstraire le comportement d'un modèle microscopique dans le but de paramétrer un modèle macroscopique. Ce couplage est parfois à double sens dans la mesure où l'échelle microscopique n'est pas simulée pour l'ensemble du domaine, du fait de son coût. Ainsi, il est nécessaire de définir des conditions aux limites acceptables pour les sous-domaines considérés pour cette échelle. À défaut d'uniformiser ces méthodes, Weinan et Engquist [Weinan et al., 2003] ont cherché, au cours de leurs différents travaux, à proposer une méthodologie générale pour la modélisation de systèmes tirant avantage de la multiplicité des niveaux de description.

La méthode multi-échelles hétérogènes (HMM, pour *Heterogeneous Multiscale Method*) est à la fois très proche de l'approche « sans équation » mais se revendique également fondamentalement différente. En particulier, elles s'opposent au niveau de l'échelle macroscopique. L'approche « sans équation » s'inscrit par définition dans une démarche dite ascendante, ou *bottom-up*, dans le sens où son objectif est de synthétiser le modèle microscopique pour faire émerger la dynamique du niveau de description supérieur, que nous ne connaissons pas *a priori*. Au contraire, la HMM propose une approche descendante, ou *top-down*, dont le but est de compléter un modèle macroscopique sur lequel nous avons des informations, en prenant pour base un modèle équivalent plus détaillé. Nous constatons donc que la HMM est une méthode « basée sur une équation » par opposition à l'approche « sans équation ». Ainsi, cette nuance impacte l'algorithme sur deux points.

Tout d'abord, nous disposons d'un vrai modèle macroscopique qui nous épargne la construction d'un time-stepper uniquement basé sur l'extrapolation des méthodes d'intégration projectives. Ensuite, puisque nous disposons de plus d'informations sur ce modèle macroscopique, et notamment sur sa dynamique, nous pouvons plus facilement définir les opérateurs de *lifting* et de *restriction*, rebaptisés opérateurs de *compression* et de *reconstruction* respectivement, en nous appuyant sur l'expertise métier du domaine. L'emploi de ces opérateurs est illustré par la figure 3.7.

Cependant, l'analyse des auteurs de la HMM tend à confirmer ce que nous avons pu remarquer au travers des exemples précédents [Abdulle et al., 2012]. Si l'expertise métier permet d'optimiser le traitement des problématiques multi-échelles de nos modèles, que sont le couplage de phénomènes et le temps de calcul, elle implique de formuler des hypothèses

ad hoc pour chaque nouveau cas. De plus, nous nous appuyons désormais sur un modèle à l'échelle macroscopique qui est, par définition, imparfait. Or ce modèle sert de base pour la construction et l'initialisation d'un modèle microscopique. Une erreur de modélisation vient donc s'ajouter aux erreurs liées aux opérateurs inter-échelles. De ce fait, si cette erreur est trop importante, nous ne pouvons pas espérer obtenir des résultats exploitables.



FIGURE 3.7 – **Représentation schématique de l’Heterogeneous Multiscale Method** – Deux niveaux de description d’un système sont associés. Des opérateurs de reconstruction et de compression permettent de faire la traduction de l’un à l’autre, l’un définissant des contraintes pour raffiner la description tandis que l’autre résume un ensemble de données en des variables macroscopiques, d’après [Weinan et al., 2003]

3.2.6 Bilan

Tout au long de cette section, nous avons étudié plusieurs familles de méthodes numériques multi-échelles. Nous avons constaté que dans certains cas la séparation des échelles permet d’examiner un système sous différents angles, avec différents niveaux de détails, et donc différents modèles. Les méthodes présentées suggèrent une utilisation conjointe de ces modèles au sein d’une même simulation et ainsi une mise à profit de l’ensemble des connaissances dont nous disposons. La section suivante a pour objet l’implémentation de ce couplage grâce à des approches de multi-modélisation.

3.3 Approches de multi-modélisation

3.3.1 Paradigmes de modélisation

Comme nous l’avons vu, la diversité des niveaux de description pour décrire un même système nous conduit à considérer différents types de modèles. Ces modèles peuvent être issus de différents champs disciplinaires ou de communautés scientifiques aux points de

vue divergents. À ce titre, le paradigme de modélisation adopté va souvent de pair avec la catégorie du modèle. Ainsi, un modèle de population reposera le plus souvent sur des équations différentielles, stochastiques ou non, tandis que les modèles individus-centrés seront construits à l'aide de système multi-agents. Le tableau 3.8 donne une liste non exhaustive de paradigmes couramment utilisés en simulation avec des informations sur leur façon de traiter les dimensions spatiale et temporelle.

paradigme	temps	espace	niveau d'abstraction
équation différentielle ordinaire	continu	non pris en compte	global : variables moyennées sur l'ensemble des éléments du système
équation différentielle ordinaire par compartiment	continu	pris en compte (compartimentation de l'espace)	intermédiaire : variables moyennées restreintes aux populations des compartiments
équation différentielle partielle	continu	pris en compte (expressions analytiques simples)	global : variables moyennées sur l'ensemble des éléments du système
modélisation stochastique	généralement continu	généralement non pris en compte	intermédiaire : variables aléatoires correspondantes aux densités de probabilité des éléments du système
modélisation logique	généralement discrétisé (transition d'un état à un autre)	a priori non pris en compte (mais des sous-classes de molécules peuvent être associées à des régions de l'espace)	intermédiaire : regroupement des éléments de même type
réseau de Petri	continu ou discrétisé (transitions)	pris en compte (au travers des places)	intermédiaire : un jeton par élément mais temps (transitions) et espace (places) sont discrétisés
automate cellulaire	généralement discrétisé (mais peut être simulé continu)	nécessairement discrétisé (quadrillage fixe)	intermédiaire : autant d'attributs par cellule (tant que ceux-ci ne dépendent pas de l'espace)
système multi-agents	généralement discrétisé (mais peut être simulé continu)	généralement continu	fin : toutes les singularités des éléments prises en compte

FIGURE 3.8 – **Paradigmes de modélisation et leurs relation au temps, à l'espace et aux éléments modélisés** – Il est possible de positionner les différents paradigmes les uns par rapport aux autres en se basant sur le niveau d'abstraction des éléments modélisés, depuis le plus global jusqu'au plus local.

Dans le cadre de notre étude, nous souhaitons développer des simulations multi-échelles redondantes, c'est-à-dire qui intègrent plusieurs niveaux de description pour un même phénomène. Cette démarche pose la question de l'association des différents paradigmes de modélisation utilisés pour chacun des modèles.

Cette association peut être réalisée grâce à des méthodes de multi-modélisation. De nombreuses méthodes existent, telles que la transformation de modèles [de Lara et Vangheluwe, 2002], la composition de modèles hétérogènes [Hardebolle et Boulanger, 2007], l'adaptation de composants [Yellin et Strom, 1997] [Morel et Alexander, 2003] ou encore les méga-modèles [Bézivin et al., 2004] [Favre, 2006]. Le lecteur intéressé trouvera une revue détaillée de ces méthodes dans [Hardebolle, 2008]. Pour notre part, nous nous sommes concentrés sur deux approches qui nous semblent les plus pertinentes dans le contexte de la modélisation multi-paradigmes : l'intégration formelle et le couplage de simulateurs.

3.3.2 Intégration formelle

L'intégration formelle consiste à utiliser un formalisme unique pour tous les modèles. Deux solutions sont envisageables pour obtenir cette homogénéité. La première consiste à traduire chaque modèle dans un même formalisme, qui peut être différent de tous ceux déjà utilisés ou non. La seconde approche consiste à encapsuler les formalismes d'origine [Ramat, 2006] en développant des interfaces de traduction dans un formalisme commun. Cette opération est illustrée par la figure 3.9.

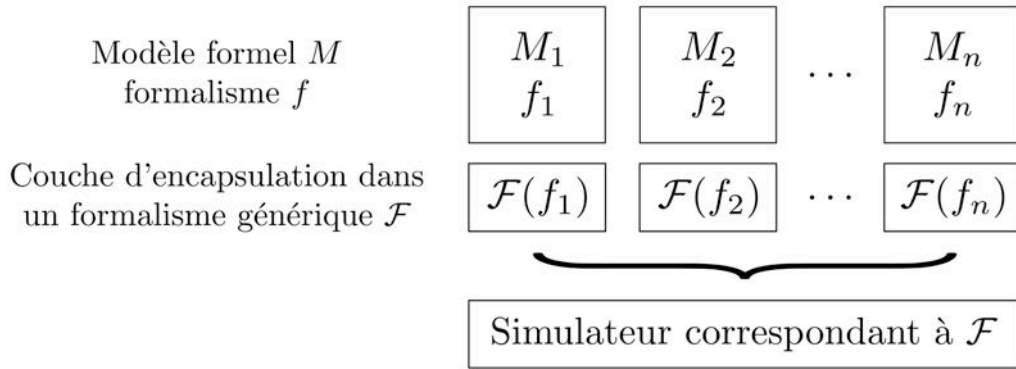


FIGURE 3.9 – **Encapsulation de formalismes** – Pour permettre la simulation d'un système multi-modèle, l'intégration formelle par encapsulation de formalisme consiste à utiliser une interface de traduction propre à chaque formalisme originel pour les rendre inter-opérationnels. Source : [Siebert, 2011]

Cette approche est notamment celle qui a été retenue pour le développement de la plateforme de multi-modélisation ModHel'X [Hardebolle et Boulanger, 2007].

Le choix du formalisme commun est alors dicté par un besoin de généricité. En effet, il doit être en mesure de faire abstraction de toutes les spécificités des paradigmes de modélisation, ce qui tend à le rapprocher d'une description en « boîte noire ». Le formalisme de multi-modélisation DEVS est celui qui se dégage dans la littérature [Fishwick et Zeigler, 1992] [Zeigler et al., 2000] [Quesnel et al., 2004] [Soulié, 2012].

Un modèle DEVS (*Discrete Event System Specification*) est représenté par un n-uplet $(X, Y, S, t_a, \delta_{int}, \delta_{ext}, \lambda, t_a)$ dans lequel X et Y représentent respectivement les ports d'entrée et de sortie du modèle, S représente l'ensemble des états du modèle, δ_{int} définit comment les états interne du modèle changent au cours du temps et δ_{ext} comment les événements externes changent les états internes du modèle et t_a définit la fonction d'avancement du temps de simulation. La figure 3.10a donne une vue schématique d'un modèle DEVS.

Plusieurs modèles DEVS peuvent être assemblés pour former un nouveau « modèle couplé » [Zeigler et al., 2000] avec ses ports d'entrées/sorties, sa dynamique interne et ses ports d'observation, comme le montre la figure 3.10b. Les travaux de [Duboz, 2004] et [Quesnel et al., 2009] donnent des exemples d'application du formalisme DEVS à travers le développement de la plateforme VLE (*Virtual Laboratory Environment*).

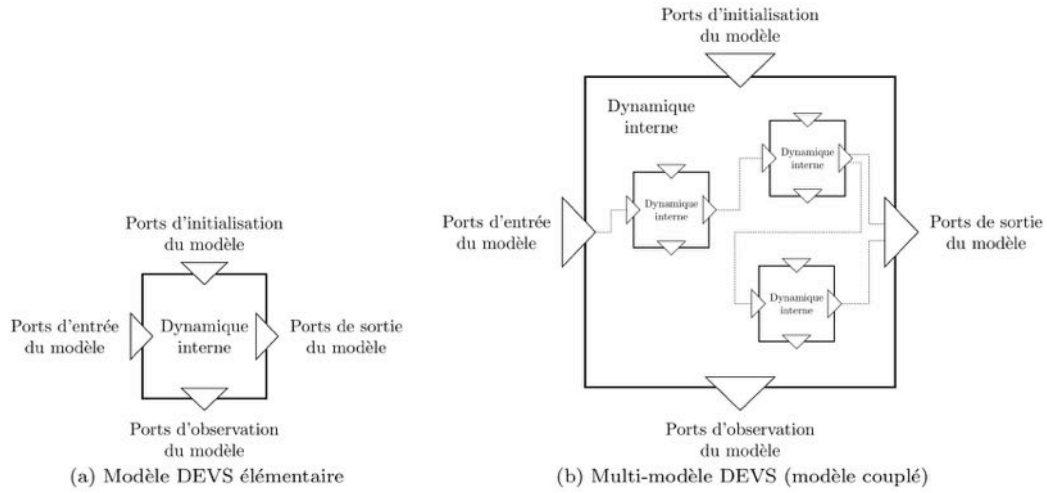


FIGURE 3.10 – **Formalisme DEVS et couplage** – a) Un modèle DEVS est une boîte qui possède des ports d’entrée et de sortie, une dynamique interne observable par des ports dédiés. b) Plusieurs modèles peuvent être couplés pour former un nouveau modèle DEVS englobant. Source : [Siebert, 2011]

3.3.3 Couplage de simulateurs et co-simulation

Une seconde approche de multi-modélisation consiste à conserver la structure d’origine de chaque modèle et de les simuler conjointement. Cette méthode est appelée la co-simulation (*co-operative simulation*). Elle est illustrée par la figure 3.11.

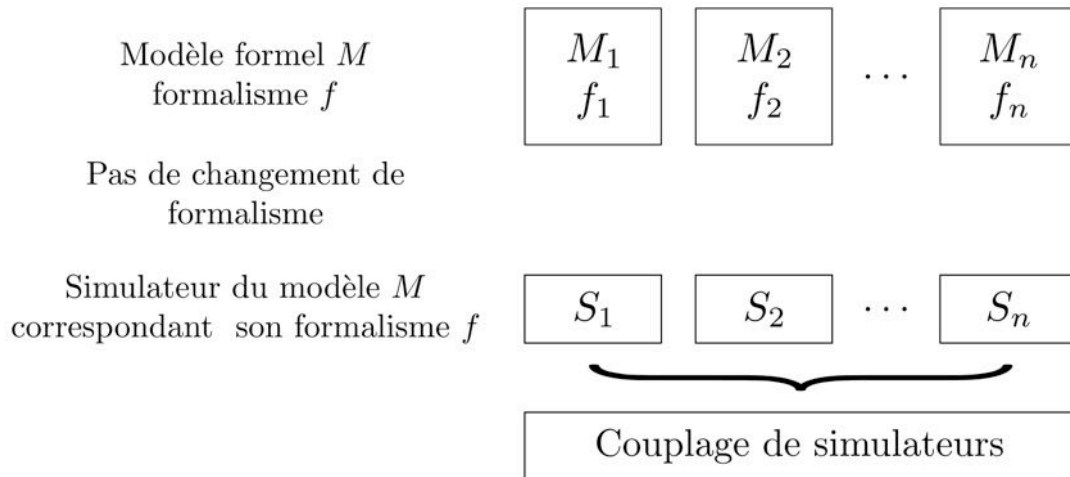


FIGURE 3.11 – **Co-simulation de modèles hétérogènes** – Chaque modèle M est exprimé dans son propre formalisme f . Le couplage de ces modèles intervient au niveau de leurs simulateurs S sans modification du modèle ou du formalisme. Source : [Siebert, 2011]

La co-simulation est une notion très utilisée en électronique [Lemarrec, 2000] [Ruelland, 2002]. Les méthodologies modernes de conception de systèmes électroniques peuvent induire l’utilisation de plusieurs langages dans la spécification des différentes parties du système.

Chaque sous-système, parfois un même sous-système, est décrit dans un langage qui lui est propre, par des équipes de conception séparées et à des niveaux d'abstraction matériels et logiciels différents. La co-simulation d'un système consiste en la simulation conjointe de l'ensemble de ses sous-systèmes. Elle apporte au concepteur la possibilité de construire son circuit en considérant les autres composants comme des boîtes noires communicantes. Nous pouvons d'ailleurs noter que la co-simulation est tout à fait compatible avec le formalisme DEVS présenté dans la section précédente, comme le montrent les travaux autour de la plateforme AA4MM (*Agents and Artifacts for Multi-Modelling*) [Siebert, 2011] [Camus et al., 2012].

Une interface de co-simulation fait le pont entre les différents composants en coordonnant les échanges de données calculées et interprétées par chaque sous-système. Cette interface recueille les données calculées par chacun des simulateurs et gère les échanges entre ces derniers. Les données échangées doivent être vérifiées pour qu'elles respectent les contraintes de types et de taille par exemple. Par ailleurs, comme nous le verrons dans la section suivante, la coordination des échanges de données conditionne fortement l'efficacité du couplage des différents simulateurs.

La principale différence entre la co-simulation et l'intégration formelle réside dans le fait que le couplage n'intervient plus au niveau des modèles mais au niveau des simulateurs. Ce point est particulièrement intéressant dans un contexte industriel où on doit souvent composer avec des simulateurs existants. La co-simulation présente également l'avantage de considérer chaque simulateur comme une sorte d'agent autonome du point de vue de sa conception et de son exécution. Cette méthode nous semble donc la méthode la plus adaptée pour mener à bien la simulation conjointe d'échelles de modélisation hétérogènes.

Dans les sections suivantes, nous détaillons les différents concepts liés à l'implémentation d'une plateforme de co-simulation, à commencer par les modèles de synchronisation des échanges entre simulateurs.

3.3.4 Synchronisation de co-simulation

La synchronisation est un élément essentiel de la co-simulation. En effet, elle doit définir la manière dont les données vont s'échanger entre les modèles. Nous allons nous intéresser aux deux principaux modèles de synchronisation : le modèle maître-esclave et le modèle distribué.

3.3.4.1 Modèle maître-esclave

Le modèle maître-esclave comprend un simulateur maître et un ou plusieurs simulateurs esclaves. Dans ce cas de figure, les simulateurs esclaves sont exécutés par le simulateur maître à la manière d'une fonction. Un simulateur esclave fonctionne alors comme une boîte noire et doit pour cela présenter des interfaces de programmation qui définissent ses entrées et sorties.

Bien que généralement assez facile à mettre en place, la co-simulation maître-esclave présente des inconvénients. Les échanges en lecture/écriture entre les deux simulateurs se font de façon alternée. Ce mode de synchronisation ne permet pas d'exécuter les simulateurs

de façon parallèle. Lorsque le simulateur maître s'exécute, le simulateur esclave est stoppé, et vice-versa. La figure 3.12a donne un exemple de cette procédure. Ce type de co-simulation est bien adapté pour des échanges synchronisés de données alors qu'il n'est pas très fonctionnel pour des applications basées sur le contrôle d'exécution.

Par ailleurs, le mode maître-esclave interdit le partage de simulateurs esclaves entre plusieurs simulateurs maîtres. La figure 3.13 illustre cet écueil. Quand le simulateur esclave S produit une même donnée D pour les deux simulateurs maîtres $M1$ et $M2$, la valeur de D ne pourra jamais être égale pour $M1$ et $M2$. En outre, en fonction de l'ordre d'appel de S par $M1$ et $M2$, ces derniers peuvent se retrouver dans un état indésirable et même indéfini si $M1$ et $M2$ sont exécutés en parallèle à des vitesses différentes.

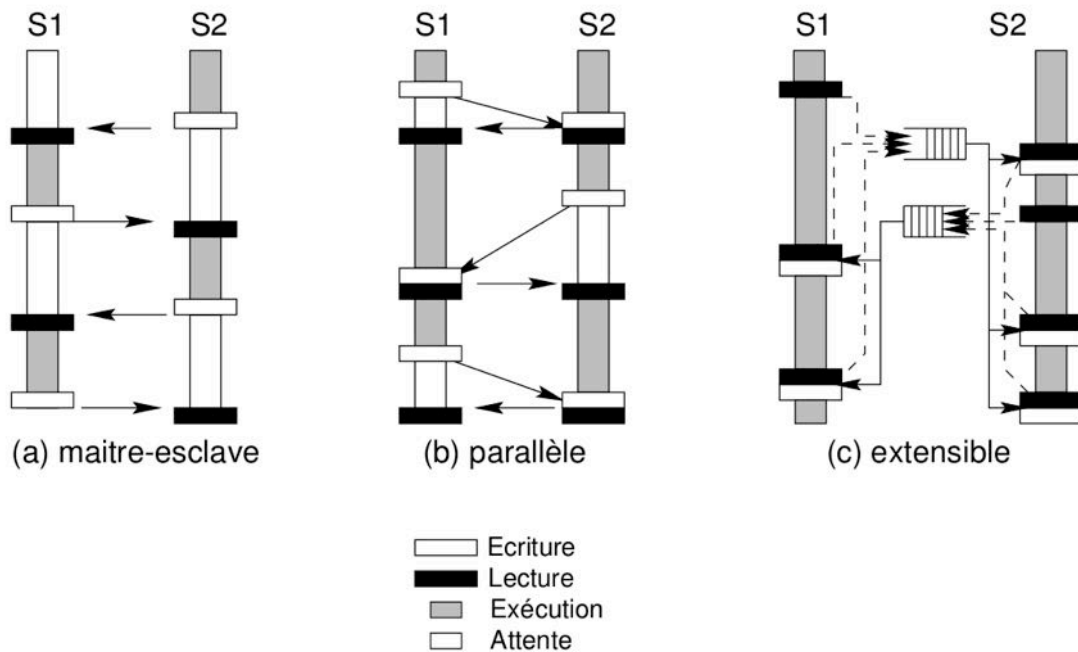


FIGURE 3.12 – **Diagrammes de communication des différents modèles de co-simulation** – Dans le mode « maître-esclave », les phases d'exécution et d'attente sont alternées entre les deux simulateurs $S1$ et $S2$. Le mode « parallèle » en (b) permet l'exécution des deux simulateurs simultanément mais un simulateur peut se mettre en attente de données de la part du second simulateur. Enfin, le mode « extensible » en (c) minimise le temps d'attente en mettant en place une communication asynchrone à travers un médium partagé.

3.3.4.2 Modèle parallèle ou distribué

Contrairement au mode maître-esclave, le mode de synchronisation distribué permet une exécution simultanée des simulateurs. Les échanges de données n'ont lieu que sur l'appel d'une procédure d'entrée/sortie (E/S) par un simulateur. Ainsi on s'assure que l'échange n'est effectué que si cela est nécessaire. Chaque simulateur envoie ses données et reste en attente pour la réception des données provenant de l'autre côté, comme le montre la figure 3.12b.

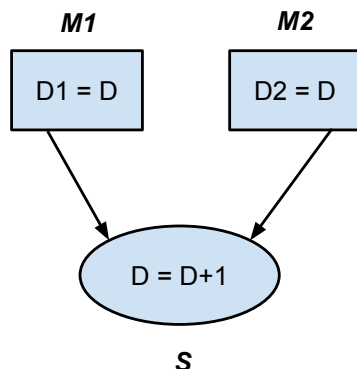


FIGURE 3.13 – **Illustration du problème de partage de données en mode « maître-esclave »** – Si à l’instant t , l’ordre d’appel est M1 - M2 et qu’à l’instant $t + 1$ l’ordre est inversé, les valeurs de D1 et D2 seront incrémentées deux fois au lieu d’une alternativement.

L’approche distribuée repose sur un bus de co-simulation qui est utilisé comme protocole de communication. Chaque simulateur dispose de procédures d’E/S qui se chargent d’extraire les données du simulateur et les communiquent au bus de co-simulation, plutôt que d’appeler directement un autre simulateur en tant qu’esclave. Ainsi chaque simulateur agit en tant que maître.

Ce modèle de co-simulation comporte plusieurs avantages. Tout d’abord, le modèle distribué autorise à la fois une flexibilité dans la conception de chacun des composants mais également une certaine modularité du simulateur global. En effet, en gardant la possibilité d’interchanger ou d’ajouter de nouveaux éléments, il n’est pas nécessaire de remettre en cause le développement de chacun des sous-simulateurs.

Ensuite, l’exécution parallèle des simulateurs peut offrir un gain de performance, notamment en distribuant le calcul sur plusieurs processeurs ou plusieurs machines. Bien entendu, ce gain ne peut être significatif que si les simulateurs sont faiblement couplés. Dans certains cas, le modèle extensible, une variante du modèle distribué, pourra être utilisé. Ce mode, illustré par la figure 3.12c, permet d’optimiser le parallélisme par le biais d’un moyen de stockage des données échangées. Ainsi, les simulateurs lisent et écrivent à leur convenance, tout au long de leur exécution, dans le bus de co-simulation et ne restent pas en attente de données. Cependant, ce mode de synchronisation introduit de nouvelles contraintes au niveau du protocole de communication.

Ceci met en évidence que la communication entre les simulateurs est la clé d’une co-simulation efficace. Le bus de co-simulation est à la fois le support de cette communication et le chef d’orchestre de ces échanges. Nous allons présenter ses principales caractéristiques.

3.3.5 Bus de co-simulation

En terme d’architecture informatique, un bus est un système de communication partagé entre différents modules, qu’ils soient matériels ou logiciels, par opposition à une liaison

point à point. Il existe une multitude d'implémentations de bus de communication, que l'on retrouve dans tous les systèmes numériques qui nous entourent, de l'automobile au système d'exploitation. Cette expression couvre à la fois la composante matérielle de la connexion (câbles, etc.) et la composante logicielle qui définit entre autres le protocole de communication. Dans le cadre de nos travaux, nous nous limiterons à l'étude de cette dernière.

Dans le contexte de la co-simulation, le bus est en charge de faire le lien entre les différents simulateurs et d'organiser les échanges de données entre ces derniers, comme le montre la figure 3.14. Ces deux rôles sont assurés par deux éléments : le support de communication et le contrôleur de données.

Support de communication

Le support de communication désigne ici l'élément logiciel qui va permettre de faire transiter les données entre les simulateurs. Plusieurs solutions existent pour mettre en place cette communication, chacune disposant de ses avantages et de ses inconvénients. Par exemple, une mémoire locale partagée a l'avantage d'être très efficace au regard du temps de transfert mais interdit la distribution des simulateurs sur plusieurs machines. Pour une communication réseau, la solution la plus largement répandue est le *Socket*. Il s'agit d'une interface logicielle avec les services réseau du système d'exploitation sur laquelle vient se brancher le contrôleur de données.

Contrôleur de données

Le contrôleur de données coordonne les échanges d'informations entre les simulateurs de l'environnement. En fonction du modèle de synchronisation employé au niveau des simulateurs, le contrôleur devra router les données de l'expéditeur au destinataire, principalement dans le cas d'une communication sur *socket*. Pour ce faire, des services de communication peuvent être utilisés à différents niveaux.

Les services de communication haut niveau, comme CORBA [Gransart et Geib, 1999] ou DDS [Pardo-Castellote, 2003], ont l'avantage de simplifier l'échange de données entre modules car ils offrent une abstraction de la communication. Ainsi, un simulateur pourra être déployé sur une machine distante sans que son implémentation soit remise en cause. Cependant, cette abstraction entraîne un surcoût des temps de communication puisqu'ils reposent sur des services de plus bas niveaux, comme TCP ou UDP, autrement plus rapides, bien que beaucoup moins flexibles. Le choix du mode de communication se fera donc en fonction des besoins en termes de performance et de flexibilité.

Ensuite, par définition, le destinataire n'a pas besoin d'être spécifié. La donnée est donc disponible à tous les simulateurs connectés au bus. Par ailleurs, la communication peut ne pas être synchronisée entre les différents simulateurs. En effet, dans le cas d'une exécution parallèle des simulateurs connectés au bus, chacun d'eux peut à tout moment amorcer une procédure d'E/S. Le bus de co-simulation doit alors également gérer les conflits d'accès aux données au sein de la mémoire. Si deux simulateurs accèdent simultanément à une donnée en lecture, il n'y a pas de problème. En revanche, si l'un lit une donnée dans le bus de communication pendant que l'autre remplace la valeur de cette donnée, le premier simulateur pourra éventuellement hériter d'une donnée corrompue. Des mécanismes d'exclusion mutuelle résolvent généralement ce genre de problème en programmation parallèle [Crépin, 2013].

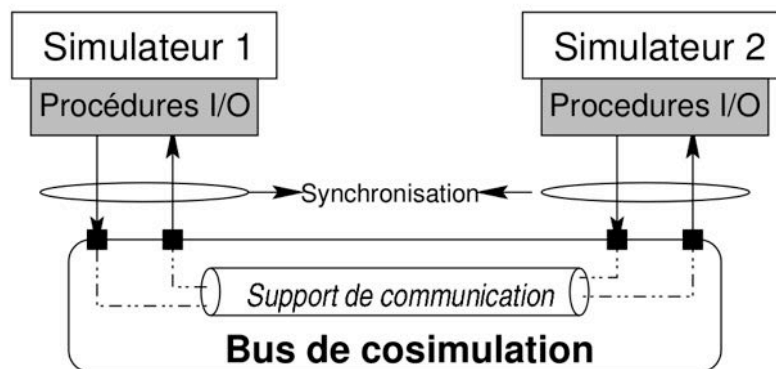


FIGURE 3.14 – **Modèle générique de co-simulation distribuée** – Chaque simulateur communique avec le bus de co-simulation à travers des primitives d’entrées-sorties. Ces primitives sont des appels externes au simulateur et sont réalisées par chaque simulateur de façon spécifique. Le bus de co-simulation est en charge du transfert des données entre les différents simulateurs.

Nous avons présenté dans cette section la co-simulation telle qu’elle est utilisée pour la conception de systèmes électroniques. Nous avons vu que cette technique permet d’interconnecter des simulateurs autonomes en termes d’exécution qui gardent la possibilité d’échanger des données. De plus, sans être nommée de la sorte, cette approche a déjà été appliquée dans le cadre de simulations multi-niveaux hybrides de flux de trafic routier [El Hmam et al., 2008] [Néron et al., 2012]. Elle nous semble donc appropriée pour la méthode de simulation redondante de modèles hétérogènes que nous proposons.

3.4 Synthèse

La simulation interactive de systèmes complexes passe par l’utilisation de modèles phénoménologiques. Ceux-ci nous imposent de nombreux paramètres pour lesquels il est parfois difficile de fixer des valeurs. Nous avons montré au début de ce chapitre que ces paramètres sont en fait les traces de l’abstraction de phénomènes sous-jacents. En d’autres termes, il est possible de calculer ces paramètres en utilisant des modèles de bas niveau que nous avons écarté auparavant à cause de leur coût computationnel élevé. Il s’agit donc de combiner les forces des deux points de vue au moyen de méthodes multi-échelles.

Nous avons décrit plusieurs méthodes similaires qui proposent de combiner différents niveaux de description pour un même système. Particulièrement, les méthodes multi-échelles hétérogènes introduisent le concept de simulations microscopiques redondantes pour appuyer un modèle macroscopique incomplet. Cependant, ces méthodes de multi-modélisation partagent le même défaut : il n’existe à ce jour pas d’implémentation générique qui permette d’appliquer directement ces méthodes à tous types de problèmes.

Ce manque est principalement dû à l’élément clé de ces méthodes qui consiste à définir des opérateurs de traduction entre les niveaux de description que l’on souhaite coupler. Or, si

nous disposons souvent de plusieurs modèles pour un même phénomène ou un même système, il n'est pas toujours facile de faire correspondre des données de l'un à l'autre. Il l'est encore moins de formaliser ce transfert sans introduire des connaissances spécifiques au domaine considéré. Le développement de ces opérateurs doit donc nécessairement être fait de manière *ad hoc*, application par application.

À défaut de pouvoir proposer une implémentation générique, nous souhaitons nous concentrer sur la démarche qui, elle, est commune à toutes ces méthodes et à tous les cas d'application. Nous cherchons donc dans la suite de ce mémoire à développer des outils qui favorisent l'implémentation de simulations redondantes.

Pour ce faire, nous nous sommes intéressés à la co-simulation qui nous offre deux avantages. Premièrement, cette technique nous autorise à coupler des modèles d'origines diverses, exprimés dans des formalismes différents. Ainsi, la transformation de simulateurs existants se limite à la définition d'interfaces de communication. Nous verrons cependant que ce point reste la difficulté principale dans le contexte de simulations multi-échelles redondantes. Deuxièmement, la co-simulation se rapproche en quelque sorte de l'expérimentation *in virtuo* en garantissant les autonomies de conception et d'exécution des différents simulateurs.

Nous proposons donc dans le chapitre suivant d'adapter la technique de la co-simulation à travers une architecture logicielle qui permet la simulation redondante d'échelles de modélisation hétérogènes.

Chapitre 4

Simulations redondantes d'échelles de modélisation hétérogènes



Philippe Geluck - LE TOUR DU CHAT EN 365 JOURS

Nous avons proposé dans le chapitre 2 une approche phénoménologique pour la simulation des systèmes complexes à l'aide des systèmes multi-agents qui implémentent des lois descriptives. Celles-ci tendent à décrire, avec plus ou moins d'empirisme, un système constitué de nombreux éléments par des variables agrégées (des concentrations, des densités de population, etc.) et s'intéressent à la dynamique de ces variables. Nous obtenons ainsi une représentation très compacte et très générale de l'évolution du système, autrement plus rapide à simuler qu'un modèle individu-centré qui considère chaque entité comme autonome. Le problème est que l'identification des paramètres agrégés des modèles macroscopiques peut parfois s'avérer délicate.

Dans le chapitre précédent, nous avons exposé un ensemble de méthodes qui s'inscrivent dans une démarche de simulations redondantes. Le principe général est d'utiliser plusieurs niveaux d'abstraction d'un même système dans le but de profiter des avantages de chacun d'eux, à savoir la vitesse d'un modèle macroscopique et la précision d'un modèle microscopique. Nous nous proposons d'adopter cette méthodologie dans le but de corriger les écueils de notre approche multi-agents phénoménologique.

Pour illustrer concrètement cette idée, nous considérons d'abord un exemple simple : le phénomène de diffusion d'une espèce chimique dans un fluide. Nous proposerons ensuite

une architecture de simulations redondantes qui s'inspire des techniques de co-simulation présentées dans le chapitre précédent.

4.1 Exemple : le phénomène de diffusion

Le phénomène de diffusion est l'un de ceux qu'on trouve généralement dans la littérature pour démontrer la possibilité de modéliser un même phénomène à des échelles différentes. À notre tour, nous l'utilisons dans cette section pour mettre en avant l'intérêt de la démarche de co-simulation de niveaux de descriptions différents que nous proposons. L'exemple développé ici est une adaptation de celui utilisé par R. Duboz pour illustrer le formalisme DEVS [Duboz, 2004].

Nous commençons par donner quelques informations théoriques et nos choix d'implémentations par les systèmes multi-agents, pour chacun des deux modèles que nous considérons. Nous expliquons ensuite comment tirer avantage de cette double modélisation d'un même phénomène, notamment pour l'identification hors-ligne, c'est-à-dire en amont de la simulation, du coefficient de diffusion nécessaire au modèle macroscopique. Enfin, nous montrons qu'il est également judicieux de suivre une telle procédure en cours de simulation, pour assurer un paramétrage adéquat du modèle macroscopique alors que son niveau d'abstraction ne permet pas d'expliciter certains phénomènes individuels.

4.1.1 Modèle macroscopique

Le phénomène de diffusion dans un milieu désigne le transport de matière qui tend à uniformiser les concentrations des espèces chimiques en présence. À l'échelle macroscopique, l'étude de ce phénomène par la théorie cinétique a permis d'établir les lois de Fick [Fick, 1855] qui constituent la base de notre modèle agrégé.

4.1.1.1 Éléments théoriques

Première loi de Fick

La première loi de Fick indique qu'un gradient de concentration d'une substance dissoute dans une solution induit un flux de matière proportionnel mais de sens opposé à celui du gradient, le flux étant la quantité de matière qui se déplace par unité de surface et par unité de temps. La constante de proportionnalité est donnée par le coefficient de diffusion de la substance. L'expression de la première loi de Fick, pour la dimension 1, est donnée par l'équation aux dérivées partielles

$$J_x(x, t) = -D \frac{\partial C}{\partial x}(x, t) \quad (4.1)$$

où :

- J_x est le flux de matière selon l'axe des x en $\text{mol} \cdot \text{m}^{-2} \cdot \text{s}^{-1}$,
- D est le coefficient de diffusion de la substance en $\text{m}^2 \cdot \text{s}^{-1}$,

- $C(x, t)$ la concentration de la substance en l'abscisse x à l'instant t en $\text{mol} \cdot \text{m}^{-3}$.

Seconde loi de Fick

La seconde loi de Fick s'obtient à partir de la précédente, sur l'hypothèse de la conservation de la matière. Considérons un volume délimité par deux surfaces d'aires S et séparées d'une distance δx . Le flux de matière se propage de la surface à l'abscisse x vers la surface à l'abscisse $x + \delta x$. S'il y a conservation de la matière, la variation de la concentration au sein du volume durant l'intervalle de temps δt est donnée par la différence entre le flux entrant et le flux sortant, soit

$$\frac{C(t + \delta t) - C(t)}{\delta t} = - \frac{J_x(x + \delta x) - J_x(x)}{\delta x} \quad (4.2)$$

A la limite $\delta t \rightarrow 0$ et $\delta x \rightarrow 0$, l'équation devient

$$\frac{\partial C}{\partial t} = - \frac{\partial J_x}{\partial x} \quad (4.3)$$

En insérant l'équation (4.1) dans la précédente, nous obtenons

$$\frac{\partial C}{\partial t}(x, t) = D \frac{\partial^2 C}{\partial x^2}(x, t) \quad (4.4)$$

qui est l'équation de diffusion.

Appliquée à la diffusion membranaire et approximée linéairement, nous obtenons la loi suivante

$$\frac{\partial C}{\partial t}(x, t) = \frac{D \cdot S}{L} \cdot \Delta C \quad (4.5)$$

avec

- D le coefficient de diffusion de l'espèce chimique dans le milieu donné en $\text{m}^2 \cdot \text{s}^{-1}$,
- S la surface d'échange en m^2 ,
- ΔC le gradient de concentration en mol ,
- L l'épaisseur de la membrane en m .

4.1.1.2 Implémentation

Le phénomène de diffusion représente le transfert de concentration de particules entre deux milieux. Par conséquent, nous choisissons de nous baser sur la première loi de Fick et de le modéliser par un agent-interaction qui considère deux compartiments. A chaque pas de temps, cet agent est chargé de scruter les concentrations de particules au sein de ceux-ci. S'il constate un déséquilibre, il calcule la quantité de matière à déplacer d'un compartiment à l'autre en fonction du coefficient de diffusion.

Si l'on considère deux compartiments A et B et la diffusion de C entre eux, les variations de concentrations seront les suivantes (en considérant un schéma d'intégration d'Euler explicite) :

$$\begin{cases} [C]_A^{n+1} = [C]_A^n - h_n d \\ [C]_B^{n+1} = [C]_B^n + h_n d \end{cases} \quad (4.6)$$

où d est calculé de la manière suivante par l'agent-interaction :

$$d = \frac{D \cdot S \cdot ([C]_A^n - [C]_B^n)}{L} \quad (4.7)$$

Ainsi, il est possible de chaîner des compartiments entre lesquels on place des agents-interaction. La figure 4.1a montre la discrétisation spatiale en tranche d'un cube d'eau au centre duquel nous plaçons une mole de caféine à l'instant initial.

4.1.2 Modèle microscopique

Le phénomène de diffusion à l'échelle microscopique est généralement décrit à l'aide d'un processus stochastique, le mouvement brownien [Brown, 1902]. Le modèle que nous présentons ici est volontairement très simplifié dans le but de nous concentrer sur le principe de co-simulation.

4.1.2.1 Éléments théoriques

Le mouvement brownien caractérise le déplacement aléatoire de particules en suspension dans un fluide. Il est provoqué par les collisions des particules considérées avec les molécules du fluide exposées à l'excitation thermique.

La succession des déplacements aléatoires conduit à un processus de diffusion dont le coefficient est donné par la loi de Stokes-Einstein, dans le cas de particules sphériques :

$$D = \frac{k_B T}{6\pi\eta R} \quad (4.8)$$

où :

- k_B est la constante de Boltzmann,
- T est la température,
- η est la viscosité du fluide,
- R est le rayon de la particule.

Cette équation nous donne donc un coefficient de diffusion dans des conditions « idéales », c'est-à-dire qui ne prend potentiellement pas en compte certains éléments de l'environnement comme les interactions avec les autres molécules en présence ou la géométrie de l'enceinte.

Une autre approche pour la modélisation du mouvement brownien consiste à se dire qu'à chaque instant, une particule donnée possède une vitesse instantanée bornée entre 0 m/s et sa vitesse terminale. La vitesse terminale d'une particule dans un fluide résulte du bilan des forces de gravité, d'Archimède et de traînée et peut être calculée par

$$v_{max} = \sqrt{\frac{2mg}{\rho AC_w}} \quad (4.9)$$

où

- m est le poids de la particule,
- g est l'accélération de la pesanteur,
- ρ est la viscosité du fluide,
- A est la surface projetée de la particule,
- C_w est le coefficient de traînée.

4.1.2.2 Implémentation

Pour implémenter le mouvement brownien, nous considérons un ensemble d'agents-entité qui représentent les particules se déplaçant sur un axe. Les agents sont tous placés initialement en $x = 0$. À chaque pas de temps, un agent donné prend une vitesse aléatoire, entre $-v_{max}$ et v_{max} , qu'il intègre pour déterminer sa nouvelle position. Ce modèle donne la position de toutes les particules à chaque instant, comme le montre la figure 4.1b. Il permet de calculer la concentration en particules en toute partie de l'espace, mais aussi éventuellement d'autres grandeurs, comme la distance moyenne parcourue par les particules par exemple.

4.1.3 Détermination du coefficient de diffusion

Grâce aux informations de la simulation microscopique, nous pouvons calculer une estimation du coefficient de diffusion. Pour cela, nous mesurons à chaque instant t l'écart quadratique moyen $\bar{\sigma}$ de l'ensemble des particules à l'aide de l'équation suivante :

$$\bar{\sigma} = \frac{1}{n} \sqrt{\sum_{i=1}^n x_i^2} \quad (4.10)$$

Nous pouvons relier $\bar{\sigma}$ et D :

$$\bar{\sigma} = \sqrt{2Dt} \quad (4.11)$$

Connaissant $\bar{\sigma}$ à un instant donné, nous pouvons donc en déduire une constante D instantanée :

$$D = \frac{1}{2t} \bar{\sigma}^2 \quad (4.12)$$

Pour notre expérience, nous considérons la diffusion de molécules de caféine, pour laquelle le coefficient D est connu, dans l'eau. De cette manière, nous pouvons évaluer la justesse

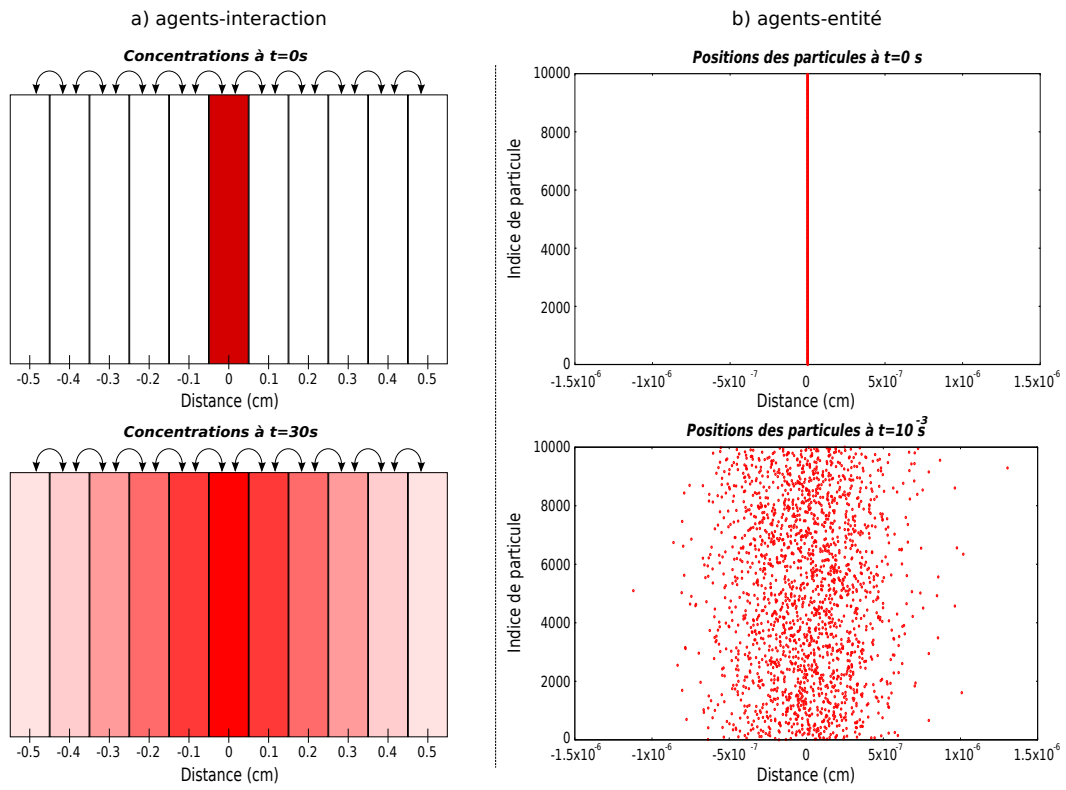


FIGURE 4.1 – **Simulation de deux niveaux de description du phénomène de diffusion** – À gauche, le modèle phénoménologique à base d'agents-interaction manipule des concentrations entre des compartiments homogènes. À droite, le modèle microscopique, à base d'agents-entité, calcule le mouvement brownien de toutes les particules.

de notre simulation individus-centrée. La simulation microscopique est réalisée avec 10^5 particules pour une durée simulée de 30s sur un axe mesurant 1cm.

La figure 4.2 montre l'évolution de D en fonction du temps. En calculant la pente de cette droite, nous obtenons le coefficient de diffusion. Ce coefficient est injecté dans l'équation (4.7) du modèle macroscopique. La figure 4.3 permet de comparer visuellement les deux niveaux de description et nous constatons que la simulation individus-centrée s'ajuste très bien avec la loi de Fick.

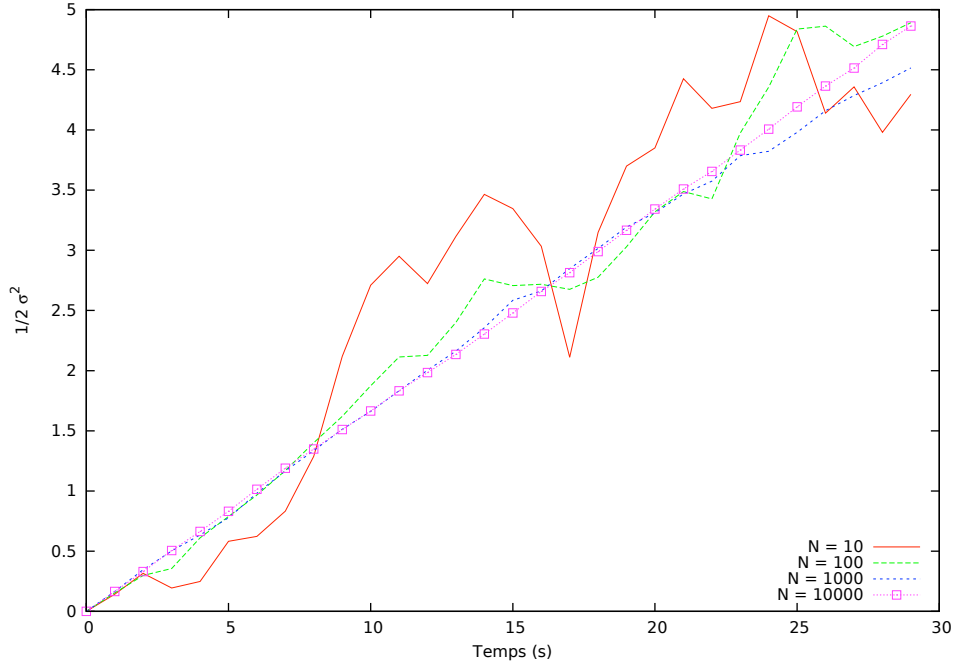


FIGURE 4.2 – **Mesure du coefficient de diffusion de la caféine** – L'évolution de l'écart quadratique de l'ensemble des particules forme globalement une droite dès lors que le nombre de particules simulées est suffisamment grand. La pente de cette droite nous donne une estimation du coefficient de diffusion.

Ce premier exemple nous a permis de montrer qu'il est possible de déterminer *a priori* la valeur de paramètres d'un modèle agrégé grâce à la simulation d'un niveau de description inférieur. Le modèle individus-centré sert ici de laboratoire virtuel dans lequel nous pouvons manipuler finement l'ensemble des propriétés du système et évaluer leurs impacts sur son comportement global. Nous allons montrer dans la section suivante qu'il est également possible d'utiliser cette technique de co-simulation pour le paramétrage en ligne de lois descriptives.

4.1.4 Détermination de paramètres en cours de simulation

Puisque les lois descriptives sont formulées à partir d'observations, nous disposons généralement de valeurs, ou de tables, pour les paramétrer. Cependant, ces valeurs relèvent le plus souvent d'estimations ou de mesures faites pour des cas particuliers. Or nous savons que

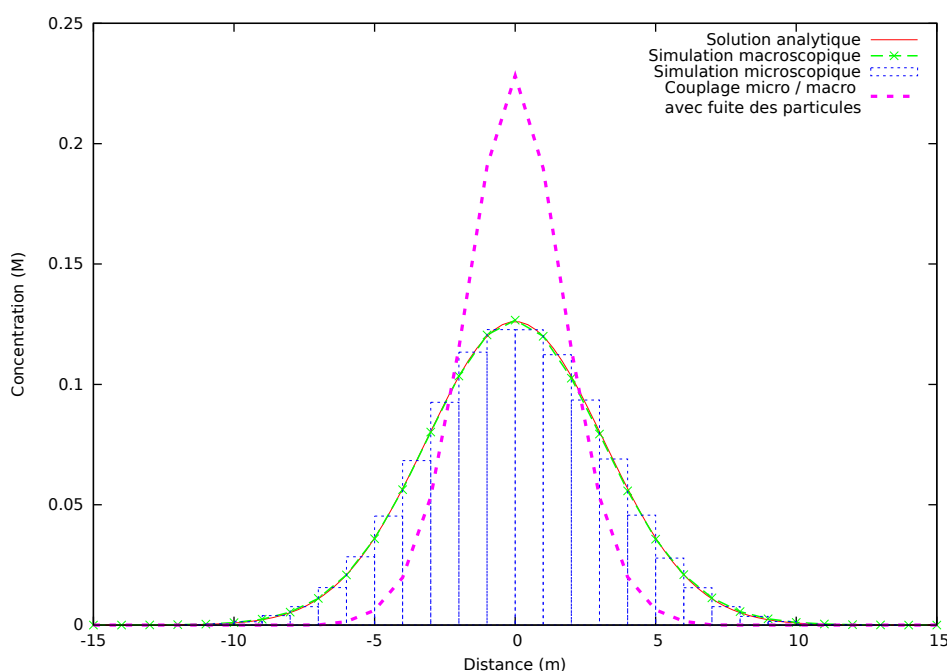


FIGURE 4.3 – **Comparaison des distributions de concentration** – Le phénomène de diffusion est simulé à deux niveaux de description différent. Nous constatons une bonne correspondance de la distribution des concentrations en fin de simulation entre les deux échelles. Par ailleurs, nous observons un changement de dynamique important du modèle macroscopique avec l’ajout d’une réaction de fuite entre particules au niveau microscopique.

les systèmes que nous cherchons à simuler peuvent s’éloigner de ces situations bien définies en raison de l’intrication des différents phénomènes.

Par exemple, l’équation (4.8) nous permet de calculer le coefficient de diffusion théorique d’une particule, en ne prenant en compte l’environnement que par la viscosité du fluide dans lequel elle est immergée. Toute influence des autres particules potentiellement présentes, de même nature ou non, est négligée. En réalité, certains phénomènes peuvent venir perturber localement la diffusion des particules et ainsi rendre le coefficient D variable dans le temps et l’espace. De par son niveau d’abstraction, notre modèle macroscopique n’est pas en mesure de prendre en compte ces phénomènes qui agissent à l’échelle individuelle. En revanche, ces effets spatiaux peuvent être intégrés à notre modèle microscopique. Il s’agit donc de coupler les modèles, non plus préalablement à la simulation du modèle macroscopique, mais tout au long de cette dernière.

Nous illustrons cette problématique en ajoutant à notre modèle un phénomène simple de fuite des particules les unes par rapport aux autres. Cette réaction à la concentration peut être modélisée par une équation de Monod qui fait augmenter la vitesse des particules en fonction de la concentration jusqu’à un maximum fixe :

$$v = v_{max} - v_{max} \exp^{-bC(x,t)} \quad (4.13)$$

où v est la vitesse des particules, v_{max} la vitesse maximale et $C(x,t)$ la concentration

courante.

Nous simulons comme précédemment la diffusion à l'aide du modèle individus-centré pour déterminer, à chaque pas de temps, et pour chaque tranche du cube, la valeur locale de D . À la différence de notre première expérience, le paramétrage de la simulation microscopique est imposé par l'état courant du modèle macroscopique et notamment par la valeur de la concentration locale. Les valeurs des coefficients sont remontées au modèle macroscopique et les agents-interaction mettent à jour les concentrations dans chaque compartiment. Ces nouvelles concentrations seront utilisées au pas de temps suivant pour paramétrer une nouvelle simulation microscopique.

La figure 4.3 montre qu'avec ce nouveau phénomène les particules n'atteignent pas le bord du cube pour une même durée de simulation. Il n'est pas question ici de faire une étude précise de ce phénomène mais de montrer que la co-simulation que nous proposons permet l'utilisation de modèles descriptifs dont le paramétrage est dynamique et peut s'adapter à la complexité du système.

Tout au long de cette section, nous avons illustré les grandes lignes de la démarche que nous souhaitons mettre en œuvre pour faciliter la simulation de systèmes complexes par des lois descriptives. Pour ce faire, nous réalisons des simulations redondantes du même système à un niveau de description inférieur pour calibrer la simulation macroscopique. Ce procédé s'inspire des méthodes multi-niveaux présentées dans le chapitre 3, et plus particulièrement des *Heterogeneous Multiscale Methods* dans la mesure où nous disposons d'un modèle macroscopique. Nous avons vu qu'il n'existe à ce jour pas d'implémentation générique de ces méthodes, notamment en raison de la nature inévitablement *ad hoc* des opérateurs de traduction entre échelles. En revanche, nous constatons que la démarche reste, elle, identique quelque soit le système étudié. À défaut de pouvoir développer une implémentation universelle, nous proposons dans la section suivante une architecture logicielle qui vise à automatiser la conception de simulateurs multi-modèles redondants.

4.2 Architecture pour la co-simulation redondante d'échelles de modélisation hétérogènes

Dans l'optique de pallier les problèmes liés au paramétrage des modèles descriptifs, nous suggérons de simuler conjointement, et de manière redondante, différentes échelles de modélisation d'un même système. L'exemple du phénomène de diffusion que nous avons développé dans la section précédente nous a permis d'illustrer cette démarche que nous désirons généraliser. Nous allons à présent exposer étape par étape l'implémentation de la méthode que nous proposons, en commençant par donner une description générale de son architecture.

4.2.1 Vue générale

Nous souhaitons construire une architecture logicielle qui généralise la méthode de simulation redondante d'échelles de modélisation hétérogènes. Ce type d'opération tend généralement à imposer un formalisme commun à tous les modèles, ou à définir un nouveau méta-modèle. Ces approches présentent indéniablement des avantages, notamment pour la compréhension du modèle et son contrôle. Pour autant, nous nous plaçons dans un contexte industriel et, de ce fait, nous sommes confrontés à la réalité des modèles, d'une grande variété de formalismes et issus de domaines très différents. Notre ligne directrice consiste donc plutôt à greffer de nouveaux outils autour de ces modèles existants pour ne pas remettre en cause leur structure générale.

Notre objectif est d'obtenir une simulation d'un système complexe qui allie précision et temps de calcul raisonnable. Nous avons vu que pour y parvenir, nous devons recourir à des modèles descriptifs de haut niveau qui nécessitent un paramétrage rigoureux et adaptatif. La qualité du paramétrage est une problématique récurrente dans le domaine de la simulation des systèmes complexes du fait de l'intrication des phénomènes et des mécanismes de rétro-action aux conséquences imprévisibles. En effet, en sus de la difficulté de choisir *a priori* des paramètres qui rendent compte de plusieurs phénomènes sous-jacents, il est parfois nécessaire d'adapter leurs valeurs en fonction de l'évolution générale du système. La première étape consiste donc à intégrer dans nos modèles phénoménologiques un outil en charge de détecter les défauts de paramétrage en cours de simulation. Nous décrirons donc dans un premier temps quels peuvent être ces défauts et comment nous proposons de les monitorer à l'aide d'un agent-*watchdog*.

Pour corriger ces défauts, notre approche consiste à prendre appui sur des modèles auxiliaires dont les résultats de simulation nous permettent de déterminer de nouvelles valeurs de paramètres. Cette opération constitue donc un couplage que nous souhaitons le plus faible possible. Dans cette idée, la technique de la co-simulation présentée dans le chapitre 3 nous semble particulièrement adaptée. Elle consiste à mettre en relation des sous-systèmes, qui sont conçus, développés et exécutés de manière autonome et distribuée avec une approche dite « boîte noire », au moyen d'un bus de communication, comme illustré par la figure 3.14. Ce bus, dit de co-simulation, sert de support aux échanges de données entre les modèles. Les données acheminées par le bus sont, d'une part, des requêtes de paramétrage en provenance de l'agent-*watchdog*, d'autre part, les réponses des modèles auxiliaires à ces requêtes. Nous détaillerons dans la suite ce protocole de communication et comment les données sont traitées par chacun des modèles.

Pour terminer, nous discuterons des conditions d'efficacité de notre méthode, notamment en termes de précision et de temps de calcul. Plus particulièrement, nous identifierons différentes règles pour la simulation du modèle auxiliaire.

4.2.2 Mécanisme de contrôle

La première étape de l'implémentation de notre architecture logicielle consiste à ajouter à nos modèles phénoménologiques un mécanisme de contrôle capable de détecter les éventuels

problèmes de paramétrage. Nous considérons la simulation d'un modèle phénoménologique, tel que nous l'avons décrit dans le chapitre 2 et représenté sur la figure 2.15. Dans ce contexte, nous distinguons quatre situations particulières dans lesquelles un défaut de paramétrage peut être capturé.

Absence de modèle/de données

Tout d'abord, nous ne disposons pas toujours de modèle pour un phénomène à l'échelle macroscopique. C'est typiquement le genre de problèmes traités à l'aide de l'approche « sans-équation » décrite dans la section 3.2.4. Ainsi, la démarche consiste à utiliser, à intervalles réguliers ou non, des données de simulations microscopiques pour identifier une trajectoire du modèle microscopique. Cette trajectoire est ensuite intégrée temporellement par le schéma d'intégration numérique du modèle macroscopique.

Conditions aux limites dynamiques

La deuxième situation se présente quand les conditions aux limites du problème varient au cours de la simulation. En effet, pour ce qui est de la simulation de systèmes complexes, il est fréquent d'observer que la réalisation d'un phénomène affecte le périmètre d'activité d'un autre phénomène. Cette caractéristique se manifeste d'autant plus que nous souhaitons mettre en œuvre des simulations interactives, c'est-à-dire qu'elles font intervenir un utilisateur qui peut agir de manière imprévisible. De ce fait, un phénomène doit être en mesure d'inspecter les données qu'il manipule pour vérifier qu'elles restent compatibles avec son exécution, faute de quoi il devra s'adapter pour tenir compte des nouvelles conditions aux limites.

Divergence

Ensuite, l'état de l'art a montré que les modèles à base d'agents-interaction sont le plus souvent simulés à l'aide de schémas d'intégration numérique. Ces derniers peuvent induire des erreurs qui se propagent d'un pas de temps sur l'autre et qui font s'éloigner la simulation de la solution exacte. Ces erreurs peuvent avoir deux origines : soit elles sont causées par un défaut de précision de la méthode numérique utilisée, soit elles proviennent d'inexactitudes dans le modèle. Nous faisons l'hypothèse suivante : si une simulation ne donne pas de résultats satisfaisants alors que sa convergence a été démontrée pour un jeu de paramètres donné, il est raisonnable d'imputer cette divergence aux paramètres du modèle qu'il convient de réévaluer. Nous supposons également que la variation des paramètres est relativement faible pour ne pas sortir du domaine de convergence du schéma numérique. Par exemple, la convergence d'une simulation numérique d'un système masse-ressort est fortement dépendante du couple de paramètres que forment le pas de temps de résolution et la raideur du système. Pour des variations faibles de la raideur, le schéma restera stable tandis que des variations fortes causeront sa divergence.

Expertise métier

Enfin, la dernière situation intervient lorsque le modèle sort de son domaine de validité. Il arrive que l'expression d'un phénomène à l'échelle macroscopique lui autorise à être simulable mathématiquement alors que les conditions expérimentales sont irréalistes. Cette situation se produit particulièrement quand un phénomène n'est que partiellement connu et est décrit à l'aide d'une loi plus ou moins empirique. Dans ces conditions, seule l'expertise métier permet de déterminer si les paramètres d'un phénomène sont bien situés dans une plage acceptable.

Ces quatre situations sont autant de points de contrôle que nous pouvons analyser tout au long de la simulation du modèle macroscopique.

Dans l'optique d'une moindre modification des modèles existants et d'une intégration aisée et modulaire, le mécanisme de contrôle que nous proposons prend la forme d'un agent-*watchdog*, chargé de scruter la simulation du modèle.

En informatique, un *watchdog*, ou chien de garde, est un processus généralement exécuté en tâche de fond. Son but est de s'assurer du bon fonctionnement d'un système et de prendre des mesures le cas échéant, comme par exemple son redémarrage.

Dans notre cas, le *watchdog* est donc un agent autonome qui a une perception globale du système et qui est capable d'analyser son exécution. En effet, ayant connaissance de la totalité du système, le contrôleur est en mesure de vérifier la cohérence des interactions entre les agents, tout en gardant la possibilité d'analyser chacun des agents individuellement.

L'agent-*watchdog* intègre donc un ensemble de connaissances et de règles formulées *a priori* qui lui permettent de contrôler l'évolution du système et de ses composants. Par exemple, il peut vérifier la pente des dérivées de chaque agent-interaction, la cohérence des attributs des agents-entité par rapport au contexte de la simulation, ou encore la fréquence des changements structurels imposés par un agent-énaction.

Sa perspective globale lui permet également de cibler les sources des défauts de paramétrage. Car si un agent se retrouve en défaut, il est possible que ce défaut soit causé par le dysfonctionnement d'un autre agent avec lequel il interagit. Encore une fois, ce sont souvent les connaissances que nous avons du système qui nous guideront dans la hiérarchisation du contrôle du paramétrage.

Dès lors qu'au moins un des points de contrôle devient invalide, l'agent-*watchdog* signale le défaut en formulant une requête de paramétrage.

4.2.3 Requêtes de paramétrage

En fonction de ses observations, l'agent-*watchdog* peut demander la réévaluation d'un paramètre en formulant une **requête**. Cette requête est envoyée au bus de co-simulation qui va sélectionner les modèles susceptibles d'y répondre, comme nous le détaillerons plus tard.

Notre but est d'utiliser des simulations externes pour calibrer le modèle. Une requête regroupe un ensemble de données qu'il est nécessaire de partager pour la mise en œuvre de la co-simulation. En premier lieu, elle fournit une copie de l'état courant de la simulation macroscopique qui servira de base pour le changement de niveau de description. En second lieu, elle contient le nom du paramètre à évaluer car jugé défaillant par l'agent-*watchdog*. En dernier lieu, la requête se voit attribuer une date d'expiration au delà de laquelle elle sera annulée et un statut, qui prend alternativement la valeur « validée » ou « non validée ».

L'intérêt de la date d'expiration d'une requête est de garantir que cette dernière est toujours cohérente avec l'état courant de la simulation. En effet, nous souhaitons éviter d'interrompre la simulation macroscopique afin de conserver son interactivité. L'état enregistré

dans la requête n'est donc pas figé durant le processus de recalibration du paramètre. Or si une requête reste en suspens trop longtemps au regard de la vitesse d'exécution du simulateur macroscopique, la valeur recalculée du paramètre sera potentiellement en décalage avec le nouvel état de la simulation, au moment de son application.

À chaque pas de temps, l'*agent-watchdog* vérifie l'état de ses requêtes. Si elles sont validées, il juge de la pertinence du résultat obtenu en le comparant par exemple à une plage prédéfinie, et l'applique ou non sur la simulation macroscopique toujours en cours d'exécution. Si la date d'expiration d'une requête est dépassée et qu'elle n'est toujours pas validée, elle est automatiquement supprimée car nous considérons qu'elle n'est plus cohérente vis-à-vis de l'état courant de la simulation.

Par ailleurs, nous fixons comme règle à l'*agent-watchdog* de n'avoir qu'une requête active à la fois pour un paramètre donné. De cette manière, nous évitons d'engorger le bus de co-simulation avec des requêtes qui seraient vraisemblablement similaires, si l'on considère une durée de vie assez courte pour les requêtes.

La figure 4.4 montre l'intégration de l'*agent-watchdog* et de ses requêtes dans notre architecture.

Nous allons à présent détailler le traitement des requêtes par le bus de co-simulation.

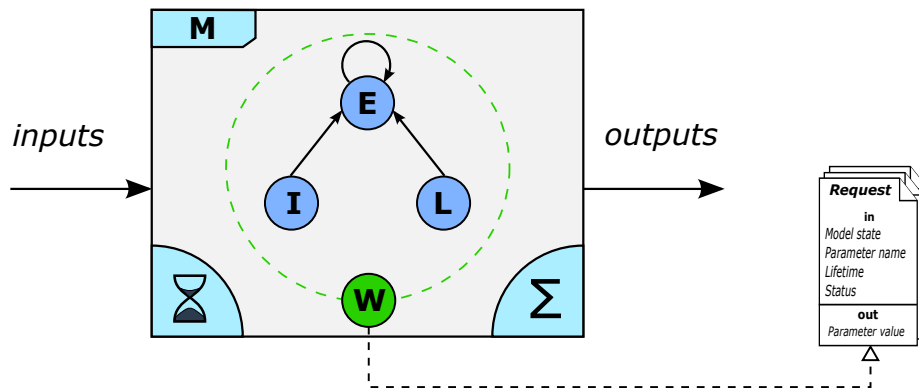


FIGURE 4.4 – **Intégration d'un agent-watchdog** – Le rôle de l'*agent-watchdog* est de détecter les défauts de paramétrage de la simulation. Il génère des requêtes de paramétrage qui contiennent l'état courant de la simulation, le nom du paramètre à recalibrer et une durée de vie au delà de laquelle la requête est considérée comme nulle par l'*agent-watchdog*.

4.2.4 Bus de co-simulation et modèles auxiliaires

Comme nous l'avons vu dans le chapitre 3, la co-simulation est une technique qui permet de lier dynamiquement des simulateurs hétérogènes en termes de paradigmes et de niveaux d'abstraction. De cette manière, chaque simulateur, ou composant, peut être

conçu indépendamment des autres, avec son propre formalisme. De plus, les simulateurs sont exécutés en parallèle, que ce soit sur une même machine ou de manière distribuée, ce qui offre plus de flexibilité et de performance.

Pour communiquer, les différents composants partagent un medium commun qu'est le bus de co-simulation. Plus qu'une simple mémoire partagée, le bus de co-simulation est un composant à part entière de notre architecture et joue le rôle d'un contrôleur de données. Ainsi, grâce à des primitives d'entrées/sorties propres à chaque composant, le bus fait transiter les données aux potentiels destinataires qui ne sont pas connus *a priori* par l'expéditeur.

Dans notre cas, ce sont les requêtes de l'agent-*watchdog* que le bus doit, dans un premier temps, faire transiter jusqu'aux composants susceptibles d'y répondre. Pour ce qui est de ses entrées, chaque composant doit donc être en mesure de recevoir une requête de paramétrage telle que nous l'avons décrite dans la section précédente. La sélection des destinataires se fait d'après la liste des sorties que chaque composant communique au bus lorsqu'il s'y connecte. Ainsi, le bus transmettra une requête pour un paramètre donné à tous les composants qui exposent ce même paramètre comme une de leurs sorties. Cette séquence d'opérations est résumée par la figure 4.5.

Les composants peuvent être de trois types, représentés sur la figure 4.6.

Modèles auxiliaires

La démarche de simulation redondante que nous suivons suppose de disposer de modèles plus fins pour un même système, que nous nommerons dans la suite les modèles auxiliaires. Ces modèles constituent un catalogue et sont susceptibles de prendre plusieurs formes.

Un modèle auxiliaire peut par exemple consister en une copie exacte du modèle macroscopique à la différence près qu'il possède une résolution spatiale et/ou temporelle plus fine, en conséquence de quoi il assure une meilleure précision des résultats dans une plage donnée. D'un autre côté, le modèle peut ne concerner qu'une sous-partie du système ou seulement un phénomène en particulier et adopter un formalisme différent. De cette manière, il est possible de composer avec des modèles de tous horizons : modèles multi-agents (phénoménologiques ou non), modèles stochastiques, simulateurs commerciaux (fonctionnement en boîte noire), etc.

Il est à noter que dans le cas d'un modèle auxiliaire multi-agents tel que nous les concevons, notre méthode a la particularité d'être récursive dans le sens où ce modèle peut à son tour générer des requêtes pour le bus. C'est cette situation que nous avons choisi d'illustrer sur la figure 4.6.

Les modèles auxiliaires n'ont pas vocation à être simulés en permanence. En effet, leur but est de répondre à des requêtes de paramétrage qui correspondent au contexte de la simulation macroscopique. Ainsi, lorsque le bus de co-simulation choisit un modèle auxiliaire, il l'instancie et ce dernier doit s'initialiser en accord avec l'état de la simulation macroscopique contenu dans la requête de paramétrage. Le modèle auxiliaire est ensuite simulé en parallèle de la simulation macroscopique.

Pour finir, le modèle auxiliaire inscrit son résultat de simulation dans la requête et il la renvoie au bus de co-simulation.

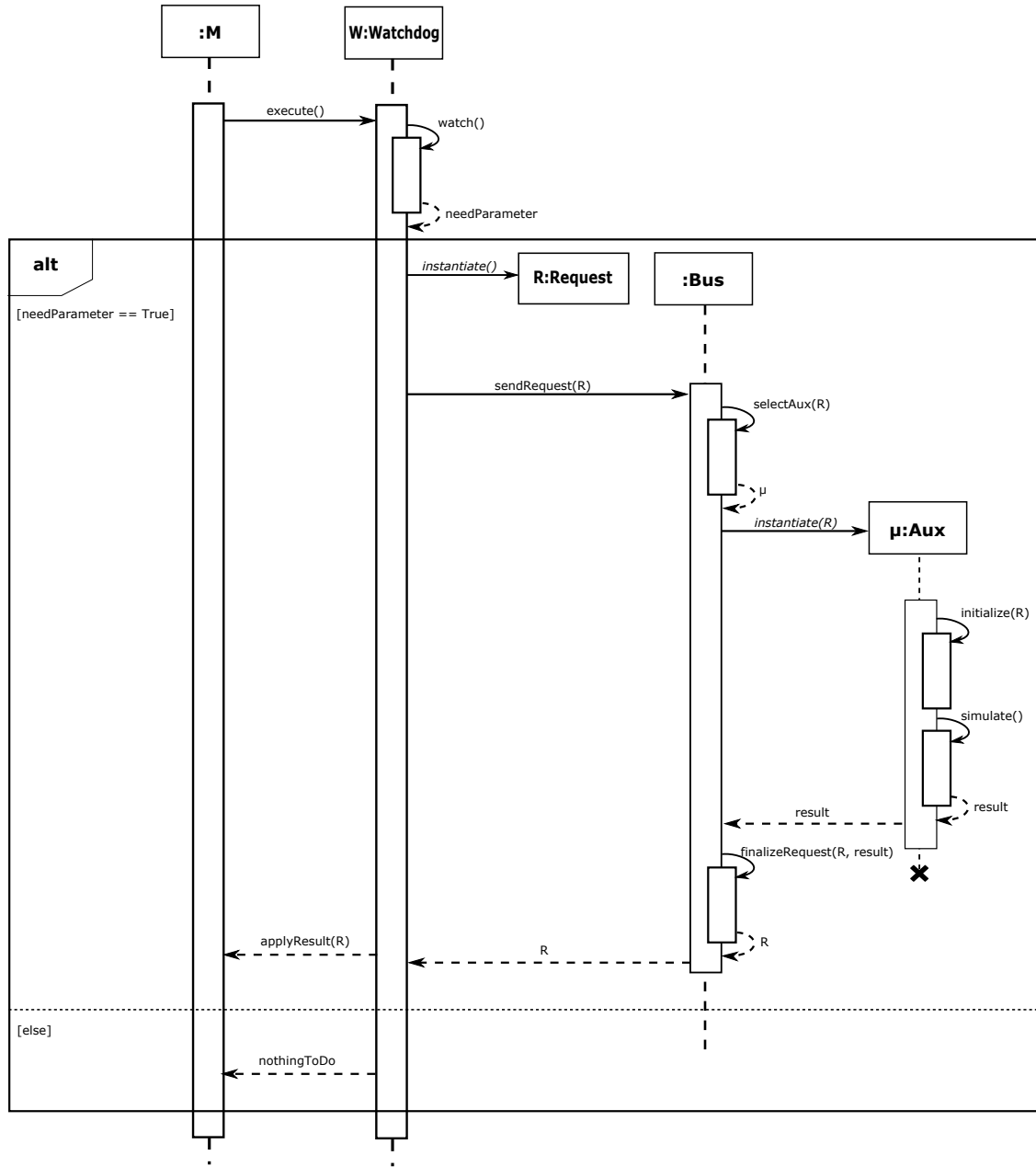


FIGURE 4.5 – **Diagramme de séquence du processus par co-simulation** – Le processus de co-simulation est initié par l'agent-*watchdog*. Il génère des requêtes de paramétrage qu'il envoie au bus de co-simulation. Le bus sélectionne le meilleur candidat pour répondre à la requête et instancie le modèle auxiliaire en utilisant l'état macroscopique stocké dans la requête. Après simulation du modèle auxiliaire, les données sont retournées à l'agent-*watchdog* qui les applique sur le modèle macroscopique.

Base de données

Le paramétrage initial d'un modèle phénoménologique se fait généralement par analogie avec des cas connus. Ces cas composent une base de données qui est alimentée par des données de terrain, par des experts métiers ou par des résultats expérimentaux, qu'ils soient issus de tests *in vitro* ou de simulations précédentes. Si cette base ne peut évidemment pas être exhaustive, elle contient potentiellement des cas similaires à l'état du modèle macrosocopique contenu dans la requête. À défaut de pouvoir proposer un simulateur microscopique pour répondre à une requête, nous pourrions interroger la base et nous appuyer sur ses données pour extrapoler une valeur adéquate pour le paramètre à calibrer.

Utilisateur

Notre approche de modélisation place l'utilisateur au cœur de l'expérimentation *in virtuo*. Ainsi, de même que nous avons vu qu'un utilisateur a la possibilité de prendre la place d'un agent, nous pouvons imaginer qu'il réponde lui-même aux requêtes pour combler dynamiquement les manques du modèle.

La connexion d'un composant au bus de co-simulation suppose qu'il dispose d'une interface de communication qui lui permet de recevoir des requêtes et de transmettre des résultats.

Pour une base de données ou un utilisateur, le développement de cette interface est minime. Dans le premier cas, c'est la définition-même d'une base de données que de répondre à des requêtes sur son contenu. Nous devons seulement définir une méthode de tri des résultats s'ils sont multiples afin de sélectionner la meilleure valeur unique pour le paramètre. Pour un utilisateur, nous considérons comme acquis les moyens d'interaction avec la simulation dans le contexte d'expérimentation *in virtuo*. En revanche, l'utilisation d'un modèle auxiliaire requiert plus d'attention.

4.2.5 Discussions

Nous nous intéressons dans cette section aux particularités liées à l'utilisation de modèles auxiliaires.

Généricité des opérateurs de traduction

La simulation d'un modèle auxiliaire prend, au moins en partie, comme état initial l'état courant de la simulation macrosocopique inscrit dans la requête au moment de sa génération. Il s'agit donc de faire correspondre des variables macrosocopiques et leurs équivalents microscopiques. C'est l'objet de l'opérateur de compression tel qu'il est évoqué dans notre revue des méthodes multi-échelles au chapitre 3. Nous avons remarqué que, de même qu'il n'existe pas de théorie générale pour formaliser les systèmes complexes, il n'est pas non plus possible de donner une expression mathématique universelle pour la traduction d'échelles. Ce processus est donc toujours dépendant du type de problème considéré et il faut à chaque fois mettre en œuvre des outils *ad hoc*.

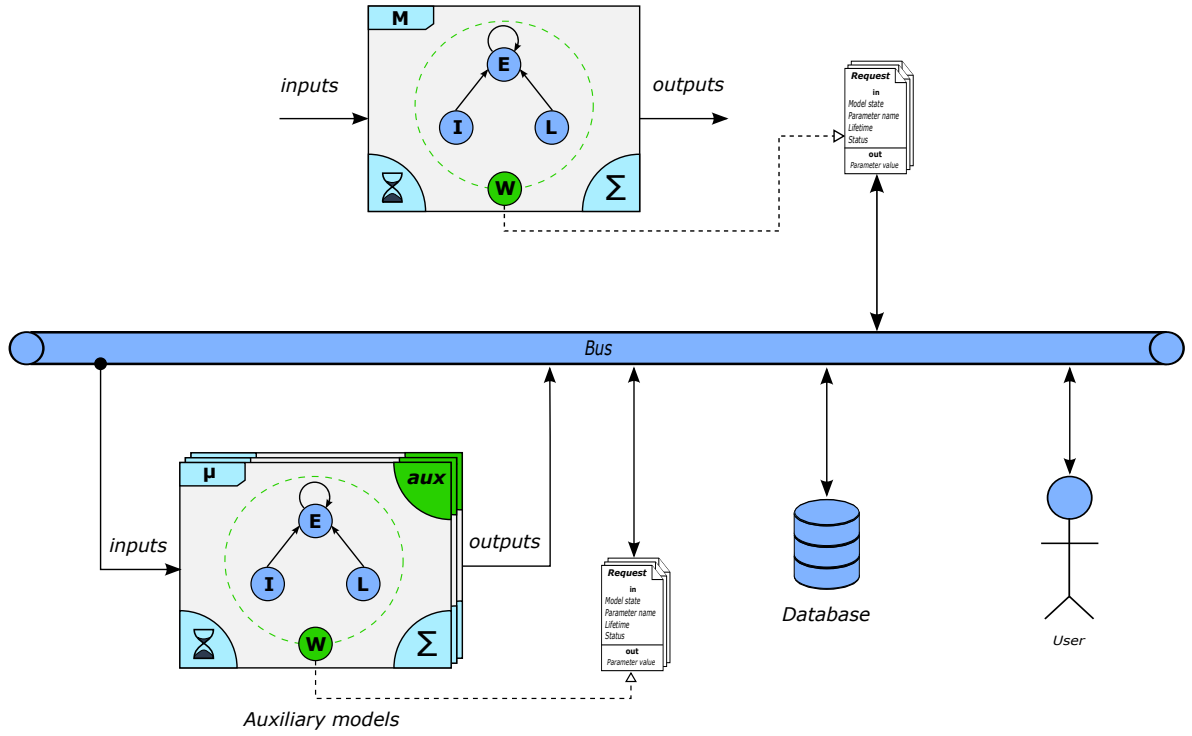


FIGURE 4.6 – **Architecture de co-simulation** – Pour améliorer nos simulations de modèles phénoménologiques, nous proposons une architecture de co-simulation redondante. La co-simulation est une technique qui permet de lier dynamiquement des simulateurs hétérogènes en termes de paradigmes et de niveaux d'abstraction. Chaque simulateur peut être conçu indépendamment des autres, avec son propre formalisme. Pour communiquer, les différents composants partagent un medium commun qu'est le bus de co-simulation. Le bus de co-simulation joue le rôle d'un contrôleur de données et fait transiter les données aux potentiels destinataires qui peuvent être des modèles auxiliaires, des bases de données ou un utilisateur. Ces composants ont vocation à calibrer à la demande les paramètres du modèle macroscopique.

Nous estimons que la majeure partie de ce travail est faite en amont du développement du simple fait de la démarche de simulation redondante. En effet, nous partons du principe que les potentielles faiblesses du modèle macroscopique sont connues *a priori* et que les modèles auxiliaires sont sélectionnés corollairement. En choisissant les modèles, nous faisons plus ou moins implicitement le lien entre les niveaux de description des différentes variables. Aussi, dans notre exemple du phénomène de diffusion, nous savons qu'une concentration correspond à une population de particules.

Domaine de simulation

Une autre difficulté réside dans la définition du périmètre de simulation microscopique. Car si nous misons sur la finesse du modèle microscopique pour calibrer le modèle macroscopique, nous sommes en réalité soumis à plusieurs contraintes concernant son exécution.

Tout d'abord, le temps de simulation du modèle microscopique ne doit pas être trop important au regard du pas de temps macroscopique. Ceci aurait pour conséquence d'introduire des erreurs dans la simulation macroscopique, comme le montre la figure 4.7. En effet, la simulation macroscopique n'est pas mise en pause pendant le processus de co-simulation et son état continue d'évoluer alors que la simulation microscopique se base sur un état figé au moment de la génération de la requête. La date d'expiration d'une requête permet de régler ce problème en refusant des résultats qui ont été calculés à partir d'un état trop ancien.

Pour autant, il est nécessaire de prendre en considération que la mesure de paramètres macroscopiques sur le modèle microscopique ne peut être fiable que lorsque ce dernier atteint un état stationnaire. Le temps de simulation ne doit donc pas non plus être trop réduit au risque d'introduire une erreur supplémentaire dans la simulation macroscopique.

Le temps de simulation d'un modèle microscopique, qui est souvent de type individus-centré, est fortement lié au nombre d'entités qui le constituent. Aussi, une piste pour réduire ce temps est de limiter au maximum le périmètre de la simulation microscopique tout en conservant une certaine représentativité du phénomène macroscopique. C'est ce procédé qu'illustre la figure 4.2 sur laquelle nous avons constaté que la précision de la mesure du coefficient de diffusion dépend en grande partie du nombre de particules simulées.

4.2.6 Bilan

Nous avons présenté dans cette section une architecture de co-simulation qui nous permet de garantir la validité des paramètres d'un modèle phénoménologique tout au long de sa simulation. Des modèles auxiliaires, généralement d'un niveau d'abstraction inférieur, sont simulés en parallèle pour nourrir le modèle de plus haut niveau. L'implémentation de cette redondance est assurée par un ensemble d'outils issus des techniques de co-simulation décrites au chapitre 3 qui nous assurent une grande flexibilité, en terme de formalismes des modèles utilisés et de couplage.

Pour autant, la mise en place de notre architecture peut présenter quelques difficultés. Nous avons par exemple montré que la traduction entre échelles de description n'est pas

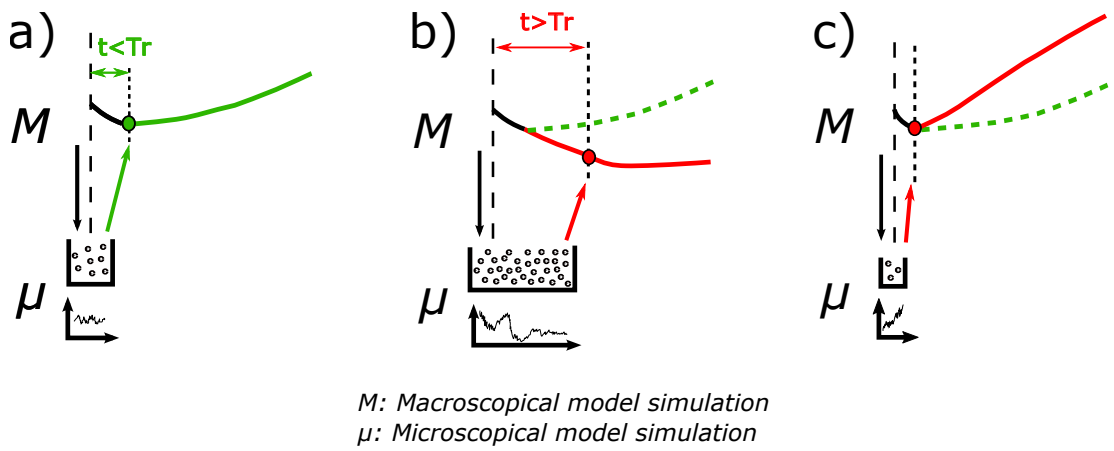


FIGURE 4.7 – **Temps de résolution d'une requête** – L'efficacité du paramétrage par co-simulation est fortement liée au temps de résolution d'une requête, la simulation macroscopique M n'étant pas interrompue pendant le processus. Ce temps de résolution (t) dépend en grande partie du domaine de simulation du modèle microscopique μ . Si le domaine est trop grand, la nouvelle valeur du paramètre sera en décalage avec l'état courant de la simulation qui aura trop divergé (cas b). Le délai d'expiration de la requête (T_r) permet d'éviter ce problème. Par ailleurs, un modèle microscopique non représentatif ou simulé pour une durée trop courte peut générer une erreur qui se propage au niveau macroscopique (cas c). Il s'agit donc de trouver un compromis entre la représentativité du modèle microscopique et son temps de simulation (cas a).

toujours chose aisée et qu'elle doit généralement être faite de manière *ad hoc*. Le but est alors de faire correspondre explicitement un paramètre macroscopique à un paramètre microscopique. Cette procédure est parfois difficile à réaliser à cause de la différence de niveau d'abstraction et n'a d'ailleurs pas toujours de sens.

Pour ces raisons, nous proposons dans la section suivante une seconde stratégie de co-simulation dans laquelle nous considérons les simulations microscopiques comme des outils de validation pour la simulation macroscopique.

4.3 Détermination implicite d'un jeu de paramètres

Pour répondre aux problèmes de paramétrage des modèles phénoménologiques, nous avons proposé dans la section précédente une architecture logicielle qui repose sur la co-simulation redondante d'échelles de modélisation hétérogènes. Celle-ci permet en particulier d'utiliser le maximum de connaissances dont nous disposons sur le système étudié et de combiner les avantages des modèles microscopiques et des modèles macroscopiques, à savoir la précision et l'efficacité respectivement, au sein d'une même simulation.

En dépit de ces nouvelles possibilités, nous avons constaté les limitations de cette première implémentation qui repose sur un calcul explicite des paramètres. Cependant, ce calcul explicite n'est pas toujours possible, et ce pour deux raisons principales. Premièrement, la particularité des systèmes complexes est d'avoir des paramètres fortement dépendants les uns des autres. Ainsi, le calcul d'un paramètre indépendamment des autres n'a pas forcément de sens. Le second problème tient au fait que nous ne disposons pas toujours de modèle plus fin pour un système. Et quand-bien-même nous en aurions, il n'est pas toujours possible de faire le lien entre des paramètres macroscopiques et microscopiques.

Ces cas particuliers nous poussent à envisager une seconde stratégie de co-simulation. Nous allons développer dans ce chapitre une nouvelle méthode qui parvient à conserver les avancées proposées par notre approche de simulations redondantes, tout en garantissant l'interactivité de nos simulations.

4.3.1 Stratégie de co-simulation

La stratégie de co-simulation que nous proposons ici s'apparente aux méthodes d'intégration projective présentées dans le chapitre 3. L'idée générale est, cette fois, d'utiliser la simulation microscopique non plus comme un estimateur de paramètres mais plutôt comme un outil de validation de la simulation macroscopique.

Notre proposition repose sur le postulat suivant : si les modèles microscopique et macroscopique partagent au moins une sortie commune (force, concentration, etc.), il est possible d'évaluer la justesse, et donc intrinsèquement le paramétrage, du modèle macroscopique grâce aux données du modèle microscopique. Cette hypothèse peut se généraliser de la manière suivante : si nous considérons plusieurs simulations microscopiques, partant d'un même état initial mais avec des jeux de paramètres différents, nous sommes en mesure de choisir le plus

pertinent d'entre eux au regard des résultats d'une simulation microscopique initialisée pareillement. Il s'agit alors d'une détermination implicite du jeu de paramètres « idéal » pour le modèle macroscopique.

Par ailleurs, afin de rendre compte des variations éventuelles des paramètres, notre idée est de répéter cette opération de sélection plusieurs fois au cours de la simulation macroscopique et donc de découper celle-ci en plusieurs périodes. La juxtaposition de ces sous-parties de simulation sélectionnées forme alors une simulation complète et interactive qui est rejouée avec un décalage temporel. L'objectif est de rendre l'opération transparente pour l'utilisateur.

Bien entendu, nous ne pouvons pas appliquer cette méthode directement. En effet, ceci impliquerait de simuler le modèle microscopique sur la même durée que le(s) modèle(s) macroscopique(s), ce qui réduirait fortement l'intérêt de nos modèles phénoménologiques. Pour implémenter cette stratégie, nous recourons donc à deux artifices.

Parallélisation de l'exploration

Sélectionner un jeu de paramètres suppose évidemment d'en tester plusieurs. Notre objectif est donc de simuler plusieurs fois le modèle macroscopique avec des jeux de paramètres différents.

Si les modèles phénoménologiques sont bien moins gourmands en temps de calcul que les modèles individus-centrés, il paraît pourtant difficile de conserver l'interactivité de la simulation affichée à l'utilisateur en exécutant les simulations macroscopiques exploratoires de manière séquentielle. Aussi, nous nous appuyons sur les architectures modernes de calcul parallèle. Nous exécutons les simulations exploratoires en parallèle, chacune sur un cœur différent, pour un temps donné. La simulation microscopique est elle-même exécutée sur un autre cœur, de même que l'affichage successif des simulations exploratoires sélectionnées.

Estimation de trajectoire

Comme nous l'avons indiqué, il n'est pas possible de simuler le modèle microscopique pour une même durée que les simulations macroscopiques exploratoires. C'est même le principal intérêt des modèles phénoménologiques que d'améliorer les temps de calcul. Ainsi, notre objectif est plus de faire correspondre les temps de simulation plutôt que les temps simulés, afin que la simulation affichée à l'utilisateur ne souffre pas de temps d'arrêt.

Par exemple, pour un temps de simulation de 10 secondes, nous pourrions simuler 100 secondes de modèle macroscopique, tandis que nous ne réussirions à simuler qu'une seconde de modèle microscopique. Nous devrions donc nous contenter des résultats générés par le modèle microscopique à l'issue de cette seconde simulée pour sélectionner la simulation macroscopique exploratoire la plus juste. De ce fait, nous analysons ces données sous deux angles.

Nous commençons par comparer « temps pour temps » les résultats de la simulation microscopique à ceux des simulations exploratoires, comme sur la figure 4.8.a. Cependant, nous avons vu que l'utilisation de l'opérateur de reconstruction, pour passer de l'état courant

de la simulation macroscopique à un état initial pour la simulation microscopique, implique un temps de relaxation du système dû à l'erreur de reconstruction. Les données des premiers pas de temps de la simulation microscopique peuvent donc être altérées et nous ne pouvons leur accorder qu'une confiance mesurée.

Les résultats à l'issue du temps de relaxation sont quant à eux plus fiables. Nous pouvons les utiliser pour tenter d'estimer la trajectoire que la simulation microscopique aurait suivi si nous l'avions exécutée jusqu'à atteindre le même temps simulé que les simulations exploratoires. Nous comparons alors le résultat au bout de cette trajectoire aux données des simulations macroscopiques à la fin de leur exécution. Cette procédure est illustrée par la figure 4.8.b. La linéarisation de trajectoire n'est toutefois valable que pour une durée de simulation exploratoire limitée.

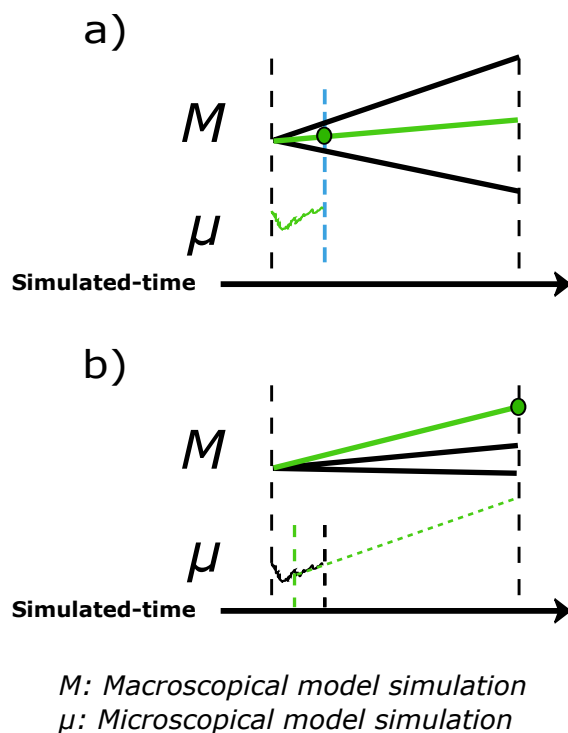


FIGURE 4.8 – **Sélection d'un jeu de paramètre** – La sélection d'une simulation exploratoire peut être faite selon deux critères. Sur la figure a, nous comparons les données de toutes les simulations macroscopiques M sur la base du temps simulé qu'elles ont en commun, c'est-à-dire celui de la simulation microscopique μ . Sur la figure b, nous estimons la trajectoire qu'aurait suivi la simulation microscopique μ si nous l'avions simulée plus longtemps et nous sélectionnons la simulation exploratoire M la plus proche de cette trajectoire.

La figure 4.9 donne une vue générale de la stratégie de co-simulation que nous venons de décrire.

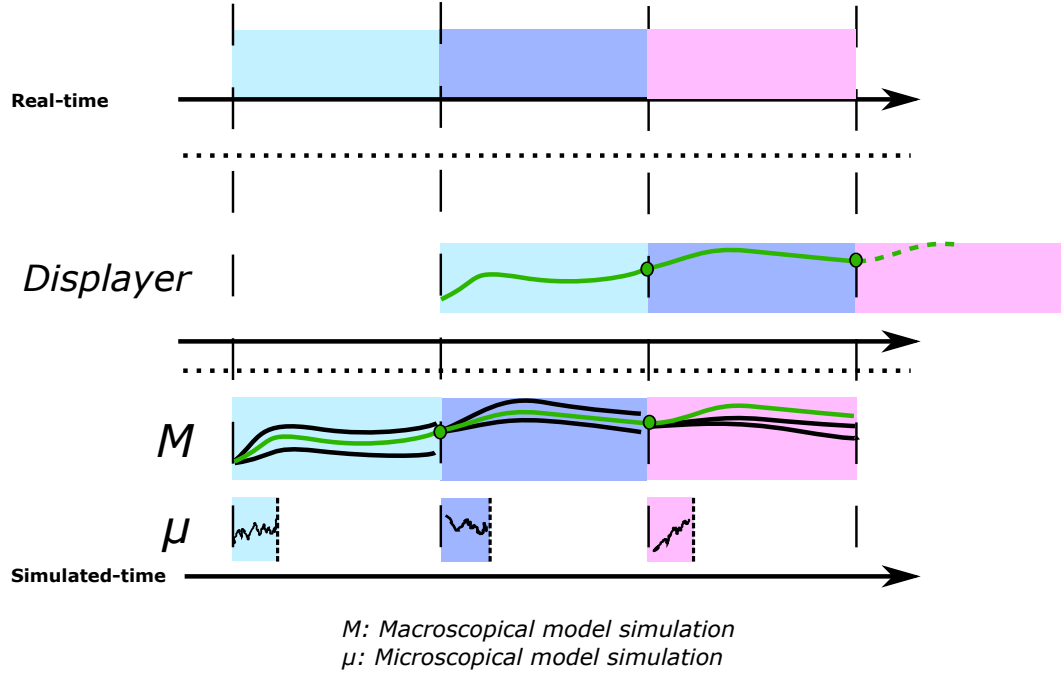


FIGURE 4.9 – **Construction d'une simulation interactive par co-simulation** – Au démarrage du simulateur, le *Displayer* n'a pas de données à afficher. La première étape consiste donc à exécuter en parallèle des simulations exploratoires (M) et une simulation microscopique (μ). Notons que les durées des simulations sont égales alors que le temps simulé n'est évidemment pas le même, puisque fortement dépendant du niveau de description. Les tronçons de courbes en vert sont ceux qui ont été sélectionnés pour affichage. Pendant l'affichage de ces données, une nouvelle co-simulation est lancée à partir de l'état final du dernier tronçon sélectionné. De cette manière, le *Displayer* ne souffre d'aucune interruption. Les plages de couleurs résument les correspondances temporelles des itérations de co-simulation ainsi que leur succession.

4.3.2 Implémentation

L'intégration de cette stratégie dans notre architecture de co-simulation passe par une redéfinition des rôles des différents composants. Leurs liens sont illustrés par la figure 4.10.

Nous ajoutons un nouveau type de simulateur, le *Displayer*, qui forme une interface interactive pour l'utilisateur. Ce simulateur est capable de jouer les données de modèles phénoménologiques qui ont été sélectionnées par son agent-*watchdog*. Le rôle de cet agent reste similaire à celui qu'il tient pour la détermination explicite de paramètres, à savoir la surveillance du déroulement de la simulation. Quand il détecte une intervention de l'utilisateur ou la fin de lecture des données, il génère une requête qui contient l'état courant de la simulation et les entrées de l'utilisateur.

Dans cette seconde stratégie, la requête n'est pas transmise au bus de co-simulation car le fonctionnement du *Displayer* est évidemment plus fortement lié au modèle phénoménologique dont il doit afficher les données. C'est donc à l'agent-*watchdog* qu'incombe l'instanciation des différentes simulations.

Les simulations exploratoires sont toutes initialisées d'après le contenu de la requête : l'état initial correspond à l'état courant du *Displayer*, éventuellement modifié par les entrées de l'utilisateur. Ensuite, chacune d'elles se voit attribuer un jeu de paramètres différents qui varient dans un périmètre imposé dans la requête. Ces variations sont définies d'après les connaissances des experts du domaine du système étudié qui, le plus souvent, disposent de plages de valeurs plutôt que de valeurs fixes pour les paramètres du modèle phénoménologique. Les simulations sont ensuite exécutées en parallèle pour un temps donné. Dans le même temps, une simulation microscopique est initialisée avec les mêmes informations, en suivant le mécanisme de traduction d'échelle de description décrit précédemment. Nous pouvons noter que le déroulement des différentes simulations, quel que soit le modèle, reste parfaitement compatible avec la première stratégie de co-simulation. Si leurs agents-*watchdog* détectent des défauts de paramétrage, ils peuvent transmettre des requêtes au bus qui les traitera en conséquence.

Enfin, l'agent-*watchdog* collecte l'ensemble des données de toutes les simulations et sélectionne la simulation exploratoire la plus proche de la simulation microscopique, pour lecture par le *Displayer*.

4.3.3 Exemple d'application : diffusion de deux molécules antagonistes

Pour illustrer le fonctionnement de notre stratégie de co-simulation pour la détermination implicite de paramètres, nous reprenons l'exemple de la diffusion développé dans la section 4.1.4.

Lors de notre première implémentation, nous avons montré l'intérêt de mettre à jour le coefficient de diffusion en cours de simulation pour rendre compte du phénomène de fuite des molécules d'une espèce chimique lié à la concentration de cette même espèce. Dans ce cas,

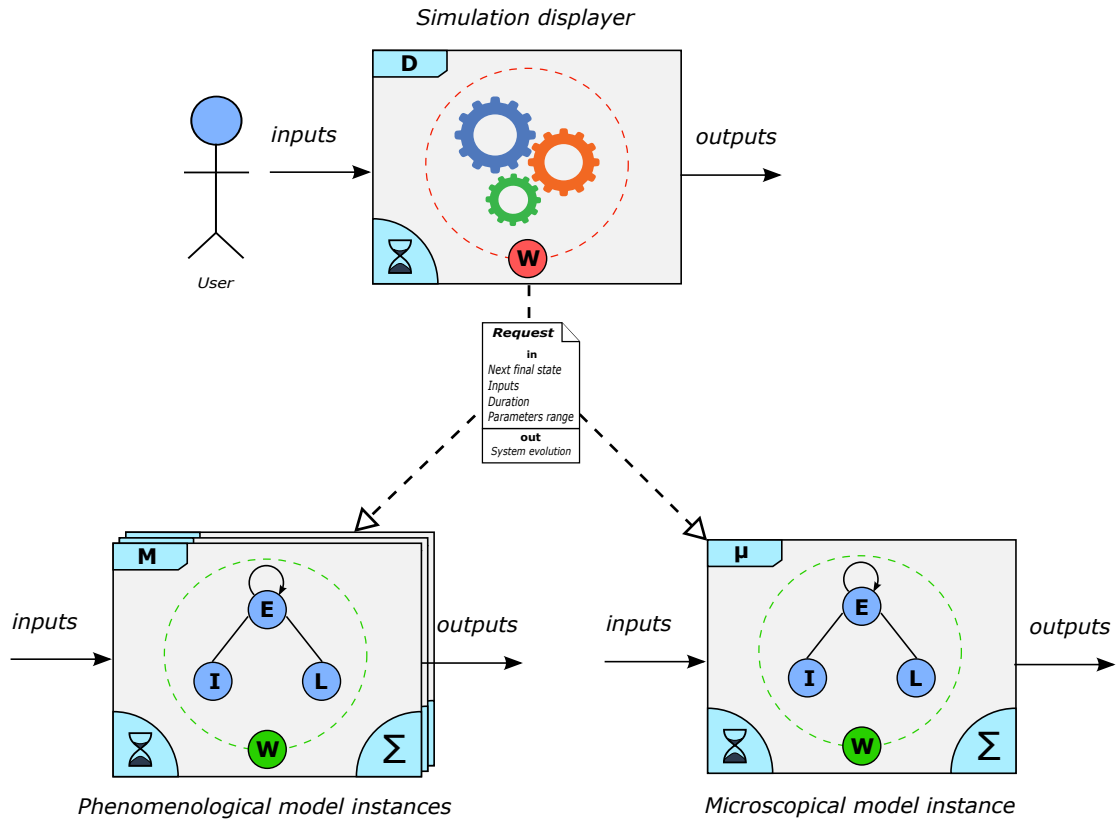


FIGURE 4.10 – **Architecture de co-simulation exploratoire** – Notre seconde stratégie de co-simulation vise à construire une simulation interactive. L’agent-*watchdog* (en rouge) a maintenant pour rôle de sélectionner des données à afficher par le *Displayer*. Ces données proviennent de simulations macroscopiques exploratoires exécutées en parallèle à partir d’un même état initial mais avec des jeux de paramètres différents. Une simulation microscopique du même système sert d’étalon pour mesurer la justesse des simulations exploratoires.

nous avons pu déterminer ce coefficient de manière explicite grâce à un modèle microscopique car un seul paramètre était inconnu.

Cette fois, nous considérons la diffusion de deux espèces chimiques qui ont une réaction de fuite l'une par rapport à l'autre. Les coefficients de diffusion de chacune des espèces dépendent donc de la concentration de l'autre.

Nous ne pouvons pas calculer chaque paramètre indépendamment sans introduire de biais de causalité dans la simulation. Nous pourrions synchroniser l'opération, c'est-à-dire calculer tous les paramètres à chaque itération et appliquer les nouveaux paramètres en une seule fois. Cela impliquerait de mettre en pause le simulateur macroscopique et d'attendre le modèle microscopique le plus lent.

Nous ne pouvons pas non plus calculer explicitement un couple de valeurs pour les coefficients car nous serions alors contraint de doubler le nombre de particules simulées par le modèle microscopique afin qu'il soit représentatif. Le temps de simulation du modèle microscopique deviendrait trop important et les valeurs obtenues seraient incohérentes au vu de l'état courant de la simulation macroscopique, à moins de d'interrompre cette dernière le temps de l'opération.

Pour ce type de problème, la stratégie de détermination implicite de paramètres semble donc plus adaptée.

Nous initialisons un modèle macroscopique avec deux concentrations identiques d'espèces chimiques A et B . Nous donnons pour chacune des valeurs estimées des coefficients de diffusion. Cet état initial est le point de départ du *Displayer*. Celui-ci n'a pas encore de données de simulation à afficher donc son agent-*watchdog* démarre le processus d'acquisition.

Trois simulations exploratoires du modèle macroscopique sont lancées avec l'état courant du *Displayer* et des variations des coefficients de diffusion. Nous avons fait le choix ici de générer des valeurs aléatoires qui s'écartent au maximum de 10% des dernières valeurs sélectionnées, qui dans le cas de l'état initial correspondent aux coefficients estimés.

Dans un quatrième *thread*, l'agent-*watchdog* instancie un modèle microscopique. Les particules des espèces A et B sont toutes placées en $x = 0$. Pour ce second exemple, l'équation 4.13 est adaptée pour calculer la vitesse des particules en fonction de leur espèce :

$$\begin{aligned} v^A &= v_{max}^A - v_{max}^A \exp^{-bC^B(x,t)} \\ v^B &= v_{max}^B - v_{max}^B \exp^{-bC^A(x,t)} \end{aligned} \tag{4.14}$$

Les quatre simulations sont exécutées en parallèle. Les trois simulations exploratoires s'arrêtent quand elles atteignent un temps simulé de $t = 1s$ et envoient leurs résultats à l'agent-*watchdog* qui interrompt alors la simulation microscopique.

Pour sélectionner la meilleure simulation exploratoire et donc le meilleur couple de coefficients de diffusion (d_A , d_B), nous comparons la concentration de l'espèce A en $x = 0$ qui est une sortie commune des simulations macroscopique et microscopique. Les données de la simulation exploratoire la plus proche du résultat du modèle microscopique sont conservées par l'agent-*watchdog* et jouées par le *Displayer*. Un nouveau cycle d'exploration de paramètres

et de simulation microscopique est relancé à partir de l'état final de la dernière simulation macroscopique sélectionnée, en parallèle de l'affichage des résultats associés. De cette manière, le *Displayer* n'est jamais interrompu et le mécanisme de co-simulation est transparent pour l'utilisateur.

4.3.4 Discussions

Coût et précision

Notre seconde stratégie de co-simulation repose principalement sur la multiplication des simulations. Son implémentation est rendue possible par la parallélisation des calculs. Nous estimons que cette répartition des simulations sur plusieurs unités de calcul n'affecte pas les performances globales du simulateur. En effet, nos simulations multi-agents sont la plupart du temps difficilement parallélisables en raison de l'intrication des phénomènes. Nous disposons donc d'unités de calcul libres.

Par ailleurs, nous pouvons noter que si la parallélisation n'offre pas ici un gain de temps de calcul, la simulation affichée par le *Displayer* sera d'autant plus juste que le nombre d'unités de calcul disponible est grand. En effet, de cette manière nous pourrions multiplier les simulations exploratoires et donc tester plus de jeux de paramètres avec une granularité plus fine.

Erreurs de sélection

À chaque itération de la co-simulation, l'agent-*watchdog* doit choisir la meilleure solution parmi les résultats des simulations exploratoires. Comme nous l'avons vu, cette sélection se base sur de multiples critères, le but étant de trouver des correspondances entre les simulations exploratoires et la simulation microscopique. Cette opération n'est évidemment pas triviale et s'apparente à une méthode d'optimisation. Au delà du fait qu'il est nécessaire de définir une fonction de *fitness* adéquate pour chaque problème, nous devons prendre en considération certaines erreurs susceptibles de perturber le processus et nous conduire à effectuer un mauvais choix.

Nous avons vu dans la section 3.2.4.2 que la traduction d'un état macroscopique à un état microscopique introduit une erreur de *lifting* qui peut être maîtrisée en simulant le modèle microscopique suffisamment longtemps pour qu'il se relaxe. Or, dans le cadre de notre stratégie de co-simulation, le temps de simulation du modèle microscopique est limité. Cette situation peut conduire à des données de validation non fiables et ainsi causer une mauvaise convergence de la simulation affichée par le *Displayer*.

Ensuite, nous devons prêter une attention particulière à l'inégalité des plages de temps de simulation que nous comparons. En effet, nous avons deux options pour choisir la meilleure des simulations exploratoires : la comparaison « temps pour temps » et l'estimation de trajectoire. Pour la première, nous pouvons nous retrouver dans une situation où les simulations exploratoires ne diffèrent qu'à partir d'un temps supérieur à celui simulé microscopiquement.

Cette situation est illustrée par la figure 4.11.

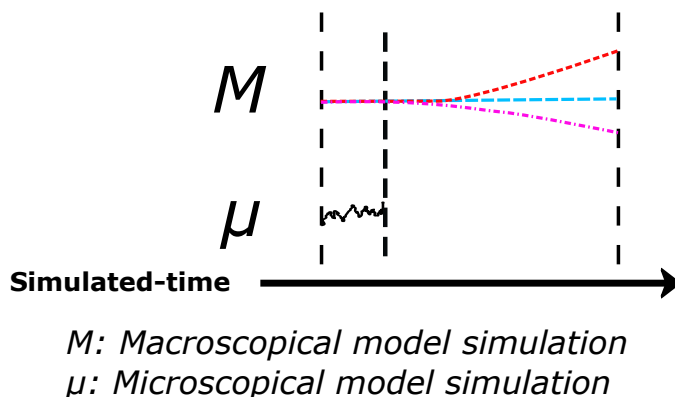


FIGURE 4.11 – **Erreur de comparaison temporelle** – Le temps simulé par le modèle microscopique peut ne pas être suffisant pour discriminer des simulations exploratoires quand celles-ci ont un début d'exécution similaire. Cette situation peut conduire à un mauvais choix de paramètres et donc à faire diverger la simulation.

Pour la seconde option, nous estimons une trajectoire sur une durée qui peut ne pas être significative au regard de la dynamique du modèle macroscopique. Sur la figure 4.12.a, nous montrons qu'une trajectoire linéaire est clairement insuffisante pour sélectionner une simulation macroscopique oscillante.

En pratique, nous remarquons donc que le bon fonctionnement de notre stratégie de co-simulation s'accompagne de la mise en œuvre d'un certain nombre de bonnes pratiques.

Pour nous assurer que le temps de simulation du modèle microscopique soit supérieur au temps de relaxation, nous pouvons réduire le domaine de simulation. Comme pour la détermination explicite de paramètres, cette réduction devra toutefois être relative afin de garder une représentativité satisfaisante.

Nous pouvons également affiner nos estimations de trajectoire en utilisant des sources extérieures aux simulations telles que l'expertise métier ou des données expérimentales. Ce cas est illustré en rouge sur la figure 4.12.b.

Fréquence de co-simulation

Par définition, les données affichées par le *Displayer* sont en décalage temporel permanent par rapport au temps en cours de co-simulation. Lorsque l'utilisateur interagit avec le *Displayer*, il manipule en fait des données qui ont été calculées lors de l'itération de co-simulation précédente. Pour que le simulateur puisse vraiment être qualifié d'interactif, l'action de l'utilisateur doit se ressentir presque immédiatement sur les données qui lui sont affichées. Ceci ne peut être garanti que si la fréquence des itérations de co-simulation est élevée au regard de la fréquence des interventions de l'utilisateur. Aussi, sur la figure 4.13,

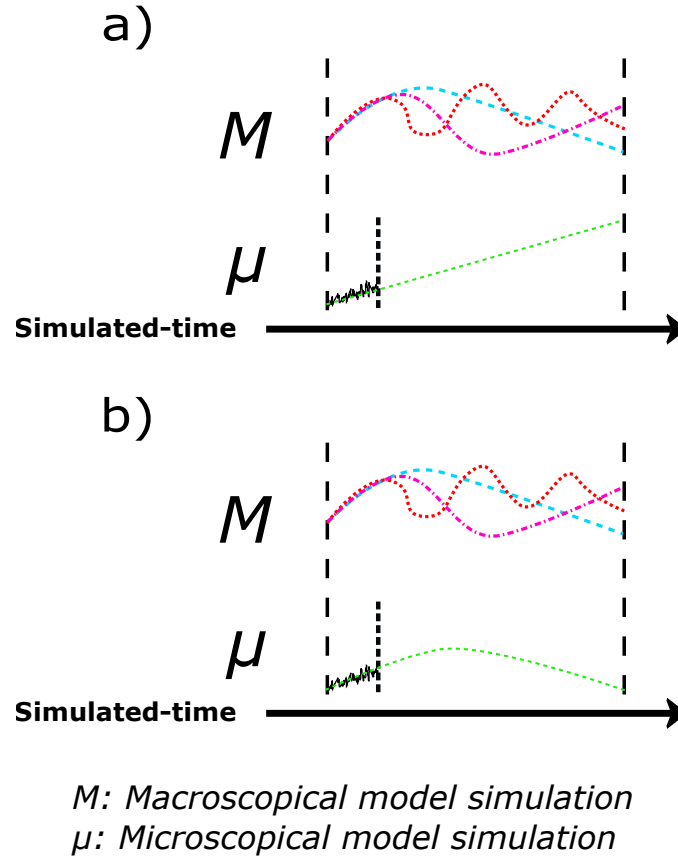


FIGURE 4.12 – **Erreur d'estimation de trajectoire** – Sur la figure a, nous observons que quand les simulations exploratoires ont une dynamique changeante sur la tranche de temps simulée, il devient difficile d'effectuer un choix fiable sur la base d'une estimation de trajectoire linéaire de la simulation microscopique. Aussi, nous préconisons d'utiliser des sources de connaissances externes (expertise métier, base de données, simulations antérieures) pour estimer des trajectoires plus conformes à la dynamique du système (figure b).

nous observons que le doublement de la fréquence de co-simulation permet de prendre en compte plus vite l'intervention de l'utilisateur.

L'augmentation de la fréquence de co-simulation peut avoir deux autres effets ambivalents. D'un côté, elle peut permettre de réduire l'erreur d'estimation de trajectoire. D'un autre côté, si la fréquence est trop élevée, la simulation microscopique n'aura jamais le temps de se relaxer et l'estimation sera biaisée.

Encore une fois, il s'agit donc de trouver un compromis acceptable entre l'interactivité du simulateur et la qualité de ses résultats.

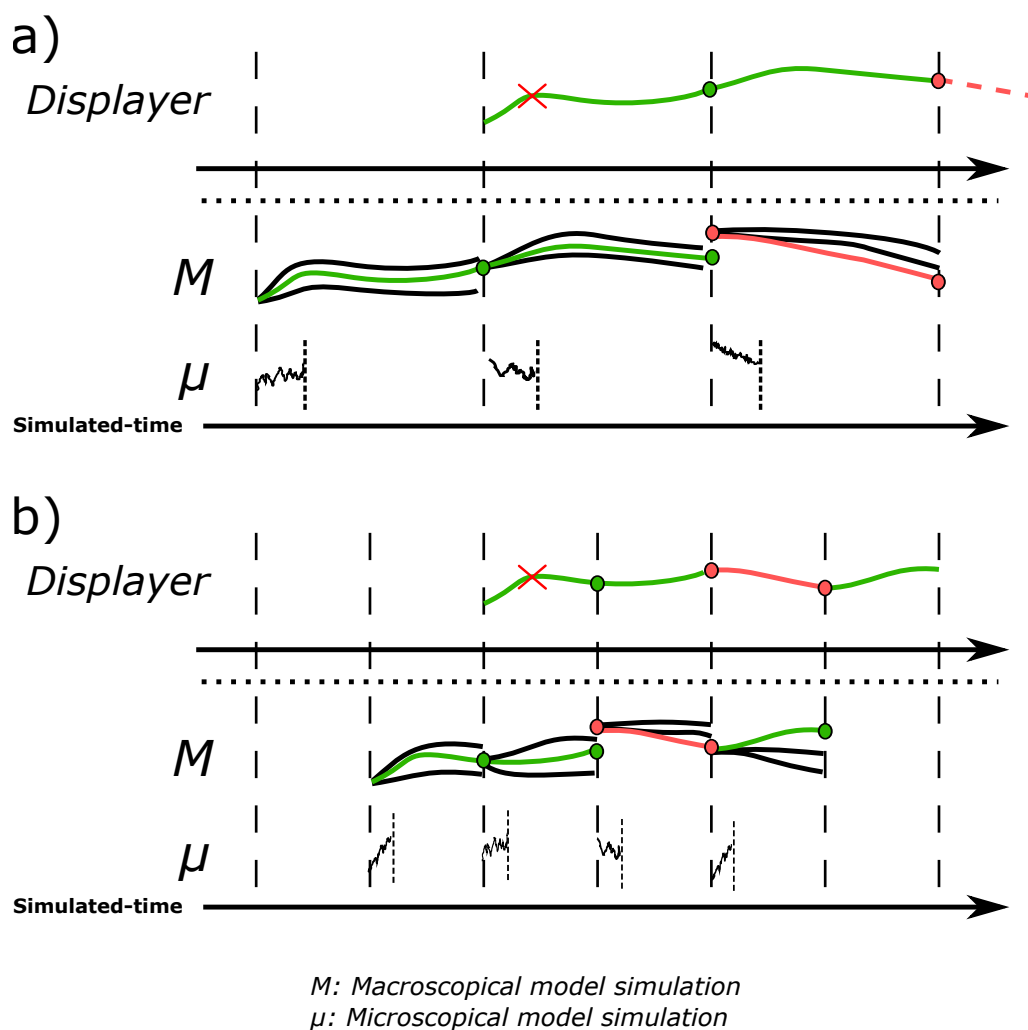


FIGURE 4.13 – **Interaction en cours de co-simulation** – Quand l'utilisateur interagit avec le *Displayer*, ses entrées ne sont prises en compte par la co-simulation qu'au prochain lancement de simulations exploratoires. Si la fréquence de co-simulation est trop basse, le *Displayer* perd en interactivité, comme le montre la figure a. Nous voyons sur la figure b que le fait de doubler la fréquence de co-simulation améliore l'interactivité. En revanche, cette augmentation de fréquence réduit le temps simulé du modèle microscopique et réduit les chances de sélectionner la bonne simulation exploratoire.

4.4 Synthèse

Dans ce chapitre, nous avons présenté une architecture de simulation qui corrige les failles des modèles phénoménologiques fortement paramétrés. Notre approche repose sur la simulation redondante d'échelles de modélisation hétérogènes, et notamment sur l'utilisation de modèles de bas niveau qui nourrissent les modèles plus descriptifs, comme le préconisent les différentes méthodes multi-modèles citées dans le chapitre 3. Mais à défaut de pouvoir proposer une implémentation générique de ces méthodes, nous nous sommes concentrés sur le développement d'outils qui permettent d'automatiser la démarche de simulation redondante qu'elles prônent.

Nous avons distingué deux stratégies de co-simulation qui, par ailleurs, peuvent tout à fait être utilisées conjointement. Dans un premier temps, nous avons implémenté une stratégie qui autorise la détermination explicite de paramètres. Celle-ci suppose d'être capable d'établir un lien direct entre un paramètre d'une loi descriptive et un modèle plus fin du même phénomène. La valeur du paramètre peut être calculée et ajustée plusieurs fois en cours de simulation, pour s'adapter à la dynamique du système et aux interventions de l'utilisateur.

Dans un second temps, nous avons proposé une nouvelle stratégie pour la détermination implicite d'un jeu de paramètres. Celle-ci intervient dans les cas où nous ne pouvons pas calculer explicitement un paramètre, soit par manque de modèle, soit dans souci de cohérence. En effet, les paramètres d'un système complexe sont souvent interdépendants et ils nécessitent dans ce cas une évaluation globale. Notre solution pour ce type de problèmes est alors de composer par morceaux une simulation interactive. Nous répétons une alternance de simulations macroscopiques exploratoires et un processus de sélection par comparaison avec une simulation microscopique exécutée en parallèle. De cette manière, nous adaptons dynamiquement et implicitement le jeu de paramètres du modèle macroscopique qui permet de rester au plus proche de la précision du modèle microscopique.

Pour illustrer notre architecture de simulation, nous nous sommes appuyés dans ce chapitre sur l'exemple relativement simple du phénomène de diffusion. Nous souhaitons à présent démontrer la pertinence de notre proposition dans le cadre d'une application industrielle. C'est l'objet du chapitre suivant de ce mémoire qui traite de la simulation autrement plus complexe de structures offshore soumises à des conditions polaires.

Chapitre 5

Application à la conception de structures offshore



Philippe Geluck - LE TOUR DU CHAT EN 365
JOURS

Nous avons proposé dans le chapitre précédent une architecture de co-simulation qui met en œuvre de manière redondante des modèles de phénomènes à des niveaux de description différents. Nous avons jusqu'ici utilisé des exemples « jouets » pour expliquer notre démarche. Nous souhaitons à présent montrer l'intérêt d'une telle approche dans un cadre industriel.

Pour ce faire, nous avons choisi de nous appuyer sur une application que nous avons développée en association avec [TECHNIP](http://www.technip.com)¹ et [BUREAU VERITAS](http://www.bureauveritas.com)², dans le domaine de l'aide au design de structures offshore pour les milieux polaires. En premier lieu, nous préciserons le contexte de ce travail et particulièrement nous détaillerons les aspects qui contraignent le design de ces structures. Ensuite, nous présenterons l'application ICE-MAS, Ice Multi-Agent Simulator, et l'approche multi-modèles retenue pour sa conception. Enfin, nous proposerons des voies d'optimisation du temps de calcul de ICE-MAS grâce à l'abstraction de

1. <http://www.technip.com>

2. <http://www.bureauveritas.com>

certains phénomènes. Nous verrons alors comment notre méthode de co-simulation autorise la simulation de ces modèles de haut-niveau grâce à un paramétrage dynamique.

5.1 Structures offshore et milieux polaires

5.1.1 Contexte

Pour l'exploitation des gisements d'hydrocarbures et le transport maritime dans les milieux polaires et sub-polaires, et notamment dans l'Arctique, l'industrie pétrolière construit depuis de nombreuses années des structures fixes ou mobiles. Dans le domaine de l'exploitation pétrolière, il s'agit de structures capables de supporter des installations de forage. Elles sont de différents types en fonction de la situation géographique, de la profondeur de l'eau et des contraintes climatiques et économiques. On distingue trois grands types de structures : les îles artificielles, les structures reposant sur le fond et les structures flottantes. Dans le domaine du transport maritime, il s'agit essentiellement de navires brise-glace, destinés à l'ouverture des routes maritimes dans les mers gelées, mais également de ponts, tel le Pont de la Confédération qui enjambe l'embouchure du Saint Laurent au Canada, et de toutes les structures destinées à la sécurité de la navigation : phares et balises notamment. Nous pouvons ajouter à cette liste les nouvelles éoliennes offshore.



FIGURE 5.1 – **Structures offshore en milieux polaires** – Le dimensionnement est un point critique de la conception des structures offshore pour les milieux polaires. En effet, ces dernières doivent supporter d'importants efforts dus aux glaces dérivantes, comme illustré à gauche. L'industrie pétrolière se base généralement sur des recommandations d'experts qui relèvent d'observations de terrains. Des essais en bassin (à droite) viennent compléter ces lois empiriques, mais ils ne peuvent être exhaustifs en raison du grand nombre de paramètres en jeu. Source : [Barkov, 2011]

Les différentes structures construites pour les milieux polaires sont exposées à des conditions d'englacement qui peuvent être préjudiciables à leur bon fonctionnement, voire à leur stabilité. Afin de dimensionner convenablement ces structures pour qu'elles résistent

à ces conditions extrêmes, il importe de connaître les efforts auxquels elles sont soumises. Certaines normes et recommandations ont été proposées pour aider au dimensionnement des structures, essentiellement basées sur des observations de terrain [Fransson et Bergdahl, 2009]. Cependant, des études récentes [Bjerka et al., 2010] ont montré que les efforts dus à la glace sont souvent mal estimés par ces normes, car elles peinent à rendre compte de la complexité des interactions glace-structure. Or le surdimensionnement d'une structure entraîne un surcoût considérable de sa construction, étant donné l'échelle de ces édifices, de même que le sous-dimensionnement peut avoir des conséquences sur la pérennité de la structure. En complément des normes de conception, les structures sont également dimensionnées sur la base d'essais en bassin de modèles réduits. Ces essais sont généralement très coûteux, en raison de la combinatoire des paramètres à tester.

Pour toutes ces raisons, le secteur de l'offshore a mené de nombreuses études ces dernières années pour tenter de mieux comprendre les phénomènes auxquels sont soumises les structures dans les milieux polaires. Plus particulièrement, ces travaux ont proposés différentes approches pour modéliser les interactions glace-structure.

5.1.2 Interactions glace-structure

Les interactions glace-structure sont des phénomènes complexes, du fait de la grande variété de conditions dans lesquelles elles se produisent. Elles ont fait, et font encore, l'objet de nombreuses études qui ont pour objectif principal la détermination des efforts maximaux que subit une structure soumise à la poussée des glaces dérivantes. En effet, dans les mers polaires, des fragments de glace de toutes tailles sont entraînés par les courants marins et par le vent. Nous nous intéressons particulièrement aux plaques de glaces qui peuvent avoir une surface de plusieurs kilomètres carrés, car ce sont celles qui induisent les charges les plus importantes sur les structures.

Le passage de la structure au travers de ces plaques de glace se fait de manière macroscopiquement continue. Si on s'intéresse de plus près à ce qui se passe au niveau du contact, on constate que l'interaction comporte des mécanismes de rupture plus complexes.

5.1.2.1 Modes de rupture

En fonction de certains paramètres de la glace (densité, épaisseur, vitesse de dérive) et de la structure (dimension, inclinaison, ancrage), différents modes de rupture peuvent survenir lors d'une collision. Nous nous concentrons dans notre étude sur les ruptures par écrasement et par flexion qui sont les plus fréquemment observées dans les conditions auxquelles sont soumises les structures en zone arctique.

Rupture par écrasement

Lors du contact entre la plaque de glace et la structure, le front de la plaque se retrouve comprimé par la structure. La glace est alors broyée ou pulvérisée suite à cette compression.

Les particules de glace ainsi formées sont éliminées par extrusion entre la paroi de la structure et la glace non encore pulvérisée.

Par ailleurs, des contraintes de cisaillement conduisent à la propagation de fissures parallèlement à la surface de la plaque. Ces fissures s'incurvent ensuite pour déboucher à la surface de la glace. Ceci conduit à la formation d'écailles de glace, d'où le nom de rupture par écaillage. L'épaisseur des écailles est inférieure à la moitié de l'épaisseur de la plaque, leur longueur étant quant à elle de l'ordre d'une à deux fois l'épaisseur de la plaque. Cette suite d'événements est illustrée par la figure 5.2a.

Ce phénomène, local dans le cas des structures larges, se produit sur toute la zone de contact dans le cas des structures étroites. La fréquence de ce mode de rupture varie entre 0.5 et 50 Hz. Lors d'une telle interaction, les efforts forment généralement des cycles composés d'une augmentation rapide suivie d'un déchargement brutal. Ceci peut s'expliquer par la variation de la surface de contact entre la plaque et la structure.

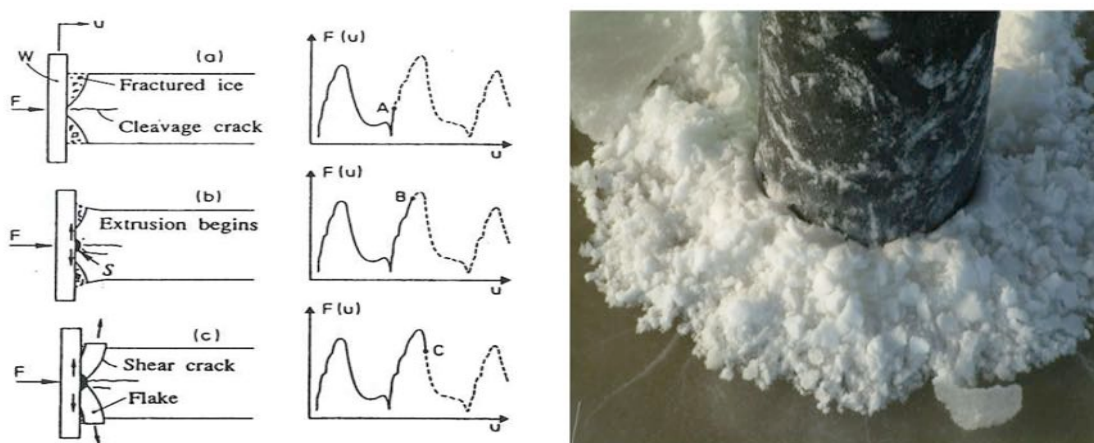


FIGURE 5.2 – **Phénomène de rupture par écrasement d'une plaque de glace** – La glace est pulvérisée et des fissures se propagent sur l'axe horizontal de la plaque. Ces fissures s'incurvent et conduisent à l'extrusion de plus grands éclats de glace. Nous observons sur le schéma de gauche que les efforts sur la structure suivent ces cycles et sont d'autant plus importants que la surface de contact est grande. À droite, nous pouvons voir que la glace pulvérisée et extrudée s'accumule sur la plaque de glace en amont de la structure. Source : [Kärnä et al., 1999]

Rupture par flexion

Dans le cas d'une structure inclinée, illustré sur la figure 5.3, le mode de rupture prédominant est la rupture par flexion. En effet, la plaque de glace a alors tendance à suivre la pente de la structure en avançant, ce qui induit une contrainte verticale importante et à terme une rupture de la plaque.

D'après les observations de terrain et en laboratoire, la ligne de fracture a une forme circulaire qui a pour origine le point de contact entre la plaque et la structure. Le rayon de cette rupture circonférentielle est variable. Le rapport B/D se situe généralement entre 0.2

et 1.2 [Kärnä et Jochmann, 2003] (où D est le diamètre de la structure à la ligne d'eau). De ce point partent également des lignes de fracture radiales qui divisent le bloc semi-circulaire extrudé de la plaque en plusieurs sous-blocs comme le montre la figure 5.3.

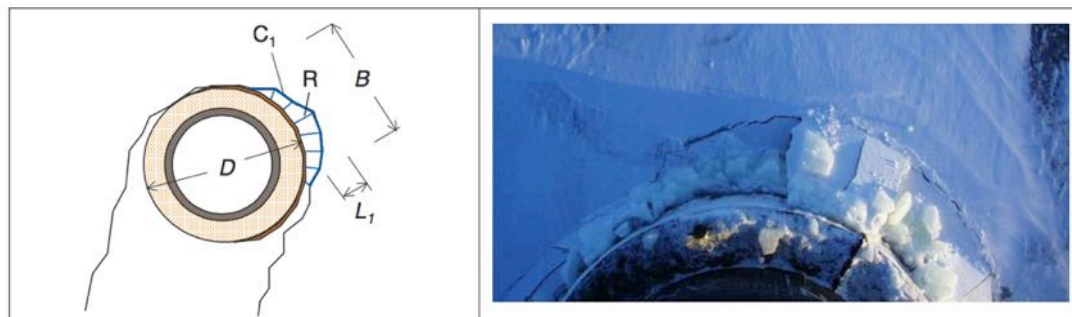


FIGURE 5.3 – **Phénomène de rupture par flexion d'une plaque de glace** – Lors du contact entre plaque de glace et une structure inclinée (de diamètre D), une ligne de fracture ($C1$) circumférentielle de rayon (B) est créée en raison de la contrainte verticale, à une distance $L1$ du point de contact. Des lignes de rupture radiales (R) causées par la contrainte horizontale viennent subdiviser le bloc initialement extrudé de la plaque. Source : [Kärnä et Jochmann, 2003]

Ces blocs sont ensuite poussés contre la structure par la plaque et/ou le courant, suivant l'inclinaison de la structure, tandis qu'une partie d'entre eux sont évacués de la zone de contact, sur les côtés de la structure. La figure 5.4 illustre l'enchaînement de ces étapes. Cette figure nous permet également de constater que, plus que la contrainte ponctuelle lors du contact, c'est surtout le poids des blocs qui s'accumulent sur la structure qui augmente au cours du temps.

Les blocs créés suite à une rupture de la plaque par flexion sont eux-mêmes soumis à ce mode de rupture. Les chocs avec les autres blocs et le frottement contre la structure notamment, peuvent causer des ruptures secondaires, voire tertiaires, des blocs tant qu'ils ont une taille supérieure à environ 5 fois l'épaisseur de la plaque. Ce comportement est surtout observé pour des faibles vitesses de dérive et des structures larges, cas dans lesquels les blocs sont moins facilement évacués du front de la structure.

5.1.2.2 Phénomène d'empilement

Que ce soit par écrasement ou par flexion, les modes de rupture que nous avons décrits conduisent à la formation de blocs de glace plus ou moins gros, qui sont poussés contre la structure et s'empilent contre elle. Ces empilements de blocs, qui peuvent atteindre quelques dizaines de mètres, contribuent à modifier les efforts sur la plaque de glace au voisinage de la structure. Ainsi, deux contraintes verticales, liées aux blocs, s'exercent sur la plaque. La première est dirigée vers le bas et est due aux blocs qui pèsent sur la plaque, qui forment ce qui est appelé la **voile** de l'amas. La seconde est dirigée vers le haut et provient des efforts générés par la flottabilité des blocs qui se retrouvent piégés sous la plaque. Ceux-ci forment la **quille** de l'amas. L'analogie avec ces termes nautiques est évidente dès lors qu'on observe le phénomène latéralement, comme sur la figure 5.5.

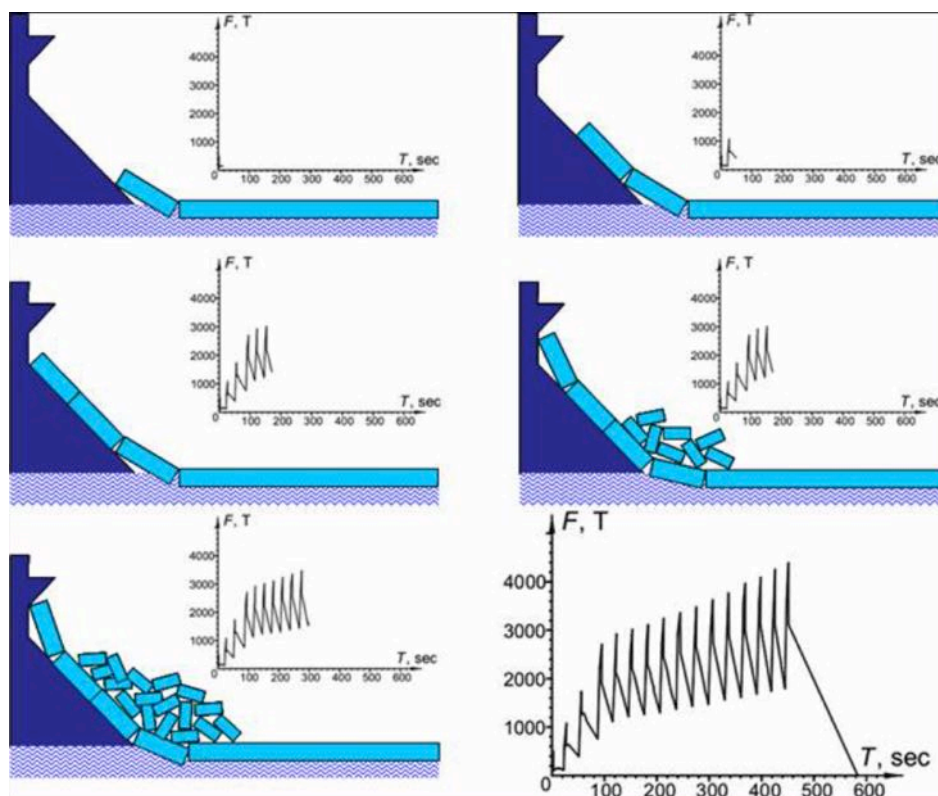


FIGURE 5.4 – **Évolution de la charge associée à la rupture par flexion** – La surface de la structure impose une contrainte verticale au front de la plaque de glace qui se fracture en créant des blocs. Ces blocs sont poussés par la plaque et s'accumulent sur la structure au rythme des ruptures successives de la plaque. La charge sur la structure connaît des pics au moment de la collision avec la plaque, tandis que la valeur moyenne augmente à mesure que l'amas de blocs grossit. Source : [Loiset, 2006]

Cette figure illustre la rupture de la plaque par flexion qui survient au moment où la contrainte exercée par la voile devient trop importante. Nous assistons alors à un transfert d'une partie des blocs de la voile vers la quille, et donc à une forte chute de la charge exercée sur la structure lorsque celle-ci est inclinée.

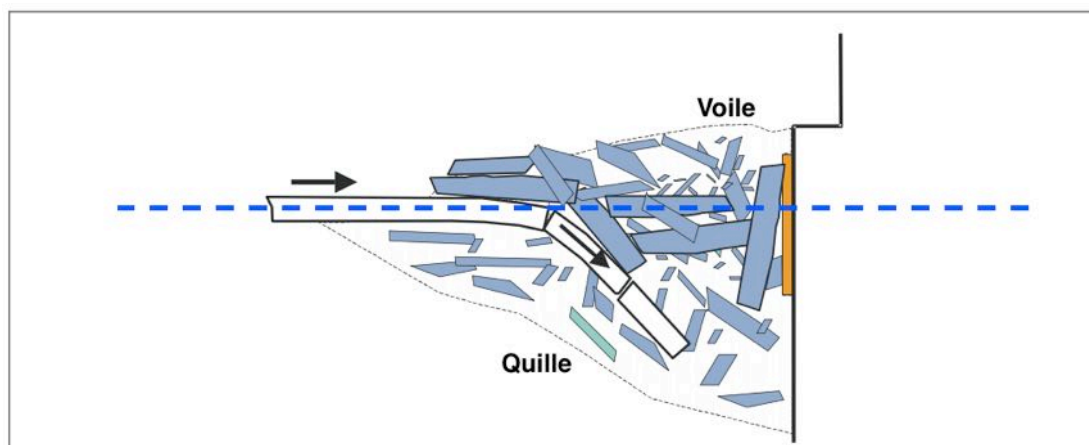


FIGURE 5.5 – **Rupture par empilement** – Les blocs de glace provenant des précédentes ruptures par flexion ou par écrasement s'accumulent sur la partie supérieure de la plaque qui finit par céder sous ce poids. Une partie des blocs chute alors dans l'eau et se retrouve sous la plaque. Source : [Kärnä et Jochmann, 2003]

5.1.3 Synthèse

Dans cette section, nous avons présenté le contexte de notre étude. Le juste dimensionnement des structures offshore pour les milieux polaires est un enjeu majeur dans la mesure où les contraintes climatiques et économiques vont croissantes. Nous avons vu que les interactions glace-structure sont des phénomènes complexes qui soumettent les structures à d'importants efforts. Il semble difficile d'estimer ces efforts *a priori* tant ils dépendent de nombreux paramètres et surtout de la dynamique des interactions.

Nous avons décrit les modes de rupture les plus fréquemment observés pour les structures que nous souhaitons étudier. Chacun d'eux implique la création de blocs de glace de tailles différentes qui jouent un rôle important dans l'évolution de la charge que supporte une structure. Il est d'ailleurs important de noter que ces différents modes de rupture ne s'excluent pas les uns les autres, mais coexistent souvent au sein du contact glace-structure, ce qui complexifie leur étude.

Ainsi, pour estimer au mieux les forces qui s'exercent sur une structure donnée, quelle que soit sa forme, mais également la dynamique de la glace aux abords de cette structure, nous proposons dans la section suivante un nouvel outil de simulation pour le design de structures offshore.

5.2 Ice-MAS : Ice Multi-Agent Simulator

La conception de structures offshore pour des conditions arctiques est une tâche complexe en raison de la difficulté de prédire les effets des phénomènes liés à la glace. C’est pourquoi de nombreux travaux ont été menés aux cours des dernières décennies pour tenter de modéliser et de simuler ces phénomènes, poursuivant ainsi deux objectifs. Tout d’abord, il s’agit de simuler l’écoulement de la glace aux abords d’une structure, qu’elle soit fixe ou flottante, pour s’assurer que l’empilement des blocs ne menacera pas les installations situées à son sommet. Ensuite, en estimant correctement les efforts imposés par la glace, il devient possible de dimensionner la structure au plus juste de manière à optimiser les coûts de production et d’améliorer la sécurité.

Les recherches ont débouché sur un ensemble de modèles analytiques, empiriques et numériques. La grande majorité des méthodes proposées ne se focalise cependant que sur un problème particulier à la fois, tant l’intrication des phénomènes rend impossible la formulation d’un modèle unique.

En juin 2012, les sociétés TECHNIP, CERVVAL et BUREAU VERITAS se sont associées pour développer un simulateur d’interactions glace-structure appelé ICE-MAS. L’intérêt de cet outil est d’adopter une approche multi-agents et multi-modèles et de superposer des méthodes de formalismes différents afin de profiter des avantages de chacune. Il a déjà fait l’objet de plusieurs publications aux conférences majeures du secteur de l’offshore [Septseault et al., 2014], [Septseault et al., 2015], [Dudal et al., 2015].

Le propos de cette section est de donner une vue d’ensemble de l’implémentation du simulateur et des résultats qui sont les siens, dans l’optique d’identifier des pistes d’optimisation du temps de simulation.

5.2.1 Approche multi-modèles

Les interactions entre la glace et les structures offshore de toutes formes font l’objet de nombreuses études, avec autant d’approches que de problèmes et d’échelles considérés. Malgré tout, prédire les efforts exercés sur une structure reste un défi car nous avons affaire à une combinaison de phénomènes complexes qui dépend de nombreux paramètres tels que la géométrie de la structure, les propriétés de la glace ou encore des courants et fonds marins. Les méthodes proposées jusqu’ici pour modéliser les interactions glace-structure, aussi performantes soient-elles, ne se concentrent chacune que sur un aspect précis du problème, et sont de ce fait incapables d’avoir une vue d’ensemble du système glace-structure-environnement. Par ailleurs, certains comportements observés restent encore aujourd’hui mal compris par les experts qui disposent seulement de données de terrains pour les décrire.

Dans le chapitre 2, nous avons montré que les systèmes multi-agents nous permettent d’utiliser conjointement des modèles d’origines diverses. Ainsi, des modèles numériques peuvent compléter des modèles analytiques par exemple, de même que des lois empiriques peuvent combler les manques de connaissances précises sur le système étudié. De plus, en adoptant une approche phénoménologique, nous pouvons avoir un point de vue centré sur la

dynamique du système.

Nous avons choisi cette démarche pour construire le simulateur ICE-MAS : nous avons couplé un ensemble de modèles spécifiques des phénomènes qui entrent en jeu dans les interactions glace-structure. De cette manière, ICE-MAS permet de simuler les actions réciproques entre la structure, la plaque de glace, les blocs issus des ruptures de la plaque, l'hydrodynamique et le fond marin. La figure 5.6 donne une vue d'ensemble de ces interactions que nous allons à présent détailler.

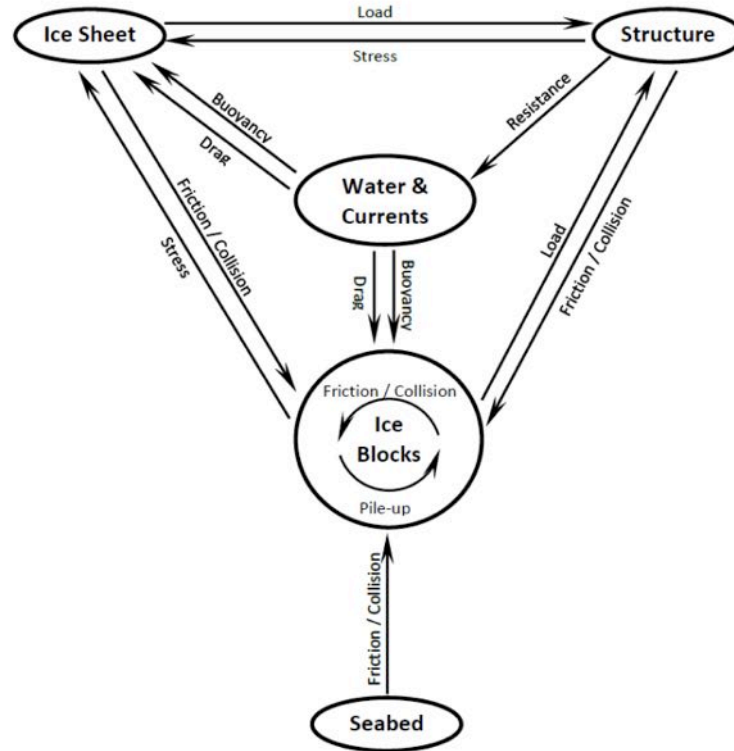


FIGURE 5.6 – **Diagramme d'interactions de Ice-MAS** – ICE-MAS est construit en suivant une approche phénoménologique. De ce fait, les flèches qui représentent les interactions entre les entités en présence sont chacune implémentée par un agent-interaction. Certaines des entités disposent d'un comportement propre (par exemple la plaque de glace qui dérive) et sont modélisées par un agent-entité.

5.2.1.1 Interactions solide/solide

La détection des collisions entre les différentes entités qui composent notre modèle est réalisée par la bibliothèque libre BULLET PHYSICS [Coumans, 2012]. Celle-ci permet la simulation des contacts de plusieurs milliers de solides en temps réel.

La structure, la plaque de glace et les blocs sont considérés comme des solides indéformables avec six degrés de liberté. Chaque collision génère une impulsion au point de contact entre les deux solides. La somme des impulsions pour un solide donné est intégrée à la fin du pas de temps pour modifier ses vitesses linéaire et angulaire.

Cette opération est réalisée par le solveur interne de BULLET auquel nous transmettons la contribution de chacun des phénomènes, hormis la force de gravité déjà calculée par BULLET.

Friction

Nous considérons une force de friction sèche qui résiste au mouvement de deux solides qui glissent l'un contre l'autre :

$$F_{friction} = \mu \cdot N \quad (5.1)$$

où

- μ est le coefficient de friction,
- N est la force normale à l'aire de contact, comme illustré par la figure 5.7.

La friction sèche opère selon deux modes : la friction statique dès lors que les vitesses relatives des deux solides en interaction sont faibles, la friction cinétique pour une plus grande vitesse relative. La différence entre ces deux modes intervient dans l'équation (5.1) au niveau de la valeur du coefficient de friction. Pour chaque couple de solides en frottement, nous adaptons donc la valeur du coefficient μ d'après les données expérimentales [Frederking et Barker, 2002] :

- $\mu = 0.45$ pour une vitesse relative nulle,
- $\mu = 0.15$ sinon.

Pour ce qui est de la friction entre deux blocs de glace, ou d'un bloc avec la plaque, nous avons choisi des valeurs communément admises pour notre cas d'étude [Fransson et al., 2011] :

- $\mu = 0.8$ pour une vitesse relative nulle,
- $\mu = 0.35$ sinon.

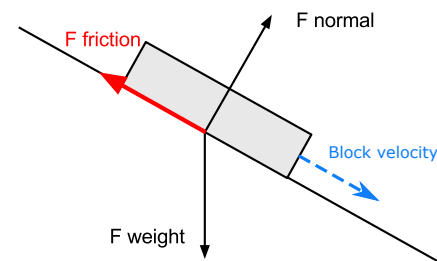


FIGURE 5.7 – **Interaction solide/solide** – L'interaction entre deux solides, qu'ils soient des blocs de glace, la structure ou la plaque de glace, génère une force résistive de friction à leur point de contact.

5.2.1.2 Hydrodynamique

Nous avons vu qu'en fonction de la géométrie de la structure, les blocs ont tendance à s'accumuler sur le front de la structure, à la fois au-dessus de la plaque en formant une voile, mais également sous l'eau en formant une quille. Ce comportement est causé par les courants marins qui poussent les blocs. Nous devons donc modéliser les effets de la mer sur les différents types de solides présents dans notre modèle.

Flottabilité

La flottabilité est une force qui s'exerce sur des corps immergés dans un fluide, et donc soumis à la poussée d'Archimède.

La flottabilité de chaque bloc est calculée en fonction de son volume immergé et elle est appliquée au centre de flottabilité du solide, comme montré par la figure 5.8. Pour un bloc de glace donné, la force de flottabilité $\vec{F}_b$ est calculée de la manière suivante :

$$\vec{F}_b = -\rho_{water} \cdot V \cdot \vec{g} \quad (5.2)$$

où :

- ρ est la densité de l'eau,
- V est le volume immergé,
- g est l'accélération.

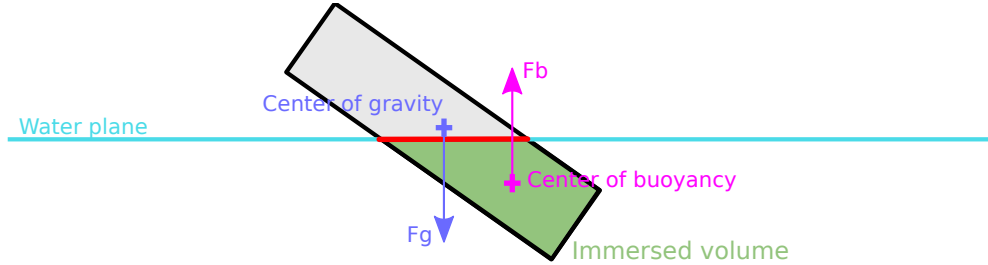


FIGURE 5.8 – **Flottabilité d'un bloc** – Pour chaque bloc de glace, nous calculons la force d'archimède $\vec{F}_b$ associée à son volume immergé. Cette force est appliquée au centre de flottabilité.

Nous reviendrons sur la flottabilité de la plaque de glace lorsque nous discuterons de la modélisation des ruptures.

Traînée

En mécanique des fluides, la traînée est la force d'inertie qui s'oppose au mouvement d'un corps dans un fluide. Dans notre modèle, cette force s'applique d'une part sur les blocs immergés dans l'eau et d'autre part sur la structure.

Nous considérons le courant marin comme un champ de vitesses uniformes, perturbé localement aux abords de la structure. Ces perturbations jouent un rôle important dans l'écoulement des blocs autour de la structure. Pour les prendre en compte, nous calculons le courant autour de la structure grâce à l'outil de calcul de dynamique des fluides OPEN-FOAM³ qui repose sur la méthode des éléments finis. Cette méthode permet la résolution numérique des équations de Navier-Stokes sur la base d'un maillage statique. Le temps de calcul de la dynamique des fluides autour de ce maillage est généralement de l'ordre de plusieurs minutes. Il est donc inenvisageable d'utiliser de calculer la forme du courant à chaque pas de temps de simulation de ICE-MAS ou même à chaque changement de forme de la quille. Pour cette raison, la forme du courant est uniquement calculée au début de la simulation, pour une structure donnée. Nous reviendrons dans la section 5.3 sur cet aspect qui nous empêche de prendre en compte les effets locaux de perturbation du courant par les blocs de glace générés au cours de la simulation. Pour autant, nous devons prendre en compte l'accumulation des blocs devant la structure qui diminue de fait son exposition au courant. Ainsi, lors de la phase de pré-calcul, nous déterminons la force maximale du courant sur la structure sur laquelle nous appliquerons un ratio adaptatif au cours de la simulation, en fonction de l'occultation des blocs. Cette opération est illustrée par la figure 5.9.

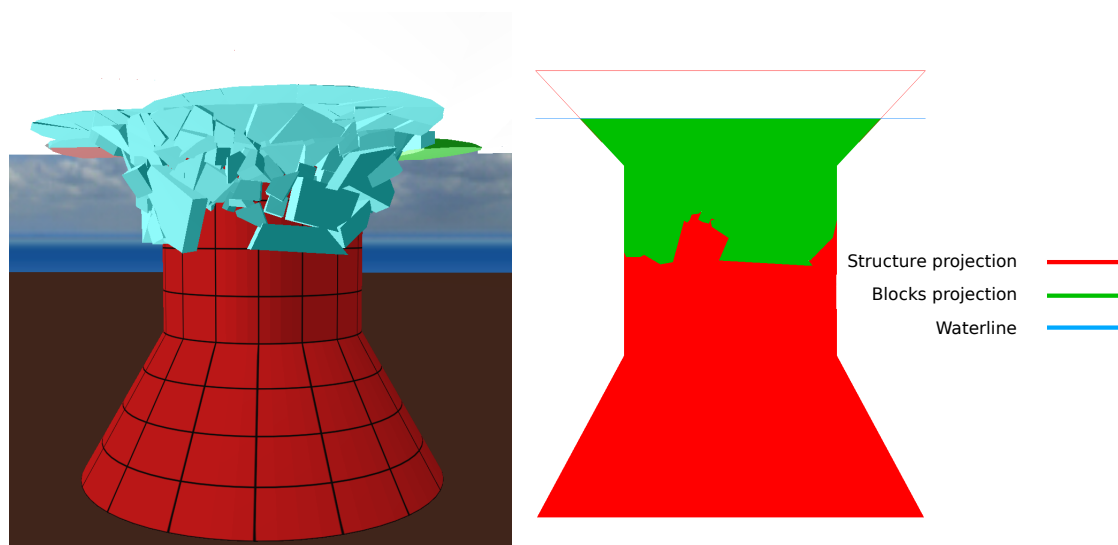


FIGURE 5.9 – **Calcul de la traînée de la structure** – Le courant exerce une force sur la structure que nous calculons au début de la simulation. Puisque le courant est pré-calculé sans glace, nous devons prendre en compte l'occultation progressive de la structure par les blocs qui s'accumulent devant la structure. Nous réalisons cette opération en calculant une projection en 2D qui nous permet de définir un ratio à appliquer à la force sans blocs.

Nous obtenons un maillage du courant autour de la structure, comme nous pouvons le voir sur la figure 5.10. Pour chaque bloc de glace, nous calculons la surface de référence, c'est-à-dire la projection en deux dimensions de l'objet par rapport à la direction du vecteur vitesse qui correspond à la maille dans laquelle il se trouve. Cette surface est utilisée pour le calcul de la force de traînée qui s'exprime :

3. Open Field Operation and Manipulation, <http://www.openfoam.com>

$$F_d = \frac{1}{2} \cdot \rho_{water} \cdot v^2 \cdot C_d \cdot A \quad (5.3)$$

où :

- ρ_{water} est la densité de l'eau,
- v est la vitesse relative de l'objet par rapport à celle du fluide,
- C_d est le coefficient de traînée,
- A est la surface de référence.

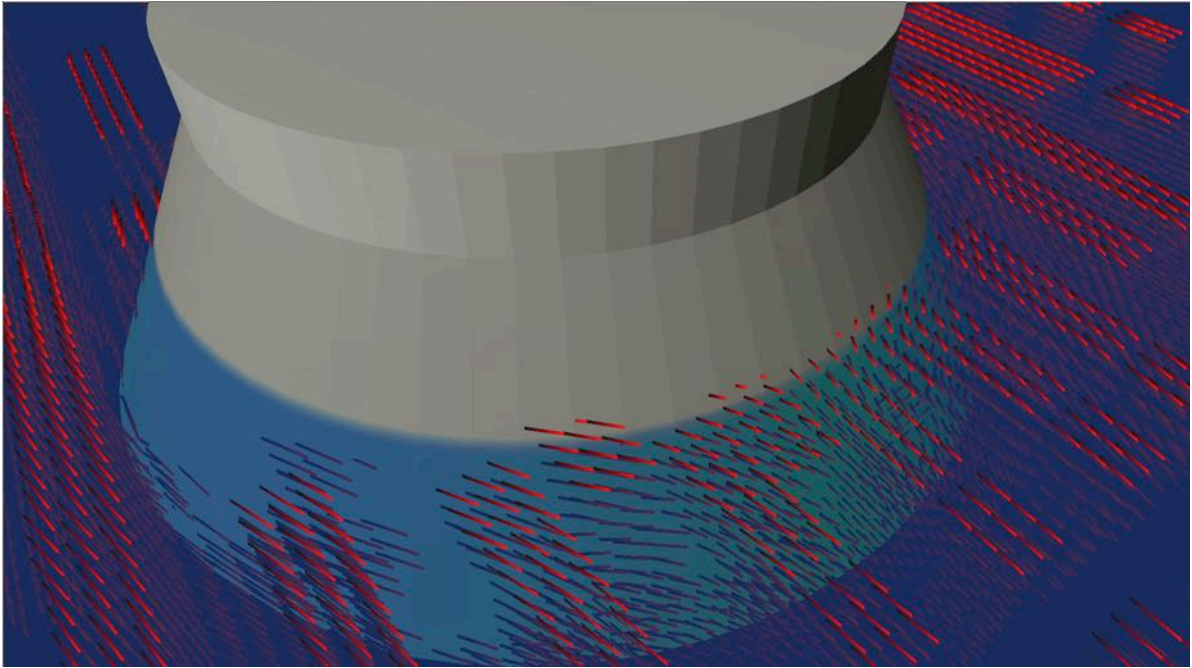


FIGURE 5.10 – **Dynamique du courant autour de la structure** – Le courant contribue à évacuer les blocs de glace qui s'accumulent devant la structure. Il est donc nécessaire de prendre en compte l'impact de la structure sur sa dynamique. Nous pouvons voir que le champs de vecteurs utilisé pour discrétiser spatialement le courant est perturbé aux abords de la structure et la contourne.

5.2.1.3 Rupture de la plaque de glace

Nous avons vu que l'interaction entre une structure et une plaque de glace dérivante conduit à une rupture de cette dernière selon deux modes principaux : la rupture par écrasement et la rupture par flexion.

Rupture par flexion

Les observations de terrain et en laboratoire ont montré que la rupture par flexion est généralement le mode de rupture dominant, en ce qui concerne l'interaction d'une plaque de

glace avec une structure inclinée. Sur la base de ces données, Nevel a proposé de modéliser la plaque comme semi-infinie et reposant sur des fondations élastiques [Nevel, 1965] :

$$D \cdot \frac{d^4 w}{dx^4} + \rho_{water} \cdot g \cdot w = q \quad (5.4)$$

où :

- w est la déflexion de la plaque,
- q est la pression exercée sur la plaque,

et D est la raideur en flexion de la plaque que l'on peut exprimer :

$$D = \frac{E \cdot h^3}{12 \cdot (1 - \nu^2)} \quad (5.5)$$

où E et ν sont respectivement le module de Young et le coefficient de Poisson classiques de l'élasticité.

La solution de l'équation 5.4 nous permet de calculer les contraintes critiques exercées sur la plaque de glace. Ces contraintes sont comparées à des valeurs de référence de la résistance en flexion de la glace pour simuler l'apparition de lignes de fractures radiales. Pour autant, le nombre de ces fractures, de même que leur propagation, ne sont pas décrits formellement. Nous devons donc nous référer aux connaissances des experts métier pour ce qui est de la formation des fractures circonférentielles et du nombre de fractures radiales. Ainsi, en nous appuyant sur les travaux de [Nevel, 1972], nous modifions le modèle de la plaque par un modèle de poutre appliqué à des quartiers adjacents, comme illustré sur la figure 5.11, pour calculer le rayon de la fracture circonférentielle, tandis que le nombre des quartiers est défini de manière empirique [Nevel, 1992].

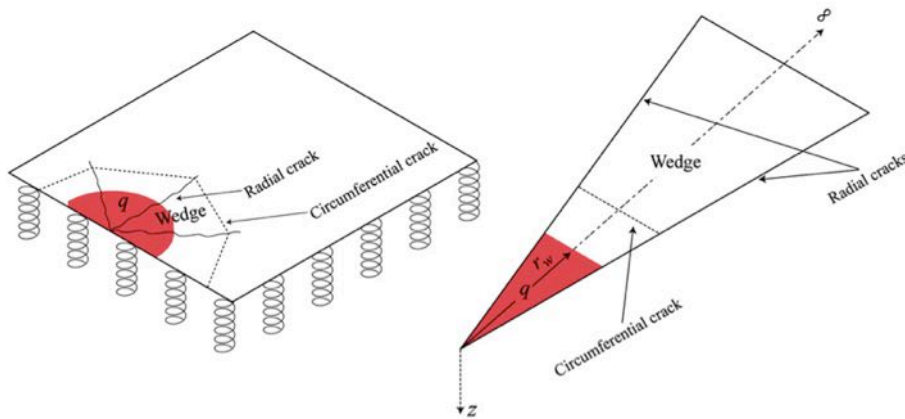


FIGURE 5.11 – **Détermination du rayon de fracture circonférentielle** – Le modèle de plaque est remplacé par un modèle de poutre sur fondation élastique et appliqué sur des quartiers. r_w désigne le rayon de la fracture circonférentielle et est égal à l'abscisse du moment de flexion maximal. Source : [Lubbad et Loset, 2011].

Comme nous l'avons vu dans la section 5.2.1.3, la rupture par flexion d'une plaque de glace peut survenir dans une seconde situation. Suite aux ruptures successives de la plaque contre la structure, les blocs de glace s'accumulent sur le front de la structure, au-dessus et au-dessous de la plaque. Le poids de ces blocs, ainsi que leur flottabilité, exercent une contrainte verticale sur la plaque qui peut dépasser la contrainte critique de flexion. Nous prenons en compte ce phénomène en distribuant les forces dues aux blocs sur la surface de plaque qu'ils couvrent [Marchenko et Karulin, 2005], comme illustré par la figure 5.12. Ceci a pour conséquence d'augmenter l'abscisse du moment de flexion maximal et donc le rayon de fracture circonférentielle.

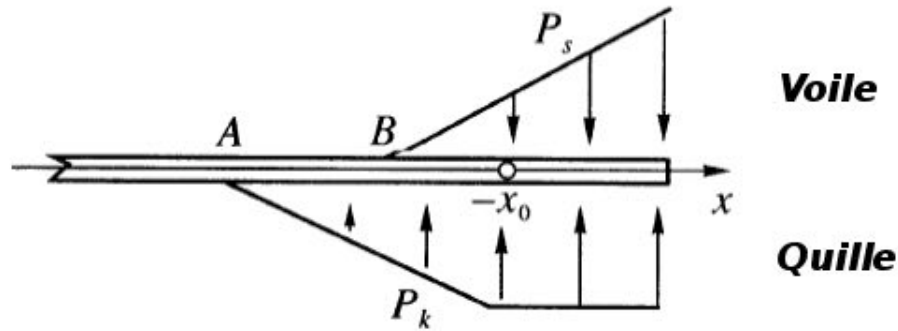


FIGURE 5.12 – **Distribution du poids et de la flottabilité des blocs sur la plaque** – Les efforts dus aux blocs sont distribués sur la surface de la plaque et modifient le rayon de fracture circonférentielle. Source : [Marchenko et Karulin, 2005].

La figure 5.13 montre la simulation de la rupture par flexion, suite d'une part à la collision initiale d'une plaque de glace avec la structure, et d'autre part, au poids des blocs accumulés sur la structure et la plaque.

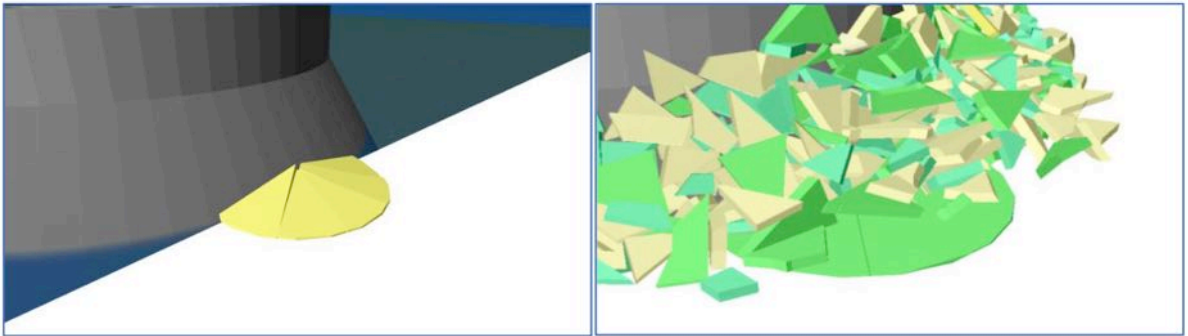


FIGURE 5.13 – **Simulation de ruptures par flexion** – La rupture de la plaque de glace par flexion intervient dans deux situations : suite à la collision avec la structure (à gauche) ou sous le poids des blocs accumulés sur le front de la structure (à droite).

Enfin, le calcul du moment de flexion est également effectué pour chaque bloc, dans la mesure où ceux-ci peuvent se rompre tant que leur taille est supérieure à la longueur caractéristique de la glace [Nevel, 1992] :

$$L = 2\sqrt{\frac{\sigma h}{3\gamma}} \quad (5.6)$$

où :

- σ est la résistance en flexion de la glace,
- h est l'épaisseur du bloc,
- γ est la densité de la glace.

Les blocs sont soumis aux forces provenant de leurs contacts avec la structure, la plaque, le fond marin et les autres blocs. La contrainte critique est calculée en prenant en compte le centre de gravité du bloc et des points de contact où les forces sont appliquées. La valeur est ensuite comparée à la résistance en flexion de la glace pour déterminer si le bloc doit se briser, selon des axes de fragilité pré-calculés qui dépendent de sa géométrie.

Rupture par écrasement

L'interaction de la plaque de glace avec une structure cylindrique ou un mur vertical se traduit généralement par une rupture de la plaque par écrasement.

Pour reproduire ce phénomène, nous utilisons le modèle de C. Daley [Daley, 1991]. La base de ce modèle est de considérer que la force mesurée sur la structure est directement liée à la formation des éclats de glace. Ainsi, lorsque les éclats se détachent de la plaque de glace, la surface de contact entre la plaque et la structure se réduit, ce qui cause une diminution de la force. L'idée générale est donc de calculer la pression exercée par la structure sur l'extrémité de la plaque. Lorsque cette pression dépasse le critère de Coulomb, une ligne de fracture apparaît avec un angle qui dépend de la surface de contact. La figure 5.14 illustre la succession de détachements des éclats sous l'avancée de la structure.

Le modèle de Daley se concentre sur la simulation de l'écrasement de la plaque et ne considère pas les éclats après leur détachement de la plaque. Pourtant, dans le cadre de notre approche multi-modèles, nous nous intéressons à l'influence de ces débris de glace. Les éclats de glace, calculés en deux dimensions, sont transformés en petits blocs de glace à part entière, idéalisés sous forme de prismes triangulaires. Cette représentation autorise la simulation de deux phénomènes importants :

- la circulation de la glace broyée autour de la structure par le courant ou par la formation d'amas de fragments devant la structure,
- l'accumulation des débris de glace au-dessus de la plaque de glace et la rupture de celle-ci par flexion sous le poids des fragments, comme illustré par la figure 5.15

5.2.2 Bilan

Nous avons développé un simulateur complet capable de prédire le comportement de la glace et les efforts qu'elle peut exercer sur les structures offshore. ICE-MAS est aujourd'hui utilisé par des ingénieurs structure pour tester *in virtuo* de nouveaux prototypes.

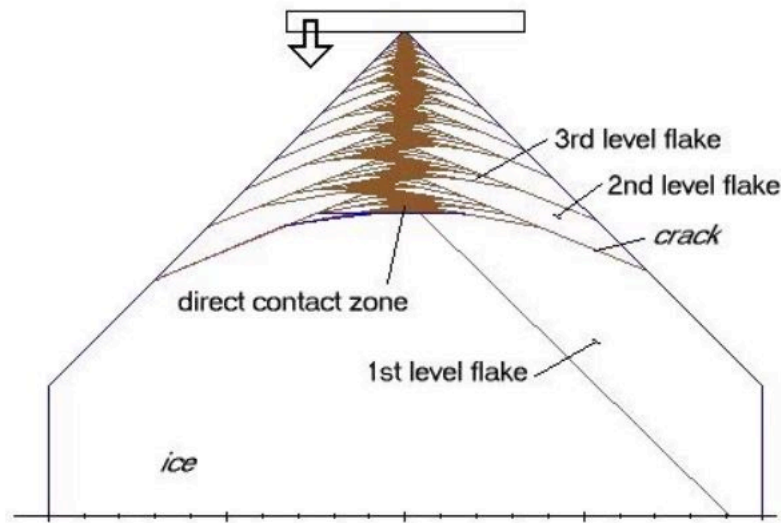


FIGURE 5.14 – **Détachement hiérarchique d'éclats de glace par écrasement** – Lorsque la structure avance, le front de la plaque de glace se rompt par écrasement et voit la formation successive d'éclats de glace selon un schéma hiérarchique, modifiant ainsi la surface de contact plaque-structure. Source : [Daley, 1991]

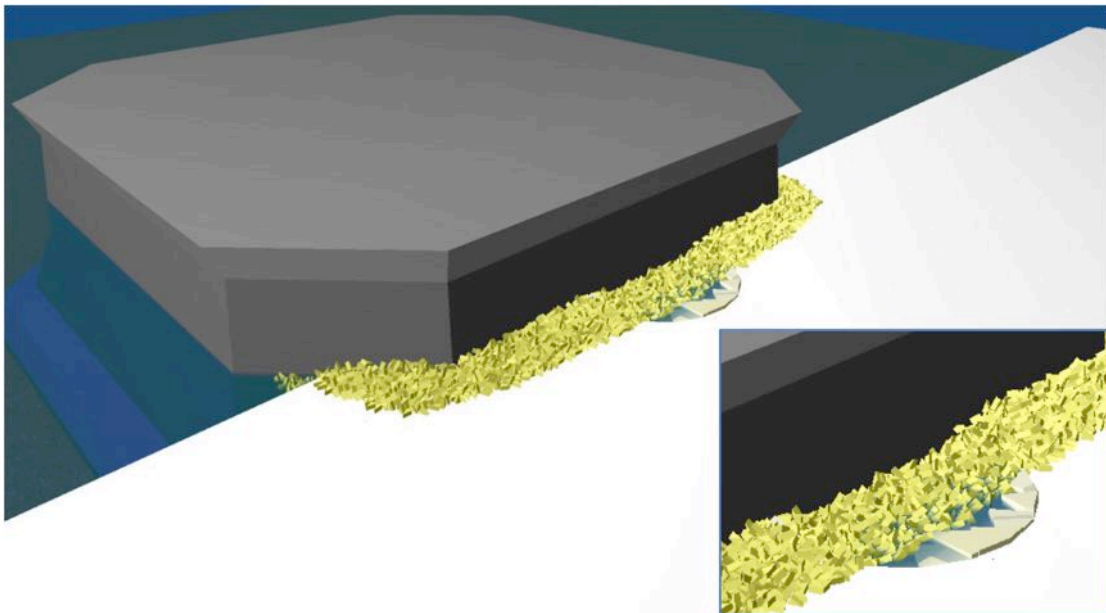


FIGURE 5.15 – **Combinaison des ruptures par flexion et par écrasement** – Sous le poids des débris de glace broyée issus de la rupture par écrasement de la plaque contre une structure verticale, la plaque de glace subit une rupture par flexion.

Notre approche a été de décomposer la modélisation de ce système complexe en ses phénomènes constitutifs. Ainsi, nous avons superposé des modèles spécifiques et de natures différentes grâce aux systèmes multi-agents. Le simulateur intègre actuellement une dizaine de phénomènes calculés pour les structures, la plaque de glace et chacun des milliers de blocs⁴. Notre propos n'est pas ici de discuter des résultats de ICE-MAS mais le lecteur intéressé pourra se référer aux publications dédiées [Septseault et al., 2014] [Dudal et al., 2015], qui montrent une bonne concordance avec des cas réels et les essais en bassin. Pour autant, nous pouvons imaginer plusieurs pistes d'optimisation pour notre outil.

La première consisterait à mieux représenter les phénomènes hydrodynamiques et notamment la forme du courant aux abords de la structure. Nous avons vu que celui-ci est pour le moment pré-calculé avant la simulation sur une structure seule et est utilisé pour évacuer les blocs de glace. Or il est évident que si le courant est perturbé par la structure, il l'est également par les blocs de glace générés au cours du temps. Nous allons donc appliquer la technique de co-simulation présentée dans la section 4.2 pour prendre en compte la présence des blocs dans le calcul du courant.

Notre seconde piste d'optimisation concerne le temps de calcul. Le simulateur tel que nous l'avons décrit nous permet de calculer les efforts sur la structure de manière précise avec un point de vue centré sur les blocs, qui nous fournit par ailleurs d'autres données dynamiques telles que la hauteur maximale des blocs ou les pressions locales. Cependant, la simulation demande de plus en plus de temps à mesure que le nombre de blocs, qui atteint rapidement plusieurs milliers, augmente. Dans une optique de prototypage rapide, nous proposons de changer de niveau de description afin d'obtenir une estimation plus grossière des efforts sur la structure en peu de temps de calcul. Nous décrirons dans la section 5.4 le modèle envisagé ainsi que ses limites qui nous pousseront à mettre en œuvre notre seconde technique de co-simulation pour déterminer de manière implicite un jeu de paramètres.

5.3 Co-simulation explicite de la dynamique du courant

Nous avons vu que le courant marin joue un rôle important dans l'activité des blocs de glace immergés aux abords de la structure. Une partie de ces blocs est poussée contre la structure, participant ainsi à la charge mesurée, tandis que les blocs les plus à l'extérieur sont évacués sur ses côtés.

Actuellement, ces phénomènes ne sont que partiellement reproduits dans ICE-MAS. En effet, l'interaction entre un fluide et un corps solide est réciproque. Or, l'impact des blocs de glaces sur la dynamique du courant est pour le moment négligée car le calcul de la dynamique de fluides (CFD⁵) est réalisé une seule fois en amont de la simulation. Ce choix est à imputer au coût en temps de calcul relativement important de l'opération.

4. Ce nombre est fortement dépendant de l'inclinaison de la structure, de sa largeur et de l'épaisseur de la plaque de glace. On dénombre en moyenne 6 000 blocs indépendants en régime établi, pour une structure inclinée à 60° et une plaque d'un mètre d'épaisseur.

5. *Computational Fluid Dynamics*

Pour autant, nous soupçonnons que la présence des blocs devant la structure perturbe la circulation du fluide de façon de moins en moins négligeable au fur et à mesure de leur accumulation. Nous souhaitons donc améliorer le réalisme de la simulation en prenant en compte les effets des blocs sur le courant.

Dans cette section, nous proposons de mettre en œuvre notre architecture de co-simulation pour mettre à jour le champ de vitesses associé au courant régulièrement en fonction du contexte.

5.3.1 Rôle de l'agent-*watchdog*

L'implémentation de notre méthode commence par l'ajout d'un agent-*watchdog* dans ICE-MAS. Son rôle est de détecter quand recalculer la dynamique du courant. En effet, si nous souhaitons améliorer notre modèle en prenant en compte les blocs accumulés devant la structure, nous ne pouvons pas raisonnablement envisager de recalculer le courant à chaque pas de simulation.

Un compromis acceptable est de dire que le flux n'est en fait modifié que petit à petit au fil des interactions entre la structure et la plaque de glace. Ainsi, nous pouvons considérer le champ de vitesses inchangé pendant plusieurs pas de temps. Plusieurs critères peuvent alors être examinés pour déterminer si le courant a besoin d'être mis à jour.

Nous pouvons par exemple choisir une fréquence constante de co-simulation. Cette solution a l'avantage d'être très simple à mettre en œuvre mais il est difficile de choisir arbitrairement la fréquence. Comme nous l'avons vu, un agent-*watchdog* ne peut avoir qu'une seule requête active pour un paramètre donné. Une fréquence arbitraire ne s'adapte pas à cette contrainte car, par définition de la co-simulation, nous ne connaissons pas à l'avance le délai de réponse, potentiellement variable, d'une requête. Une fréquence trop faible engendrera des calculs superflus et une fréquence trop élevée peut laisser diverger la simulation trop longtemps.

Nous préférons une seconde option qui consiste à évaluer un critère physique. L'agent-*watchdog* peut évaluer l'évolution du volume occupé par l'ensemble structure-blocs au cours du temps. Nous définissons un seuil sur la différence de ce volume entre deux itérations de co-simulation. Ce critère permet de ne calculer le courant que quand c'est réellement nécessaire et de s'adapter à la dynamique de la simulation. Nous pouvons notamment voir, sur la figure 5.16, que la fréquence des itérations de co-simulation n'est pas régulière et augmente au fur et à mesure que les blocs s'accumulent. Quand le seuil est dépassé, l'agent-*watchdog* génère une requête pour le paramétrage du champ de vitesses avec une copie de l'état de la simulation et de ses entrées.

La requête est ensuite envoyée au bus de co-simulation qui la transmet aux composants capables de fournir des données de dynamique des fluides.

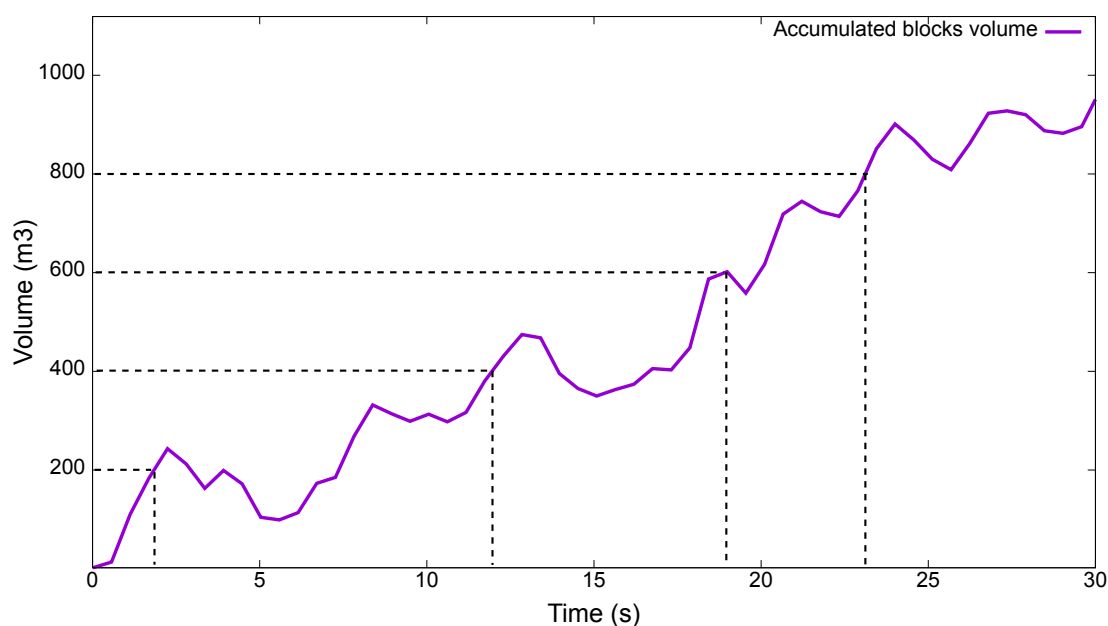


FIGURE 5.16 – **Fréquence de co-simulation et volume** – Pour fixer la fréquence des itérations de co-simulation, nous considérons un critère physique plutôt qu’une constante. Ainsi, l’agent-*watchdog* calcule le volume des blocs accumulés devant la structure à chaque pas de temps car ce sont ces blocs qui perturbent le courant. Quand la variation de volume entre deux itérations dépasse le seuil de $200m^3$, une requête de paramétrage est générée. Nous observons que la fréquence de co-simulation a tendance à augmenter au fur et à mesure que les blocs s’accumulent.

5.3.2 Traitement de la requête

Comme nous l'avons vu, plusieurs types de composants peuvent être connectés au bus de co-simulation pour répondre à la requête. Dans le cas présent, nous avons identifié deux potentiels fournisseurs de données de CFD : une base de cas et un modèle auxiliaire, qui n'est autre qu'une interface pour le logiciel OPENFOAM.

Base de cas

Pour déterminer la dynamique du courant autour de la structure et des blocs, nous pouvons nous référer à une base de données nourrie par les simulations précédentes. La sélection du cas le plus proche se fait à partir des données contenues dans la requête.

Cette opération se rapproche du domaine de la recherche multicritères qui propose de nombreuses méthodes de calcul de solution optimale. Nous n'approfondirons pas ces approches dans le cadre de ce mémoire. Le lecteur intéressé pourra se tourner vers des travaux en lien avec notre architecture de co-simulation [Beauchesne et al., 2015].

Pour notre application, le parcours de la base est simplifié car celle-ci a été construite spécifiquement. Deux données sont utilisées comme clés : la boîte englobante des blocs accumulés devant la structure et la vitesse du courant en eau libre⁶ qui fait partie des paramètres d'entrée de ICE-MAS. La boîte englobante est construite en ne prenant en compte que les blocs dont la vitesse linéaire est faible en comparaison de la norme du vecteur vitesse du courant. Nous considérons ces données comme suffisantes pour discriminer efficacement les entrées de la base de cas car nous savons que le comportement général de l'accumulation des blocs reste similaire, quelle que soit la structure, et est équitablement réparti sur toute sa largeur.

Ainsi, à chaque ensemble $\{L_b, W_b, H_b, v_b\}$ de la base (respectivement les trois dimensions de la boîte englobante et la vitesse du courant) est associé une grille de vecteurs vitesse. La sélection du meilleur résultat consiste à trouver le jeu de données qui minimise la distance avec le même ensemble issu de la requête que nous calculons de la manière suivante :

$$d = \sqrt{(L_r - L_b)^2 + (W_r - W_b)^2 + (H_r - H_b)^2 + (v_r - v_b)^2} \quad (5.7)$$

où les indices r et b désignent les données de la requête et de la base respectivement.

Le meilleur résultat est inscrit dans la requête qui est retournée à l'agent-*watchdog*.

Modèle auxiliaire

Pour calculer un champ de vitesses qui s'adapte au contexte de la simulation, nous utilisons la même technique que pour son initialisation en faisant appel au logiciel OPENFOAM.

6. La vitesse du courant en eau libre est toujours donnée linéairement car elle est considérée colinéaire au déplacement uniaxial de la plaque de glace.

Le principe est de calculer les perturbations du courant sur un objet 3D qui représente l'ensemble structure-blocs. Cet objet est généré à partir de l'état de la simulation contenu dans la requête. La première étape du calcul de CFD consiste à mailler le domaine à simuler. Le résultat de cette opération pour un ensemble structure-blocs extrait de ICE-MAS est représenté sur la figure 5.17. Les conditions aux limites (vitesses initiales, viscosités, etc.) sont ensuite initialisées, toujours grâce aux données de la requête.

Nous démarrons ensuite une simulation en spécifiant un seuil de stabilité comme critère d'arrêt primaire et un nombre d'itérations maximum en critère secondaire. Ainsi, en fonction du contexte, nous obtiendrons les résultats au mieux le plus tôt possible et au pire au bout d'un temps borné.

La grille de vecteurs vitesse est une sortie d'OPENFOAM obtenue par post-traitement, c'est-à-dire par transformation du maillage de simulation en un maillage cubique, utilisable par ICE-MAS. Ces données sont à la fois inscrites dans la requête qui est retournée à l'agent-*watchdog*, mais elles sont également enregistrées dans la base de cas.

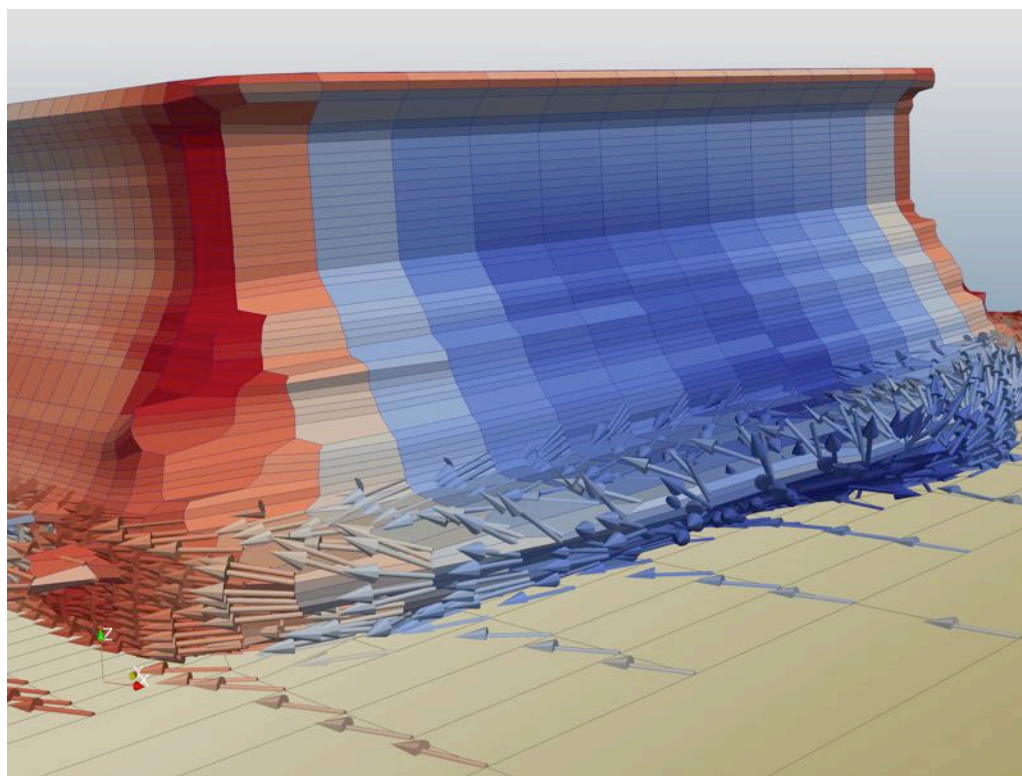


FIGURE 5.17 – **Maillage et simulation OpenFOAM** – La géométrie 3D de l'ensemble structure-blocs est extraite de la simulation ICE-MAS et maillée pour être simulée par OPENFOAM. À l'issue de la simulation nous récupérons les vecteurs vitesses du courant dont les normes et les orientations prennent désormais compte de la présence des blocs.

5.3.3 Intégration de la requête

Une fois le traitement de la requête terminé par les composants, l'agent-*watchdog* évalue la pertinence du résultat retourné.

Nous vérifions en premier lieu la norme des vecteurs vitesses de la grille calculée. Si celle-ci dépasse déraisonnablement la vitesse du courant en eau libre, nous pouvons estimer que le calcul de CFD n'a pas convergé.

En second lieu, nous mesurons la différence entre l'état courant de la simulation et celui inscrit dans la requête. Plus particulièrement, nous vérifions que les dimensions de la boîte englobante des blocs n'ont pas trop changé.

Si le résultat de la requête respecte ces deux conditions, la nouvelle grille de courant remplace celle de ICE-MAS. Dans le cas contraire, la requête est supprimée et une nouvelle co-simulation est lancée d'après le nouvel état courant de la simulation.

5.3.4 Résultats et discussion

Fréquence de co-simulation

Comme nous l'avons vu dans la section 5.2.1.2, la simulation OPENFOAM est une opération coûteuse au regard du pas de temps de simulation qui ne peut être faite de manière synchrone sur un pas de temps. Toutefois, en considérant maintenant un calcul asynchrone et effectué en parallèle grâce à la co-simulation, le temps de simulation de OPENFOAM est acceptable au regard de la dynamique de ICE-MAS. En effet, l'accumulation des blocs est un processus assez lent et, de ce fait, l'état de la simulation ICE-MAS n'a pas le temps de varier de façon significative le temps de la co-simulation. Ainsi, la nouvelle grille de courant calculée est cohérente avec l'état de la simulation lors de son application.

Impact sur la charge globale

La figure 5.18 montre la comparaison de la charge globale appliquée sur la structure, avec et sans la mise à jour du courant par co-simulation. Nous observons une différence de l'ordre de 10% sur la charge totale, ce qui peut être expliqué par le fait que les blocs ne sont plus poussés artificiellement sur la structure. En effet, en utilisant une grille de courant calculée uniquement au début de la simulation, les blocs accumulés subissaient une force de traînée colinéaire au vecteur vitesse du courant en eau libre. Désormais, la présence des blocs est répercutée sur le courant.

Par ailleurs, le fait de prendre en compte la présence des blocs affecte également la manière dont ils sont évacués sur les côtés de la structure. Nous remarquons que les blocs les plus proches de la structure ne sont plus autant détournés qu'ils l'étaient par le courant pré-calculé. Cela implique que pour un courant suffisamment faible, ils ont maintenant tendance à stagner de plus en plus devant la structure. Leur présence affecte les autres phénomènes.

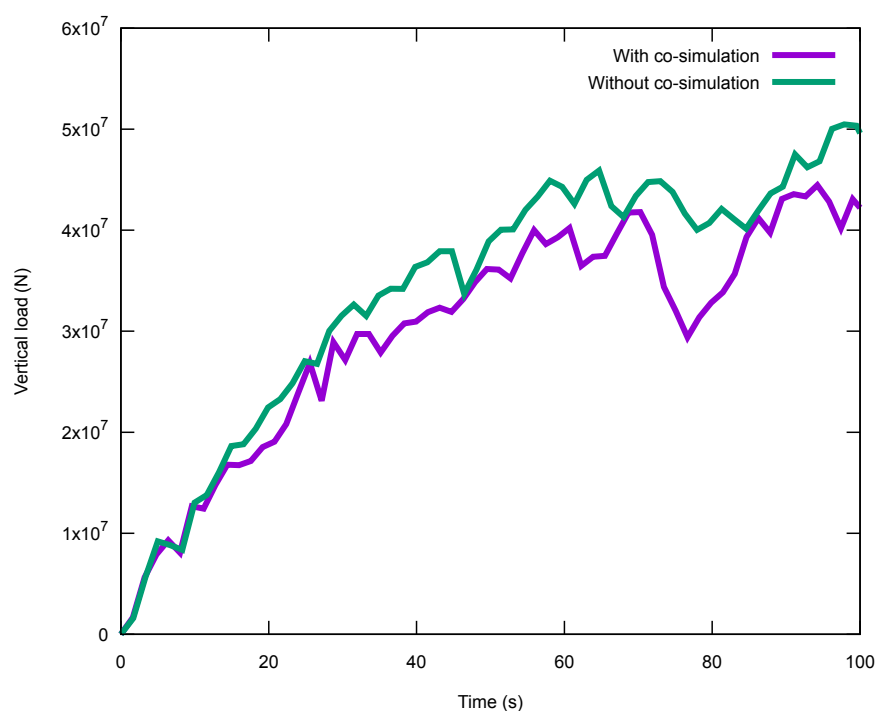


FIGURE 5.18 – **Influence de la co-simulation du courant sur la charge** – Nous comparons les efforts exercés sur la structure par les blocs, avec et sans co-simulation du courant. Nous remarquons que les efforts sont moindres en mettant à jour la dynamique du courant régulièrement. Cette différence s’explique principalement par deux phénomènes qui se compensent : d’une part les blocs immergés ne sont plus poussés artificiellement sur la structure et d’autre part ils ont tendance à plus s’accumuler et à retenir les blocs qui s’empilent sur la structure au-dessus de la plaque.

En particulier, les nouveaux blocs issus de ruptures de la plaque sont soutenus par les blocs immergés et ne plongent plus. Une partie d'entre eux est toujours évacuée sur les côtés de la structure, non plus par le courant, mais au-dessus du niveau de l'eau. L'autre partie reste accumulée sur la structure et exerce une charge supplémentaire.

Au delà d'être la parfaite illustration de l'intrication des phénomènes qui définit les systèmes complexes, ces résultats nous montrent que la détermination explicite d'un paramètre seul est une opération potentiellement délicate dont les conséquences générales sont difficiles à prédire.

5.3.5 Bilan

Dans cette section, nous avons procédé à une première passe d'optimisation de ICE-MAS. Celle-ci a consisté à rendre compte de l'impact des blocs de glace, accumulés sur le front de la structure au cours du temps, sur la dynamique du courant.

L'agent-*watchdog* ajouté à ICE-MAS détecte les variations significatives du volume des blocs accumulés et demande une mise à jour de la grille du courant. Soit les données associées à la requête nous permettent de faire le lien avec une simulation antérieure enregistrée dans une base de cas, et ce cas peut être réutilisé directement. Soit le cas est inédit et nous instancions un modèle auxiliaire pour calculer la nouvelle dynamique du courant à l'aide d'OPENFOAM.

Grâce à cette stratégie de détermination explicite du courant en cours de simulation, les phénomènes hydrodynamiques qui contribuent à l'évacuation des blocs s'adaptent maintenant au contexte de la simulation. Nous avons observé une différence significative de la mesure de la charge exercée sur la structure par rapport à la version du simulateur qui considérait le courant constant au cours du temps.

Ce résultat n'est pas encore le cas idéal qui consisterait à calculer la dynamique du fluide, de manière synchrone, à chaque pas de temps. Cependant, l'opération présentée ici a pu être effectuée sans augmenter le temps de calcul et sans interrompre l'exécution de ICE-MAS.

Nous visons à présent une optimisation du temps de calcul grâce à l'utilisation d'un modèle de plus haut niveau du phénomène d'empilement des blocs. Nous avons vu dans les chapitres précédents que ce type d'abstraction implique la synthèse de phénomènes microscopiques par des paramètres qu'il est parfois difficile d'évaluer. C'est pourquoi nous proposons de nous appuyer sur notre seconde stratégie de co-simulation pour déterminer ces jeux de paramètres implicitement et dynamiquement.

5.4 Co-simulation implicite du phénomène d'empilement

Pour construire le simulateur ICE-MAS, nous avons suivi une approche multi-agents mixte. De cette manière, nous avons pu mêler des lois empiriques à des agents-interaction et

des agents-entité. Ces derniers représentent les milliers de blocs de glace dont la dynamique individuelle est calculée par la bibliothèque BULLET PHYSICS.

Nous avons vu que la multiplication des agents-entité constitue généralement un frein à l'interactivité des simulations. Pour remédier à ce problème dans le cas de ICE-MAS, nous souhaitons employer un modèle de plus haut niveau qui abstrait le phénomène d'empilement des blocs. Ce changement de point de vue nous privera évidemment de certaines informations liées à la dynamique des blocs, mais il nous permettra d'effectuer du prototypage rapide pour cibler les meilleures conception de structures.

5.4.1 Abstraction de l'empilement des blocs

Comme nous l'avons montré dans la section 5.1.2.2, les études expérimentales des interactions glace-structure par les experts du domaine concluent toutes à un même scénario. Quand une structure posée sur le fond marin rencontre une plaque de glace dérivante, celle-ci se rompt et des débris de glace s'accumulent sur le front de la structure et forment deux amas :

- la voile, en appui sur la structure et le dessus de la plaque,
- la quille, en appui sur la structure et le dessous de la plaque.

Le modèle de A. Marchenko [Marchenko et Karulin, 2005] fait abstraction des blocs de glace en tant qu'« individu » et se concentre sur le processus qui conduit à la formation de ces deux éléments. Ce processus est illustré par la figure 5.19. La voile et la quille sont approximées respectivement par les polygones $BS_4S_3S_2$ et AS_2S_3 . L'approche de Marchenko est d'exprimer ces formes par le vecteur

$$\chi = \{\phi_s, \phi_k, \gamma_s, \gamma_k\} \quad (5.8)$$

où

- ϕ_s et ϕ_k sont les angles que forment la voile et la quille avec la plaque de glace.
- γ_s et γ_k sont les porosités (ratio air/blocs) de la voile et de la quille respectivement.

L'interaction de la glace avec la structure suit un cycle constitué de deux étapes. Dans un premier temps, les blocs de glace issus de la rupture de la plaque alimentent la voile en suivant la trajectoire $ABS_2S_3S_4$. Les fragments de glace retombent ensuite sur la surface libre BS_4 de la voile. Dans un second temps, la plaque de glace cède sous le poids de la voile. Une partie des blocs de la voile est alors transférée dans la quille pour atteindre un équilibre hydrostatique entre les deux amas.

Logiquement, la charge sur la structure a une dynamique similaire. Lors de la première phase, la charge croît linéairement à cause de l'apport continu de blocs dans la voile. Au moment du transfert de volume entre la voile et la quille, nous observons une forte chute de la charge. Cette dynamique est illustrée par la figure 5.20.

Marchenko utilise des lois de conservation de la masse, des moments et de l'énergie que nous ne détaillerons pas ici pour calculer l'évolution du vecteur χ . La rupture de la plaque due au poids de la voile est calculée selon la même méthode d'évaluation du moment de flexion présentée dans la section 5.2.1.3.

Le niveau d'abstraction choisi par Marchenko pour décrire la formation des amas nous prive d'informations précises sur la dynamique des blocs de glace. En contrepartie, les résultats de ce modèle dégradé donnent une approximation de la charge exercée sur la structure dans un temps de simulation bien plus court. L'auteur a toutefois remarqué que son modèle avait tendance à surestimer cette charge par rapport aux données expérimentales. Cette erreur est en grande partie causée par le fait que les différents paramètres du vecteur χ sont supposés constants au cours de la simulation. Par ailleurs, ce modèle 2D peut être extrudé pour donner la charge sur toute la largeur d'une structure mais il ne prend pas en compte l'échappement des blocs le long de ses côtés.

Pour profiter des avantages de ce modèle de haut niveau, tout en gommant ses écueils, nous proposons d'utiliser notre seconde stratégie de co-simulation.

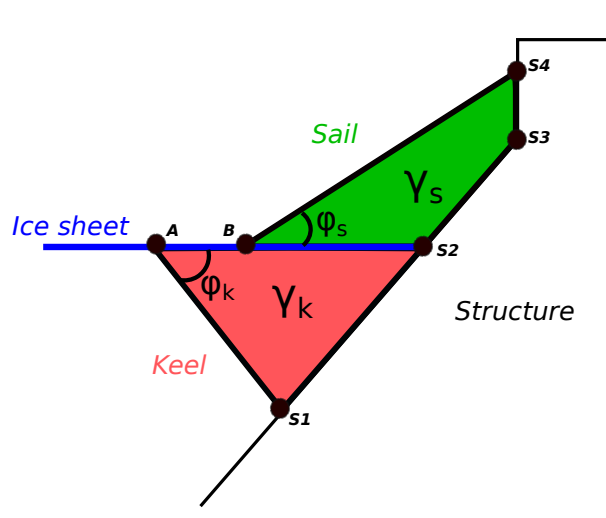


FIGURE 5.19 – **Abstraction du phénomène d'empilement des blocs** – Le modèle de Marchenko [Marchenko, 2006] formule une description haut niveau de la formation des amas de blocs de glace devant la structure. Les blocs générés par la rupture de la plaque sont poussés le long de la ligne $ABS_2S_3S_4$ pour former la voile (*Sail*). La plaque finit par céder sous le poids de la voile et une partie d'entre eux est transférée dans la quille (*Keel*), décrite par la ligne ABS_2S_1 , pour obtenir un nouvel équilibre hydrostatique. Les paramètres γ_s et γ_k correspondent aux porosités respectives de la voile et de la quille.

5.4.2 Rôle de l'agent-*watchdog*

Comme nous l'avons vu, l'apport principal de notre architecture de co-simulation consiste à adapter le paramétrage des modèles de haut niveau en cours de simulation, sur la base de résultats issus de simulations plus fines. Dans le cas présent, les paramètres à calibrer sont

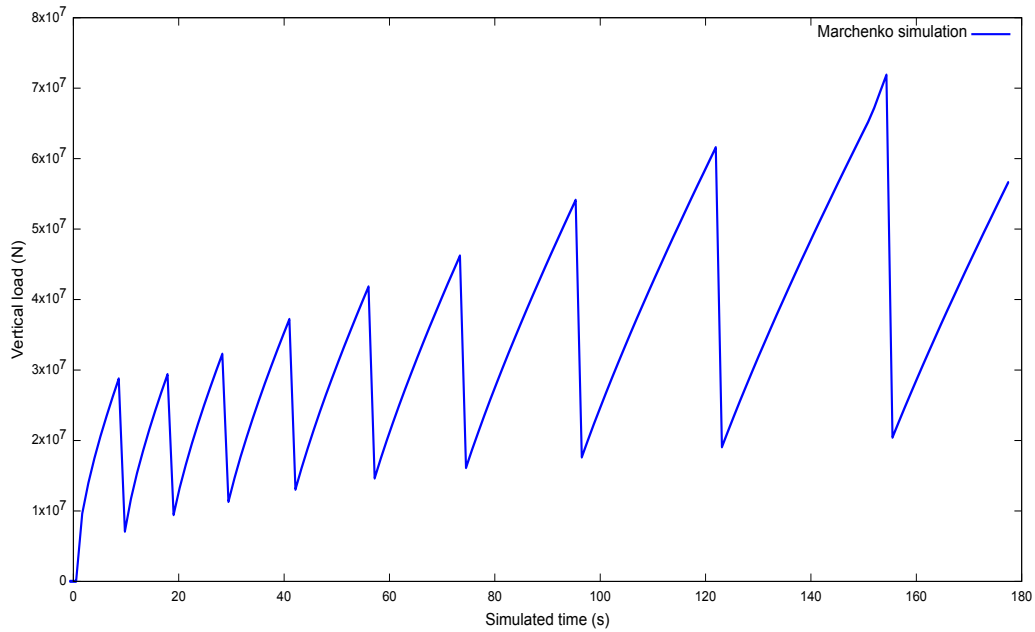


FIGURE 5.20 – **Simulation du modèle de Marchenko** – Nous observons que la charge globale sur la structure suit bien le même cycle de croissance/décroissance que la formation de la voile et de la quille. Lorsque la plaque cède sous le poids de la voile, la force chute jusqu'à ce que l'équilibre hydrostatique soit retrouvé.

les composantes du vecteur $\chi = \{\phi_s, \phi_k, \gamma_s, \gamma_k\}$. Du fait de leur nombre et de leur interdépendance, il paraît difficile de les calculer tous indépendamment les uns des autres. Aussi, nous préférons cette fois appliquer notre seconde stratégie pour déterminer implicitement un nouveau jeu de paramètres à chaque itération de co-simulation.

Notre seconde stratégie repose sur l'hypothèse que le modèle macroscopique et le modèle microscopique partagent une sortie commune. C'est le cas pour le modèle de Marchenko et ICE-MAS qui calculent tous les deux la charge globale sur la structure au cours du temps. Notre idée est donc de composer une simulation interactive (le *Displayer*) qui affichera une succession de simulations du modèle Marchenko dont les résultats seront validés par des instances de ICE-MAS.

Au début de la simulation, l'agent-*watchdog* du *Displayer* détecte qu'il n'y a pas de données à afficher à l'utilisateur. Il génère donc une requête qui contient l'état initial de simulation et une consigne de temps à simuler. Pour les itérations de co-simulation qui suivent, l'état inscrit dans la requête correspond à l'état final des données en cours de lecture.

Pour répondre à cette requête, il instancie ensuite plusieurs modèles de Marchenko avec des combinaisons variées des paramètres identifiés comme sensibles. Le premier jeu de paramètres référence est logiquement « nul » ($\chi = \{0^\circ, 0^\circ, 100.0\%, 100.0\%\}$), dans le sens où il n'y a pas encore de blocs, et nous tirons aléatoirement deux combinaisons autour de ces valeurs, avec un écart maximum de 10% pour chacun des paramètres. Ce choix est motivé par l'expertise métier qui nous indique qu'encore une fois, nous considérons un processus relativement lent pour lequel nous n'observons généralement pas de variations brusques pour

deux instants suffisamment proches. De plus, avec un écart plus grand, nous risquerions de faire diverger la simulation comme nous le montrerons plus loin. Dans un souci de clarté et de présentation des résultats, nous nous sommes contentés de trois instances simulées en parallèle mais nous aurions pu multiplier les tests de combinaisons pour obtenir une meilleure précision. Enfin, une simulation de ICE-MAS est instanciée dans un dernier *thread*.

Chacune des simulations doit ensuite traiter la requête et s'initialiser en conséquence.

5.4.3 Traitement de la requête

Chaque instance du modèle de Marchenko est initialisée avec un jeu de paramètres différent. Les trois simulations sont exécutées jusqu'à atteindre le temps fixé dans la requête. Ce temps détermine en quelque sorte la fréquence des itérations de co-simulation. Nous verrons par la suite que, tout comme pour la première stratégie de co-simulation, cette fréquence doit être choisie avec soin.

En ce qui concerne l'instance de ICE-MAS, nous devons implémenter un opérateur de *lifting* pour traduire l'état du modèle macroscopique en un état microscopique. Nous utilisons les données géométriques de la voile et de la quille pour construire une enveloppe 3D que l'on remplit avec des blocs de glace. Leurs formes et leurs tailles sont basées sur les données empiriques exprimées par [Nevel, 1972] et l'équation (5.6).

La porosité de chaque enveloppe peut être traitée de deux façons. La première consiste à générer les blocs un par un jusqu'à atteindre le taux de remplissage souhaité. C'est l'approche retenue par [Shafrova, 2007] pour reproduire des crêtes de glace. De notre côté, nous préférons calculer la répartition des blocs en 2D dans le triangle qui représente la voile ou la quille dans le modèle de Marchenko. Ces polygones sont ensuite extrudés sur une profondeur L et le motif est répété sur toute la largeur de la structure. Sur la figure 5.21, nous pouvons voir la construction de la voile et de la quille en 2D et leur passage en 3D. L'avantage de cette seconde méthode est double. D'une part, elle permet une génération beaucoup plus rapide des amas, car dans le cas de la première méthode, les blocs doivent être générés un par un et une simulation dédiée est utilisée pour déterminer leur placement. D'autre part, notre technique nous assure un temps de relaxation beaucoup plus court car les blocs sont déjà stabilisés.

Une fois les amas générés, nous démarrons la simulation ICE-MAS. Celle-ci est arrêtée par l'agent-*watchdog* une fois que les simulations du modèle macroscopique ont atteint leur objectif de temps simulé. La trajectoire de la charge exercée sur la structure est estimée par l'agent-*watchdog* qui sélectionne la simulation exploratoire qui en est la plus proche. Ce processus est illustré par la figure 5.22.

Les résultats de la simulation sélectionnée sont inscrits en tant que réponse de la requête qui est ensuite rejouée par le *Displayer*. Le jeu de paramètres qui a donné le résultat sélectionné est conservé et devient la nouvelle référence pour la prochaine itération de co-simulation.

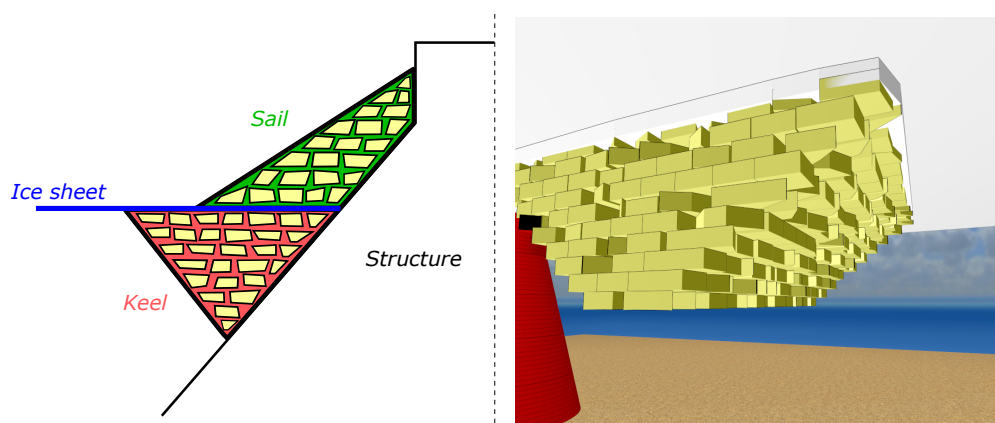


FIGURE 5.21 – **Traduction du modèle de Marchenko à Ice-MAS** – L’opérateur de *lifting* entre le modèle de Marchenko et ICE-MAS consiste à générer des amas de blocs de glace qui respectent les géométries de la voile et de la quille, de même que leurs porosités respectives. La génération des blocs est d’abord réalisée en 2D (à gauche) et s’appuie sur des connaissances d’experts qui nous indiquent des tailles et des formes de référence. Ces polygones sont ensuite extrudés pour obtenir des amas aussi larges que la structure (à droite).

5.4.4 Résultats et discussion

Pilotage du modèle de Marchenko

La figure 5.23 compare les charges globales calculées d’une part par ICE-MAS et par la co-simulation du modèle de Marchenko d’autre part. Nous remarquons que la co-simulation donne une bonne estimation de la charge au cours du temps. En particulier, elle reste pertinente dans la phase transitoire du processus d’accumulation des blocs. Nous observons de plus une convergence vers un jeu de paramètres unique en phase stationnaire.

Il est important de noter que cette comparaison est faite à temps simulé égal. Le temps de simulation de ICE-MAS est de l’ordre de la journée pour obtenir ces résultats alors que la co-simulation est réalisée en temps réel. Le gain de temps est donc très important et nous permet d’obtenir une simulation interactive d’une version auto-adaptative du modèle de Marchenko. Le temps associé aux étapes de la co-simulation (génération de requêtes, instanciation des simulations, reconstruction de l’état microscopique et sélection du résultat) est quant à lui négligeable.

Fréquence de co-simulation, systèmes oscillants et exploration des paramètres

Nous avons discuté dans le chapitre 4 de l’importance de la fréquence de co-simulation pour les systèmes oscillants. Pour ces cas particuliers, dont fait partie le modèle de Marchenko, la sélection de la meilleure simulation exploratoire sera d’autant plus juste que nous aurons adapté la fréquence de co-simulation par rapport à la fréquence d’oscillation du système.

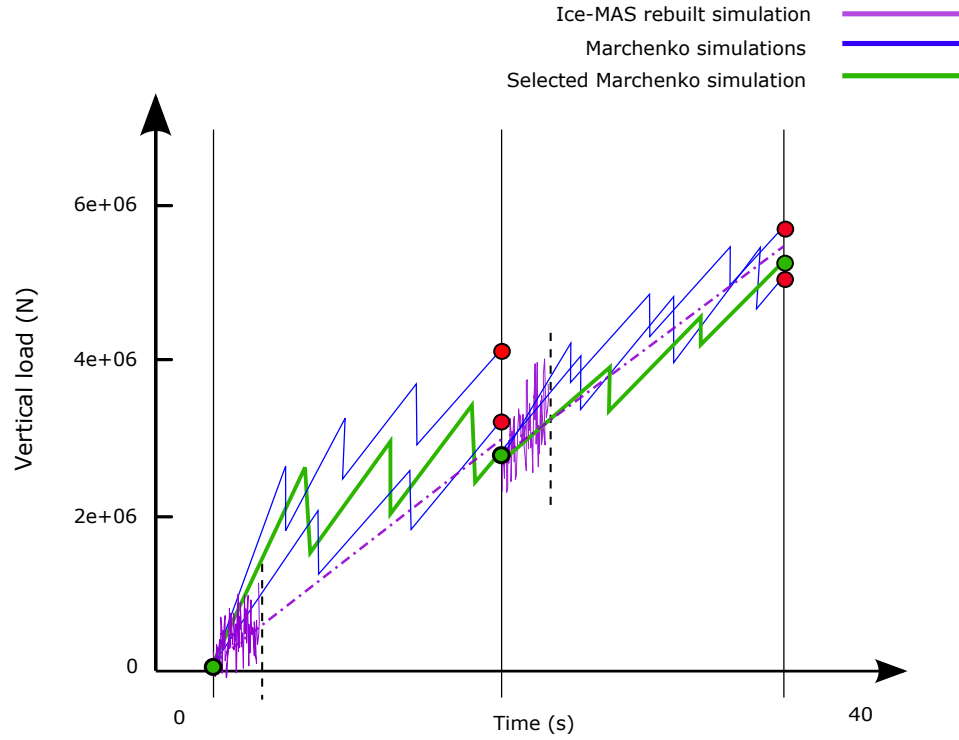


FIGURE 5.22 – **Itérations de co-simulation du modèle de Marchenko** – À chaque itération de co-simulation, trois instances du modèle de Marchenko sont simulées en parallèle avec différents jeux de paramètres qui conditionnent leurs résultats. Sur un quatrième cœur, nous initialisons le simulateur ICE-MAS pour qu'il corresponde à l'état final du dernier modèle de Marchenko sélectionné. Cette instance est simulée jusqu'à ce que les trois instances de Marchenko aient atteint la barrière de synchronisation. Son temps simulé est évidemment bien inférieur mais ses résultats nous permettent d'évaluer la trajectoire que ICE-MAS aurait suivi si nous l'avions simulé plus longtemps. Nous sélectionnons alors la simulation de Marchenko qui donne les meilleurs résultats au regard de la trajectoire ainsi estimée.

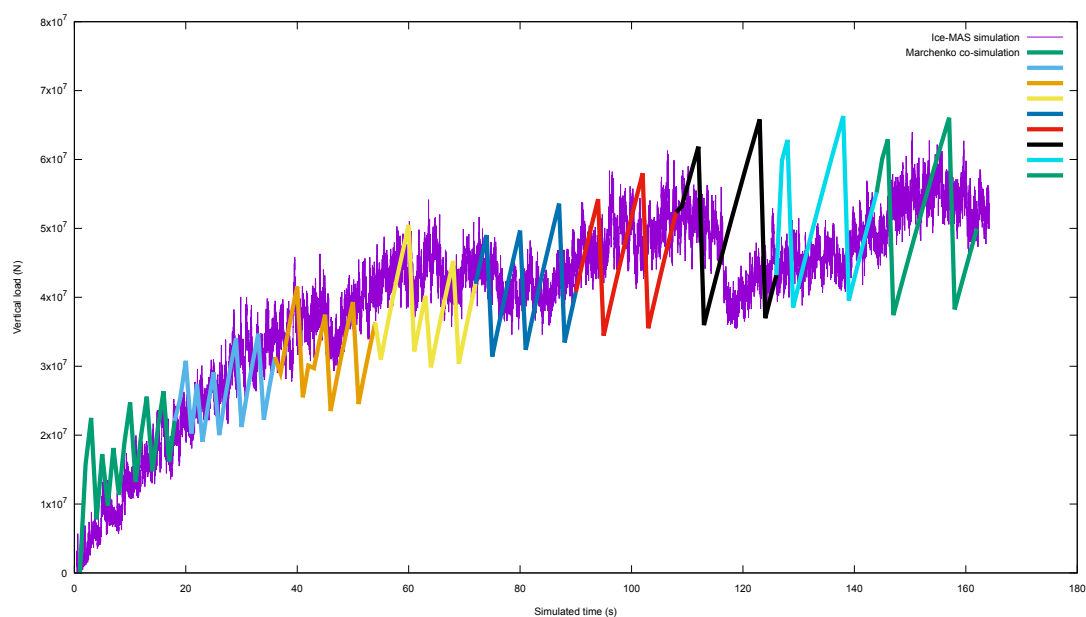


FIGURE 5.23 – **Comparaison de Ice-MAS et du modèle de Marchenko piloté** – La simulation auto-adaptative du modèle de Marchenko montre une bonne corrélation avec la simulation ICE-MAS originale, à temps simulé égal. Nous voyons que la détermination implicite des paramètres nous permet de reproduire la phase transitoire de l’accumulation des blocs, ce que ne pouvait proposer un jeu de paramètres constant. Puisque ICE-MAS n’est plus simulé à plein temps mais comme outil de validation ponctuel, nous gagnons énormément en terme de temps de simulation.

Cependant, dans le cas présent, il est difficile de choisir une fréquence de co-simulation qui corresponde à toutes les simulations exploratoires car la fréquence d'apparition des ruptures de plaque dépend fortement du jeu de paramètres testé.

Par ailleurs, nous observons sur la figure 5.24 que la sélection d'un jeu de paramètres différent de celui de l'itération de co-simulation précédente entraîne des perturbations absentes du modèle de base. Du fait de sa nature oscillante et suivant le temps de simulation accordé aux simulations exploratoires, une simulation de Marchenko peut être interrompue en phase de croissance de la voile ou juste après la rupture de la plaque. Si elle est sélectionnée, son état est conservé pour l'itération suivante et de nouvelles simulations sont démarrées, dont deux avec des jeux de paramètres légèrement altérés. Le fait est que si ces variations sont trop grandes, elles peuvent causer des instabilités et provoquer une nouvelle rupture de la plaque. Ceci correspond en quelque sorte à un temps de relaxation pour le modèle macroscopique. Ces perturbations sont acceptables dans la mesure où elles ne perturbent pas le processus de sélection, auquel cas une simulation exploratoire peut être invalidée et son jeu de paramètres éliminé d'office des candidats.

D'autres solutions envisageables pour ajuster la fréquence de co-simulation seront proposées en guise de perspectives en fin de mémoire.

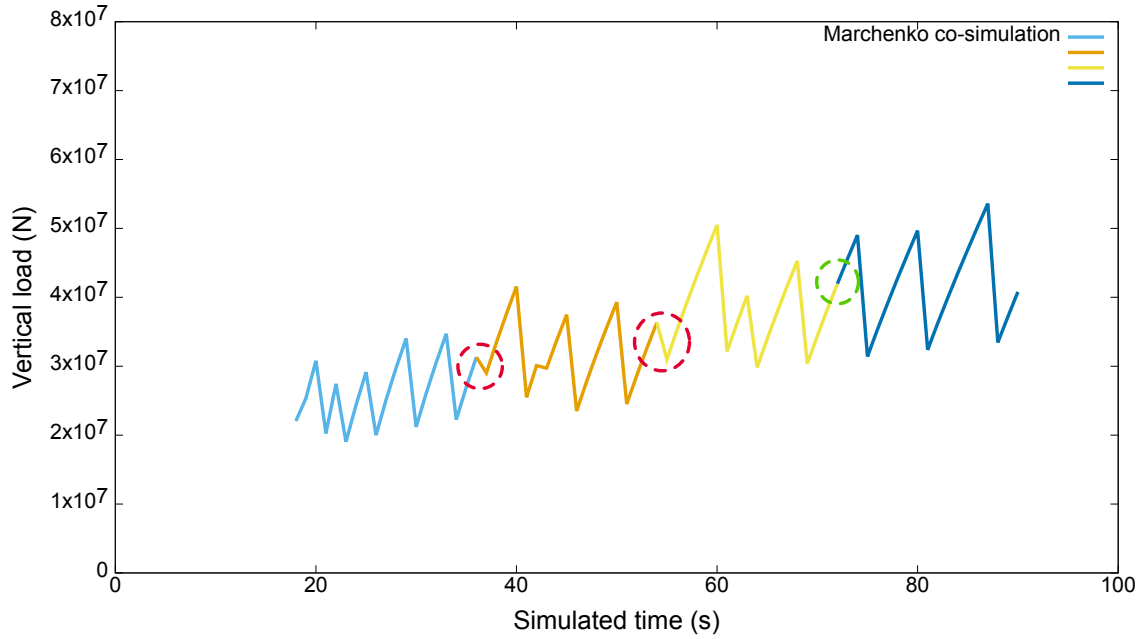


FIGURE 5.24 – **Paramétrage du processus de co-simulation** – Les cercles rouges soulignent des perturbations dans l'évolution de la charge qui surviennent après un changement de jeu de paramètres. Elles trouvent leur origine dans une combinaison de plusieurs facteurs. En premier lieu, la nature oscillante du modèle de Marchenko fait qu'avec une fréquence de co-simulation trop faible, une simulation exploratoire peut s'arrêter pendant une phase de croissance de la voile. L'état du modèle à cet instant peut alors être sélectionné comme étant la nouvelle référence. Un jeu de paramètres trop éloigné du précédent peut causer des instabilités au début de la nouvelle itération. Le cercle vert nous montre que le problème n'apparaît pas quand le jeu de paramètres reste inchangé entre deux itérations.

5.4.5 Bilan

Dans cette section, nous avons apporté une solution au problème du temps de calcul de ICE-MAS. Notre travail a consisté à développer un outil de prototypage rapide qui peut être utilisé comme une première passe pour le design d’une structure offshore. Pour ce faire, nous nous sommes appuyés sur un modèle de plus haut niveau qui se focalise sur le cycle de formation des amas de glace, la voile et la quille, sur le front de la structure.

Notre seconde stratégie de co-simulation nous a permis de construire un simulateur interactif, dans le sens où son temps de calcul nous permet d’obtenir des résultats dans un temps bien plus intéressant. Ce simulateur joue une succession de périodes de simulations macroscopiques sélectionnées au regard de leur justesse par rapport à ICE-MAS. En agissant de la sorte, nous répondons à la contrainte du paramétrage du modèle de Marchenko. Les données affichées sont en décalage temporel permanent par rapport à celles qui sont en train d’être simulées, mais le processus est totalement transparent pour l’utilisateur.

Nous notons que cette seconde stratégie est généralement plus simple à mettre en œuvre que la détermination explicite de paramètres. Nous nous apercevons en effet que l’estimation d’un paramètre est souvent réalisée grâce à un modèle microscopique qui représente le système en entier avec une granularité plus fine. Ce modèle pourrait donc être utilisé dans le cadre de la détermination implicite de paramètres.

Par ailleurs, la détermination implicite des paramètres telle qu’elle est réalisée par notre stratégie de co-simulation nous dispense de définir un opérateur de *restriction* pour traduire les données microscopiques en données macroscopiques.

5.5 Synthèse

Ce chapitre nous a permis de mettre en application nos différentes propositions.

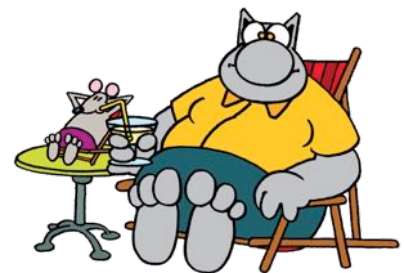
Tout d’abord, nous avons construit un simulateur d’interactions glace-structure, ICE-MAS, grâce à une approche multi-agents phénoménologique. Nous avons décrit chacun des phénomènes un à un en intégrant des informations d’origines diverses telles que des lois analytiques, des lois empiriques et des données expérimentales. Cette première implémentation a montré son efficacité pour estimer la charge globale exercée sur la structure. Cependant, nous avons identifié deux pistes d’optimisation pour améliorer d’une part sa précision et d’autre part son interactivité.

La première étape a consisté à mieux représenter les phénomènes hydrodynamiques. Les calculs de dynamique du courant sont trop coûteux pour être réalisés à chaque pas de temps ou même de manière synchrone. De ce fait, pour prendre en compte l’impact des blocs de glace qui occultent la structure au cours de la simulation, nous utilisons notre architecture de co-simulation. Les données 3D sont extraites régulièrement de ICE-MAS pour calculer une nouvelle grille de vecteurs vitesses à l’aide d’une base de cas ou du logiciel OPENFOAM. Cette stratégie nous permet d’obtenir des résultats plus réalistes pour un coût de calcul supplémentaire insignifiant. En effet, ce calcul explicite est réalisé en parallèle et son résultat

est appliqué dès que possible de manière asynchrone, ce qui est rendu possible à la fois par la lenteur du processus d'accumulation des blocs et par la lenteur du simulateur lui-même.

Le temps de simulation de ICE-MAS dépend fortement du nombre de blocs simulés individuellement. Ainsi, la seconde étape d'optimisation nous a conduit à changer de niveau de description pour le phénomène d'empilement des blocs sur le front de la structure. Cependant, comme nous avons pu le voir tout au long de ce mémoire, l'abstraction des phénomènes conduit généralement à des difficultés de paramétrage. Notre seconde stratégie de co-simulation nous a permis de répondre à ce problème en déterminant implicitement les paramètres sensibles du modèle de Marchenko. Nous avons construit un simulateur dont la tâche est de jouer des données de simulations du modèle macroscopique sélectionnées par ICE-MAS. Bien entendu, cet outil est moins expressif que le simulateur original mais il peut être utile pour effectuer un premier prototypage rapide des structures offshore pour lesquelles une étude plus poussée pourra être réalisée.

Conclusion



Philippe Geluck - LE TOUR DU CHAT EN 365
JOURS

Le travail de recherche exposé dans ce mémoire a consisté à développer des outils pour la modélisation et la simulation de systèmes complexes. Dans le contexte industriel dans lequel nous nous plaçons, nous mesurons la qualité d'un modèle de système complexe selon deux critères. En premier lieu, il doit être en mesure de donner une vue du fonctionnement global du système et pas seulement d'échantillons peu représentatifs. En second lieu, la simulation de ce modèle doit être interactive, dans le sens où son temps de calcul ne doit pas gêner la boucle d'interaction entre le modèle et son expérimentateur.

Pour développer de tels outils, nous devons donc élever notre point de vue sur le système et nous focaliser sur les interactions qui forment sa dynamique. Cette approche est qualifiée de phénoménologique et consiste à décrire, séparément et un à un, les différents phénomènes qui agissent dans ou sur le système. Nous avons vu que les systèmes multi-agents se prêtent plutôt bien à cet exercice de décomposition. Plus encore, une approche multi-agents mixte autorise la superposition de tous types de modèles, qu'ils soient analytiques, numériques ou empiriques. Nous avons démontré toute la pertinence de l'utilisation conjointe de ces deux approches à travers le développement du modèle MACMA pour l'étude des mécanismes de refroidissement du manteau terrestre.

Toutefois, nous nous sommes aperçus que l'implémentation des modèles de haut niveau, le plus souvent par des équations différentielles, devient rapidement dépendante d'une multitude de paramètres qu'il faut alors ajuster. Cette calibration doit même généralement être dynamique pour s'adapter aux changements de conditions aux limites par exemple. Ce point constitue un véritable obstacle au moment de traiter des modèles abstraits dont les paramètres peuvent avoir perdu de leur sens.

Pour traiter ce problème, nous avons en quelque sorte rebroussé chemin. Notre hypothèse de travail est la suivante : nous disposons souvent de plusieurs modèles pour un même système ou un même phénomène. Ces modèles sont organisés en termes de niveau de description, et les paramètres effectifs des modèles macroscopiques résument les phénomènes sous-jacents des modèles microscopiques. Notre revue des méthodes multi-échelles nous a permis de découvrir toute une catégorie de méthodes qui cherche à exploiter cette redondance des modèles. Particulièrement, l'objectif est d'utiliser les modèles microscopiques pour nourrir les modèles macroscopiques. Si ce concept est séduisant dans le cadre de notre problématique, nous avons pu nous rendre compte du manque de généralité de ces méthodes. En effet, elles nécessitent de définir des opérateurs de traduction inter-échelles fortement dépendants du domaine étudié.

À défaut de pouvoir exprimer une méthode générique de simulations redondantes, nous nous sommes fixés comme objectif d'automatiser au maximum l'implémentation de simulations qui suivent cette démarche, qui elle reste identique quel que soit le problème. Pour ce faire, nous avons développé une architecture logicielle qui est capable de mettre en relation différents modèles d'échelles hétérogènes. La diversité des formalismes des modèles que nous devons coupler nous a fait nous tourner vers la technique de la co-simulation.

Pour mener à bien le pilotage automatique d'une simulation macroscopique, nous avons conçu deux stratégies de co-simulation. La première est une implémentation directe des méthodes multi-échelles que nous avons présentées. Elle consiste à simuler un modèle microscopique sur une courte durée pour extrapoler les paramètres d'un modèle macroscopique. La seconde stratégie est plus subtile et fait intervenir des notions plus éloignées, comme les méthodes de tirs ou les méthodes d'optimisation. L'idée est cette fois d'explorer des jeux de paramètres et d'utiliser le modèle microscopique comme un outil de validation et de sélection du meilleur d'entre eux.

L'application de ces propositions au problème concret du design de structures offshore nous a permis de mettre en évidence les particularités de chacune de ces stratégies. Tout d'abord, d'une manière générale, nous avons pu remarquer que notre architecture de co-simulation s'accommode bien des contraintes industrielles auxquelles nous sommes soumis. En particulier, le couplage des simulateurs existants et de formalismes différents que sont ICE-MAS et OPENFOAM s'est fait sans heurt et a nécessité assez peu de développement. De plus, la démarche incrémentale offerte par la co-simulation s'inscrit très bien dans un processus de développement efficace.

Ensuite, nous distinguons deux usages cibles pour nos stratégies de co-simulation. La détermination explicite d'un paramètre intervient généralement dans un souci de précision de la simulation macroscopique. Quant à la détermination implicite de jeux de paramètres, nous pouvons dire qu'elle est une généralisation de la première stratégie à plusieurs dimensions. Cet aspect lui permet de sélectionner des tronçons valides de simulation qui deviennent la

nouvelle discrétisation temporelle de la simulation et assurent une certaine continuité des résultats.

Nos deux stratégies de co-simulation se recouvrent sur un point : elles n'ont pas de conséquence sur le temps de calcul du modèle macroscopique. Hormis l'exécution de l'agent-*watchdog*, l'ensemble des opérations de co-simulation est supportée par d'autres cœurs de processeur.

Enfin, nous avons contribué au développement et à l'amélioration du simulateur ICE-MAS qui est aujourd'hui utilisé par des experts du domaine de l'offshore pour concevoir et dimensionner de futures structures. Le site internet <http://www.ice-mas.com> donne plusieurs exemples de simulations qui illustrent les différents phénomènes implémentés. De nouveaux travaux ont débuté pour prendre en compte les interactions entre des icebergs et des structures.

Perspectives

Tout au long de ce mémoire, nous nous sommes attachés à relever les défauts de nos propositions. Nous dressons ici trois pistes de travaux complémentaires qui pourraient les corriger.

Automatisation du contrôle

L'activité de contrôle de l'agent-*watchdog* est régie par des points de contrôle. Nous sommes partis du principe que nous connaissions à l'avance les faiblesses de notre modèle et de ce fait nous avons défini à chaque fois des règles *ad hoc* comme des seuils sur le volume de glace accumulé ou le dépassement d'un délai préfixé. Nous pensons toutefois que la détection de défauts dans le modèle hôte de l'agent-*watchdog* pourrait être partiellement automatisée. Plusieurs travaux [Klein, 2009] [Boes, 2014] ont abouti à des méthodes de contrôle des systèmes multi-agents. La décomposition du système en agents autonomes peut permettre d'isoler les variations numériques ou les changements structurels éventuels. Il faudra cependant garder à l'esprit que, dans le cadre de systèmes complexes, le défaut d'un agent peut ne pas être de son fait mais provenir d'un autre agent duquel il subit une influence.

Estimation de trajectoires complexes

Nous avons discuté plusieurs fois dans ce mémoire de l'estimation de trajectoires d'une simulation microscopique. Celle-ci est à la base du processus de sélection de notre seconde stratégie de co-simulation. Pour bon nombre de problèmes, une trajectoire linéaire extrapolée des résultats du modèle microscopique sera suffisante pour effectuer le bon choix, surtout si la fréquence de co-simulation est élevée. Dans le cas de systèmes oscillants ou d'une fréquence de co-simulation trop faible, il est nécessaire d'estimer plus finement la trajectoire. Comme nous

l'avons vu, cette estimation peut être faite en prenant en compte des connaissances métier. Mais plus simplement, puisqu'au moment de la sélection, les simulations exploratoires sont déjà simulées, nous pouvons peut-être extraire un *pattern* commun susceptible de guider notre choix.

Fréquence de co-simulation adaptative

Cette dernière perspective est directement liée à la précédente. Nous avons vu que la fréquence de co-simulation était d'une grande importance pour les systèmes oscillants. L'exemple du modèle de Marchenko ne fait peut-être pas école, mais nous avons le sentiment qu'une fréquence de co-simulation calquée sur la fréquence des événements majeurs du modèle macroscopique (ex : rupture de la plaque dans le modèle de Marchenko) améliorerait la justesse des sélections. Nous pourrions encore plus souvent nous contenter d'une estimation de trajectoire linéaire. L'implémentation de cette approche passe soit par l'incorporation de connaissances métier, soit par une analyse fréquentielle d'une simulation antérieure. Ainsi, la fréquence de co-simulation pourra même s'adapter à des oscillations non régulières.

Co-simulation à N niveaux

Plus qu'une correction de notre technique de co-simulation, cette dernière perspective vise à généraliser l'approche multi-niveaux en prenant appui sur la récursivité de notre architecture logicielle. En effet, nous avons vu que le simulateur Ice-MAS est tour à tour le modèle macroscopique paramétré par OpenFOAM, et l'outil de validation microscopique du paramétrage du modèle de Marchenko. Nous pouvons donc facilement imaginer une co-simulation à trois niveaux qui exploiterait un simulateur Ice-MAS optimisé. Cependant, ce cas précis ne fonctionnerait pas en raison du temps de calcul du simulateur OpenFOAM qui est bien supérieur à la période de co-simulation du modèle de Marchenko. Nous pouvons donc voir ici que la durée de vie de la requête de paramétrage d'un niveau de co-simulation est contrainte par celle de la requête provenant du niveau supérieur. Il conviendra alors de soit choisir des modèles microscopiques capables de satisfaire cette contrainte, soit d'optimiser la taille du domaine microscopique simulé et ainsi trouver le meilleur compromis entre la vitesse d'exécution et la qualité du résultat estimé pour reparamétriser le modèle macroscopique.

Pour finir, nous souhaitons souligner que nos travaux rejoignent d'une certaine manière ceux réalisés au CERV sur la notion de « la simulation dans la simulation » appliquée à la modélisation de comportements humains [Buche, 2012] [Polceanu, 2015]. La validité de modèles comportementaux est souvent difficile à démontrer. On considère dans ces cas que la crédibilité de la simulation fait office de preuve. Au contraire, nos méthodes de co-simulation explicite et implicite laissent entrevoir la possibilité d'une validation mathématique de ce concept dans le contexte de la simulation de phénomènes physiques. Ces perspectives théoriques devraient faire l'objet d'une future thèse.

Publications scientifiques

- Béal, P.-A., Septseault, C., Le Yaouanq, S. L., Dudal, A., Cahay, M., Roberts, B. (2016). Simulation of Iceberg Impacts using ICE-MAS : A Design Tool for Simulating Ice-structure Interactions. *Artic Technology Conference 2016*. Publication acceptée, présentation à venir.
Cité page 3
- Combes, M., Grigné, C., Husson, L., Conrad, C. P., Le Yaouanq, S., Parenthoën, M., Tisseau, C., et Tisseau, J. (2012). Multiagent simulation of evolutive plate tectonics applied to the thermal evolution of the Earth. *Geochemistry, Geophysics, Geosystems*. vol. 13, no 5.
Cité pages 2 et 51
- Combes, M., Grigné, C., Husson, L., Le Yaouanq, S., Parenthoën, M., Tisseau, C., et Tisseau, J. (2011). MACMA : Mantle cooling mechanisms simulated by agents. *Geophysical Research Abstracts* vol. 13
Cité pages 2, 32, 38 et 51
- Grigné, Combes, M., C., Tisseau, C., Le Yaouanq, S., Parenthoën, M., et Tisseau, J. (2012). Multi-agent modelling of earth's dynamics : Towards a virtual laboratory of plate tectonics. *AGU Fall Meeting Abstracts*.
Cité page 2
- Dudal, A., Sepseault, C., Béal, P.-A., Le Yaouanq, S., et Roberts, B. (2015). A new arctic platform design tool for simulating ice-structure interaction. *Proceedings of the 23rd International Conference on Port and Ocean Engineering under Arctic Conditions*.
Cité pages 3, 124 et 134
- Grigne, C., Combes, M., Le Yaouanq, S., Husson, L., Conrad, C. P., et Tisseau, C. (2011). Thermal Evolution of the Earth from a Plate Tectonics Point of View. *AGU Fall Meeting Abstracts*
Cité page 2
- Le Yaouanq, S., Le Gal, C., Redou, P., et Tisseau, J. (2015). Co-simulation of redundant and heterogeneous modelling scales for a phenomenological approach. *APSAC, Applied Physics, Simulation and Computers, Vienna*. pages 109-115.
Cité page 3

Le Yaouanq, S., Redou, P., Le Gal, C., Abgrall, J. F., et Tisseau, J. (2011a). Multi-agent systems and heterogeneous scales interactions. application to pharmacokinetics of vitamin k antagonists. *Advances in Artificial Life, ECAL 2011*. pages 447-454.

Cité page [2](#)

Le Yaouanq, S., Redou, P., Le Gal, C., Abgrall, J. F., et Tisseau, J. (2011b). Redundant modelling of multiscale interactions with multi-agent systems. application to pharmacokinetic/pharmacodynamic of vitamin k antagonists. *European Conference of Complex Systems, ECCS 2011*. page 137.

Cité page [2](#)

Septseault, C., Béal, P.-A., Le Yaouanq, S., Dudal, A., et Roberts, B. (2014). A New Ice SimulationTool Using a Multi-Model Program. *Artic Technology Conference 2014*.

Cité pages [3](#), [124](#) et [134](#)

Septseault, C., Béal, P.-A., Le Yaouanq, S. L., Dudal, A., Roberts, B., et al. (2015). Update on a new ice simulation tool using a multi-model program. *OTC Arctic Technology Conference, 23-25 March, Copenhagen, Denmark*.

Cité pages [3](#) et [124](#)

Références bibliographiques

- Abdulle, A., Weinan, E., Engquist, B. A., et Eijnden, E. V. (2012). The heterogeneous multiscale method. *Acta Numerica*, 21 :pages 1–87.
Cité page 73
- Ascher, U. M. et Petzold, L. R. (1998). *Computer methods for ordinary differential equations and differential-algebraic equations*. Society for Industrial Mathematics.
Cité pages 26 et 172
- Asselin, G. (2013). *Une approche multi-agents pour le développement d'un jeu vidéo*. Thèse de doctorat, Université de Toulouse, Université Toulouse III-Paul Sabatier.
Cité page 15
- Barkov, A. (2011). Research and design solutions for development of oil fields in the northern caspian sea. Dans *Proceedings of RAO conference*.
Cité page 118
- Bates, J., Loyall, A. B., et Reilly, W. S. (1992). Integrating reactivity, goals, and emotion in a broad agent. Dans *Proceedings of the Fourteenth Annual Conference of the Cognitive Science Society*.
Cité page 16
- Béal, P.-A., Le Gal, C., et Tisseau, J. (2008). Approche multi-phénomènes appliquée à la simulation des systèmes mécaniques. Dans *Actes du séminaire du LISyC*, pages 25–36.
Cité pages 27 et 171
- Beauchesne, C., Parenteau, T., Septseault, C., et Béal, P.-A. (2015). Development & large scale validation of a transient flow assurance model for the design & monitoring of large particles transportation in two phase (liquid-solid) riser systems. Dans *Offshore Technology Conference*.
Cité page 137
- Becker, T. W. et Faccenna, C. (2009). A review of the role of subduction dynamics for regional and global plate motions. Dans *Subduction Zone Geodynamics*, pages 3–34. Springer.
Cité page 48

- Beni, G. (2005). From swarm intelligence to swarm robotics. Dans *Proceedings of the 2004 International Conference on Swarm Robotics*, SAB'04, pages 1–9, Berlin, Heidelberg. Springer-Verlag.
Cité page 15
- Bernard, C. (1865). *Introduction à l'étude de la médecine expérimentale*. Flammarion.
Cité page 11
- Bézivin, J., Jouault, F., et Valduriez, P. (2004). On the need for megamodels. Dans *Proceedings of the OOPSLA/GPCE : Best Practices for Model-Driven Software Development workshop, 19th Annual ACM Conference on Object-Oriented Programming, Systems, Languages, and Applications*.
Cité page 75
- Bjerka, M., Albrektsen, A., et Gurtner, A. (2010). Static and dynamic ice actions in the light of new design codes. Dans *ASME 2010 29th International Conference on Ocean, Offshore and Arctic Engineering*, pages 733–739. American Society of Mechanical Engineers.
Cité page 119
- Boes, J. (2014). *Apprentissage du contrôle de systèmes complexes par l'auto-organisation coopérative d'un système multi-agent : application à la calibration de moteurs à combustion*. Thèse de doctorat, Université de Toulouse, Université Toulouse III-Paul Sabatier.
Cité page 155
- Bommel, P. (2009). *Définition d'un cadre méthodologique pour la conception de modèles multi-agents adaptée à la gestion des ressources renouvelables*. Thèse de doctorat, Université Montpellier II-Sciences et Techniques du Languedoc.
Cité page 11
- Bonabeau, E., Dorigo, M., et Theraulaz, G. (1999). *Swarm intelligence : from natural to artificial systems*. Number 1. Oxford university press.
Cité page 15
- Brown, A. J. (1902). Enzyme action. *Journal of the chemical society, transactions*, 81 :373–388.
Cité page 88
- Buche, C. (2012). *Adaptive behaviors for virtual entities in participatory virtual environments*. Habilitation à diriger des recherches, Université de Bretagne occidentale - Brest.
Cité page 156
- Campin, J.-M., Hill, C., Jones, H., et Marshall, J. (2011). Super-parameterization in ocean modeling : Application to deep convection. *Ocean Modelling*, 36(1) :90–101.
Cité page 65
- Camus, B., Siebert, J., Bourjot, C., et Chevrier, V. (2012). Modélisation multi-niveaux dans AA4MM. *arXiv preprint arXiv :1210.5936*.
Cité page 78
- Chevrier, V. et St Dizier, A. (2005). L'intelligence en essaim ou comment faire complexe avec du simple ? *Interstices*.
Cité page 15

- Combes, M. (2011). *Earth's mantle cooling mechanisms studied by multiagent systems*. These de doctorat, Université de Bretagne occidentale - Brest.
Cité pages 2, 32, 38 et 51
- Combes, M., Buin, B., Parenthoën, M., et Tisseau, J. (2010). Multiscale multiagent architecture validation by virtual instruments in molecular dynamics experiments. *Procedia Computer Science*, 1(1) :761–770.
Cité page 31
- Comete, C. et Abib, M. (1998). *Codesign : conception conjointe logiciel-matériel*. Eyrolles.
Cité page 2
- Conrad, C. P. et Lithgow-Bertelloni, C. (2002). How mantle slabs drive plate tectonics. *Science*, 298(5591) :207–209.
Cité page 35
- Coumans, E. (2012). *Bullet 2.80 Physics SDK Manual*. <http://www.bulletphysics.org>.
Cité page 125
- Crépin, L. (2013). *Couplage de modèles population et individu-centrés pour la simulation parallélisée des systèmes biologiques. Application à la coagulation du sang*. Thèse de doctorat, Université de Bretagne occidentale-Brest.
Cité pages 12 et 81
- Daley, C. (1991). *Ice edge contact : a brittle failure process model*. Acta polytechnica Scandinavica : Mechanical engineering series. Finnish Academy of Technical Sciences.
Cité pages 132 et 133
- de Lara, J. et Vangheluwe, H. (2002). Atom³ : A tool for multi-formalism modelling and meta-modelling. Dans *IN PROC. FASE 02, SPRINGER LNCS 2306*. Citeseer.
Cité page 75
- Delaux, S., Stevens, C. L., et Popinet, S. (2010). High-resolution computational fluid dynamics modelling of suspended shellfish structures. *Environmental Fluid Mechanics*, 11(4) :405–425.
Cité page 63
- Demazeau, Y. (1995). From interactions to collective behaviour in agent-based systems. Dans *In : Proceedings of the 1st. European Conference on Cognitive Science. Saint-Malo*.
Cité page 15
- Desmeulles, G. (2006). *Réification des interactions pour l'expérience in virtuo de systèmes biologiques multi-modèles*. Thèse de doctorat, Université de Bretagne occidentale-Brest.
Cité pages 20 et 21
- Duboz, R. (2004). *Intégration de modeles hétérogenes pour la modélisation et la simulation de systemes complexes*. Thèse de doctorat, Université du Littoral - Côte d'Opale.
Cité pages 76 et 86
- Dudal, A., Sepseault, C., Béal, P.-A., Le Yaouanq, S., et Roberts, B. (2015). A new arctic platform design tool for simulating ice-structure interaction. *Proceedings of the 23rd International Conference on Port and Ocean Engineering under Arctic Conditions*.
Cité pages 3, 124 et 134

- Eddy, F. et Gooi, H. (2011). Multi-agent system for optimization of microgrids. Dans *Power Electronics and ECCE Asia (ICPE ECCE), 2011 IEEE 8th International Conference on*, pages 2374–2381.
Cité page 15
- El Hmam, M. S., Jolly, D., Abouaissa, H., et Benasser, A. (2008). Modelisation hybride du flux de trafic. *Méthodologies et Heuristiques pour l'Optimisation des Systèmes Industriels (MHOSI'08)*, pages 193–198.
Cité page 82
- Erickson, A., Von Herzen, R., Sclater, J., Girdler, R., Marshall, B., et Hyndman, R. (1975). Geothermal measurements in deep-sea drill holes. *Journal of Geophysical Research*, 80(17) :2515–2528.
Cité page 36
- Etzioni, O., Levy, H. M., Segal, R. B., et Thekkath, C. A. (1993). OS agents : Using AI techniques in the operating system environment. Rapport technique.
Cité page 15
- Favre, J.-M. (2006). Megamodelling and etymology. Dans *Dagstuhl Seminar Proceedings*. Schloss Dagstuhl-Leibniz-Zentrum für Informatik.
Cité page 75
- Fick, A. (1855). Ueber diffusion. *Annalen der Physik*, 170(1) :59–86.
Cité page 86
- Fishwick, P. A. et Zeigler, B. P. (1992). A multimodel methodology for qualitative model engineering. *ACM Transactions on Modeling and Computer Simulation (TOMACS)*, 2(1) :52–81.
Cité page 76
- Fransson, L. et Bergdahl, L. (2009). Recommendations for design of offshore foundations exposed to ice loads. Rapport technique, Elforsk.
Cité pages 1 et 119
- Fransson, L., PAtil, A., et Andren, H. (2011). Experimental investigation of friction coefficient of laboratory ice. Dans *Proc. 21th International Conference on Port and Ocean Engineering under Arctic Conditions*.
Cité page 126
- Frederking, R. et Barker, A. (2002). Friction of sea ice on steel for condition of varying speeds. Dans *Proc. 12th International Offshore and Polar Engineering Conference*, pages 26–31.
Cité page 126
- Fuchs, P., Arnaldi, B., et Tisseau, J. (2003). *La réalité virtuelle et ses applications*, volume 1 of *Le Traité de la Réalité Virtuelle*, chapitre 1. Presses de l'Ecole des Mines de Paris, 2^e édition.
Cité page 13
- Gear, C. W., Li, J., et Kevrekidis, I. G. (2003). The gap-tooth method in particle simulations. *Physics Letters A*, 316(3) :190–195.
Cité page 71

- Grabowski, W. W. et Smolarkiewicz, P. K. (1999). Crcp : A cloud resolving convection parameterization for modeling the tropical convecting atmosphere. *Physica D : Nonlinear Phenomena*, 133(1) :171–178.
Cité page [64](#)
- Gransart, C. et Geib, J.-M. (1999). *CORBA : Des concepts à la pratique*. Dunod.
Cité page [81](#)
- Grigné, C. (2003). *Etude de l'effet des continents sur la convection du manteau*. Thèse de doctorat, Paris, Institut de physique du globe.
Cité page [33](#)
- Grigné, C., Combes, M., Le Yaouanq, S., Husson, L., Conrad, C. P., et Tisseau, C. (2011). Thermal Evolution of the Earth from a Plate Tectonics Point of View. *AGU Fall Meeting Abstracts*.
Cité page [2](#)
- Grigné, C., Combes, M., Tisseau, C., Le Yaouanq, S., Parenthoën, M., et Tisseau, J. (2012). Multi-agent simulations of earth's dynamics : Towards a virtual laboratory for plate tectonics. Poster.
Cité page [35](#)
- Grigné, C., Labrosse, S., et Tackley, P. J. (2007). Convection under a lid of finite conductivity : Heat flux scaling and application to continents. *Journal of Geophysical Research : Solid Earth (1978–2012)*, 112(B8).
Cité page [36](#)
- Grimm, V. (1999). Ten years of individual-based modelling in ecology : what have we learned and what could we learn in the future ? *Ecological modelling*, 115(2) :129–148.
Cité page [11](#)
- Hairer, E., Norsett, S., et Wanner, G. (1993). *Solving Ordinary Differential Equations I, Nonstiff Problems*. Springer.
Cité pages [25](#) et [171](#)
- Hardebolle, C. (2008). *Composition de modèles pour la modélisation multi-paradigme du comportement des systèmes*. Thèse de doctorat, Université Paris Sud-Paris XI.
Cité page [75](#)
- Hardebolle, C. et Boulanger, F. (2007). Modhel'x : A component-oriented approach to multi-formalism modeling. Dans *Models in Software Engineering*, pages 247–258. Springer.
Cité pages [75](#) et [76](#)
- Harrouet, F. (2000). *oRis : s'immerger par le langage pour le prototypage d'univers virtuels à base d'entités autonomes*. Thèse de doctorat, ENIB.
Cité page [23](#)
- Herviou, D. (2006). *La perception visuelle des entités autonomes en réalité virtuelle : application à la simulation de trafic routier*. Thèse de doctorat, Université de Bretagne Occidentale.
Cité page [19](#)

- Heuret, A. et Lallemand, S. (2005). Plate motions, slab dynamics and back-arc deformation. *Physics of the Earth and Planetary Interiors*, 149(1) :31–51.
Cité page 42
- Jaupart, C., Labrosse, S., et Mareschal, J. (2007). Temperatures, heat and energy in the mantle of the earth. *Treatise on geophysics*, 7 :253–303.
Cité pages 36 et 49
- Jaupart, C. et Mareschal, J.-C. (2011). *Heat generation and transport in the Earth*. Cambridge university press Cambridge, UK.
Cité page 48
- Kärnä, T. et Jochmann, P. (2003). Field observations on ice failure modes. Dans *Proceedings of the 17th International Conference on Port and Ocean Engineering under Arctic Conditions, Trondheim, Norway, June*, pages 16–19.
Cité pages 121 et 123
- Kavenoky, A. (1978). The SPH homogeneization method. Rapport technique, CEA Centre d'Etudes Nucleaires de Cadarache, 13-Saint-Paul-les-Durance (France). Dept. des Reacteurs a Eau.
Cité page 69
- Kerdélo, S. (2006). *Méthodes informatiques pour l'expérimentation in virtuo de la cinétique biochimique. Application à la coagulation du sang*. Thèse de doctorat, Université de Bretagne Occidentale.
Cité page 20
- Kevrekidis, I. G., Gear, C. W., Hyman, J. M., Kevrekidid, P. G., Runborg, O., Theodoropoulos, C., et al. (2003). Equation-free, coarse-grained multiscale computation : Enabling mocosmic simulators to perform system-level analysis. *Communications in Mathematical Sciences*, 1(4) :715–762.
Cité pages 68 et 71
- Klein, F. (2009). *Contrôle d'un Système Multi-Agents Réactif par Modélisation et Apprentissage de sa Dynamique Globale*. Thèse de doctorat, Université Nancy II.
Cité page 155
- Kubera, Y. (2010). *Simulations orientées-interaction des systèmes complexes*. Thèse de doctorat, Université des Sciences et Technologie de Lille-Lille I.
Cité page 20
- Kwon, S., Lee, Y., Park, J. Y., Sohn, D., Lim, J. H., et Im, S. (2009). An efficient three-dimensional adaptive quasicontinuum method using variable-node elements. *Journal of Computational Physics*, 228(13) :4789–4810.
Cité page 67
- Kärnä, T., Kamesaki, K., et Tsukuda, H. (1999). A numerical model for dynamic ice-structure interaction. *Computers & Structures*, 72(4-5) :645–658.
Cité page 120

- Labrosse, S. et Jaupart, C. (2007). Thermal evolution of the earth : Secular changes and fluctuations of plate characteristics. *Earth and Planetary Science Letters*, 260(3) :465–481.
Cité pages 36 et 48
- Lales, C. (2007). *Modélisation gros grains et simulation multi-agents-Application à la membrane interne mitochondriale*. Thèse de doctorat, Université Sciences et Technologies-Bordeaux I.
Cité page 60
- Lallemand, S., Heuret, A., Faccenna, C., et Funiciello, F. (2008). Subduction dynamics as revealed by trench migration. *Tectonics*, 27(3).
Cité page 42
- Lau, M. (2001). *Ice Forces on a Faceted Cone due to the Passage of a Level Ice Field*. Thèse de doctorat, Memorial University of Newfoundland.
Cité page 1
- Le Bris, C. (2005). *Systèmes multi-échelles : modélisation et simulation*, volume 47. Springer.
Cité page 27
- Le Gal, C., Olagnon, M., Parenthoen, M., Beal, P.-A., et Tisseau, J. (2007). Comparison of sea state statistics between a phenomenological model and field data. Dans *OCEANS 2007 - Europe*, pages 1–6. IEEE.
Cité page 20
- Le Moigne, J.-L. (1993). *La modélisation des systèmes complexes*. Dunod.
Cité page 7
- Le Yaouanq, S., Redou, P., Le Gal, C., Abgrall, J. F., et Tisseau, J. (2011a). Multi-agent systems and heterogeneous scales interactions. application to pharmacokinetics of vitamin k antagonists. *Advances in Artificial Life, ECAL 2011*, pages 447–454.
Cité page 2
- Le Yaouanq, S., Redou, P., Le Gal, C., Abgrall, J. F., et Tisseau, J. (2011b). Redundant modelling of multiscale interactions with multi-agent systems. application to pharmacokinetic/pharmacodynamic of vitamin k antagonists. Dans *European Conference of Complex Systems, ECCS 2011*, page 137.
Cité page 2
- Lemarrec, P. (2000). *Cosimulation multiniveaux dans un flot de conception multilangage*. Thèse de doctorat, Institut National Polytechnique de Grenoble-INPG.
Cité page 77
- Lesne, A. (2003). *Approches multi-échelles en physique et en biologie*. Thèse de doctorat, Université Pierre et Marie Curie.
Cité page 69
- Li, Z.-X., Bogdanova, S., Collins, A. S., Davidson, A., De Waele, B., Ernst, R., Fitzsimons, I. C., Fuck, R., Gladkochub, D., Jacobs, J., et al. (2008). Assembly, configuration, and break-up history of rodonia : a synthesis. *Precambrian research*, 160(1) :179–210.
Cité page 49

- Lind, J. (2001). *Iterative software engineering for multiagent systems : the MASSIVE method*. Springer-Verlag.
Cité page 15
- Loset, S. (2006). Ice actions on sloping-sided structures. Course material, University of Science and Technology, Trondheim.
Cité page 122
- Lubbad, R. et Loset, S. (2011). A numerical model for real-time simulation of ship-ice interaction. *Cold Regions Science and Technology*, 65(2) :111–127.
Cité page 130
- Marchenko, A. (2006). A method of calculating ice loads when ice piles up on a fixed wall. *Journal of Applied Mathematics and Mechanics*, 70(3) :387 – 398.
Cité page 143
- Marchenko, A. et Karulin, E. (2005). A method for the calculation of ice loads by ice motion against wide structure. Dans *Proceedings of POAC 2005*, volume 1, pages 441–452.
Cité pages 131 et 142
- Mariano, P., Pereira, A., Correia, L., Ribeiro, R., Abramov, V., Szirbik, N., Goossenaerts, J., Marwala, T., et De Wilde, P. (2001). Simulation of a trading multi-agent system. Dans *Systems, Man, and Cybernetics, 2001 IEEE International Conference on*, volume 5, pages 3378–3384 vol.5.
Cité page 15
- Ming, P. et Yang, J. Z. (2009). Analysis of a one-dimensional nonlocal quasi-continuum method. *Multiscale Modeling & Simulation*, 7(4) :1838–1875.
Cité page 66
- Minsky, M. (1965). Matter, mind and models. *International Federation of Information Processing Congress*.
Cité page 9
- Morel, B. et Alexander, P. (2003). Automating component adaptation for reuse. Dans *Automated Software Engineering, 2003. Proceedings. 18th IEEE International Conference on*, pages 142–151. IEEE.
Cité page 75
- Navarrete Gutierrez, T. (2012). *A control architecture for complex systems, based on multi-agent simulation*. Thèse de doctorat, Université de Lorraine.
Cité page 15
- Néron, E., Audin, F., Ndiaye, I., Boukebab, K., Serrhini, K., Maïzia, M., Gralepois, M., Gasnier, B., et Desramault, N. (2012). Système d’aide à la décision multicritère pour la logistique de l’évacuation à grande échelle. Dans *7ème Workshop Interdisciplinaire sur la Sécurité Globale*.
Cité pages 19 et 82
- Nevel (1992). Ice forces on cones from floes. Dans *IAHR-92*, volume 3, pages 1391 – 1404.
Cité pages 130 et 131

- Nevel, D. E. (1965). A semi-infinite plate on an elastic foundation. Rapport technique, DTIC Document.
Cité page [130](#)
- Nevel, D. E. (1972). The ultimate failure of a floating ice sheet. Dans *las Actas del International Association for Hydraulic Research, Ice Symposium*, pages 17–22.
Cité pages [130](#) et [145](#)
- Nikolaeva, K., Gerya, T., et Marques, F. (2010). Subduction initiation at passive margins : numerical modeling. *Journal of Geophysical Research : Solid Earth (1978–2012)*, 115(B3).
Cité page [48](#)
- Pardo-Castellote, G. (2003). Omg data-distribution service : Architectural overview. Dans *Distributed Computing Systems Workshops, 2003. Proceedings. 23rd International Conference on*, pages 200–206. IEEE.
Cité page [81](#)
- Parenthoën, M. (2004). *Animation phénoménologique de la mer - une approche énaïve*. Thèse de doctorat, Université de Bretagne occidentale - Brest.
Cité page [20](#)
- Pavliotis, G. et Stuart, A. (2008). *Multiscale methods : averaging and homogenization*, volume 53. Springer.
Cité page [69](#)
- Polceanu, M. (2015). *ORPHEUS : Reasoning and Prediction with Heterogeneous rEpresentations Using Simulation*. Thèse de doctorat, Université de Bretagne Occidentale (UBO).
Cité page [156](#)
- Quesnel, G., Duboz, R., et Ramat, E. (2004). Wrapping into devs simulator : A case study. Dans *Proceedings of the Conference on Conceptual Modeling and Simulation, Genoa, Italy*, pages 28–30.
Cité page [76](#)
- Quesnel, G., Duboz, R., et Ramat, É. (2009). The virtual laboratory environment—an operational framework for multi-modelling, simulation and analysis of complex dynamical systems. *Simulation Modelling Practice and Theory*, 17(4) :641–653.
Cité page [76](#)
- Ramat, E. (2006). Introduction à la modélisation et à la simulation à événements discrets. *Modélisation et simulation multi-agents : Applications pour les sciences de l’homme et de la société*, pages 50–73.
Cité page [76](#)
- Randall, D., Khairoutdinov, M., Arakawa, A., et Grabowski, W. (2003). Breaking the cloud parameterization deadlock. *Bulletin of the American Meteorological Society*, 84(11) :1547–1564.
Cité page [64](#)
- Redou, P., Gaubert, L., Desmeulles, G., Béal, P. A., Le Gal, C., et Rodin, V. (2010). Absolute stability of chaotic asynchronous multi-interactions schemes for solving ODE. *Computer*

- Modeling in Engineering and Sciences*, 70(1) :11.
Cité pages 28 et 181
- Redou, P., Kerdelo, S., Le Gal, C., Querrec, G., Rodin, V., Abgrall, J. F., et Tisseau, J. (2005). Reaction-agents : first mathematical validation of a multi-agent system for dynamical biochemical kinetics. *Progress in Artificial Intelligence*, pages 156–166.
Cité page 20
- Redou, P., Sébastien, K., Desmeulles, G., Abgrall, J.-F., Rodin, V., et Tisseau, J. (2007). Formal validation of asynchronous interaction-agents algorithms for reaction-diffusion problems. Dans *PADS'07, 21st International Workshop on Principles of Advanced and Distributed Simulation*, pages 26–34.
Cité pages 20 et 28
- Rosen, R. (1969). Hierarchical organization in automata theoretic models of the central nervous system. Dans *Information processing in the nervous system*, pages 21–35. Springer.
Cité page 62
- Ross, R. B. et Mohanty, S. (2008). *Multiscale simulation methods for nanomaterials*. John Wiley & Sons.
Cité page 69
- Ruelland, R. (2002). *Apport de la co-simulation dans la conception de l'architecture des dispositifs de commande numérique pour les systèmes électriques*. Thèse de doctorat, Institut National Polytechnique de Toulouse.
Cité page 77
- Samaey, G., Kevrekidis, I. G., et Roose, D. (2007). Patch dynamics : Macroscopic simulation of multiscale systems. *PAMM*, 7(1) :1025803–1025804.
Cité pages 71 et 72
- Seghrouchni, A. E. F., Florea, A. M., et Olaru, A. (2010). Multi-agent systems : A paradigm to design ambient intelligent applications. Dans Essaaidi, M., Malgeri, M., et Badica, C., éditeurs, *Intelligent Distributed Computing IV*, volume 315 of *Studies in Computational Intelligence*, pages 3–9. Springer Berlin Heidelberg.
Cité pages 1 et 15
- Septseault, C., Béal, P.-A., Le Yaouanq, S., Dudal, A., et Roberts, B. (2014). A New Ice SimulationTool Using a Multi-Model Program. *Artic Technology Conference 2014*.
Cité pages 3, 124 et 134
- Septseault, C., Béal, P.-A., Yaouanq, S. L., Dudal, A., Roberts, B., et al. (2015). Update on a new ice simulation tool using a multi-model program. Dans *OTC Arctic Technology Conference*. Offshore Technology Conference.
Cité pages 3 et 124
- Shafrova, S. (2007). *First-year sea ice features. Investigation of ice field strength heterogeneity and modelling of ice rubble behaviour*. Thèse de doctorat, Norwegian University of Science and Technology, Trondheim.
Cité page 145

- Shoham, Y. (1993). Agent-oriented programming. *Artificial intelligence*, 60(1) :51–92.
Cité page 16
- Siebert, J. (2011). *Approche multi-agent pour la multi-modélisation et le couplage de simulations. Application à l'étude des influences entre le fonctionnement des réseaux ambiants et le comportement de leurs utilisateurs*. Thèse de doctorat, Université Henri Poincaré-Nancy I.
Cité pages 76, 77 et 78
- Sieburg, H. (1990). Physiological studies in silico. *Sciences of Complexity*.
Cité page 13
- Silver, P. G. et Behn, M. D. (2008). Intermittent plate tectonics? *science*, 319(5859) :85–88.
Cité page 51
- Soulié, J. (2012). *Multi-Modélisation et Simulation de Système Complexes : de la théorie à l'application*. Thèse de doctorat, Université du Littoral Côte d'Opale.
Cité page 76
- Tadmor, E. B., Ortiz, M., et Phillips, R. (1996). Quasicontinuum analysis of defects in solids. *Philosophical Magazine A*, 73(6) :1529–1563.
Cité page 66
- Timco, G. et Croasdale, K. (2006). How well can we predict ice loads? Dans *Proceedings 18th International Symposium on Ice, IAHR'06*, volume 1, pages 167–174.
Cité page 1
- Tisseau, J. (2001). *Réalité virtuelle : autonomie in virtuo*. Habilitation à Diriger des Recherches, Université de Rennes 1.
Cité pages 7, 14 et 17
- Varela, F. (1979). *Principles of biological autonomy*. North Holland New York.
Cité page 20
- Weinan, E., Engquist, B., et Huang, Z. (2003). Heterogeneous multiscale method : a general methodology for multiscale modeling. *Physical Review B*, 67(9) :092101.
Cité pages 73 et 74
- Wiener, N. (1950). *Cybernétique et société (1950)*. Paris : Union Générale d'Editions (1971).
Cité page 7
- Williams, R., Konev, B., et Coenen, F. (2014). Multi-agent environment exploration with ar.drones. Dans *Advances in Autonomous Robotics Systems*, volume 8717 of *Lecture Notes in Computer Science*, pages 60–71. Springer International Publishing.
Cité page 15
- Wilson, J. T. (1966). Did the atlantic close and then re-open? volume 211 :676–681.
Cité pages 37 et 49
- Wooldridge, M. (2001). Intelligent agents : Theory and practice. *Knowledge Engineering Review*.
Cité page 16

- Wooldridge, M. et Jennings, N. R. (1994). Agent theories, architectures, and languages : A survey. Dans Wooldridge, M. J. et Jennings, N. R., éditeurs, *Intelligent Agents : ECAI-94 Workshop on Agent Theories, Architectures, and Languages*, pages 1–39. Springer, Berlin,.
Cité page [16](#)
- Yellin, D. M. et Strom, R. E. (1997). Protocol specifications and component adaptors. *ACM Transactions on Programming Languages and Systems (TOPLAS)*, 19(2) :292–333.
Cité page [75](#)
- Zeigler, B. P., Praehofer, H., et Kim, T. G. (2000). *Theory of modeling and simulation : integrating discrete event and continuous complex dynamic systems*. Academic press.
Cité page [76](#)
- Zeigler, B. P., Praehofer, H., Kim, T. G., et al. (1976). *Theory of modeling and simulation*, volume 19. John Wiley New York.
Cité page [10](#)

Méthodes numériques multi-agents

Les phénomènes que nous souhaitons expérimenter sont classiquement modélisés par des équations différentielles, organisées en système. Celui-ci ne peut cependant que très rarement être résolu analytiquement dans le contexte. Dans de tels cas, il est possible d'utiliser une méthode numérique pour tenter de déterminer une approximation de la solution. Une telle méthode démarre depuis un état initial connu ou estimé et trouve des approximations successives qui devraient converger vers la solution sous certaines conditions.

Dans le cas général de la résolution de systèmes d'équations différentielles ordinaires, les *problèmes de Cauchy* s'expriment classiquement sous la forme [Hairer et al., 1993] :

$$\begin{cases} y'(t)=f(t, y(t)) \\ y(0)=y_0 \end{cases} \quad (\text{A.1})$$

Nous supposons que le problème de Cauchy (A.1) admet une unique solution sur un intervalle $[0, T]$ (T fini). Pour résoudre numériquement (A.1), la méthode naturelle est de découper l'intervalle $[0, T]$ en N intervalles, de longueurs non nécessairement identiques.

$$0 = t_0 < t_1 < \dots < t_n < \dots < t_{N-1} < t_N = t_0 + T \quad 0 \leq n \leq N$$

Le principe de la résolution numérique d'un système d'équation différentielles consiste à approcher les valeurs des solutions exactes $y(t_n)$ par des valeurs numériques y_n , au moyen d'itérations successives de pas d'intégration. Chaque pas consiste à passer de l'état y_n à l'état y_{n+1} en utilisant une méthode numérique. Ces méthodes sont appelées méthodes à un pas. Pour les présenter, nous prendrons comme exemple les méthodes d'Euler qui n'ont qu'un intérêt pédagogique de par leur faible précision.

Par soucis de clarté, la suite de section s'appuiera sur un exemple scolaire, développé dans [Béal et al., 2008], impliquant la simulation de plusieurs phénomènes. Nous choisissons de développer un exemple à la fois simple et révélateur des forces et des faiblesses de l'utilisation de méthodes numériques multi-agents : un système masse/ressort.

Le système est considéré à un seul degré de liberté (on pose x , la position de la masse) et les trois phénomènes qui interviennent sont :

- le poids modélisé par $F_p = -m \times g$;
- la réaction du ressort $F_r = -k \times x$;
- l'amortissement $F_a = -\gamma \times x'$.

Nous allons maintenant détailler la simulation de ce système en utilisant le principe fondamental de la dynamique $\sum \vec{F} = m \times x''$ grâce à différentes méthodes d'intégration.

Méthodes explicites

Les méthodes explicites sont les méthodes de résolution numérique de (A.1) qui peuvent s'écrire sous la forme

$$y_{n+1} = y_n + h_n \Phi_f(t_n, y_n, h_n) \quad (\text{A.2})$$

où la longueur h_n du $n_i\grave{e}me$ pas est

$$h_n = t_{n+1} - t_n \quad (\text{A.3})$$

et $\Phi_f : [t_0, t_0 + T] \times \mathbb{R} \times \mathbb{R} \rightarrow \mathbb{R}$ est une fonction que l'on supposera continue.

Dans le cas de la méthode d'Euler explicite nous avons :

$$y_{n+1} = y_n + h_n f(t_n, y_n) \quad (\text{A.4})$$

Comme indiqué précédemment, la résolution numérique d'un système d'équations différentielles engendre forcément une erreur numérique. Pour valider les calculs, celle-ci doit être estimée. A l'instant t_{n+1} , cette erreur peut-être décomposée en deux parties :

- **l'erreur de consistance** ϵ_n , correspondant à l'erreur qui vient d'être commise sur le pas de temps $[t_n, t_{n+1}]$

$$\epsilon_n = y(t_{n+1}) - [y(t_n) + h_n f(t_n, y(t_n))] \quad (\text{A.5})$$

- **l'erreur globale** e_n qui provient de tous les pas de temps antérieurs

$$e_n = y(t_n) - y_n \quad (\text{A.6})$$

Malgré ces erreurs, on peut montrer que le schéma est :

- **consistant** si $\sum_{n=0}^N |e_n|$ tend vers 0 quand h_n tend vers 0
- **stable** si e_n reste borné pour tout ϵ_n

Si une méthode est consistante et stable, on peut en déduire qu'elle est **convergente** [Ascher et Petzold, 1998], c'est à dire que

$$\lim_{h \rightarrow 0} \max_{0 \leq n \leq N} |e_n| = 0 \quad (\text{A.7})$$

Voyons à présent comment la méthode d'Euler explicite s'applique à notre exemple.

Le système masse/ressort est un système du second ordre, on se ramène donc à un système du premier ordre en posant :

$$X = \begin{pmatrix} x \\ x' \end{pmatrix} \Rightarrow X' = \begin{pmatrix} x' \\ -g - \frac{k}{m} \times x - \frac{\gamma}{m} \times x' \end{pmatrix} \text{ soit : } X' = \begin{pmatrix} 0 & 1 \\ -\frac{k}{m} & -\frac{\gamma}{m} \end{pmatrix} \times X + \begin{pmatrix} 0 \\ g \end{pmatrix}$$

En appliquant (A.2) à ce système, on obtient :

$$\begin{cases} x(t+dt) = x(t) + dt \times x'(t) \\ x'(t+dt) = x'(t) + dt \times \left(-g - \frac{k}{m} \times x(t) - \frac{\gamma}{m} \times x'(t) \right) \end{cases} \quad (\text{A.8})$$

On simule ce système avec les paramètres $k = 10$, $m = 1$, $\gamma = 1$.



FIGURE A.1 – **Méthode d'Euler explicite** – Le graphique de gauche montre plusieurs résultats convergents avec un pas de temps inférieur à $\frac{1}{k}$, le graphique du milieu montre un résultat non convergent pour un pas de temps égal à $\frac{1}{k}$, et le graphique de droite, un résultat divergeant avec un pas de temps supérieur à $\frac{1}{k}$

Sur le graphique de gauche de la figure A.1, on observe trois résultats obtenus avec un pas de temps inférieur à $\frac{1}{k}$, plus le pas de temps est faible plus le résultat est précis et le système converge vers le même état d'équilibre. Sur celui du centre, on a choisi un pas de temps égal à $\frac{1}{k}$, le système ne trouve plus son point d'équilibre. Enfin, sur celui de droite, avec un pas de temps supérieur à $\frac{1}{k}$, le système diverge.

Lors de la résolution numérique d'un système d'équations différentielles, la valeur du pas d'intégration est dictée à la fois par la précision numérique souhaitée mais aussi par la stabilité. On dira que le système est **raide** si la stabilité de la méthode numérique employée induit une contrainte sur le pas de temps plus forte que l'exigence de précision. Notre exemple du système masse/ressort en est la parfaite illustration. La raideur k , physique cette fois, du ressort nous impose un pas de temps inférieur à $\frac{1}{k}$ pour conserver la stabilité de la simulation.

Le coût du schéma d'Euler explicite est déterminé par les évaluations de la fonction f à chaque pas de temps. Economiser du temps de calcul revient donc à réduire N , à augmenter h_n . Cependant, les résultats présentés ci-dessus nous montrent qu'il est impossible d'augmenter considérablement le pas de temps d'une méthode explicite sans mettre à mal la stabilité du système. Pour palier ce problème, il est possible d'utiliser une seconde catégorie de méthodes numériques, les méthodes implicites.

Méthodes implicites

Nous avons vu que les méthodes explicites permettent d'approximer la valeur de y_{n+1} à partir de la valeur de y_n à t_n et de l'intégration sur le pas de temps h_n de la fonction Φ_f ,

elle-même étant la dérivée de la fonction $f(t_n, y_n, h_n)$. Dans certains cas, nous ne disposons pas de méthode directe pour calculer y_{n+1} en fonction de y_n . Nous recourons alors à des méthodes dites « implicites » qui utilisent en plus une approximation de y_{n+1} à l'instant t_{n+1} pour le calcul de Φ_f .

Nous avons donc :

$$y_{n+1} = y_n + h_n \Phi_f(t_n, h_n, t_{n+1}, y_{n+1}, h_{n+1}) \quad 0 \leq n \leq N-1 \quad (\text{A.9})$$

La A -stabilité est une propriété imposée lorsque sur un certain intervalle on désire choisir un grand pas d'intégration h pour « gagner du temps ». On dit qu'une méthode est A -stable si

$$\forall h > 0, \lim_{n \rightarrow \infty} y_n = 0 \quad (\text{A.10})$$

pour toute équation (dite "test")

$$\begin{cases} y' = -\lambda y(t) & \text{pour } \operatorname{Re}(\lambda) > 0 \\ y(0) = y_0 \end{cases} \quad (\text{A.11})$$

Ainsi, le schéma d'Euler implicite

$$y_{n+1} = y_n + h_n f(t_{n+1}, y_{n+1}) \quad 0 \leq n \leq N-1, \quad (\text{A.12})$$

peut s'écrire :

$$y_{n+1} = \prod_{k=0}^n \frac{1}{1 + \lambda h_k} y_0 \quad (\text{A.13})$$

On peut montrer que $y_n \rightarrow 0$ quand $n \rightarrow +\infty$ pour tout n indépendamment de h_n .

La méthode d'Euler implicite est donc A -stable, ce qui n'est pas le cas d'Euler explicite. On pourra ainsi prendre un grand pas de temps pour accélérer la résolution. En contrepartie, le calcul de y_{n+1} à partir de y_n revient à résoudre un système, linéaire ou non, à chaque itération. Dans l'évaluation du coût global de la méthode implicite, il faudra tenir compte du celui de chaque pas de temps en plus du nombre de pas.

Dans le cas de notre système masse/ressort, en utilisant la transformation en variable vectorielle comme précédemment et en appliquant (A.12) on obtient :

$$\begin{cases} x(t+dt) = x(t) + dt \times x'(t+dt) \\ x'(t+dt) = x'(t) + dt \times \left(-g - \frac{k}{m} \times x(t+dt) - \frac{\gamma}{m} \times x'(t+dt) \right) \end{cases} \quad (\text{A.14})$$

En simulant le système avec les mêmes paramètres, on obtient les résultats présentés figure A.2.

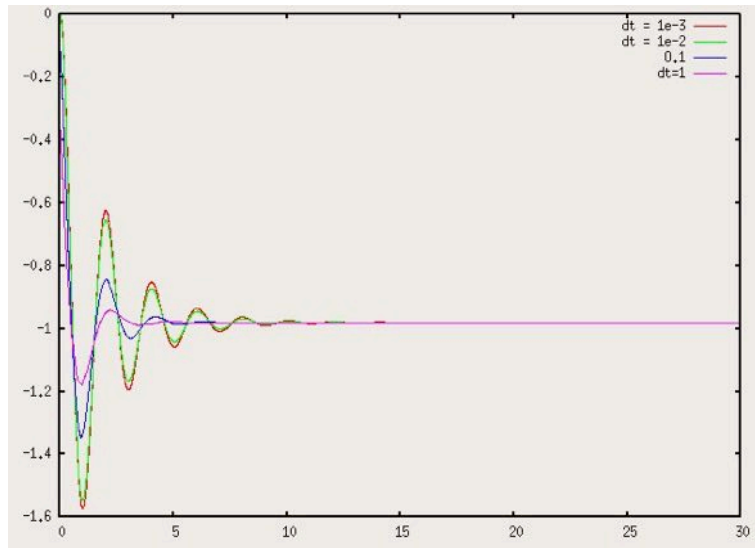


FIGURE A.2 – Méthode d'Euler implicite – Les résultats convergent quelque soit le pas de temps.

Avec la méthode d'Euler implicite, le système converge vers son point d'équilibre quel que soit le pas de temps choisi. Si le pas de temps est trop faible, la convergence est trop rapide mais ne fera pas diverger le reste de la simulation. De plus, si on change le pas de temps en cours de simulation en choisissant une fréquence plus adaptée, le système peut retrouver un fonctionnement correct.

Méthodes numériques et systèmes multi-agents

Dans le cas d'une simulation multi-agents, le problème global peut être décomposé en sous-systèmes. Les agents résolvent chacun une partie du calcul, à des fréquences adaptables et hétérogènes. Les agents sont localement itérés puis leurs résultats sont superposés.

Dans une optique de comparaison des résultats, nous détaillons tout d'abord l'application d'un schéma d'intégration d'Euler explicite sur notre exemple du système masse/ressort en le modélisant par un système multi-agents.

Simulation explicite

Chaque agent représente un phénomène physique indépendant du reste du système et possède sa propre fréquence d'exécution. L'ensemble des agents travaillent sur l'état du système représenté par le vecteur $X = \begin{pmatrix} x \\ x' \end{pmatrix}$.

En appliquant (A.2) à chaque phénomène, on obtient les trois algorithmes itératifs autonomes suivants :

Poids

$$\begin{aligned} x &\leftarrow x + dtPoids \times x' \\ x' &\leftarrow x' + dtPoids \times g \end{aligned}$$

Ressort

$$\begin{aligned} x &\leftarrow x + dtRessort \times x' \\ x' &\leftarrow x' + dtRessort \times \frac{k}{m} \times x \end{aligned}$$

Amortissement

$$\begin{aligned} x &\leftarrow x + dtAmortissement \times x' \\ x' &\leftarrow x' + dtAmortissement \times -\frac{\gamma}{m} \times x' \end{aligned}$$

En simulant le système avec les mêmes paramètres que précédemment, on obtient les résultats présentés figure A.3.

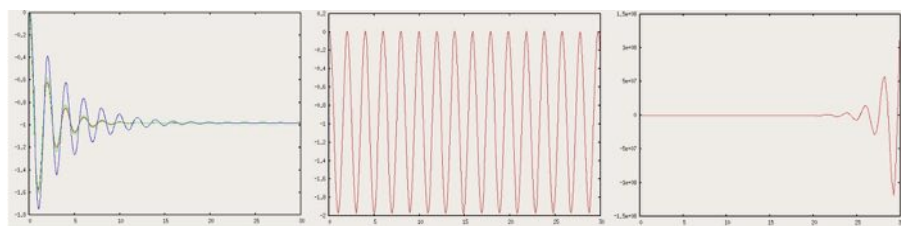


FIGURE A.3 – **Méthode d'Euler explicite multi-agents à pas de temps égaux** – Les pas de temps d'exécution sont égaux pour tous les phénomènes. On retrouve des résultats similaires à ceux de la figure A.1.

Si on choisit une valeur identique pour tous les pas de temps des phénomènes, on obtient des résultats très similaires à ceux du système modélisé de façon dite « monolithique », c'est-à-dire en un seul bloc, ce qui n'est pas surprenant.

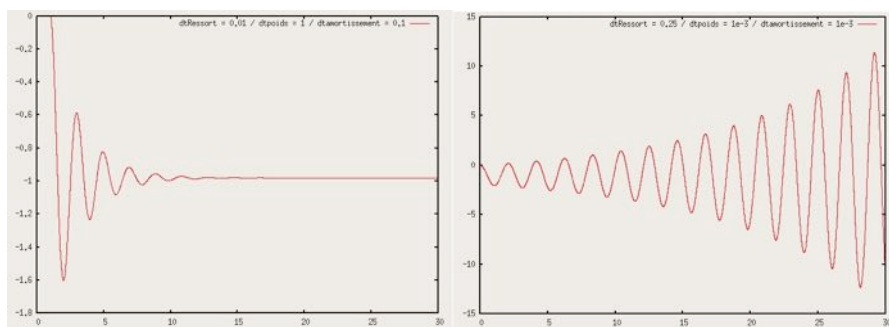


FIGURE A.4 – **Méthode d'Euler explicite multi-agents à pas de temps différents** – Les pas de temps d'exécution sont propres à chaque phénomène. On observe que la stabilité dépend du pas de temps choisi pour simuler l'élongation du ressort.

En revanche, avec des pas de temps propres à chaque phénomène, la stabilité du système dépend du pas de temps choisi pour les phénomènes qui nécessiteraient un pas de temps faible

pour être stable avec une intégration d'Euler explicite (dans notre exemple, le ressort). En d'autres termes les phénomènes les plus raides dirigent la stabilité de l'ensemble, comme le montrent les résultats de la figure A.4.

Simulation implicite

Nous allons maintenant décrire l'utilisation d'une méthode d'intégration d'Euler implicite dans une simulation multi-agents. Pour cela, on suppose connue de chaque phénomène la force totale appliquée au système (notée σ) ainsi que l'effort qu'il a produit au pas de temps précédent. On applique (A.12) à chacun des phénomènes en considérant les efforts extérieurs constants. On note par l'indice $n - 1$ la valeur des variables telles qu'elles étaient avant l'activation du phénomène et on cherche la valeur de ces variables à l'indice n .

- Le comportement du phénomène associé au poids s'écrit alors :

1. calcul de l'effort : $Fp_n = -m \times g$

2. intégration : $\sigma - Fp_{n-1} - m \times g = m \times x''$
 $\Leftrightarrow x'' = \frac{\sigma - Fp_{n-1}}{m} - g$

en utilisant le vecteur $X \begin{vmatrix} x \\ x' \end{vmatrix}$ et (A.12) on obtient :

$$\begin{cases} x'_n = x'_{n-1} + \frac{\sigma - Fp_{n-1}}{m} - g \\ x_n = x_{n-1} + dt \times x'_n \end{cases}$$

3. actualisation de σ : $\sigma_n = \sigma_{n-1} - Fp_{n-1} + Fp_n$

- Celui du ressort :

1. calcul de l'effort : $Fr_n = -k \times x$

2. intégration : $\sigma - Fr_{n-1} - k \times x = m \times x''$
 $\Leftrightarrow x'' = \frac{\sigma - Fr_{n-1}}{m} - \frac{k}{m} \times x$

en utilisant le vecteur $X \begin{vmatrix} x \\ x' \end{vmatrix}$ et (A.12) on obtient :

$$\begin{cases} x'_n = x'_{n-1} + \frac{\sigma - Fr_{n-1}}{m} - \frac{k}{m} \times x_n \\ x_n = x_{n-1} + dt \times x'_n \end{cases}$$

$$\Leftrightarrow \begin{cases} x'_n = \frac{1}{1 + \frac{k \times dt^2}{m}} \times \left(x'_{n-1} + \frac{dt}{m} \times (\sigma - Fr_{n-1} - k \times x_{n-1}) \right) \\ x_n = x_{n-1} + dt \times x'_n \end{cases}$$

3. actualisation de σ : $\sigma(n) = \sigma_{n-1} - Fr_{n-1} + Fr(n)$

- et celui de l'amortissement :

1. calcul de l'effort : $Fa_n = -\gamma \times x'$

2. intégration : $\sigma - Fa_{n-1} - \gamma \times x' = m \times x''$
 $\Leftrightarrow x'' = \frac{\sigma - Fa_{n-1}}{m} - \frac{\gamma}{m} \times x'$

en utilisant le vecteur $X \begin{vmatrix} x \\ x' \end{vmatrix}$ et (A.12) on obtient :

$$\begin{cases} x'_n = x'_{n-1} + \frac{\sigma - Fa_{n-1}}{m} - \frac{\gamma}{m} \times x'_n \\ x_n = x_{n-1} + dt \times x'_n \end{cases} \Leftrightarrow \begin{cases} x'_n = \frac{1}{1 + \frac{\gamma \times dt}{m}} \times \left(x'_{n-1} + \frac{dt}{m} \times (\sigma - Fa_{n-1}) \right) \\ x_n = x_{n-1} + dt \times x'_n \end{cases}$$

3. actualisation de σ : $\sigma_n = \sigma_{n-1} - Fa_{n-1} + Fa_n$

Notons que dans notre exemple la formalisation de la méthode d'Euler implicite est simple à expliciter puisque le comportement de nos phénomènes est linéaire. Dans le cas de phénomènes non linéaires, il peut être nécessaire d'avoir recours à des procédures itératives plus coûteuses.

Lorsqu'on simule à nouveau le système ainsi obtenu, en utilisant les mêmes paramètres que précédemment, on observe une plus grande stabilité.

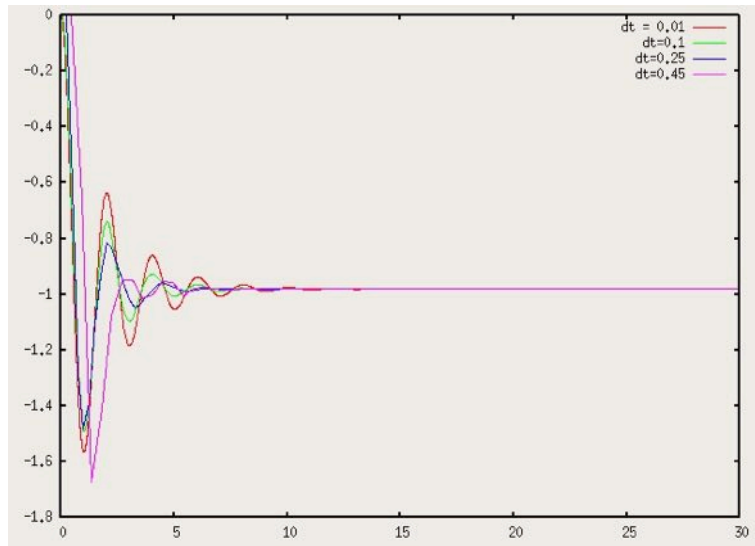


FIGURE A.5 – **Méthode d'Euler implicite multi-agents à pas de temps égaux** – Les pas de temps d'exécution sont égaux pour tous les phénomènes. Les résultats convergent et restent stables pour des pas de temps allant de 0.1s à 0.45s.

Nous constatons sur la figure A.6 que lorsque tous les phénomènes s'activent à la même fréquence, on peut observer des résultats plus stables qu'en utilisant une méthode d'intégration explicite. Cependant, lorsque la fréquence devient trop basse, on voit apparaître des valeurs s'écartant de plus en plus du point d'équilibre comme sur la figure A.7.

Pour des fréquences d'activation extrêmement basses, le système ne trouve plus son point d'équilibre et diverge. On n'a donc pas un comportement totalement similaire à celui observé sur le système monolithique qui nous permettait d'assurer la stabilité quel que soit le pas de temps fixé.

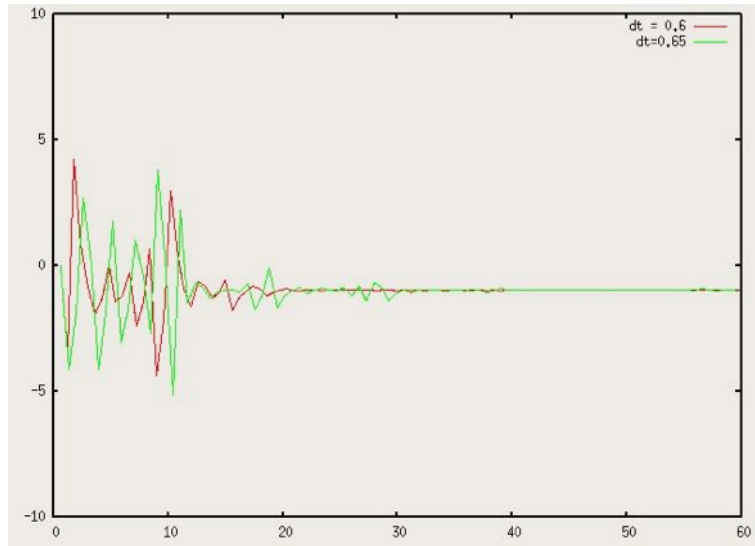


FIGURE A.6 – **Cas de stabilité de la méthode d'Euler implicite multi-agents** – Les pas de temps d'exécution sont égaux pour tous les phénomènes. Les résultats sont en dehors des plages de valeurs normales mais restent stables.

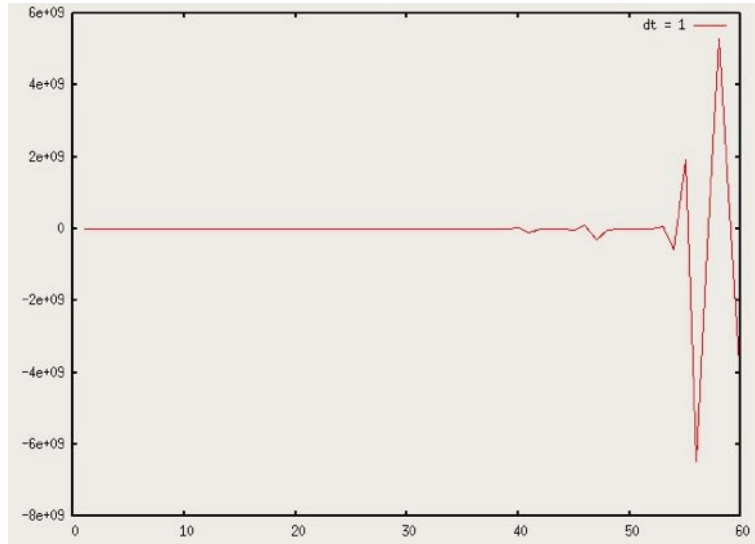


FIGURE A.7 – **Cas d'instabilité de la méthode d'Euler implicite multi-agents** – Les pas de temps d'exécution sont égaux pour tous les phénomènes. Les résultats divergent

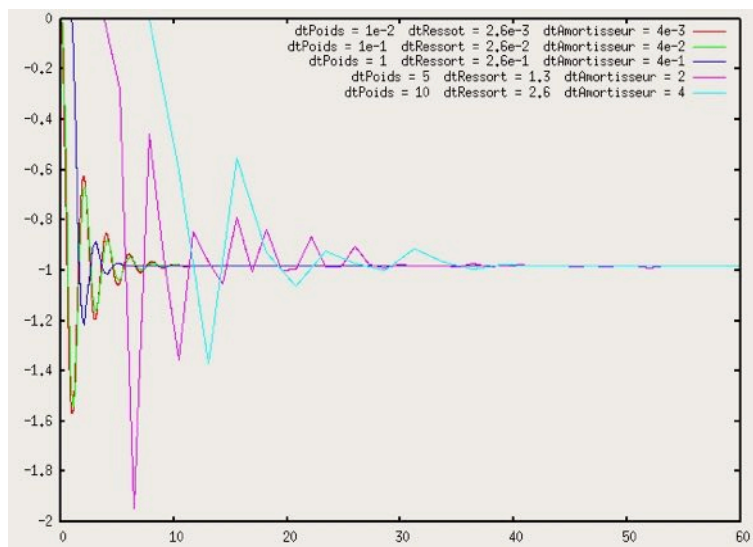


FIGURE A.8 – **Choix des différents pas de temps pour la méthode d’Euler implicite multi-agents** – Les pas de temps d’exécution sont différents pour tous les phénomènes. Le ressort est simulé à la plus haute fréquence. La simulation reste stable.

En utilisant des fréquences différentes pour chaque phénomène, les expériences montrent que le système reste stable si on fixe le pas de temps du ressort inférieur aux pas de temps des autres phénomènes. On remarque notamment sur les courbes rose et bleu clair de la figure A.8 que le système trouve son point d’équilibre alors que les pas de temps sont tous supérieurs au pas de temps choisi lors de la simulation représentée sur la figure A.7.

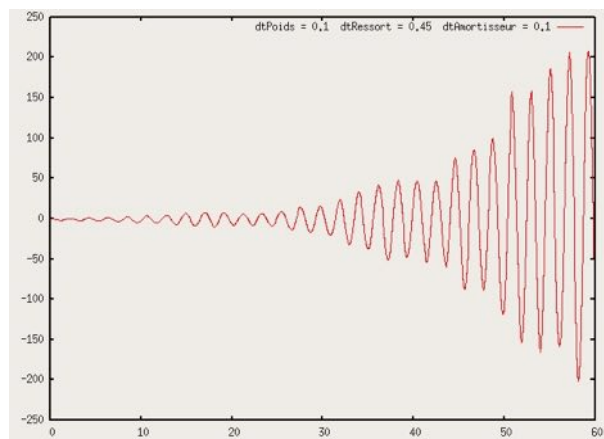


FIGURE A.9 – **Raideur et méthode d’Euler implicite multi-agents** – Les pas de temps d’exécution sont différents pour tous les phénomènes. Le ressort est simulé à la plus basse fréquence. La stabilité de la simulation n’est plus garantie.

Pour une simulation dans laquelle le ressort est simulé à la fréquence la plus basse, la stabilité du système dépend de cette fréquence. On remarque sur la figure A.9 que le système n’est pas stable alors que toutes les fréquences d’activation des phénomènes sont supérieures à celles de la simulation représentée par la courbe rose sur la figure A.5.

Pour autant, ce n'est pas l'approche agent qui est ici en cause mais plutôt l'ordonnancement asynchrone des activités. En effet, les travaux de [Redou et al., 2010] ont montré que la stabilité absolue des schémas implicites n'est plus garantie dès lors qu'ils ne sont pas résolus de manière synchrone et d'autant plus que les phénomènes simulés sont antagonistes.

Résumé de thèse

Deux points de vue sont souvent opposés dans le cadre de la modélisation des systèmes complexes. D'un côté, une modélisation microscopique cherche à reproduire précisément le comportement des entités qui composent le système en les représentant de manière individuelle, par des agents par exemple ; mais en raison du grand nombre d'entités à considérer, la simulation de tels modèles entraînent des temps de calculs souvent prohibitifs pour le passage à l'échelle de systèmes réels ; par ailleurs, ce niveau de détail n'est pas toujours souhaitable lorsqu'il s'agit d'étudier le comportement du système dans sa globalité. Une autre approche consiste à changer notre point de vue sur le système et à le décomposer, non plus en entités autonomes, mais en ses phénomènes constitutifs pour étudier l'impact de chacun d'entre eux. Cette approche phénoménologique nous permet de décrire le comportement du système tel que nous le percevons ce qui se révèle particulièrement profitable lorsque nous sommes confrontés à des phénomènes qui ne sont que partiellement connus. Ces modèles macroscopiques reposent sur des lois descriptives qui autorisent des simulations plus rapides mais impliquent l'introduction de paramètres qui peuvent être difficilement identifiables dans le contexte.

Pour répondre à ce problème, nous proposons de combiner les différents points de vue de modélisation. Les lois descriptives étant le résultat d'abstractions qui résument un ou plusieurs phénomènes sous-jacents, nous faisons l'hypothèse qu'une simulation d'un modèle microscopique est équivalente à une simulation d'un modèle macroscopique correctement paramétré, à temps simulé égal et pour un phénomène donné. Aussi, notre idée est de simuler conjointement ces niveaux de description grâce aux méthodes multi-échelles hétérogènes qui reposent sur l'utilisation redondante de simulations microscopiques pour nourrir un modèle macroscopique incomplet. L'objectif est d'obtenir une simulation descriptive rapide tout en profitant de la précision d'un modèle microscopique.

Pour mettre en œuvre cette approche, nous nous appuyons sur la technique de la co-simulation qui nous autorise à coupler des modèles d'origines diverses et exprimés dans des formalismes différents. En particulier, nous proposons une architecture logicielle qui généralise la démarche de simulation redondante d'échelles de modélisation hétérogènes. Nous distinguons deux stratégies de co-simulation qui permettent de piloter un modèle macroscopique en cours de simulation. La première consiste à estimer dynamiquement, et de manière explicite, de nouvelles valeurs pour un paramètre critique donné à l'aide d'un simulateur microscopique dédié. La seconde stratégie se rapproche de certaines méthodes d'optimisation et permet de déterminer implicitement un jeu de paramètres interdépendants sur la base d'une sortie commune des différents niveaux de description simulés.

Nous appliquons nos travaux au problème concret du design de structures offshore pour des conditions polaires. Nous détaillons dans un premier temps l'implémentation d'un simulateur phénoménologique d'interactions glace-structure. Dans un second temps, nous illustrons l'intérêt et l'intégration de nos stratégies de co-simulation pour, d'une part améliorer la précision des simulations des phénomènes hydrodynamiques, et d'autre part guider un modèle de plus haut niveau à des fins de prototypage rapide.

Mots clefs : systèmes complexes, systèmes multi-agents, méthodes multi-échelles, approche phénoménologique, co-simulation, interactions glace-structure, prototypage

— Résumé —

*Co-simulation redondante d'échelles de modélisation hétérogènes
pour une approche phénoménologique.*

Deux points de vue sont souvent opposés dans le cadre de la modélisation des systèmes complexes. D'un côté, une modélisation microscopique cherche à reproduire précisément le comportement des nombreuses entités qui composent le système, ce qui impose des temps de calculs prohibitifs pour le passage à l'échelle de systèmes réels. À l'inverse, l'approche phénoménologique se concentre sur le comportement global du système. Ces modèles macroscopiques reposent sur des lois descriptives qui autorisent des simulations plus rapides mais impliquent l'introduction de paramètres qui peuvent être difficilement identifiables dans le contexte. Pour répondre à ce problème, nous proposons de combiner les différents points de vue de modélisation et d'utiliser des simulations microscopiques pour nourrir un modèle macroscopique incomplet. L'objectif est d'obtenir une simulation descriptive rapide tout en profitant de la précision d'un modèle microscopique. Pour cela, nous proposons une architecture logicielle qui s'appuie sur la technique de la co-simulation pour généraliser la démarche de simulation redondante d'échelles de modélisation hétérogènes. Nous distinguons deux stratégies de co-simulation qui permettent de piloter un modèle macroscopique en cours de simulation. La première consiste à estimer dynamiquement, et de manière explicite, de nouvelles valeurs pour un paramètre critique donné à l'aide d'un simulateur microscopique dédié. La seconde stratégie permet de déterminer implicitement un jeu de paramètres interdépendants sur la base d'une sortie commune des différents niveaux de description simulés. Nous appliquons nos travaux au problème concret du design de structures offshore pour des conditions polaires. Nous détaillons dans un premier temps l'implémentation d'un simulateur phénoménologique d'interactions glace-structure. Dans un second temps, nous illustrons l'intérêt et l'intégration de nos stratégies de co-simulation pour, d'une part améliorer la précision des simulations des phénomènes hydrodynamiques, et d'autre part guider un modèle de plus haut niveau à des fins de prototypage rapide.

Mots clefs : systèmes complexes, systèmes multi-agents, méthodes multi-échelles, approche phénoménologique, co-simulation, interactions glace-structure, prototypage.

— Abstract —

*Co-simulation of redundant and heterogeneous modelling scales
for a phenomenological approach.*

There are usually two opposite points of view for the modelling of complex systems. First, microscopical models aim at reproducing precisely the behavior of each entity of the system. In general, their great number is a major obstacle both to simulate the model in a reasonable time and to identify global behaviors. By contrast, the phenomenological approach allows the construction of efficient models from a macroscopic point of view as a superposition of phenomena. A drawback is that we often have to set empirical parameters in these descriptive models. To respond to this problem, we want to make joint use of different levels of description and to use microscopical simulations to feed incomplete macroscopical models. We would then obtain enhanced descriptive simulations with the precision of microscopical models in this way. To this end, we propose a redundant multiscale architecture which is based on the co-simulation methodology in order to generalize the redundant multiscale approach. We suggest two specific co-simulation strategies to guide a macroscopical simulation. The first one consists in dynamically and explicitly estimating critical parameters of a macroscopical model thanks to a dedicated microscopical simulator. The second one allows to implicitly determine a full set of dependant parameters on the basis of an output shared by the different levels of description. Then we apply our works to the effective problem of the design offshore structures for arctic conditions. We first describe the implementation of an ice-structure simulation tool by means of a phenomenological and multi-model approach. In a second phase, we show the benefits of our co-simulation strategies to improve the precision of hydrodynamics simulations on the one hand, and on the other to pilot a more macroscopical model for the purpose of fast prototyping.

Key words : complex systems, multi-agent systems, multiscale methods, phenomenological approach, co-simulation, ice modelling, offshore structures, design optimisation.