



Modélisation du mouvement de foules denses : phénoménologie et couplage de modèles

Etienne Pinsard

► To cite this version:

Etienne Pinsard. Modélisation du mouvement de foules denses : phénoménologie et couplage de modèles. Analyse numérique [math.NA]. Université Paris-Saclay, 2022. Français. NNT : 2022UPASM035 . tel-03955838

HAL Id: tel-03955838

<https://theses.hal.science/tel-03955838>

Submitted on 25 Jan 2023

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

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

Modélisation du mouvement de foules denses : phénoménologie et couplage de modèles

*Modeling dense crowds movements : phenomenology and
model coupling*

Thèse de doctorat de l'université Paris-Saclay

École doctorale de Mathématique Hadamard (EDMH) n°574
Spécialité de doctorat : Mathématiques appliquées
Graduate School : Mathématiques, Référent : Facultés des sciences d'Orsay

Thèse préparée dans le **Laboratoire de Mathématiques d'Orsay (Université Paris-Saclay, CNRS)** et **Laboratoire central de la préfecture de Police de Paris (Préfecture de Police de Paris)**, sous la direction de **Bertrand MAURY**, Professeur, le co-encadrement de **Jean-Luc PAILLAT**, Ingénieur

Thèse soutenue à Paris-Saclay, le 08 décembre 2022, par

Etienne PINSARD

Composition du jury

Cécile APPERT-ROLLAND Directrice de recherche, Laboratoire de Physique des 2 infinis Irène Joliot-Curie	Présidente
Pierre DEGOND Directeur de recherche, Institut de Mathéma- tiques de Toulouse	Rapporteur & Examineur
Boris ANDREIANOV Professeur, Institut Denis Poisson	Rapporteur & Examineur
Paola GOATIN Directrice de recherche, INRIA Sophia Antipolis	Examinatrice
Julien PETTRÉ Directeur de recherche, INRIA - Rennes	Examineur
Bertrand MAURY Professeur, Laboratoire de Mathématique d'Or- say	Directeur de thèse

Titre : Mouvements de foules : phénoménologie et couplage de modèles

Mots clés : mouvements de foules ; modélisation ; couplage de modèles

Résumé : Il existe aujourd'hui une multitude de modèles pour les mouvements des foules, prenant en compte une large palette de comportements humains. Le besoin applicatif d'outils de simulation variés est de plus en plus présent dans le cadre de l'organisation de grands évènements à l'approche des Jeux Olympiques de 2024.

L'objectif de la thèse est d'étudier le couplage de modèles existants pour tirer parti des avantages de chaque approche. Un grand nombre de modèles ont été proposés ces 20 dernières années, de types microscopiques (ou lagrangiens, basés sur un suivi des individus) ou macroscopiques (ou eulériens, où

la foule est représentée par un champ de densité). Lorsque l'on s'intéresse au mouvement de foules importantes dans un bâtiment ou un domaine extérieur complexe, le niveau de description optimal peut varier selon la zone. Un premier axe de cette thèse porte sur l'analyse et l'implémentation de stratégies de couplages de modèles entre deux zones pour des modèles à une dimension. Dans un second temps, le couplage est étudié pour un modèle microscopique de dimension deux. Enfin, nous proposons un nouveau modèle macroscopique écrit par une approche numérique de la modélisation, basée sur un effet d'inhibition des individus.

Title : Crowd motion : phenomenology and model coupling

Keywords : crowd motion ; mathematical model ; model coupling

Abstract : Several crowd motion models are available today and are able to take into account a wide variety of different human behaviours. These tools are more and more used to guide the organisation of massive gatherings, such as the Olympic games of Paris in 2024.

The objective of this thesis is to study the coupling of existing models in order to benefit from the advantages of each model. During the last 20 years, several microscopic (or lagrangian, i.e. following individuals) and macroscopic (or eulerian, where indi-

viduals are represented by a continuum) models have been proposed. The optimal level of description in the context of crowd motion may vary if we consider buildings, external spaces, or complex real life situations. A first approach focuses on the analysis and implementation of model coupling strategies in one dimension. The case of a microscopic two dimensional model is then studied. Finally, we introduce a new numerical model for the simulation of macroscopic movements of individuals that takes into account an inhibitive effect between individuals.

Remerciements

Je tiens tout d'abord à remercier les personnes qui ont participé à mon encadrement : mon directeur de thèse au LMO Bertrand Maury, mon encadrant au LCPP Jean-Luc Paillat, ainsi que mon co-encadrant officieux Sylvain Faure, du LMO. J'ai eu une très grande chance de les avoir pour ma thèse, et c'est grâce à leurs qualités aussi bien humaines que scientifiques que j'ai pu aller au bout de ces trois ans. Rien ne me fera plus plaisir que de continuer à travailler avec ces personnes les prochaines années.

Je souhaite remercier mes deux rapporteurs, Boris Andreianov et Pierre Degond, pour leur relecture attentive de mon manuscrit et leur remarques constructives, ainsi que Julien Pettré, Cécile Appert-Rolland et Paola Goatin pour avoir accepté de faire partie de mon jury et pour les échanges passionnant lors de ma soutenance.

Je suis reconnaissant envers le directeur du LCPP Christophe Perzron pour m'avoir fait confiance, le directeur adjoint Aurélien Thiry pour avoir soutenu ces travaux, ainsi que le chef de la section EMP Jean-Pierre Orazy. J'ai aussi eu la chance de rencontrer de formidables personnes dans la section M2E pour qui mon affection n'a pas de limites : Adissa, Aurélien, Delphine, Eddie, Jenifer, Louis, Mathieu, Nicolas, Sylvie et Renato. Je remercie aussi Emilie, Insaf, Quentin et Sylvie pour les discussions du midi et les cafés, et Anne Thiry-Muller en particulier pour ses relectures de mon manuscrit. Je remercie enfin tous les permanents, apprentis et stagiaires que j'ai pu côtoyer.

J'ai pu me faire de nombreux amis parmi les membres du LMO, et je tiens à remercier mes cobureaux Nhi et Guillaume pour tous ces thés bus ensemble qui ont illuminé mes journées au labo, ainsi que Clément, Elise, Nir et Pierre pour les discussions mathématiques ou non. Je souhaite aussi exprimer ma gratitude envers les gestionnaires d'équipe et du laboratoire Valérie et Marie-Christine pour m'avoir aidé avec mes soucis administratifs, ainsi que Mathilde du pôle informatique pour m'avoir sauvé la mise pendant ma thèse et surtout pour ma soutenance.

Je remercie aussi Quentin Jullien du CSTB pour nos discussions que nous avons eues et pour son travail dans le groupe EVAC. Je tiens aussi par la même occasion à remercier les autres membres de ce groupe pour tous les travaux auquel j'ai pu participer.

Merci à mes amis du CCOF et du CETCA Elizabeth, Léana, John, Marion, Sylvain, Thomas et Walid avec qui je me suis énormément amusé en réalisant des vidéos entre deux nuits blanches. Merci à Adrien, Bénédicte, Cyrille et Fabien, mes amis de longue date qui m'ont soutenu depuis tant d'années et qui sont toujours aussi chers à mes yeux. Merci à ma chère mamie que j'ai été ravi de voir à ma soutenance. Merci à ma famille, papa, maman, Marine, Wandrille et Thomas, même sans être à Toulouse j'ai l'impression que vous avez été avec moi tous les jours ces trois ans.

Lou, les mots ne rendront pas justice à la place que tu as eue, mais merci pour absolument tout

1	Introduction et contexte	1
1.1	Les mouvements de foule dans la réglementation française	2
1.2	Un bref regard sur la modélisation de foules	5
2	Modélisation de mouvement de foules : état de l’art	11
2.1	Modèles microscopiques	12
2.2	Modèles macroscopiques et couplage d’échelles	20
2.3	Observables et validation des modèles	26
3	Couplage de modèles 1D	31
3.1	Couplage de modèles d’anticipation	32
3.2	Couplage de modèles d’inhibition	44
3.3	Cas d’application et extensions	51
4	Couplage de modèles d’inhibition	61
4.1	Couplage microscopique vers macroscopique de modèles d’inhibition	65
4.2	Couplage macroscopique vers microscopique de modèles d’inhibition	70
4.3	Extensions, perspectives	75
5	Modèle numérique macroscopique avec inhibition basé sur une résolution hiérarchique	79
5.1	Discretisation, hiérarchie et positions du problème	82
5.2	Différentes approches	85
5.3	Discussion, perspectives	93
6	Conclusion générale et perspectives	95
6.1	Conclusion sur les études menées	95
6.2	Perspectives - couplages et modélisation de foules	96
A	Algorithme de Fast Marching	100
	Références	102

Chapitre 1

Introduction et contexte

1.1	Les mouvements de foule dans la réglementation française . . .	2
1.2	Un bref regard sur la modélisation de foules	5

Les évènements culturels, politiques ou sportifs sont depuis toujours des occasions pour organiser des manifestations et des célébrations. Ces rassemblements sont protéiformes, allant des concerts en plein air à des défilés traversant une ville entière. Des accidents graves survenus ces dernières années soulèvent cependant des inquiétudes quant à la sécurité de ces assemblées. Les bousculades comme celles de la Love Parade en 2010 à Duisbourg ou à la Mecque en 2015 démontrent le besoin de mesures de gestion des foules. À cela s'ajoute la menace des attentats, qui peut en elle-même causer la panique comme cela s'est vu lors du mouvement de foule en 2017 à Turin. Afin de répondre à cette attente, les pouvoirs publics en France s'intéressent de plus en plus à la sécurisation de tels évènements. Aucune réglementation n'existe pourtant aujourd'hui pour encadrer les grands rassemblements.

En parallèle, l'étude de mouvements de foule a connu un regain d'intérêt rapide ces 3 dernières décennies, notamment depuis les travaux de Helbing en 1995 [HM95]. Les apports de simulations de foules de plus en plus variées et de nouvelles études sur le sujet ont remis en question les idées reçues sur le déroulement d'une évacuation.

Un des buts à long terme des travaux engagés au cours de cette thèse est d'appliquer des méthodes et d'utiliser des connaissances acquises depuis Helbing sur la simulation de mouvements de foules pour la réglementation française. Il paraît important ici de décrire plus en détails la construction de cette réglementation et ses possibles évolutions. Cela permettra de mieux comprendre la position que peut prendre la simulation dans cette situation. Nous présenterons ensuite des éléments de contexte sur la modélisation de mouvements de foules, tandis qu'un état de l'art de celle-ci plus complet sera présenté dans le chapitre 2.

1.1 Les mouvements de foule dans la réglementation française

Le traitement de mouvement de foules n'apparaît pas explicitement dans la réglementation française à l'heure actuelle. La question de l'évacuation est cependant encadrée dans le cas des bâtiments, et constitue une base sur laquelle s'appuient les pouvoirs publics. Un bref résumé de cette réglementation des bâtiments nous permettra de présenter les questionnements actuels dans le cadre des rassemblements et manifestations.

Les bâtiments. Le code de la construction et de l'habitation (CCH) fixe la réglementation concernant tous les bâtiments en France. Parmi tous les aspects traités, la question de l'évacuation est abordée dans le livre premier, titre IV : sécurité des personnes contre les risques d'incendie. Les règlements de sécurité contre les risques d'incendie et de panique qui en découlent fixent des obligations de moyens pour la construction d'établissements et classifient leur usage. L'idée essentielle est que pour qu'un bâtiment puisse être construit, il faut (entre autres conditions) que les occupants du bâtiment puissent évacuer ou se mettre dans un lieu adapté si un incendie se déclare, en ayant été exposé à une quantité de fumées jugée acceptable [Thi17]. Les textes offrent certaines dispositions pour répondre à ces exigences. Nous pouvons citer par exemple l'article CO 36 du règlement de sécurité contre les risques d'incendie pour les Établissements Recevant du Public (ERP), qui fixe le nombre et la largeur dégagements, c'est-à-dire toute partie de la construction permettant le cheminement d'évacuation des occupants : porte, sortie, issues, couloir, escalier, etc.

Parmi toutes les dispositions "clefs en main" disponibles dans la réglementation incendie, certaines ne sont parfois pas celles retenues par les commanditaires de bâtiments. Dans certains cas, il est alors possible de réaliser une étude d'ingénierie de sécurité incendie, c'est-à-dire une démonstration de sécurité fondée sur les outils de l'ingénieur, pour montrer que les solutions alternatives sont satisfaisantes. Le désenfumage est un exemple de champ technique dans lequel le recours à ces études est autorisé en France.

Le processus d'évacuation en cas d'incendie commence lorsque l'alarme incendie du bâtiment sonne. En général, les individus ont un délai de réaction de commencer un mouvement vers un lieu sûr ou une sortie du bâtiment. La durée de cette phase, appelée temps de pré-mouvement, est généralement estimée de façon empirique en fonction de l'usage du bâtiment ([GB16]). La modélisation du mouvement des individus rentre alors en jeu. Cette partie est de nos jours de plus en plus effectuée à l'aide de logiciels. En France, avant l'utilisation de logiciels cette phase pouvait être estimée dans le cas particulier des gares. Dans ce cadre, l'article GA 23 de l'arrêté du 24 décembre 2007 donne une table de valeurs de flux de passages et de vitesse de marche à utiliser pour les calculs. Le temps mis par chaque zone à se vider est ensuite estimé en divisant l'effectif théorique prévu par ces valeurs de flux.

Un bureau d'étude peut proposer une étude d'ingénierie du désenfumage, dans laquelle il va évaluer différents scénarios de propagation d'incendie à l'aide d'outils de modélisation d'incendie. Le guide du LCPP [Thi17] explique comment dimensionner ces calculs en fonction de l'usage du bâtiment et de puissances tabulées. Une simulation d'évacuation est ensuite menée, en prenant éventuellement en compte les résultats des simulations d'incendie. Dans le cadre de la répartition des foules la question des scénarios à simuler est plus complexe. On suppose souvent une répartition uniforme de l'occupation des espaces alloués, mais cela ne prend pas en compte des possibles

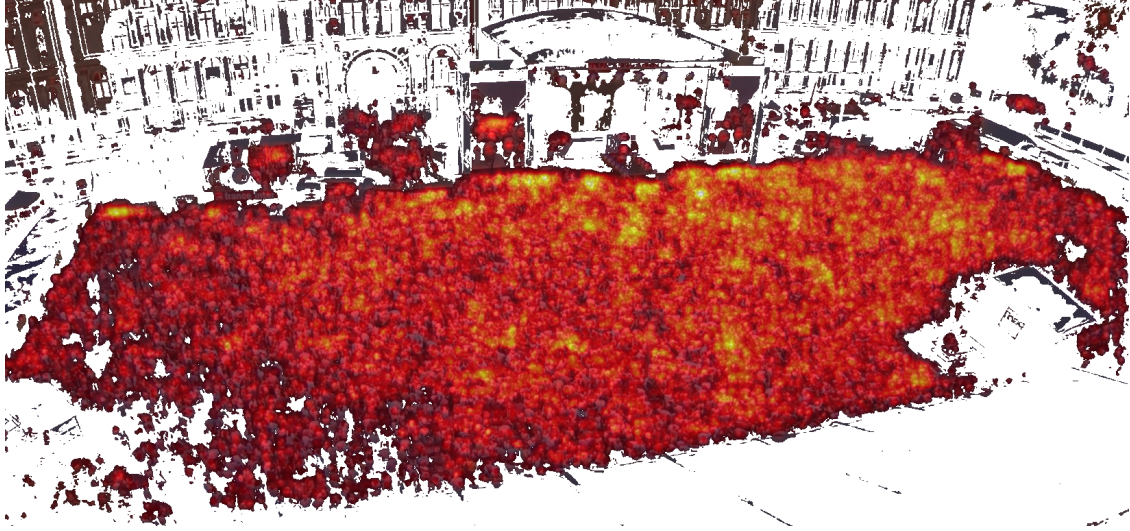


FIGURE 1.1 – Carte de densité mesurée par le LCPP lors d’un concert en extérieur dans Paris en 2019.

regroupements de personnes au sein de ces espaces (devant un point d’intérêt par exemple, comme ce fut le cas lors de l’incendie du Bazar de la Charité en 1887¹, ou bien devant une scène lors d’un concert comme montré en figure 1.1)

Un des rôles du LCPP aujourd’hui est de valider les scénarios utilisés pour les études d’ingénierie d’incendie, et d’intégrer les résultats dans l’appréciation du risque d’incendie. Pour la simulation d’évacuation, aucun guide ne cadre l’établissement de ces scénarios.

Les grands rassemblements. Les dispositions concernant les grands rassemblements répondent à des besoins différents : la réglementation ne s’étant pas développée par rapport à des questions de risques incendie, les considérations ne sont pas les mêmes. Le mémoire de PRV3 de la BSPP *La sécurité des « fanzones » et autres grands rassemblements* [GMQD20] montre les différentes phases de l’élaboration de la réglementation.

Après les attentats du 11 septembre 2001, la prise en compte de la menace terroriste va pousser à des mesures de plus en plus strictes. Les dispositifs mis en place par les responsables des sites deviennent plus importants et les aménagements du territoire plus complexes pour assurer la sûreté des événements [Vio13]. La création du code de sécurité intérieure de 2012 structure alors les textes jusque-là ajoutés au fil de l’eau pour les différents types de rassemblements.

Les attentats en France de 2015 et 2016 ont poussé à rajouter à cela de nouvelles dispositions pour l’encadrement de grands événements et à multiplier les dispositifs, le tout dans le but d’assurer la sûreté des sites en réduisant les sources de risques. Cette notion de *sûreté* (lutte contre les actions volontaires d’atteinte aux personnes) prend le pas sur la notion de *sécurité* (risques accidentels, dangers d’origine non intentionnelle). Si ces deux concepts peuvent paraître proches, les choix opérationnels qu’ils induisent peuvent être contradictoires. Un exemple notable est le dispositif de

1. Le départ de l’incendie fut causé par un non-respect de consignes de sécurité à propos du nouveau cinématographe présenté, objet qui a aussi causé un attroupement aggravant la désorganisation lors du déclenchement du feu [Hur97, p. 8-11].

sécurité mis en place au stade de France pour l'Euro 2016 : la fouille systématique à l'entrée a causé l'ouverture de seulement 4 portes d'entrée sur les 24 disponibles, ce qui a entraîné des bouchons importants devant les points de fouilles. En cas d'incendie ou de mouvement de panique à l'intérieur du stade, les 24 sorties que des mesures de *sécurité* avaient prévues n'auraient pas pu être utilisées à cause du dispositif de *sûreté*.

Ces derniers ajouts sont cependant de nature différente : le « choc de simplification administrative » de mars 2013 ayant suspendu l'ajout de nouvelles réglementations, les apports de la réglementation constituent essentiellement en la publication de nouveaux guides de recommandations d'échelles nationale, ou plus marginalement à l'échelle locale.

En parallèle, depuis 2004, des *fan zones* apparaissent en France dans le paysage urbain. Il s'agit d'espaces aménagés, originellement pensées comme zones de visionnage pour très grand public [Vio13, p. 189]. L'objectif initial d'avoir des espaces ouverts se heurte cependant au besoin d'assurer la sûreté mentionnée précédemment : la présence de contrôles des entrées et le fait de limiter les ouvertures de l'espace par exemple sont des mesures de sûreté qui peuvent aller à l'encontre des considérations sécuritaires. La variabilité des aménagements mis en place ajoute une autre difficulté au problème, en rendant cet équilibre non seulement difficile à trouver, mais aussi difficile à transcrire entre deux fan zones différentes.

Les limites, évolutions possibles et la place de la simulation. L'approche prescriptive du règlement français est critiquée sur certains aspects. Pour les ERP par exemple, l'article CO 34 impose des dégagements par étage et par effectif en fixant un nombre d'unités de passage (UP)² et un nombre de dégagements. Dans un hall d'entrée très passant, la contrainte stricte du nombre de dégagements peut conduire à la création d'une série de portes étroites telle que la somme de leurs UP corresponde au taux défini, alors qu'un nombre réduit de portes plus larges pourrait être plus bénéfique en cas d'évacuation. Dans le cas des grands rassemblements, c'est plutôt une absence de réponse à la variété des situations qui pose problème.

Une potentielle évolution à long terme réside dans l'apport de la loi ESSOC (État au service d'une société de confiance). Une des conséquences possible de cette loi serait de créer officiellement en France l'*ingénierie de l'évacuation* comme une discipline à part entière, et non plus comme une sous-partie d'une autre ingénierie. L'approche prescriptive serait complétée pour l'évacuation, en proposant la possibilité de prouver par la simulation numérique que le bâtiment aurait des performances comparables à ce que propose la réglementation (on parle alors de solution d'effet équivalent à la solution de référence).

La loi ESSOC concernant uniquement le CCH (et donc les bâtiments), cet apport ne serait pas applicable aux fan zones et grands rassemblements en extérieur. Dans ce cas, le besoin se fait sentir avant tout sur des réponses opérationnelles et des moyens de vérifier la sécurité d'une installation. Pour la Brigade de Sapeurs Pompiers de Paris (BSPP), les questions de sécurité se présentent d'une part lors de la prévention des risques sur la fan zone, mais aussi en cas d'incident sur l'accès des secours : si un camion de pompiers doit se frayer un chemin parmi une foule en déplacement, par où doit-on prévoir l'arrivée de ce camion pour faciliter son accès ? À l'heure actuelle, ces sujets sont traités de manière empirique. Les dispositifs mis en place sont testés et améliorés au fur et à mesure des événements. Le cas de l'Euro 2016 évoqué plus haut est un exemple du danger que peut

2. L'unité de passage (UP) sert à quantifier la largeur des dégagements. Une porte de 1UP aura une largeur 0.90m, 1.40m pour 2UP et à partir de 3UP chaque UP vaut 60cm.

entraîner cette démarche, et de la faible marge d'expérimentation que peuvent avoir les autorités. Le LCPP collabore avec la BSPP sur ces sujets et voit de plus l'arrivée des dossiers de sécurité de préparation aux Jeux Olympiques de 2024. Les questions de la prévention des risques et de la coordination en cas d'incident sont donc aussi des enjeux de discussion pour les années à venir.

La simulation de foule est une piste pour fournir des réponses à ces différentes questions. Dans le cadre de la construction, les acteurs privés n'ont pas attendu la loi ESSOC pour se munir d'outils commerciaux de simulations qui se sont développés ces dernières années. Pour les grands rassemblements, l'utilisation de simulation n'est pas fréquente, mais le besoin est présent.

C'est donc avec deux objectifs liés que s'est engagée cette thèse :

- préparer et cadrer l'arrivée possible de l'ingénierie de l'évacuation en comparant les différents modèles et leurs capacités,
- développer des outils adaptés aux problématiques complexes et variées que posent les fan zones et autres grands rassemblements.

1.2 Un bref regard sur la modélisation de foules

La modélisation de foules est une science jeune, mais déjà très riche. Nous allons ici faire une courte introduction de ses principes, *et nous présenterons dans le chapitre 2 les éléments précis de modélisation mathématique qui seront évoqués à titre de mise en contexte.*

Des premières études paraissent dès les années 70 sur les mouvements de foules et leur modélisation [Fru71 ; PM78]. Ces premiers travaux établissent des comparaisons entre des mouvements de foules et la mécanique des fluides. Les mouvements humains sont cependant très différents des systèmes mécaniques usuels en physique ([MF18]) :

- Les êtres vivants se meuvent en utilisant leur propre énergie interne, et non seulement à cause de forces externes.
- Chaque individu est capable de prendre une décision propre, et n'est donc pas soumis à une loi de déplacement.
- Le cône de vision propre à chaque individu implique une forme de dissymétrie dans les interactions et brise ainsi la loi de l'action et de la réaction que l'on a pour les systèmes particules.
- Chaque agent possède des caractéristiques physiques et psychologiques différentes, et ils ne sont donc pas interchangeables.

Les travaux de Helbing [HM95] sont les premiers à prendre en compte une part de ces aspects et ont eu une forte popularité au sein de la communauté scientifique. Helbing et ses co-auteurs décrivent un modèle assimilant les piétons à des disques dans un espace à deux dimensions soumis à des forces d'interaction entre elles. La somme de ces forces anime le mouvement de chaque particule x_i , de façon similaire à la seconde loi de Newton :

$$m_i \frac{d^2 x_i}{dt^2} = \frac{m_i}{\tau_i} \left(U_i - \frac{dx_i}{dt} \right) + \sum_j f_{ij},$$

où m_i est la masse d'un piéton, τ_i est un temps de réaction individuel, U_i est la vitesse que souhaite atteindre naturellement un individu seul et f_{ij} sont des *forces sociales*. Ces forces ont pour but de

prendre en compte des effets de comportements grégaires entre groupes ou bien d'évitement d'autres individus en vue d'éviter des contacts. Elles ne sont pas nécessairement symétriques et peuvent se baser sur un cône de vision individuel. Ce cadre permet aussi de donner des expressions plus générales aux forces sociales, pour prendre en compte des comportements humains plus complexes.

Modèles microscopiques et macroscopiques. De nombreux modèles ont depuis été développés dans la littérature. Différentes communautés scientifiques se sont intéressées à ces sujets et ont développé des approches très variées. On peut les catégoriser selon plusieurs critères, mais nous allons ici nous concentrer sur leur aspect microscopique ou macroscopique.

Les modèles microscopiques se concentrent sur une approche lagrangienne du mouvement, c'est-à-dire qui suit les individus. On retrouve ainsi le modèle de Helbing, mais aussi de nombreuses variations développées depuis celui-ci. D'autres modèles laissent de côté cette approche mécaniste et se concentrent sur le développement d'algorithmes visant à reproduire les systèmes de décision humains [MHT11 ; Rey+99]. Adopter un point de vue microscopique pour le modélisateur paraît assez naturel pour les mouvements de foules : en raisonnant avec des concepts psychologiques très simples, il est possible d'écrire des règles individuelles de mouvement, et de se poser la question de leur validité par rapport à sa propre expérience dans un premier temps. On peut de plus facilement prendre en compte les variations individuelles de comportement avec un jeu de paramètres propres à chaque agent.

Notons que si on se focalise sur le domaine académique et le développement de modèles, il existe aussi des solutions commerciales de simulations utilisées en ingénierie en France et dans le reste du monde [LRK20]. Ces modèles sont essentiellement de nature microscopique et peuvent se baser explicitement sur des modèles académiques (comme le module EVAC de FDS [KH09], utilisant le modèle de Helbing) mais développent aussi leurs propres modèles comme Pathfinder [Thu] ou EXODUS [GLG+15].

Les avantages d'une approche microscopique vont de pair avec des limites lors de la mise en application des simulations numériques. Tout d'abord, la volonté de faire un modèle « réaliste », ou prenant en compte la variété des comportements humains, conduit à une complexification rapide des modèles sous-jacents. L'ajout de nouveaux comportements dans le modèle de Helbing, par exemple, peut se faire en ajoutant un nouveau type de forces, lui-même dépendant d'un ou plusieurs nouveaux paramètres qu'il faudra déterminer selon la situation d'intérêt et dont il faudrait déterminer la variabilité au sein de la population humaine. S'il peut être envisagé de déterminer des caractéristiques physiques sur les populations (poids, taille, etc.), des paramètres de modélisation psychologiques ou sociologiques posent des questions bien plus profondes.

Une expérience numérique menée durant ces travaux de thèse sur le logiciel Pathfinder permet de mettre en évidence comment cette variabilité peut être difficile à prévoir. Nous avons simulé l'évacuation de gradins de stade en modifiant deux paramètres de fonctionnement interne du logiciel (le temps d'anticipation des collisions et le temps d'accélération, voir le tableau 1.1). Une variation d'un de ces deux paramètres entraîne une diminution d'environ 5% du temps d'évacuation. En revanche, en utilisant les deux valeurs extrêmes en même temps nous obtenons une baisse de 35% du temps d'évacuation. Cette utilisation certes caricaturale d'un logiciel précis montre l'importance que peuvent avoir certains paramètres de fonctionnement internes. Pour les organismes responsables de la prévention, cette variabilité rend difficile la lecture des résultats d'une simulation, car l'arrivée de chaque nouveau modèle nécessite une compréhension fine de ces phénomènes pour avoir un regard

Temps d'évacuation (s)		Temps d'anticipation de collision (s)	
		1, 5	0
Temps d'accélération (s)	1, 1	475	460
	0	465	305

TABLE 1.1 – Temps d'évacuation de gradins avec le logiciel *Pathfinder*

critique sur une simulation.

L'interprétation de résultats de simulations microscopiques pose aussi des problèmes. En admettant que le modèle de mouvement choisi est approprié à la situation que l'on cherche à simuler, que doit-on choisir comme conditions initiales ? Lorsque l'on souhaite simuler l'évacuation d'un stade, quelle valeur doit-on mettre dans les paramètres individuels décrivant les agents à évacuer ? Si le modèle de mouvement possède une variabilité stochastique, comment intégrer cette variabilité ? Des éléments de réponse ont été proposés dans le cadre d'études menées durant cette thèse avec un consortium d'entités impliquées dans ces questions en France dans [JDW+20].

Les modèles macroscopiques suivent une approche eulérienne du mouvement et raisonnent sur des densités, en analogie avec la mécanique des fluides. On peut retrouver essentiellement deux approches principales, définies selon la façon dont est structuré l'espace. Les premiers modèles de foules [Fru71] représentaient chaque bâtiment par un graphe connectant des pièces entre elles, avec des flux de déplacement entre chaque pièce³. Ce type d'approche a été plus récemment formalisé avec des outils mathématiques poussés [MFAB18]. Les autres modèles macroscopiques considèrent les bâtiments comme un espace à deux dimensions dans lequel se déplace une densité de piéton. Ces modèles s'écriront le plus souvent sous la forme d'une équation de conservation :

$$\partial_t \rho + \nabla \cdot (\rho \mathbf{v}) = 0,$$

où ρ est la densité de piétons et $\mathbf{v}$ est la vitesse effective. L'expression de v dépendra du modèle. Le modèle de Hughes [Hug02] est un exemple connu de ce type. La vitesse dépend d'une part de la densité locale, ainsi que d'un terme déterminant la direction vers la sortie la moins congestionnée. Ces équations ont connu un fort intérêt de la part de la communauté des mathématiques appliquées, le modèle de Hughes étant une équation aux dérivées partielles difficile à étudier d'un point de vue théorique ou numérique. D'autres modèles macroscopiques sont aussi étudiés activement [Rou11; TGD14], mais leur utilisation sur des cas réels est à l'heure actuelle moins fréquente.

Les modèles macroscopiques ont par essence une vision à grande échelle des mouvements et sont indiqués pour d'importantes quantités d'agents. Cela simplifie des questions de calculs numériques : un modèle calculant des forces de pression de collision peut demander d'importantes ressources de calcul lorsque le nombre d'individus atteint plusieurs milliers de personnes. Un modèle macroscopique pourra résoudre ces questions plus efficacement. Le prix à payer est une prise en compte moins fine des interactions individuelles, ainsi qu'une indistinguabilité des agents qui ne sont vus

3. À noter que, même aujourd'hui, des approches règlementaires comme dans le code de la construction en France adoptent un point de vue assez proche. Des modèles utilisés pour les gares et se basant sur des flux décrits par l'article CO 34 décrivent de façon indirecte une approche basée sur des graphes pour estimer le temps d'évacuation.

que comme un continuum. Une étude sur les modèles de contacts [MRSV11] montre comment pour des modèles microscopiques et macroscopique basés sur des principes similaires le flux au niveau de la porte peuvent se comporter très différemment. Dans ce cas-là le modèle macroscopique ne peut pas reproduire des effets de blocage au niveau d’une porte qui sont présents avec le microscopique.

Les modèles macroscopiques ont cependant un avantage important par rapport aux microscopiques : la difficulté de choisir des paramètres, des conditions initiales et un traitement statistique des résultats a une réponse plus directe. En s’intéressant à des effets moyens, on limite le nombre de paramètres nécessaires à la description du modèle.

Modèles granulaires. Un modèle microscopique nous intéresse en particulier pour la simulation d’évacuation. Il s’agit d’un modèle prenant en compte de façon explicite et directe les contacts entre individus, développé initialement pour des mouvements de particules dures ([MV11]). On qualifie ce modèle de « granulaire » du fait de sa vision des individus comme des grains de sables en contact. Ce modèle permet d’avoir une modélisation adaptée aux foules denses et qui soit sans paramètre. Les situations denses sont cruciales dans le cadre de la prévention en cas de mouvements de panique où on s’attend à avoir des bousculades.

Ce modèle microscopique représente les individus comme des disques durs de rayon fixe soumis à un mouvement spontané. Si deux disques sont en contact, il existe un ensemble de vitesses qu’ils peuvent adopter sans se chevaucher. Le modèle décrit alors la vitesse de ces disques comme étant la projection de leurs vitesses spontanées souhaitées sur cet ensemble de vitesses admissibles. Les contacts sont donc traités globalement pour un cluster de piétons. Une approche inertielle chercherait à évaluer des vitesses et des forces de collisions, mais lorsque le nombre de particules augmente les simulations numériques faites avec des modèles granulaires sont bien plus robustes. Ce modèle est particulièrement adapté pour reproduire des effets de bouchons qui peuvent apparaître devant une porte et ralentir le flux global (appelé *capacity drop* [Cep09]). On modélise ainsi le comportement d’individus paniqués cherchant à sortir à tout prix.

La thèse de F. Al Reda [Red17] présente une façon de prendre en compte des effets sociaux en plus des effets physiques. Dans ce nouveau modèle, on considère que lors d’un contact chaque individu choisit sa vitesse en fonction des personnes dans son cône de vision. Le modèle prend ainsi en compte une phase d’inhibition, où les individus vont adopter la vitesse la plus proche de leur vitesse souhaitée sans pour autant bousculer les personnes devant eux. On peut ainsi décrire le comportement d’individus polis, mais pressés et cherchant à évacuer le plus rapidement possible⁴. Des calculs numériques ont validé l’utilisation de ce modèle sur des expériences d’évacuation menées en laboratoire ([Red17, Chapitre 4]).

Il existe une version macroscopique du modèle granulaire ([Rou11]). Celle-ci décrit le mouvement d’une densité soumise à un mouvement décrit par un champ de vitesse souhaitée. On fixe pour cette densité un seuil maximal et les champs de vitesse admissibles sur les clusters de densité maximale sont ceux vérifiant cette fois une contrainte de non-concentration. La vitesse est aussi exprimée comme projection du champ de vitesse souhaitée sur ces vitesses admissibles.

Problématiques multi-échelle. Dans les cas d’application réels et de grande ampleur les problématiques rencontrées mettent en jeu de multiples échelles. On peut s’intéresser aux gradins

4. Pour les évacuations dans le cas d’incendies on pourrait s’attendre à ce que la panique s’empare des individus et incite à utiliser principalement le modèle granulaire, mais des études sur les comportements des rescapés du World Trade Center observent un comportement calme malgré la situation ([Kul16]).

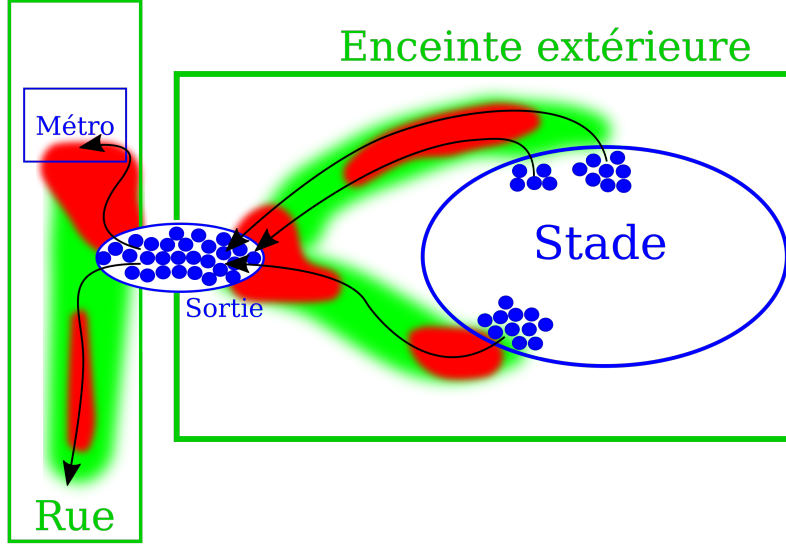


FIGURE 1.2 – Évacuation aux abords d’un stade sportif. Les points de congestion sont à traiter avec un modèle microscopique, tandis que les grands espaces sont à traiter avec un modèle macroscopique.

d’un stade en simulation microscopique, mais l’extérieur du stade sera plus propice à une modélisation macroscopique, et lorsque les individus s’approchent d’une clôture ou d’une station de métro adjacente la modélisation microscopique redevient nécessaire, comme schématisé en Figure 1.2.

Pour résoudre ce problème, nous avons fait le choix d’étudier des méthodes de *couplages aux interfaces* des modèles. Plus précisément, notre approche consiste à définir des interfaces entre les zones de chaque type où l’on pourra échanger entre un point de vue lagrangien ou eulérien. À l’heure actuelle les avantages et inconvénients de chaque approche sont connus, mais très peu de travaux cherchent à développer ces couplages pour la modélisation de mouvements de foules.

La difficulté principale dans cette approche du couplage réside dans la transmission de l’information et de la matière à l’interface : si un flux continu de piéton traverse une porte, la matière traverse l’interface toujours dans le même sens, mais si un ralentissement se produit en aval un bouchon peut se former et remonter le courant. L’information et la matière peuvent alors se transmettre dans des directions opposées. Pour les modèles granulaires qui traitent les contacts de manière globale, l’information se propage à vitesse infinie dans le cluster microscopique ou macroscopique et il faut donc traiter la notion de contact directement dans le couplage.

Nous présentons dans cette thèse des développements pour définir ces méthodes de couplage.

Le chapitre 3 porte sur le problème de couplage dans le cas de modèles à une dimension d’espace. Dans ce cadre simplifié nous regardons des modèles microscopiques et macroscopiques issus de l’étude du trafic routier et développons un couplage à l’interface passant du modèle microscopique au macroscopique, puis inversement. Les modèles d’inhibitions microscopiques et macroscopiques sont ensuite présentés, nous montrons comment ceux-ci peuvent être interprétés comme une limite

raide des modèles précédents, pour enfin adapter les procédures de couplage précédentes à ces modèles. Ces deux types de modèles et leur couplage sont comparés ensuite sur un cas applicatif.

Le chapitre 4 est dédié aux couplages entre le modèle d'inhibition microscopique et son équivalent macroscopique à une dimension. Nous écrivons le couplage à une interface pour ces modèles, en vérifiant que les contacts sont correctement pris en compte au niveau du changement de modèle et que l'information d'un bouchon se propage correctement.

Dans le chapitre 5, nous étudions enfin une version macroscopique du modèle d'inhibition. Seule la version microscopique de ce modèle existe aujourd'hui et repose sur l'existence d'un graphe hiérarchique des cônes de visions lors d'une évacuation. Nous introduisons des modèles numériques inspirés de cette notion de hiérarchie appliquée à une discrétisation cartésienne de l'espace.

Chapitre 2

Modélisation de mouvement de foules : état de l’art

2.1	Modèles microscopiques	12
2.2	Modèles macroscopiques et couplage d’échelles	20
2.3	Observables et validation des modèles	26

L’étude de l’évacuation de foules se développe à la fois sur les aspects expérimentaux et théoriques depuis plusieurs années. D’un point de vue empirique, des premières études s’intéressaient aux flux de passages au niveau d’une issue ou à la relation entre vitesse de déplacement et densité locale ([Fru71]). Les observations collectées depuis permettent d’obtenir des caractéristiques individuelles physiologiques et comportementales telles que la vitesse de marche libre par exemple ([Daa04]). D’autres études s’intéressent à des effets observables émergents pour des foules de fortes densités, comme l’apparition de turbulences ou d’ondes d’accélération et de décélération ([HJ13]). Certains de ces phénomènes ne sont cependant pas systématiquement observés : il a été mis en évidence dans la littérature qu’un obstacle en amont d’une issue pourrait augmenter le flux en cas d’évacuation ([HBJW05]), mais l’occurrence de cet effet n’est pas systématique comme peuvent le souligner des expériences ([GMP+18]).

De nombreux modèles sont développés en lien avec ces observations. On retrouve des approches de nature lagrangienne (ou microscopiques) se concentrant sur une modélisation des individus, ainsi que des modèles eulériens (ou macroscopiques) raisonnant sur un continuum d’agents. Les modèles microscopiques nécessitent de paramétrer individuellement chaque agent qui sera modélisé. La grande variabilité des agents à modéliser rend complexe l’interprétation que l’on peut faire des résultats de ces simulations. Ces modèles sont cependant les plus appropriés pour des simulations dans les cas où les interactions entre individus sont très fortes, comme en amont d’une porte ou d’un goulot d’étranglement pour une évacuation. Les modèles eulériens simulent un continuum d’agents et sont donc moins sensibles à la question de l’interprétation statistique. Ces modèles ne reproduisent pas à l’heure actuelle fidèlement la totalité des effets observés dans les évacuations de foules.

2.1 Modèles microscopiques

Nous présentons dans cette section quelques modèles microscopiques pour la modélisation de mouvements de foules. On retrouve dans cette catégorie de nombreux modèles académiques, mais aussi des logiciels commerciaux dont le rôle est de plus en plus important dans les études d'évacuation.

Vitesse souhaitée. La notion de vitesse souhaitée est transverse à plusieurs modèles de mouvements de foule. On peut la définir comme la vitesse qu'adopterait un individu seul. Selon le type de modèle considéré, elle peut prendre la forme d'un champ de vecteurs dans $\mathbb{R}^2$ unique ou dépendant des individus, et dont la norme peut être un paramètre. On supposera par la suite implicitement que l'on peut définir cette quantité dans chacun des modèles, car il s'agit rarement d'un point de difficulté mathématique. Notons toutefois que la forme de ce champ peut être une question ouverte dans certains cas : le contournement d'obstacles en particulier et le cheminement dans la géométrie des lieux peuvent être implicitement encodées dans ce champ. Selon les modèles cela peut avoir une importance plus ou moins grande et dépendre de la façon dont est géré le choix de la sortie empruntée par chaque agent, une question souvent séparée de la définition du modèle de mouvement (bien que dans le cadre des applications réelles, elle soit d'une importance majeure).

Forces sociales

Modèle de Helbing. Le modèle de Helbing ([HM95],[HFV00]) se base sur une loi de Newton pour des sphères dures représentant les individus. Pour N individus $1, \dots, N$ dont on note les positions $x_i(t) \in \mathbb{R}^2$ et $u_i = \dot{x}_i$, leur mouvement est déterminé par l'équation :

$$m_i \frac{du_i}{dt} = \frac{m_i}{\tau_i} (U_i - u_i) + \sum_{j \neq i} f_{ij}^c + \sum_{j \neq i} f_{ij}^{soc} + \eta_i, \quad (2.1)$$

où :

- m_i est la masse de l'individu i , τ_i un temps de relaxation individuel, et U_i sa vitesse souhaitée
- f_{ij}^c sont les forces physiques entre deux individus exercées sur i .
- f_{ij}^{soc} sont les forces de répulsion sociales entre individus exercées sur i .
- η_i est un terme de fluctuation aléatoire.

Nous avons ici omis les forces d'interaction avec les obstacles et murs, qui peuvent se traiter de façon similaire aux interactions avec les individus. Le premier terme traduit une relaxation vers la vitesse souhaitée en l'absence d'interaction. Les deux autres termes traduisent les effets d'interactions : les forces physiques sont des forces de contacts élastiques entre sphères, tandis que les forces sociales traduisent les tendances individuelles à rester à une certaine distance des autres individus.

Pour la suite on introduit la distance entre deux individus D_{ij} prenant en compte leurs rayons ainsi que le vecteur unitaire e_{ij} pointant de i vers j :

$$D_{ij} = |x_i - x_j| - r_i - r_j, \quad (2.2)$$

$$e_{ij} = \frac{x_j - x_i}{|x_i - x_j|}.$$

Le terme de force sociale peut prendre plusieurs formes, Helbing propose :

$$f_{ij}^{soc} = -F \exp(-D_{ij}/\delta) \left(\lambda + (1 - \lambda) \frac{1 + \cos(\varphi_{ij})}{2} \right) e_{ij}, \quad (2.3)$$

où :

- F est un paramètre fixant l'intensité de cette force, δ est un paramètre fixant la portée de cette force,
- φ_{ij} est l'angle entre u_i et e_{ij} ,
- λ quantifie l'anisotropie de cette force : pour $\lambda = 1$ on a une expression isotrope, et pour $\lambda < 1$ les agents présents de front ont plus d'influence que ceux derrière.

Notons ici que les paramètres F , δ et λ peuvent varier entre les individus. Cette force est l'outil principal permettant de reproduire une dynamique d'interaction humaine. On peut la comprendre comme un terme d'anticipation de collision, caractérisé par une distance δ . Dans le cas $\lambda < 1$, cette force viole le principe d'action-réaction entre les particules. Enfin bien que cette quantité soit dimensionnée comme une force, elle ne représente pas une interaction physique.

La force de contact est un terme de contact élastique et vaut :

$$f_{ij}^c = \kappa D_{ij}^- e_{ij}, \quad (2.4)$$

où κ est un coefficient de friction, qui peut être pris uniforme comme ici ou dépendre des individus. D_{ij}^- est la partie négative de D_{ij} , de sorte que cette force s'applique uniquement au contact direct. Contrairement à l'expression précédente, ce terme traduit donc une force physique. Le terme de fluctuations représente des variations individuelles et une forme d'indécision. Helbing [HFV02] ne donnait pas de forme précise à ce terme qu'il a introduit.

Ce modèle peut reproduire un très grand nombre d'effets observés dans la littérature : la force de friction κ permet d'observer des effets d'encombrements, les forces sociales anisotropes permettent de prendre en compte la vision des individus, etc. Pour une présentation de certains de ces aspects, voir [MF18].

Des ajouts au modèle de forces sociales. EVAC [KH09] est un module de simulation d'évacuation couplé au code de simulation de mécanique des fluides spécialisé dans l'incendie *Fire Dynamics Simulator*. Ce module se base sur le modèle de forces sociales de Helbing, mais ajoute aussi le résultat de travaux ultérieurs par d'autres équipes de recherches. La forme ellipsoïdale des êtres humains est approchée par trois disques, comme représenté sur la Figure 2.1

La distance D_{ij} doit être modifiée dans l'équation (2.2) pour être à la place la distance entre les deux cercles de i et j les plus proches, d'autre part les composantes rotationnelles des forces sociales et physiques doivent être introduites.

Un modèle de contre-flux tiré de [HKHE12] est implémenté pour améliorer la qualité de la simulation de flux de piétons se croisant et dans les cas de hautes densités. Le champ de vision est divisé en trois secteurs angulaires devant chaque individu, chaque secteur se voyant attribuer un score dépendant des obstacles et des vitesses des piétons présents. La vitesse effective est ensuite projetée sur le secteur maximisant ce score local de collisions. Le score attribué à chaque secteur permet aussi de prendre en compte les préférences de dépassement par la droite ou la gauche.

Les effets du feu simulé par FDS pour les incendies peuvent aussi être pris en compte. La concentration de fumée pénalise linéairement la vitesse souhaitée, et lorsque la dose de fumée atteint un certain seuil les individus sont considérés comme incapables.

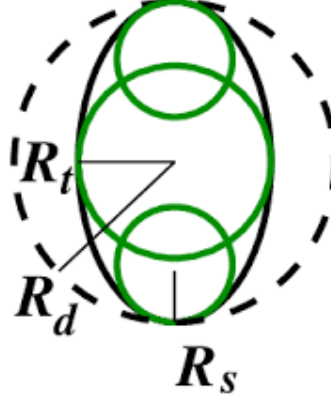


FIGURE 2.1 – Représentation d’un agent dans FDS+EVAC. La forme du corps humain est approximée par trois cercles superposés et de positions relatives fixes. Figure tirée de [KH09]

Automates cellulaires

Dans les modèles d’automates cellulaires ([BML92 ; SSNI95], ou dans le logiciel Exodus [GLG+15]), l’espace est discrétisé en un maillage cartésien avec un pas d’espace fixe Δx (typiquement de l’ordre de 40 ou 50 cm). Les individus sont représentés par des particules contraintes à occuper exactement une cellule de ce damier. Les mouvements sont alors une série de sauts stochastiques entre deux cases voisines à chaque pas de temps $k \in [0, T]$, comme une marche aléatoire. La définition du voisinage d’une cellule dépend du modèle, mais sera typiquement les cellules adjacentes et les cellules en diagonales lorsqu’on considère un voisinage dit de Moore.

Pour une salle carrée de taille $[L \times L]$ telle que $N\Delta x = L$ on note (i, j) les indices d’une case de la grille de centre $x_{ij} = (\Delta x(i - \frac{1}{2}), \Delta x(j - \frac{1}{2}))$, avec $(i, j) \in [1, N]^2$. La probabilité de transition d’une cellule à une autre va dépendre du *Static Floor Field* S_{ij} . Si on a une cible commune située en x_0 , on peut prendre :

$$S_{ij} = |x_{ij} - x_0|.$$

D’autres modèles prennent à la place la distance de Manhattan sur le graphe des liaisons entre cellules pour S_{ij} (comme EXODUS). À l’étape k on définit pour un individu les poids :

$$w_\gamma = \exp(-\kappa S_{ij+h}), \quad (2.5)$$

où κ contrôle l’intensité du biais attribué à S_{ij} et $h \in H$ l’ensemble des directions possibles, tel que :

$$H = \{(0, 0), (1, 0), (0, 1), (-1, 0), (0, -1)\}.$$

Dans le cadre du voisinage de Moore on ajoutera les directions diagonales à H . La probabilité de transition de chaque individu à l’instant k dans la direction h vaut :

$$p_h = \frac{w_h}{\sum_{h \in H} w_h}$$

Il peut arriver lors d'une transition que deux individus ou plus cherchent à se déplacer vers la cellule. On peut alors ajouter un paramètre de friction μ , tel qu'il y ait une probabilité μ qu'il n'y ait pas de vainqueur et une probabilité $1 - \mu$ que l'un des deux l'emporte (tiré aléatoirement).

Modèles d'anticipation de collisions

Les modèles présentés dans cette section ont une plus grande diversité de formalismes, mais reposent sur la supposition que les piétons sont capables d'adapter leur vitesse pour éviter les collisions avec des obstacles ou d'autres individus.

Follow-The-Leader (ou FTL). Cette famille de modèles ([GHR61]) regroupe plusieurs modèles décrivant le déplacement de N agents sur une section linéaire. On note les positions des individus $x_i(t)$ et on suppose qu'ils ont un diamètre commun $d > 0$. Les agents sont placés à des positions initiales x_i^0 telles que :

$$x_1^0 < x_2^0 < \dots < x_N^0, \quad (2.6)$$

$$x_{i+1}^0 - x_i^0 \geq d \quad \forall i < N, \quad (2.7)$$

c'est-à-dire que l'on suppose des agents ordonnés initialement et séparés par une distance d . Leurs vitesses sont décrites par la relation :

$$\dot{x}_i(t) = \varphi(x_{i+1}(t) - x_i(t)) \quad \forall i < N, \quad (2.8)$$

où φ est une fonction décrivant le comportement d'anticipation. Chaque individu ralentit lorsque la distance avec son voisin frontal diminue pour éviter la collision. φ sera donc croissante et vérifie typiquement :

$$\varphi(w) = 0 \quad \forall w \leq d, \quad \lim_{w \rightarrow \infty} \varphi(w) = U,$$

assurant une vitesse nulle de l'agent lorsque au contact avec son prédécesseur ainsi qu'une vitesse souhaitée U en marche libre. Plusieurs formes pour φ sont présentes dans la littérature, comme montré en Figure 2.2. On attribue enfin à l'entité N une fonction vitesse dépendante du temps $V(t)$ pour compléter le système d'équations (2.8).

On peut démontrer l'existence de solutions à ce système dans le cas où φ est aussi globalement Lipschitz par une application du théorème de Cauchy-Lipschitz (voir [MF18, Propositions 2.3-2.4]). L'hypothèse sur le caractère Lipschitz de φ donne aussi que les distances $x_{i+1} - x_i$ restent strictement supérieures à d , c'est-à-dire qu'il n'y aura pas de contact direct. Ce modèle peut aussi être étendu à une version de second ordre en temps, capable notamment de reproduire des phénomènes d'émergence d'instabilités créant des bouchons ([LJK+12]).

Modèles à deux dimensions. Plusieurs formulations prenant en compte des notions d'anticipation de collisions ont été proposées à deux dimensions. Ces modèles supposent que pour choisir une vitesse, chaque individu va évaluer en fonction de ses voisins et des obstacles présents une vitesse permettant d'éviter une collision dans un futur proche ([TP21]). Les premiers modèles à adopter et populariser cette approche ([PPD07; vMM08]) proposent différentes expressions de l'espace des vitesses *admissibles* (qui ne vont pas entraîner une collision).

Plusieurs auteurs proposeront par la suite des variantes dans les méthodes de résolutions et la détermination des vitesses inadmissibles ou du choix de la vitesse effectivement empruntée. Par

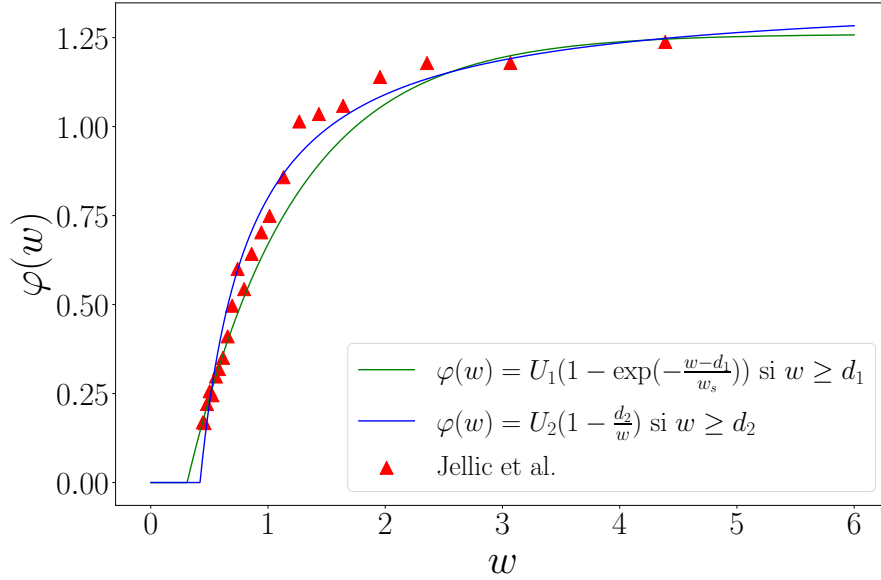


FIGURE 2.2 – Exemples de fonction vitesse, avec des données tirées de [JALP12]. Les fonctions sont tracées avec les valeurs obtenues par régression sur les données : $U_1 = 1.26 \text{ m.s}^{-1}$, $d_1 = 0.31 \text{ m}$, $w_s = 0.91 \text{ m}$, $U_2 = 1.38 \text{ m.s}^{-1}$ et $d_2 = 0.42 \text{ m}$.

exemple, Moussaid et al. ([MHT11]) décrivent un modèle pour le choix de la direction instantanée empruntée par un piéton en fonction de son angle de vision. Le modèle suppose que chaque piéton peut observer les obstacles bloquant son trajet et va choisir la direction qui minimise l'écart par rapport son objectif initial. Une autre heuristique assure qu'un piéton gardera une distance de sécurité avec les autres piétons.

Le logiciel de simulation Pathfinder ([Thu]) propose un modèle basé sur des algorithmes décrits dans [Rey+99]. Le modèle initial consiste à calculer pour chaque agent une vitesse souhaitée dépendante de la géométrie, et dont la valeur maximale est fixée par la distance avec les agents les plus proches. L'anticipation des collisions se fait ensuite localement sur la direction de la vitesse souhaitée. Le cône de vision individuel est séparé en un ensemble de directions et sur chacune est calculée une série de coûts pour la direction. Ces coûts vont dépendre de la présence d'obstacles, d'individus, de l'inclinaison de la direction, etc. La direction de la vitesse sera celle qui minimise la combinaison de ces coûts.

Modèle de Vicsek. Une branche de l'étude des mouvements d'entités actives s'intéresse aux mouvements collectifs, c'est-à-dire à des mouvements ordonnés émergents spontanément dans certains systèmes d'agents en mouvements¹. Le modèle de Vicsek ([VCB+95]) permet d'expliquer plusieurs phénomènes de mouvements collectifs avec très peu d'ingrédients.

Sous sa forme la plus simple, ce modèle décrit le mouvement de N entités itérativement avec

1. Les systèmes en question vont bien au-delà des piétons et on retrouve des mouvements collectifs chez les oiseaux ([LLE10]) ou les bancs de poissons ([DC97]).

un pas de temps Δt . Pour chaque particule i , de position $x_i(t)$ et de vitesse $v_i(t)$ d'amplitude fixée commune et d'angle $\theta_i(t)$, les équations du mouvement sont données par :

$$\begin{aligned} x_i(t + \Delta t) &= x_i(t) + \Delta t v_i(t), \\ \theta_i(t + \Delta t) &= \langle \theta_j \rangle_{j \in S_i} + \eta_i(t), \end{aligned}$$

où $\langle \theta_j \rangle_{j \in S_i}$ désigne la moyenne des θ_j pour tout j dans une sphère centrée sur $x_i(t)$ et de rayon r , et $\eta_i(t)$ est un terme de bruit tiré uniformément sur un intervalle d'amplitude $2\pi\eta$. Chaque agent du modèle va chercher à aligner sa trajectoire en fonction des individus dans son voisinage avec une perturbation aléatoire. Le modèle a trois paramètres principaux : η , contrôlant l'amplitude du bruit, la densité moyenne d'individus dans la simulation et la célérité des individus. Des simulations ont montré que des variations sur la valeur de la densité de particules et l'amplitude du bruit pouvaient faire émerger des comportements collectifs où les agents s'alignent dans des directions similaires ([CSV97; VCB+95]). Ces transitions de phase peuvent être continues ou discontinues selon les valeurs de v_0 et plusieurs variantes du modèles ont été proposées par la suite ([CGGR08]), ainsi que des versions continues ([SSG+06; DM08]).

Modèles granulaire et d'inhibition sociale

Modèle granulaire. Le modèle granulaire ([MV11]) représente N individus comme des disques durs de rayon r . On note q le vecteur contenant leurs positions

$$q = (q_1, q_2, \dots, q_N) \in \mathbb{R}^{2N}.$$

On peut aussi définir la distance entre deux individus en prenant en compte leur rayon :

$$D_{ij} = |q_j - q_i| - 2r.$$

La contrainte de disques durs implique que les individus ne peuvent se chevaucher, on peut donc définir un ensemble de configurations admissibles pour les individus :

$$K = \{q \in \mathbb{R}^{2N}, D_{ij} \geq 0\}.$$

Dans le cadre de ce modèle les individus cherchent à suivre leur vitesse souhaitée, mais leurs positions doivent rester dans l'ensemble des configurations admissibles K . En l'absence de contact (c'est-à-dire $D_{ij} \geq 0 \forall i, j < N$), chaque individu peut adopter sa vitesse souhaitée U_i . Lorsqu'il y a un contact, il existe un ensemble de vitesses qui ne brisent pas la contrainte de non-recouvrement pour les individus concernés. Cet ensemble de vitesse admissible est :

$$\mathcal{C}_q = \{v \in \mathbb{R}^{2N}, D_{ij} = 0 \implies e_{ij} \cdot (v_j - v_i) \geq 0\}, \quad (2.9)$$

où e_{ij} est un vecteur normal entre les centres de deux disques :

$$e_{ij} = \frac{q_j - q_i}{|q_i - q_j|}.$$

Le modèle décrit alors que la vitesse effective est la plus proche de la vitesse souhaitée parmi cet ensemble de vitesses admissibles² :

$$\dot{q} = P_{C_q} U,$$

où P_{C_q} est la projection euclidienne sur l'ensemble défini dans l'équation (2.9). Cette projection est bien définie, mais la fonction $q \mapsto P_{C_q} U$ est non lisse et empêche donc d'appliquer le théorème de Cauchy-Lipschitz. La démonstration du caractère bien posé de ce modèle repose sur une approche dite de *catching-up*. Cette approche a été introduite par Moreau ([Mor77]) dans le cas d'un point q contraint à rester dans un espace convexe mouvant $K(t)$. Moreau montre qu'il est possible de construire une suite d'approximations de la solution q^k à des temps $k\tau$, $k \in \mathbb{N}$, $\tau > 0$ en projetant q^k sur l'espace des contraintes à l'instant q^{k+1} . Pour appliquer cette approche au modèle de mouvement de foules, il faut d'une part prendre en compte la vitesse souhaitée U , mais aussi que l'espace de contrainte K n'est pas convexe en général. Cet espace est cependant prox-régulier, c'est-à-dire que la projection sur celui-ci est bien définie pour des points suffisamment proches de K [MV11].

Cette notion de projections successives permet aussi de construire un schéma numérique naturel pour la discrétisation de ce problème. En développant au premier ordre pour un pas de temps τ l'expression de $D_{ij}(q(t+\tau))$ on peut exprimer l'étape de projection à chaque pas de temps comme la résolution d'un problème d'optimisation sous contraintes quadratiques. Il s'agit d'une approche similaire aux schémas numériques d'advection-projection. La librairie python *cromosim* ([FM]) propose une implémentation de cet algorithme.

Dans cette formulation le contact entre deux disques est symétrique et respecte donc la loi d'action-réaction. Cette interaction n'est pas inertielle : l'équation du mouvement est de premier ordre en temps. On ne traite pas de collisions, mais bien de contacts directs. Cette approche permet de retrouver des effets parfois observés pour les évacuations de foules, comme l'effet fluidifiant d'un obstacle (voir Figure 2.3, et la section 2.3 pour une discussion de cet effet.)

L'étape de projection est un traitement purement physique des contacts, sans prise en compte de processus de décision individuel. Ce modèle traduit une vision très simplifiée du comportement humain : chaque individu est égoïste et ne prend pas en compte les autres individus. Ce modèle est donc plus adapté pour l'évacuation en situation d'urgence où on s'attend à de fortes densités.

Dans le modèle tel que présenté ici la masse individuelle n'est pas prise en compte directement, il est cependant possible de le faire en remplaçant la projection euclidienne simple par une projection pondérée par les masses individuelles.

Modèle d'inhibition. Il est proposé dans la thèse de F. Al Reda ([Red17]) d'enrichir le modèle précédent avec une étape supplémentaire. Cette étape modélise une forme d'inhibition individuelle : dans une situation de forte densité, chacun souhaite à se rapprocher de sa vitesse souhaitée sans pour autant bousculer les autres individus devant soi. Une projection granulaire est ensuite effectuée pour prendre en compte les recouvrements qui pourraient toujours advenir.

Plus précisément, on introduit un graphe d'influence entre le groupe des individus en contact. Il s'agit d'un graphe orienté où un individu j pointerait vers un individu i si j est dans le cône de

2. La notion de vitesse admissible a ici un sens différent de celui qu'elle peut avoir pour les modèles basés sur des évitements de collision : on ne parle pas ici de vitesses qui évitent la collision, mais d'une vitesse effective en cas de contact direct.

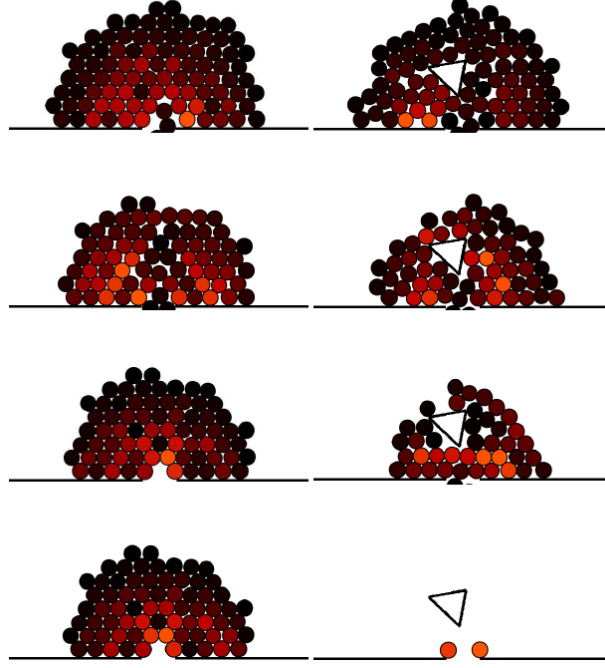


FIGURE 2.3 – Simulations du modèle granulaire, tiré de [MF18] et illustrant la capacité du modèle à reproduire l’effet fluidifiant d’un obstacle.

vision de i , que l’on note $j \in V_i$. Ce graphe est supposé acyclique dans les situations d’évacuation, où les individus cherchent à atteindre un même but. Dans le cas d’un graphe acyclique, l’algorithme de tri topologique permet de trier les individus de façon à ce que $j \in V_i \implies j > i$.

L’étape d’inhibition sociale va alors consister à parcourir la hiérarchie des individus et de déterminer une vitesse inhibée $\bar{u}$ à partir de la vitesse souhaitée initiale U . Cette étape consiste alors à résoudre une suite de problème de minimisation :

$$\bar{u}_i = P_{\mathcal{C}_i(q, \bar{u}_{-i})} U_i, \quad (2.10)$$

où $P_{\mathcal{C}_i(q, \bar{u}_{-i})}$ désigne la projection pour la distance euclidienne sur l’espace $\mathcal{C}_i(q, \bar{u}_{-i})$, $\bar{u}_{-i}$ désigne le vecteur $(\bar{u}_1, \dots, \bar{u}_{i-1}, \bar{u}_{i+1}, \dots, \bar{u}_N)$ et l’espace des contraintes $\mathcal{C}_i(q, \bar{u}_{-i})$ est défini par :

$$\mathcal{C}_i(q, \bar{u}_{-i}) = \{v \in \mathbb{R}^2, \forall j \in V_i, D_{ij}(q) = 0 \implies e_{ij} \cdot (v - \bar{u}_j) \leq 0\} \quad (2.11)$$

L’arrangement hiérarchique des individus permet notamment de pouvoir résoudre ces problèmes en commençant par l’individu avec l’indice le plus élevé et en continuant de manière décroissante. Le champ de vitesse résultant peut en revanche tout de même donner lieu à des chevauchements entre les individus. Le champ de vitesses doit donc être projeté globalement de la même façon que le modèle granulaire, de façon à ce que la vitesse effective vérifie alors :

$$\dot{q}(t) = P_{\mathcal{C}_{q(t)}} \bar{u}. \quad (2.12)$$

Ce modèle permet de prendre en compte des effets de comportement humain dans la modélisation, alors que le modèle granulaire seul se limite aux phénomènes physiques de contacts. Ainsi on

peut voir le modèle granulaire comme décrivant le comportement de piétons égoïstes, alors que ce modèle d’inhibition modélise celui de piétons plus polis.

Dans la thèse de F. Al Reda ([Red17]) ces deux modèles sont comparés à une série d’expériences sur le flux au niveau d’une porte en fonction de la compétitivité des agents ([GPP+16]). Dans ces expériences, on demande à des participants de passer à travers une sortie étroite le plus vite possible pour simuler une évacuation. Les expérimentateurs demandent à une partie des individus d’adopter un comportement égoïste, en les autorisant à pousser pour se frayer un chemin. Tous les autres participants ont pour instruction de ne pas bousculer pour avancer. Les modèles granulaires et d’inhibition retrouvent quantitativement les comportements de ces flux instantanés pour les individus respectivement égoïstes et polis. Ces résultats permettent de montrer l’intérêt de chacun de ces modèles pour représenter des comportements humains différents et difficile à reproduire avec des modèles inertiels.

2.2 Modèles macroscopiques et couplage d’échelles

Des modèles macroscopiques ont été aussi développés pour l’étude des mouvements de foules. Des observations expérimentales sur la relation entre vitesse et densité pour les piétons ([Fru71]) ont menés des premiers auteurs à considérer des modèles raisonnant sur des quantités macroscopiques. La communauté mathématique s’est depuis emparée de ce sujet et de très nombreux travaux sont menés et font appel à des notions de mathématiques très avancées. Certains de ces modèles ont de plus été couplés afin d’étudier les mouvements de foules à différentes échelles. Contrairement aux modèles microscopiques, les modèles macroscopiques ne sont pas utilisés à l’heure actuelle dans les simulations pour l’ingénierie d’incendie.

Modèles à compartiments.

Des modèles d’évacuations basés sur une organisation en réseau d’un bâtiment peuvent se retrouver dès les années 70 ([Fru71]). Ces modèles consistent à décrire un réseau de flux et de connexions entre les pièces d’un bâtiment et représenter l’évacuation par un écoulement entre les nœuds de ce graphe. Le graphe doit être déterminé à partir de la structure à évacuer, comme montré en Figure 2.4. Un modèle minimaliste de ce type est introduit dans [MFAB18] et se base sur la donnée d’une capacité pour chaque porte, représentant le flux maximal pouvant la traverser. Le flux effectif est ensuite donné comme le flux le plus proche de cette capacité en fonction du nombre de personnes disponibles.

L’attrait de ce type de modèle est dans sa simplicité d’explication ainsi que son nombre très réduit de paramètres : l’expression d’un flux de passage pour une porte ou un point d’intérêt et d’une vitesse de cheminement entre deux salles sont suffisantes pour décrire une évacuation d’un bâtiment tracé. Un défaut important de ce type de modèle réside dans leur sensibilité au choix initial du graphe. Pour des bâtiments avec un grand nombre de cloisons, il est parfois simple de donner un graphe représentant les cheminements, mais pour des zones larges et ouvertes par exemple la traduction en graphe peut parfois être moins pertinente.

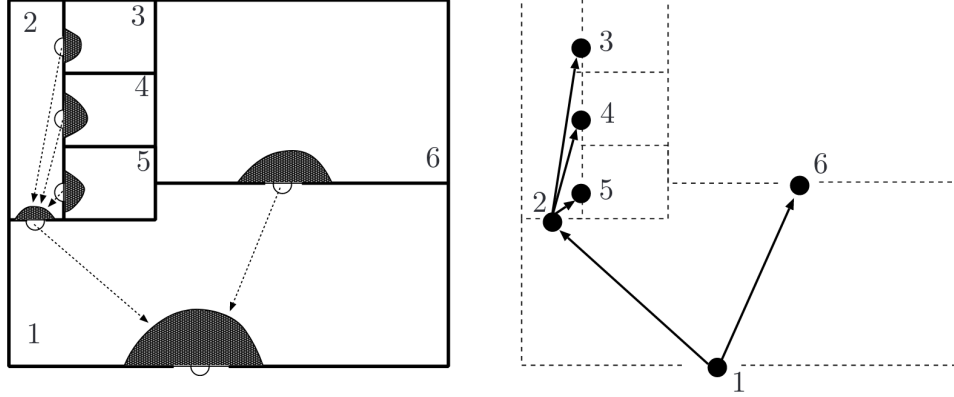


FIGURE 2.4 – Exemple de graphe représentant une évacuation. Tiré de [MF18]. Les chiffres représentent les indices des noeuds du graphe utilisé pour représenter l'évacuation.

Modèle de Lighthill-Whithams-Richards

Le modèle de Lighthill-Whithams-Richards (ou LWR) a été introduit initialement pour l'étude du trafic routier ([LW55]). Les agents sont décrits par une densité linéique $\rho(t, x)$ dont l'évolution est donnée par une équation de conservation où la vitesse d'advection dépend de la densité locale :

$$\partial_t \rho + \partial_x (\rho v(\rho)) = 0. \quad (2.13)$$

La fonction v est décroissante en ρ , pour préserver, dans le cas de véhicules, une distance de sécurité d'autant plus grande que la distance au véhicule qui précède est importante. Plusieurs formes de fonctions existent dans la littérature (voir section 2.3)

Ce modèle peut être vu comme une limite macroscopique du modèle FTL. En associant à la distance inter-agent la densité via $\rho = 1/w$ et en faisant tendre le nombre d'agents vers l'infini et leur taille vers 0 dans le modèle FTL on peut retrouver dans un certain sens des solutions au modèle FTL (voir [DR15], ou [HR17] qui montre que le modèle FTL peut être vu comme une discrétisation particulière du modèle LWR).

De nombreuses variations de ce modèle ont été proposées au cours des dernières décennies. Le modèle de Aw, Rascle ([AR00]) et Zhang ([Zha02]) est une extension de second ordre en temps, particulièrement utilisée dans la modélisation du trafic routier. Une série de travaux à porté sur la modélisation de contraintes sur le flux en un point :

$$\begin{aligned} \partial_t(\rho) + \partial_x(f(\rho)) &= 0, \\ f(\rho(t, 0)) &\leq q(t), \end{aligned}$$

où $q(t)$ est le flux maximal autorisé au point 0. Le caractère bien posé a été étudié dans [CG07], et un schéma numérique a été proposé dans [AGS10]. D'autres extensions ont ensuite permis de considérer des contraintes de flux non locales ([ADR14; AS20]) modélisant des effets similaires au *faster-is-slower*.

L'étude du trafic routier a enfin mené à étendre le modèle LWR sur des réseaux routiers, comportant des jonctions ([CP02]). L'espace modélisé est constitué de routes reliées entre elles au niveau de jonctions. Sur chaque route l'évolution de la densité est donnée par une équation du type LWR, tandis que pour les jonctions des conditions sont ajoutées pour décrire les échanges de flux. Une première condition est la conservation du flux au niveau d'une jonction. En s'intéressant à une unique jonction avec n routes entrantes et m routes sortantes et située en 0, cette condition s'écrit :

$$\sum_{i=1}^n f_i(\rho_i(t, 0-)) = \sum_{j=n+1}^{n+m} f_j(\rho_j(t, 0+)),$$

où nous avons pris la convention de noter par des indices $i = 1, \dots, n$ les grandeurs associées aux routes entrantes et $j = n+1, \dots, n+m$ celles des routes sortantes. Cette condition de conservation de la matière n'est cependant pas suffisante pour définir des solutions à ce problème. Le comportement de la jonction est usuellement décrit en choisissant un solveur de Riemann pour la jonction. Il en existe plusieurs dans la littérature ([BČG+14]), basés sur des phénomènes observés ou des considérations mathématiques.

Modèle de Hughes

Le modèle de Hughes ([Hug02]) est basé sur une équation de conservation à deux dimensions en espace. La norme de la vitesse est traitée de la même façon que pour le modèle LWR avec une dépendance en fonction de la densité locale pour traiter la congestion.

Hughes introduit aussi une modélisation sophistiquée de la direction souhaitée pour un modèle macroscopique. On suppose que tous les agents ont un même but commun, qui est une région Γ de l'espace. Chaque agent cherche alors à rejoindre Γ le plus rapidement possible, et l'on peut définir $T(x)$ comme étant le temps du trajet le plus court entre le point x et la zone Γ . Il est possible de montrer ([Set99, Partie 2]) que $T(x)$ vérifie alors une équation dite eikonale :

$$|\nabla T(x)| = \frac{1}{c(x)}, \quad (2.14)$$

où $c(x)$ est la vitesse de déplacement locale. On peut alors définir un vecteur unitaire μ qui va donner la direction vers la sortie la plus proche :

$$\mu = \frac{-\nabla T(x)}{|\nabla T(x)|}.$$

Afin de prendre en compte l'effet de la congestion dans cette estimation de trajet, Hughes propose donc que la norme de la vitesse dépende de ρ localement, et que sa direction soit donnée par une équation eikonale prenant en compte la réduction de vitesse due à la congestion :

$$\partial_t \rho + \nabla \cdot (\rho v(\rho) \mu) = 0, \quad (2.15)$$

$$\mu = \frac{-\nabla T(x)}{|\nabla T(x)|}, \quad (2.16)$$

$$|\nabla T(x)| = \frac{1}{v(\rho(x))}. \quad (2.17)$$

Avec cette prise en compte de la congestion, les individus peuvent éviter des zones de fortes densités. Dans le cadre de géométries avec plusieurs sorties une partie de la foule peut adapter sa stratégie et se diriger des zones moins encombrées (voir [TGD14] pour une implémentation numérique du modèle et des cas où l'on retrouve ce comportement).

Limites hydrodynamiques

La variété de modèles microscopique présents dans la littérature a conduit au développement de nombreux modèles macroscopiques obtenus par passage à la limite de modèles microscopiques. Cette procédure de passage à la limite peut être très complexe à obtenir et nécessite de pouvoir écrire un modèle mésoscopique intermédiaire ([Deg19]). Le modèle *Self-Organised Hydrodynamics* (SOH) fait partie de cette catégorie, étant obtenu par passage à la limite du modèle de Vicsek ([DM08] et [DH13] pour des méthodes numériques de simulation). Ce modèle est décrit par le système d'équations :

$$\partial_t \rho + \partial_x (c_1 \rho \omega) = 0, \quad (2.18)$$

$$\rho (\partial_t \omega + c_2 (\omega \cdot \nabla) \omega) + \lambda P_{\omega^\perp} (\nabla p(\rho)) = 0, \quad (2.19)$$

$$|\omega| = 1, \quad (2.20)$$

où :

- c_1 et c_2 sont deux célérités advectives du modèle et λ est un paramètre du modèle,
- $P_{\omega^\perp}$ désigne la projection sur le plan orthogonal à ω ,
- $p(\rho)$ est un terme de pression.

L'équation (2.18) traduit la conservation de la matière et l'équation (2.19) la conservation du moment. Ce système ressemble aux équations d'Euler dans le cas isentropique, mais ici la vitesse ω est de module fixé (d'où la présence du terme de projection qui permet de satisfaire cette contrainte) et les célérités c_1 et c_2 peuvent avoir des valeurs différentes. Avec plusieurs formes de $p(\rho)$ il est possible de prendre en compte différents comportements de la foule, par exemple une aversion de la congestion.

Modèle granulaire macroscopique

L'analogie macroscopique du modèle granulaire présenté en section 2.1 a été introduit durant la thèse d'Aude Roudneff-Chupin ([Rou11]). La contrainte de non-recouvrement pour des disques est traduite dans ce cas par une densité maximale ρ_{max} autorisée. On se place dans Ω un ouvert de $\mathbb{R}^2$, et on note $U(x)$ la vitesse souhaitée en tout point $x \in \Omega$ des individus, le modèle s'écrit alors :

$$\begin{cases} \partial_t \rho + \nabla \cdot (\rho u) = 0, \\ u = P_{C_\rho} U, \end{cases} \quad (2.21)$$

où P_{C_ρ} désigne la projection pour la norme L^2 sur l'espace des vitesses admissibles C_ρ , c'est-à-dire les vitesses dont la divergence est positive sur les zones où $\rho = \rho_{max}$. Cela signifie que l'on se restreint à des champs de vitesse non concentrants lorsque l'on atteint la densité maximale. Pour formaliser cette idée il est nécessaire d'introduire une formulation duale de cet espace C_ρ :

$$C_\rho = \left\{ v \in L^2(\Omega), \int_\Omega v \cdot \nabla p \leq 0, p \geq 0 \text{ p.p.}, \forall p \in H_\rho^1(\Omega) \right\}, \quad (2.22)$$

où $H_\rho^1(\Omega)$ est l'ensemble des fonctions test défini par :

$$H_\rho^1(\Omega) = \{ p \in H^1(\Omega), p(\rho_{max} - \rho) = 0 \text{ p.p.} \}. \quad (2.23)$$

Ce modèle et le modèle granulaire microscopique de la section 2.1 donnent des résultats qualitativement très différents pour le flux à une porte. Les arches qui peuvent se créer devant une sortie dans le modèle microscopiques vont ralentir, ou même bloquer le flux. Pour le modèle macroscopique la contrainte de densité va modifier l'allure du champ de vitesses en amont, mais le modèle ne va pas reproduire ces effets d'arche. Le flux ne sera donc pas stoppé à l'ouverture (voir [MRSV11] pour une discussion développée).

Jeux à champ moyen

Les jeux à champ moyens sont une classe de modèles introduits indépendamment par Lasry et Lions ([LL06a ; LL06b ; LL07]) et Caines, Huang et Malhamé [HMC06].

Ils s'intéressent à la modélisation d'un continuum d'agents rationnels cherchant à maximiser une certaine fonction dépendant de l'ensemble des agents. Chaque agent est ainsi capable d'établir un raisonnement, d'estimer l'évolution des autres et d'établir sa stratégie en fonction (avec des termes stochastiques pour prendre en compte des incertitudes individuelles). Un exemple de système de jeu à champ moyen appliqué à l'évacuation d'une foule sur un intervalle de temps $[0, T]$ et sur le tore $\mathbb{T}^d$ peut s'écrire :

$$-\partial_t u - \Delta u + H(x, m(t, x), \nabla u) = 0, \quad (2.24)$$

$$\partial_t m - \Delta m - \nabla \cdot \left(\frac{\partial H}{\partial p}(x, m(t, x), \nabla u) m \right) = 0, \quad (2.25)$$

$$u(T, x) = \phi(x, m(T, x)), \quad m(0, x) = m_0(x), \quad (2.26)$$

où :

- $u(t, x)$ est la fonction valeur pour l'équation de Hamilton-Jacobi (2.24). Cette quantité peut être vue comme un analogue de temps de trajet à la sortie prenant en compte la densité des agents dans le modèle de Hughes.
- $m(t, x)$ est la densité des agents, évoluant selon l'équation de Fokker-Planck (2.25).
- $H(x, m, p)$ est le Hamiltonien du système, une fonction décrivant l'évolution des agents et pouvant prendre en compte des pénalisations de vitesse en cas de congestion par exemple.
- $\phi(x, m)$ et $m_0(x)$ sont des conditions terminales et initiales du système.

Il est important de noter que l'équation (2.24) est dite « backward » en temps, c'est-à-dire évoluant dans le sens contraire au temps alors que l'équation (2.25) est « forward » en temps. Ces équations font l'objet de nombreux travaux théoriques et leur simulation numérique nécessite de faire appel à des outils numériques complexes afin de traiter ce caractère forward-backward ([AL20]).

Couplages d'échelle

Le constat des limites des approches microscopiques et macroscopiques amène assez naturellement à une approche combinée, au développement de *couplages*. Le principe d'une approche couplée consiste à utiliser plusieurs modèles ou plusieurs échelles pour une étude de modélisation. Cette idée a déjà été abordée dans la littérature sur les mouvements de foules et du trafic routier, montrant avant tout que le principe de couplage de modèles peut se voir de plusieurs façons très différentes.

Les modèles de Lighthill-Whithams-Richards [LW55 ; Ric56] (ou LWR) et du Follow-The-Leader (ou FTL) et leurs très nombreuses variantes constituent des blocs de base pour de nombreux couplages destinés à l'étude du trafic routier. De nombreux articles développent des méthodes pour étudier le mouvement d'un agent ponctuel dans une densité ([BCL+18]) avec une influence sur l'écoulement produit ([DG14b ; DG14a]). Les équations étudiées peuvent prendre la forme :

$$\begin{cases} \partial_t \rho(t, x) + \partial_x [f(\rho(t, x), x - y(t))] = 0, \\ \dot{y}(t) = w(\rho(t, y)), \end{cases} \quad (2.27)$$

où $f(\rho, x)$ est une fonction de flux, y est la position d'un véhicule particulier (par exemple un camion), et w une fonction de vitesse décrivant le mouvement de ce véhicule. Le couplage permet de modéliser l'influence d'un élément ponctuel sur la densité globale.

Une autre approche développée par Colombo et Marcellini dans [CM15] consiste à alterner des représentations microscopiques et macroscopiques avec des interfaces mouvantes. Plus précisément, on regarde les équations :

$$\begin{cases} \partial_t \rho + \partial_x (f(\rho)) = 0, & x \in]-\infty, x_1(t)[\cup]x_n(t), +\infty[\\ \dot{x}_i(t) = v\left(\frac{1}{x_{i+1} - x_i}\right), & \forall i \in \{1, \dots, n-1\} \\ \dot{x}_n = v(\rho(t, x_n)), \end{cases} \quad (2.28)$$

où $x_1, \dots, x_n$ sont les positions d'individus microscopiques, et v est une fonction donnant la vitesse en fonction de la densité locale. Le piéton leader x_n et le piéton en fin de file x_1 sont les interfaces entre les modèles microscopiques.

Des travaux plus anciens comme ceux de la thèse de E. Bourrel se sont intéressés à la définition d'un changement d'échelle : au lieu d'utiliser une représentation mixte des deux modèles on s'intéresse à passer de l'un à l'autre pour des portions choisies de route [BH02 ; BL03]. Ce type de couplage par interface a ensuite été abordé d'un point de vue théorique pour une variante du modèle FTL de second ordre en temps avec un modèle LWR de second ordre en temps [LP10] et de premier ordre en temps [GP17].

Ceux-ci se placent dans le cadre du couplage micro vers macro à une interface fixe en $x = 0$. Le principe consiste alors à attribuer à chaque modèle de part et d'autre de l'interface des conditions de bord appropriées. Pour le modèle FTL il va s'agir d'une vitesse du leader, tandis que pour le modèle LWR pour un ordre 1 un flux à l'interface sera nécessaire. Ces conditions au bord ne seront valables que jusqu'à ce que le leader du modèle FTL atteigne l'interface, où il doit devenir macroscopique.

Dans [GP17], les auteurs proposent ainsi qu'une solution au problème couplé doit être composée :

- d'une fonction $l(t)$ donnant l'indice du leader microscopique à l'instant t . Cette fonction est continue par morceau et croissante, ses points de discontinuité étant les temps pour lesquels un leader atteint l'interface.
- d'un ensemble de fonctions (x_i, v_i) pour le modèle microscopique, vérifiant les équations du modèle FTL dans le cas si $i < l(t)$, et dépendant du modèle macroscopique si $i = l(t)$.
- d'une densité ρ , solution pour le modèle macroscopique avec comme condition au bord en 0 un flux dépendant du modèle micro et changeant aussi à chaque fois qu'un leader atteint l'interface.

Cette définition permet de donner un sens mathématique à la modélisation du couplage, permettant de donner un cadre à ce type de changement d'échelle.

Une autre approche importante dans le domaine des foules est celle de [CPT14] : il s'agit du développement d'un modèle prenant en compte les deux échelles en même temps en calculant des contributions microscopiques et macroscopiques sur le mouvement des individus et de leur densité globale. Le mouvement individuel et global est ensuite une combinaison linéaire de ces deux contributions. Cette façon d'aborder le couplage permet de prendre en compte des effets à plusieurs échelles et dans ce paradigme le couplage est vu comme une superposition des deux échelles. En revanche cette approche ne paraît pas généralisable simplement pour d'autres modèles que ceux introduits : la nécessité de définir l'interaction d'une particule au milieu d'une densité fait intervenir un formalisme provenant de la théorie de la mesure et ne convient pas forcément à tous les modèles.

Enfin le problème a aussi été abordé d'un point de vue plus informatique, avec par exemple la mise en place de passages entre automates cellulaires, modèles microscopiques en espace continu et modèles macroscopiques [BCB21 ; XLC+10]. Ces développements montrent les difficultés pratiques à la mise en place de couplages de modèle et nécessaires à surmonter pour appliquer ces calculs à des cas réels. Assez peu d'études sur les implications et la transmission d'information lors du changement d'échelle sont cependant présentées, le couplage étant principalement abordé en tant que moyen d'optimiser des performances de calcul et pas nécessairement pour l'intérêt propre d'avoir un modèle par situation.

2.3 Observables et validation des modèles

Les études expérimentales dans le domaine des mouvements de foules se multiplient ces dernières années pour identifier et caractériser les différents phénomènes pouvant émerger des comportements individuels. Ces différentes données peuvent servir à nourrir la modélisation, en donnant des valeurs à des paramètres utilisés pour la taille des individus ou la vitesse de marche libre par exemple. Beaucoup d'expériences se concentrent en revanche sur l'identification et la caractérisation de phénomènes émergents propres aux foules en mouvement, comme le *Faster is Slower* mentionné ci-après. Ces phénomènes peuvent servir à construire de nouveaux modèles ou à comparer des calculs à la réalité dans un cadre contrôlé.

La validation de modèles. Il serait tentant de chercher à valider les modèles, en comparant des résultats de simulation à chaque phénomène observé. Il n'existe pas à notre connaissance de cas test globalement accepté par la communauté pour chaque situation. Un document fréquemment utilisé est un guide publié par l'*International Maritime Organisation* (IMO), qui fournit un ensemble de tests et de données d'entrée pour un logiciel de simulation de mouvement de foules ([IMO07]).

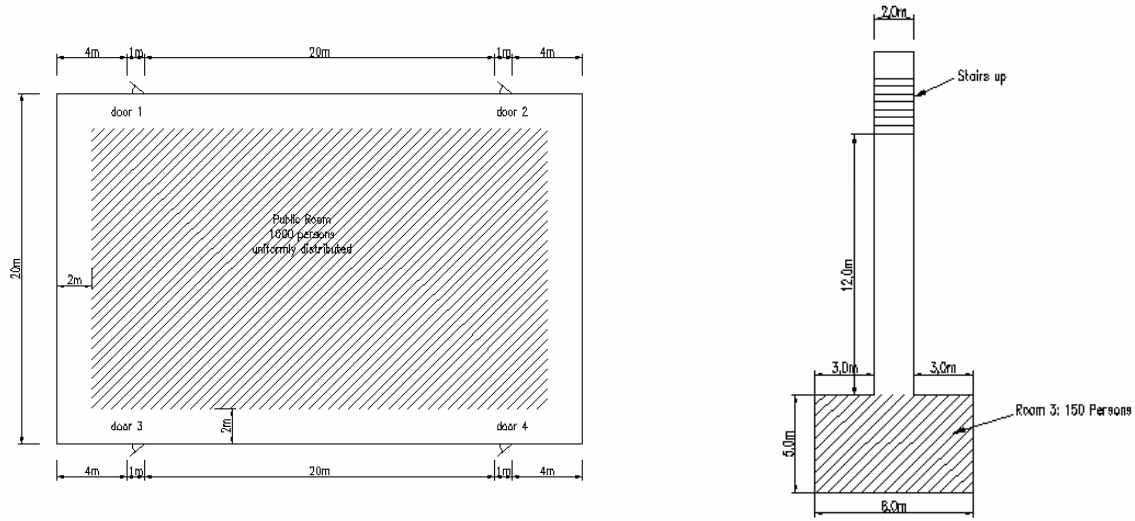


FIGURE 2.5 – Quelques exemples de géométries dans les tests IMO ([IMO07])

Ces tests permettent de vérifier un certain bon fonctionnement du logiciel (en vérifiant qu'on est capable de fixer correctement les paramètres d'entrée, ou qu'il n'y a pas de pénétration de murs solides, etc). Assez peu de critères de validation précis sont cependant donnés : le test n°9 par exemple demande de simuler l'évacuation d'une salle rectangulaire donnée en figure 2.5 et de vérifier qu'en fermant deux des quatre portes le temps d'évacuation double approximativement.

Dans le cadre de mouvements de foules cette démarche de validation n'est en général pas possible, les phénomènes observés étant très souvent par nature difficiles à observer systématiquement. Les comportements humains possèdent une part inhérente d'aléatoire qui peut être délicate à estimer et certains processus de pensées complexes peuvent être difficiles à étudier en laboratoire : lors d'un exercice d'évacuation en cas d'incendie en laboratoire, à quel point le comportement observé diffère d'une évacuation réelle ? On ne peut ainsi rejeter ou accepter des modèles selon leur capacité à reproduire des résultats expérimentaux, mais plutôt catégoriser les modèles selon les phénomènes qu'ils peuvent reproduire et décider de la démarche de modélisation la plus adaptée.

Vitesse de déplacement et densité. En cas de marche libre des individus on observe des vitesses de marche comprises entre 1.2 m/s et 1.5 m/s. D'après [Wei92], la vitesse de marche libre des individus dans le monde suit une loi normale de moyenne 1.34 m/s, avec un écart-type de 0.26m/s.

En cas de congestion en revanche il est connu que la vitesse individuelle diminue. Beaucoup d'études ont été menées afin de déterminer la relation entre vitesse et densité locale. Des expériences dans des couloirs circulaires sont proposées dans [JALP12] et [MGM+12]. Dans le cas de mouvements de panique, les vidéos d'une bousculade à la Hadj en 2006 ont été analysées dans [HJA07], et montrent des vitesses positives même à des densités bien supérieures à ce qui est généralement considéré dans la littérature. Ces différentes observations sont représentées en Figure 2.6. À partir d'observations de terrain, différentes relations ont été proposées dans [Wei92] ou [PM78].

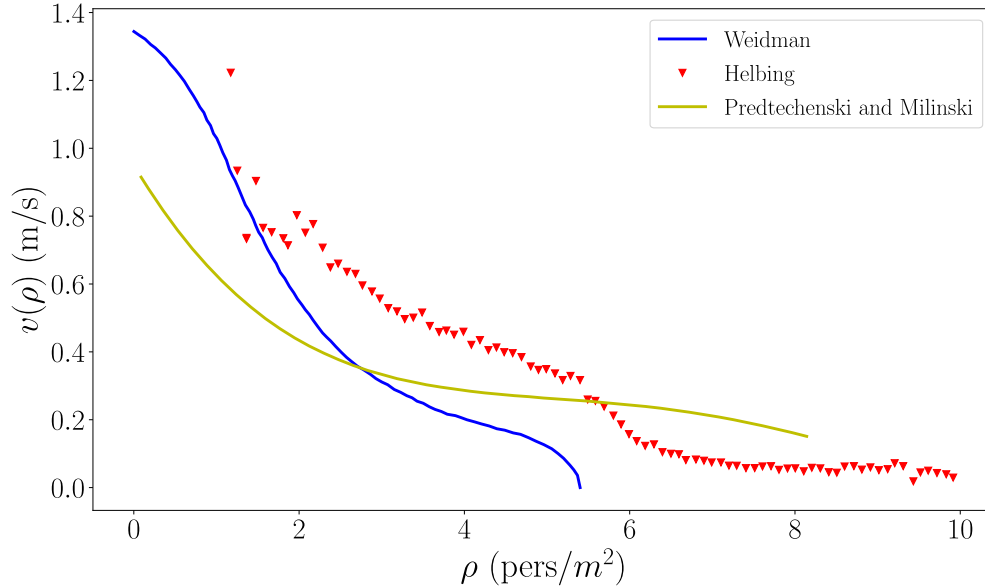


FIGURE 2.6 – Relations entre densité et vitesse dans [Wei92 ; PM78 ; HJA07]. Les traits pleins représentent des formules théoriques et les points des données expérimentales.

Flux en amont de rétrécissements

Dans le cadre d'évacuations de bâtiments, le flux à un goulot d'étranglement est un indicateur significatif de la fluidité et la vitesse de l'évacuation. Les portes par exemple sont des structures où on peut s'attendre à des congestions plus importantes. Les expériences menées dans [Cep09] montrent que, pour des fortes densités en amont d'une porte, le flux passant diminue (effet parfois appelé *capacity drop*). Plusieurs études essaient de caractériser ces effets de congestion et d'en estimer les causes lorsque cela est possible.

Effet zipper. L'effet zipper est un phénomène d'auto-organisation lors du passage de couloirs où les piétons s'organisent spontanément en rangées pour optimiser l'espace disponible. Cet effet a été observé dans [HD05] pour des flux unidirectionnels dans différentes largeurs de couloirs. Ces résultats mettent en évidence que le flux augmente par crans successifs, et non linéairement. Dans [PMLW14] la répartition des plus proches voisins dans des expériences de flux unidirectionnels a été étudiée et montre que cette organisation favorise la formation de groupes de deux piétons, qui crée une alternance au niveau des temps de passage entre deux individus en sortie de couloir (deux piétons côte à côte vont sortir pratiquement en même temps, les deux suivants seront après un intervalle plus grand, etc.). Ces corrélations ont été observées dans d'autres expériences menées sur des évacuations d'urgence à travers une porte [NBK17]. Les auteurs ont observé qu'en cas de fortes bousculades le flux pouvait devenir instable et observent que le passage par la porte se fait par clusters de deux ou trois personnes. Cette similarité phénoménologique les conduit à parler d'effet zipper généralisé, qui pourrait s'appliquer dans des situations plus variées.

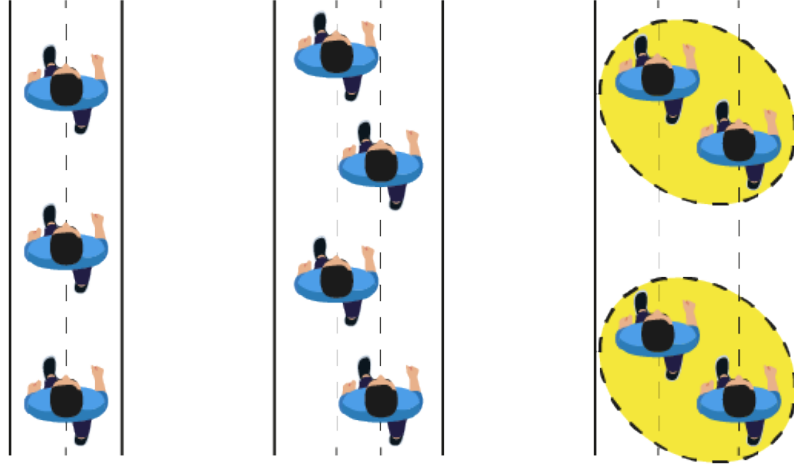


FIGURE 2.7 – Illustration de l'effet zipper.

Faster-is-Slower. On parle de *Faster is Slower* (ou FiS) pour certains systèmes qui peuvent perdre en efficacité lorsque les individus cherchent à améliorer leur performance. Il s'agit d'un effet qui peut être observé dans différents types de systèmes humains, ou même chez des animaux ou en écologie ([GH15]). Dans le cadre d'évacuation au passage d'une porte en particulier, on aura un effet de FiS si une augmentation de la vitesse souhaitée d'un individu entraîne une diminution du flux à la porte. Cette définition a été proposée dans [HFV00] et a depuis fait l'objet d'expérimentations en laboratoire. Dans [GZP+14] une évacuation de 85 personnes a été effectuée en autorisant ou non aux individus de pousser pour arriver à la sortie. Lorsque autorisés à se bousculer, les individus mettent jusqu'à 20% de temps en plus à sortir. D'autres expériences en laboratoire dans [GPP+16 ; NBK17] ont mis en évidence l'impact du comportement d'individus agressifs sur le flux global d'évacuation en favorisant l'apparition de bouchons.

D'autres expériences contredisent ces résultats : une revue de différentes méthodes expérimentales ([Hag20]) montre que les résultats expérimentaux ne permettent pas à l'heure actuelle de savoir si cet effet est systématique, ou sous quelles conditions il se manifeste.

Effet d'un obstacle. Une solution proposée au problème d'accumulation de densité en amont d'une porte a été de placer un obstacle devant la porte. L'idée derrière cette solution contre-intuitive est qu'un obstacle bien placé serait capable d'empêcher la formation de bouchons en amont et d'empêcher les forces de pression de s'accumuler sur des nombres trop importants de piétons. Appliquée à des systèmes de particules granulaires, cette stratégie est connue pour être efficace. Dans le cadre d'évacuation, dans [HBJW05] les auteurs présentent une expérience d'évacuation répétée avec et sans un panneau en bois de 45cm en amont de la sortie. Les auteurs observent un flux jusqu'à 30% plus élevé avec un obstacle.

D'autres expériences ont par la suite été menées afin de quantifier cet effet, mais ne l'ont pas toujours observé. L'effet d'un obstacle en amont d'une porte est aujourd'hui un sujet très controversé. Dans [Hag20], une revue des expériences disponibles dans la littérature entre 2016 et 2020 montre que sur les 10 expériences considérées, seules 2 observent une amélioration du temps de sortie avec la présence d'un obstacle. Dans les autres expériences, il est noté que l'effet de l'obstacle est très sensible aux conditions expérimentales, et peut parfois avoir un effet néfaste sur le flux. D'autres effets d'un obstacle sont aussi étudiés par ces expériences. Dans [FZG+20], deux séries

d'expériences d'évacuation avec des étudiants et des soldats sont présentées. Dans chaque cas un obstacle cylindrique de 1m de diamètre a été placé à différentes distances de l'issue d'évacuation (de 50 à 70cm), avec jusqu'à 200 participants. Les auteurs étudient le niveau de congestion et la stabilité du flux avec et sans obstacle et montrent que le mouvement d'évacuation est plus stable en présence d'obstacle, mais que les densités sont aussi plus élevées.

Chapitre 3

Couplage de modèles 1D

3.1	Couplage de modèles d'anticipation	32
3.2	Couplage de modèles d'inhibition	44
3.3	Cas d'application et extensions	51

Dans ce chapitre nous allons nous intéresser aux méthodes de couplage pour des modèles à une dimension. Ces situations simplifiées nous permettront d'identifier différentes approches du problème qui serviront ensuite pour le passage à deux dimensions.

Le problème du couplage entre micro et macro qui nous intéresse consiste à décrire le passage d'entités lagrangiennes en quantité eulérienne appropriée à partir d'une interface fixée en amont, comme l'illustre la figure 3.1. Lors de ce changement d'échelle la modélisation de l'effet de chacun des modèles sur l'autre n'est pas unique et aura un rôle sur la propagation d'information entre les deux modèles. Dans ce chapitre, nous écrivons des conditions préservant la stabilité du flux à l'interface pour des modèles similaires à échelle microscopique et macroscopique, ainsi que des schémas numériques adaptés.

Nous abordons le couplage à l'interface pour deux types de modèles. Pour des modèles basés sur l'anticipation individuelle, nous regarderons les différentes formulations possibles pour transmettre

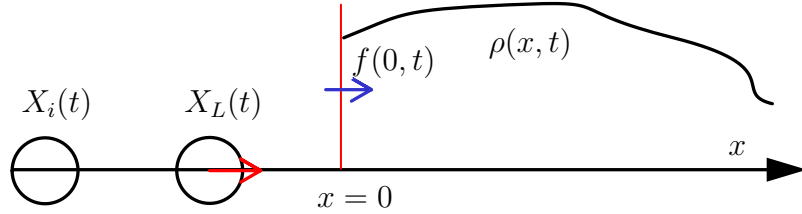


FIGURE 3.1 – Couplage microscopique vers macroscopique. Le mouvement se fait vers les x positifs et le changement de modèle à l'interface $x = 0$.

l'information au niveau de l'interface, ainsi que des approches numériques pour ce type de problèmes. Nous nous intéresserons ensuite à des modèles raides de type granulaire, en montrant le lien que l'on peut faire entre ces deux approches lorsque l'on considère une anticipation de plus en plus raide. Nous développerons ensuite les schémas numériques pour le couplage de ces modèles. L'organisation de cette section est résumée dans la figure 3.2

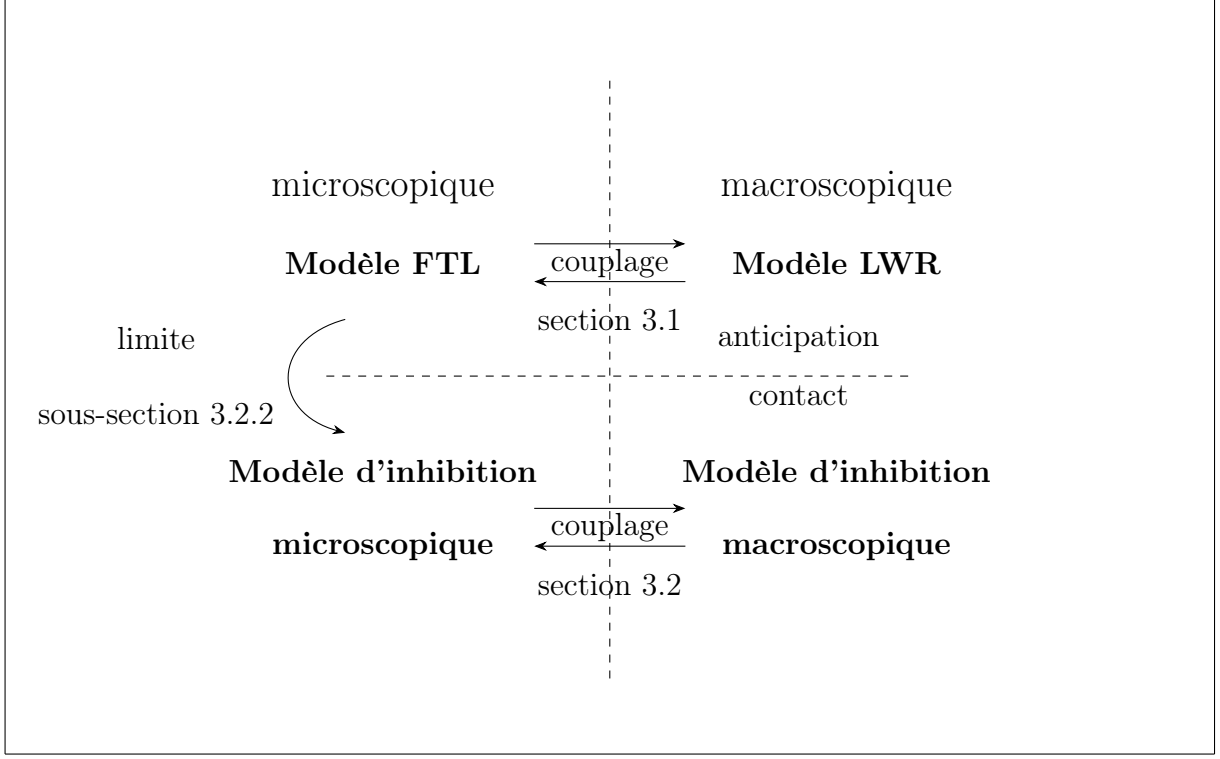


FIGURE 3.2 – Organisation des parties de cette section.

3.1 Couplage de modèles d'anticipation

Nous allons utiliser les modèles Follow-the-Leader (FTL) et Lighthill-Whithams-Richards (LWR), décrits dans le chapitre 2. Nous rappelons ici les équations principales pour faciliter la lecture du manuscrit.

Le modèle Follow-The-Leader, ou FTL, décrit le mouvement d'individus $1, \dots, L$ de positions $X_i = X_i(t)$ et de vitesses $V_i = V_i(t)$ pour tout i tel que $1 \leq i \leq L$ par le système d'équations :

$$V_i = \varphi(X_{i+1} - X_i) \quad \forall 1 \leq i < L, \quad (3.1)$$

$$V_L = V(t), \quad (3.2)$$

$$X_{i+1}(0) - X_i(0) > 2r \quad \forall 1 \leq i < L, \quad (3.3)$$

où $\varphi : \mathbb{R}^+ \rightarrow \mathbb{R}^+$ est une fonction croissante à valeurs dans $[0, U]$, U est la vitesse maximale d'un individu et $V(t)$ est une fonction donnée (typiquement continue et à valeurs dans $[0, U]$) et r

est le rayon d'un individu, supposé identique entre tous les individus. On supposera que φ est nulle sur $[0, 2r]$.

Le nombre d'individus L change au cours du temps, par conséquent dans la suite nous utiliserons $L(t)$ la fonction donnant le nombre d'individus dans le modèle microscopique à l'instant t .

Le modèle Lighthill-Whithams-Richards (LWR), est basé sur une équation de conservation du premier ordre avec une vitesse dépendant de la densité locale :

$$\partial_t \rho + \partial_x (f(\rho)) = 0, \quad (3.4)$$

où $f(\rho) = \rho v(\rho)$, $v \mapsto v(\rho)$ est une fonction donnée. Plusieurs formes de fonction ont été proposées pour $v(\rho)$ dans la littérature (voir section 2.3). On suppose généralement que v est décroissante et que pour un certain ρ_{max} on a $v(\rho_{max}) = 0$.

En choisissant $v(\rho) = \varphi(1/\rho)$, on peut interpréter le modèle FTL comme une formulation lagrangienne du modèle LWR. En effet, il est possible de prouver qu'en choisissant une limite appropriée du modèle FTL lorsque le nombre de particules tend vers $+\infty$ on retrouve le modèle LWR (voir [DR15] pour une formulation précise de ce résultat). Dans [HR17] il est aussi montré comment le modèle FTL peut être vu comme un schéma numérique lagrangien pour le modèle LWR.

Discretisation et schémas numériques. Le modèle FTL est discrétisé avec un schéma implicite en temps pour l'équation (3.1). Pour un pas de temps Δt , on note X_i^n l'approximation de $X_i(t)$ à l'instant $n\Delta t$ et on a :

$$X_i^{n+1} = X_i^n + \Delta t \varphi(X_{i+1}^{n+1} - X_i^{n+1}) \quad \forall 1 \leq i < L^n, \quad (3.5)$$

où L^n est le leader à l'instant $n\Delta t$. La résolution numérique se fait à l'aide d'une méthode de Newton qui ne pose pas de problème particulier. Ce type de schéma a l'avantage d'être stable, même pour des grands pas de temps. Ce choix permettra d'utiliser un pas de temps commun entre le modèle micro et macro.

Pour le modèle LWR, l'espace est discrétisé avec un pas Δx sur une grille $x_j = (j + \frac{1}{2})\Delta x \quad \forall j \in \mathbb{Z}$. On utilise un schéma numérique de type volumes finis avec un pas de temps Δt pour toute cellule j :

$$\frac{\rho_j^{n+1} - \rho_j^n}{\Delta t} + \frac{F_{j+1/2}^n - F_{j-1/2}^n}{\Delta x} = 0, \quad (3.6)$$

où ρ_j^n est l'approximation de $\rho(n\Delta t, x_j)$ et $F_{j+1/2}^n$ le flux à l'interface entre les cellules x_j et x_{j+1} . Pour notre schéma numérique nous prenons le flux introduit dans [FFL+16] :

$$F_{j+1/2}^n = F(\rho_j^n, \rho_{j+1}^n) = \rho_j^n v(\rho_{j+1}^n). \quad (3.7)$$

Ce schéma repose sur le fait que la matière se propage vers l'aval, alors que l'information se propage vers l'amont. On a $F(\rho, \rho) = \rho v(\rho)$ donc ce schéma est consistant, et il est monotone sous la condition CFL :

$$\Delta t < \Delta x / (U + \rho_{max} \max_{\rho} (v'(\rho))). \quad (3.8)$$

On peut le montrer en réécrivant le schéma sous la forme :

$$\rho_j^{n+1} = G(\rho_{j-1}^n, \rho_j^n, \rho_{j+1}^n),$$

la monotonicit  est v rifi e lorsque $G(\rho_1, \rho_2, \rho_3)$ est croissante par rapport   chacune de ses variables. On a :

$$G(\rho_1, \rho_2, \rho_3) = \rho_2 - \frac{\Delta t}{\Delta x}(\rho_2 v(\rho_3) - \rho_1 v(\rho_2))$$

v est d croissante, donc G est croissante par rapport   ρ_1 et ρ_3 . En d rivant par rapport   son second argument on a :

$$\partial_{\rho_2} G(\rho_1, \rho_2, \rho_3) = 1 - \frac{\Delta t}{\Delta x}(v(\rho_3) - \rho_1 v'(\rho_2)).$$

La condition CFL assure ainsi la monotonicit  du sch ma.

3.1.1 Couplage microscopique vers macroscopique : premiers sch mas heuristiques

Dans le cadre du couplage microscopique vers macroscopique, nous devons attribuer d'une part une vitesse au leader de la partie microscopique et un flux pour le mod le macroscopique. On supposera par la suite que l'interface entre les deux mod les se situe en $x = 0$. Pour la vitesse du leader un premier choix de mod lisation possible est de prendre la vitesse macroscopique   l'interface :

$$V_{L(t)} = v(\rho(t, 0^+)), \quad (3.9)$$

o  $L(t)$ donne l'indice du leader du mod le microscopique   l'instant t et $V_{L(t)} = \dot{X}_{L(t)}$. Le couplage est compl t  en attribuant un flux   l'interface pour le mod le macroscopique. Le leader ayant une vitesse donn e, en lui attribuant une densit  $\rho_b(t)$ observ e par le mod le macroscopique   l'interface le flux pourra  tre exprim  :

$$f(t, 0^+) = \rho_b(t) V_{L(t)} \quad (3.10)$$

Une premi re possibilit  de mod lisation consiste    crire :

$$\rho_b(t) = \begin{cases} 1/2r & \text{si } -2r \leq X_{L(t)} \leq 0 \\ 0 & \text{sinon,} \end{cases} \quad (3.11)$$

La densit  vue au bord est ainsi concentr e autour du leader qui s'appr te   passer dans le mod le macro.

En rassemblant toutes les  quations concern es, le syst me d' quations pour le probl me de couplage micro-macro est donn  par :

$$\partial_t \rho + \partial_x(f(\rho)) = 0 \quad \forall (x, t) \in \mathbb{R}^+ \times \mathbb{R}^+ \quad (3.12)$$

$$V_i = \varphi(X_{i+1} - X_i) \quad \forall 1 \leq i < L(t), \quad (3.13)$$

$$V_{L(t)} = v(\rho(t, 0^+)), \quad (3.14)$$

$$L(t) = \sup \{i, X_i < 0\}, \quad (3.15)$$

$$\rho_b(t) = \begin{cases} 1/2r & \text{si } -2r \leq X_{L(t)} \leq 0 \\ 0 & \text{sinon,} \end{cases} \quad (3.16)$$

$$f(t, 0^+) = \rho_b(t) V_{L(t)}. \quad (3.17)$$

Schémas numériques et cas test. Nous reprenons ici la discrétisation des équations (3.1) et (3.4) introduites avec les schémas numériques (3.5) et (3.6). Pour le modèle LWR on a un maillage $x_j = (j + \frac{1}{2})\Delta x \forall j \in \mathbb{N}$. Pour la vitesse du leader on a :

$$\frac{X_{L^n}^{n+1} - X_{L^n}^n}{\Delta t} = v(\rho_0^n). \quad (3.18)$$

Une difficulté réside ici dans l'approximation du flux $F_{-1/2}^n$ à l'interface. On peut exprimer ce flux en utilisant la conservation de la masse entre les deux modèles. À chaque pas de temps on peut déterminer la quantité intégrée de piéton de chaque côté de l'interface. On note $m[X]$ la masse intégrée à gauche de l'interface telle que :

$$m[X] = \int_{-\infty}^0 \rho_{L(t)}[X](x) dx. \quad (3.19)$$

Pour avoir une conservation de la masse entre les deux modèles, il est nécessaire que la perte de masse dans le modèle lagrangien soit compensée par un ajout équivalent dans le modèle macroscopique. La perte de masse dans le modèle microscopique est donnée par :

$$F_{-1/2}^n = \frac{m[X^n] - m[X^{n+1}]}{\Delta t}, \quad (3.20)$$

où $X^n = (X_i^n)_i$. Si les positions $X(n\Delta t)$ sont connues, $m(t)$ peut être calculée directement sans approximation supplémentaire. Entre deux pas de temps on peut ainsi avoir un ou plusieurs piétons qui passent l'interface.

Le schéma numérique global devient alors :

$$X_i^{n+1} = X_i^n + \Delta t \varphi(X_{i+1}^{n+1} - X_i^{n+1}) \quad \forall 1 \leq i < L^n, \quad (3.21)$$

$$L^n = \sup \{i, X_i^n < 0\}, \quad (3.22)$$

$$\frac{\rho_j^{n+1} - \rho_j^n}{\Delta t} + \frac{F_{j+1/2}^n - F_{j-1/2}^n}{\Delta x} = 0 \quad \forall j \in \mathbb{N}, \quad (3.23)$$

$$F_{j+1/2}^n = \rho_j^n v(\rho_{j+1}^n), \quad (3.24)$$

$$\frac{X_{L^n}^{n+1} - X_{L^n}^n}{\Delta t} = v(\rho_0^n), \quad (3.25)$$

$$F_{-1/2}^n = \frac{m[X^n] - m[X^{n+1}]}{\Delta t}. \quad (3.26)$$

Cas test. Nous présentons ici à deux cas tests :

- le premier en figure 3.3 montre une portion de route de 100 m avec 90 individus à gauche répartis uniformément sur 50 m et à droite et une densité constante à 1.8 pers/m. À 100 m la densité au bord est fixée à ρ_{max} .
- le second cas test en figure 3.4 avec 40 individus à gauche et une densité à 0.9 pers/m à droite.

Des oscillations apparaissent au niveau de l'interface entre les deux modèles. Ces oscillations sont d'autant plus marquées à basse densité par rapport aux hautes densités. Le cas test en figure 3.3 montre en revanche la capacité du couplage à faire remonter l'information à l'interface.

Ces oscillations sont dues au choix de représentation pour la densité au bord vue par le modèle macroscopique (3.11) : en choisissant une fonction aussi locale, le flux à l'interface en est impacté et oscille fortement.

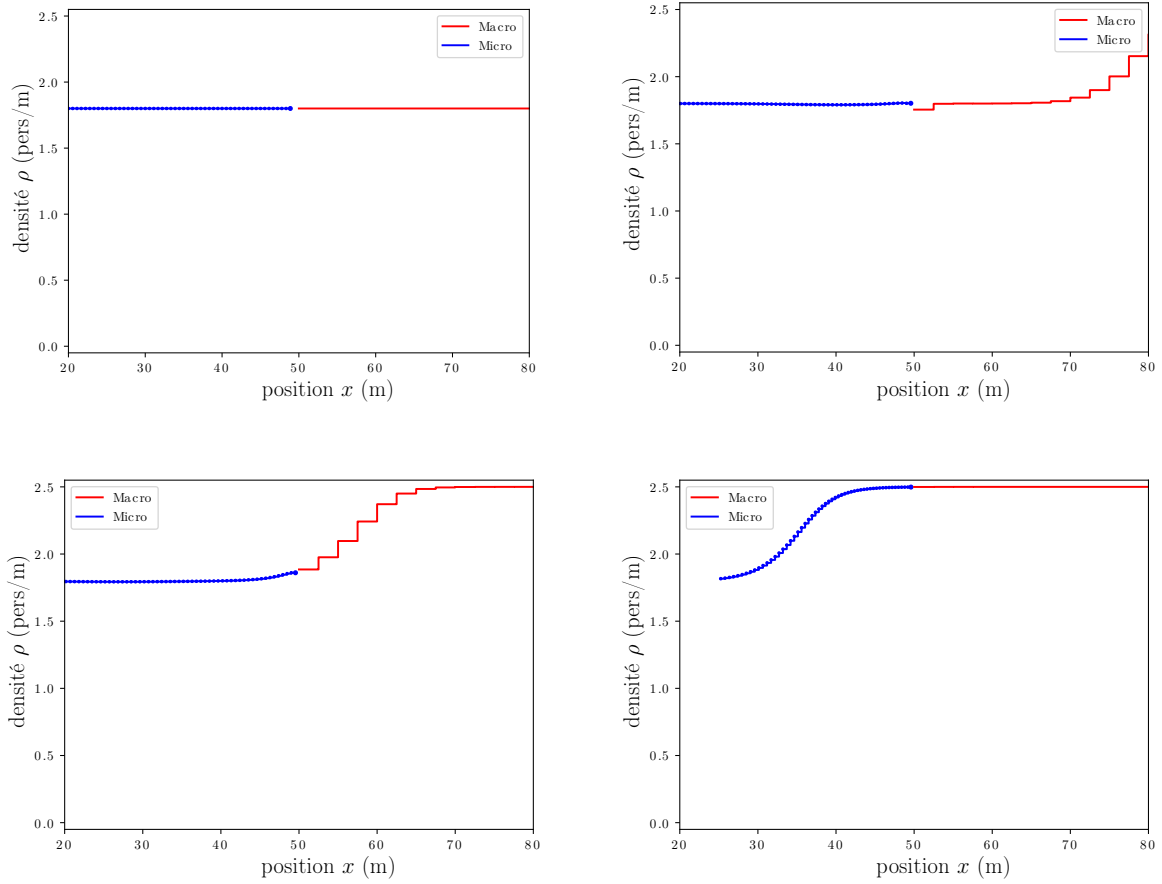


FIGURE 3.3 – Premier cas test à $t = 0, 25, 50, 75$ s, les paramètres sont donnés tableau 3.1. La figure est agrandie sur la partie entre 20 et 80m.

3.1.2 Modification de l'absorption

Dans la littérature ([DR15] par exemple) on peut voir que les outils de convergence du modèle micro vers le modèle macroscopique reposent sur la définition d'une densité pour chaque individu telle que :

$$\rho_i(t, x) = \begin{cases} \frac{1}{X_{i+1} - X_i} & \text{si } X_i \leq x < X_{i+1} \\ 0 & \text{sinon.} \end{cases} \quad (3.27)$$

Pour le leader cette définition n'est pas applicable directement : celui-ci ne perçoit aucun individu devant lui. Une première solution serait de choisir pour la densité $\rho_b(t)$:

$$\rho_b(t) = \frac{1}{|X_{L(t)}| + 2r} \quad (3.28)$$

qui correspond à la densité que voit le leader si un piéton imaginaire se tient au niveau de l'interface. Avec cette définition cependant la densité augmente toujours au fur et à mesure qu'un individu s'approche de l'interface, créant de nouvelles oscillations.

paramètre	valeur
$v(\rho)$	$U(1 - \frac{\rho}{\rho_{max}})$
U	$1.2m/s$
ρ_{max}	$2.5m^{-1}$
Δx	$6.25m$
Δt	$1.0s$

TABLE 3.1 – Paramètres pour le cas test considéré en Figures 3.3-3.4.

On choisit ainsi une autre fonction pour la densité du leader. On introduit la fonction τ donnant le temps d'intersection du précédent leader avec l'interface, c'est-à-dire :

$$\tau(t) = \begin{cases} 0 & \text{si } L(t) = L(0) \\ \sup \{t, X_{L(t)+1} < 0\} & \text{sinon.} \end{cases} \quad (3.29)$$

Dans cette définition on prend comme convention qu'à l'instant 0 un piéton imaginaire se situe à l'interface. On écrit enfin la densité vue par le modèle macroscopique comme étant :

$$\rho_b(t) = \frac{1}{|X_{L(\tau)}(\tau)|}. \quad (3.30)$$

On fixe ainsi la valeur de la densité au moment où il devient leader. En approchant de l'interface on aura alors une injection continue à la vitesse $v(\rho(t, 0^+))$. Notre nouvelle version du couplage s'obtient donc en remplaçant l'équation (3.16) par l'équation (3.30).

Modification du schéma numérique. Le schéma numérique doit prendre en compte les temps τ où le calcul de la vitesse du leader change, et qui peuvent se trouver entre deux itérations.

Nous modifions ainsi le schéma numérique en introduisant une détection des instants τ où un piéton traverse l'interface. Nous allons ajouter entre deux itérations une étape de vérification pour savoir si le leader traverse l'interface, et si c'est le cas on fait avancer le temps jusqu'au moment où le leader traverse l'interface. On obtient ainsi la densité au bord pour le leader suivant et on continue l'itération, comme illustré en figure 3.5. On aura ainsi autant de sous-pas de temps qu'il y aura de leaders qui franchissent l'interface pendant Δt .

Ce sous-pas de temps est tel que :

$$\Delta t_k = \min \left\{ \Delta t - \sum_{j < k} \Delta t_j, \frac{|X_{L^{n,k}}^{n,k}|}{v(\rho_0^n)} \right\}, \quad (3.31)$$

où $L^{n,k}$ est le leader au pas n, k et $X_i^{n,k}$ la position d'un individu i au pas n, k , et vérifie :

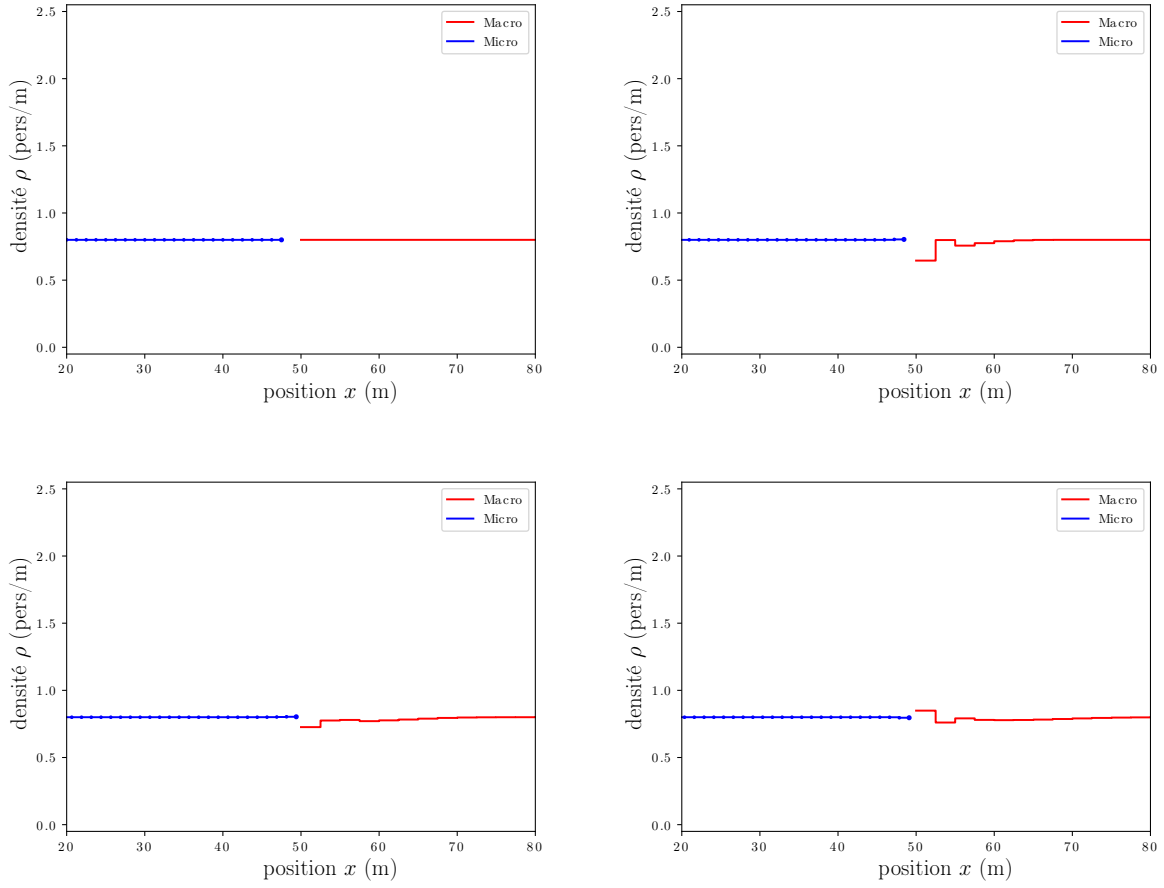


FIGURE 3.4 – Cas de simulation à basse densité à $t = 0, 25, 50, 75$ s. Les paramètres sont donnés tableau 3.1. La figure est agrandie sur la partie entre 20 et 80m.

$$X^{n,0} = X^n, \quad (3.32)$$

$$X_i^{n,k+1} = X_i^{n,k} + \Delta t^k \varphi(X_{i+1}^{n,k+1} - X_i^{n,k+1}) \quad \forall i < L^{n,k}, \quad (3.33)$$

$$X_{L^{n,k}}^{n,k+1} = X_{L^{n,k}}^{n,k} + \Delta t^k v(\rho_0^n), \quad (3.34)$$

$$X^{n+1} = X^{n,K^n}, \quad (3.35)$$

et K^n est tel que :

$$\sum_{k \leq K^n} \Delta t_k = \Delta t. \quad (3.36)$$

Avec cette formulation la densité et le flux à l'interface se réécrivent :

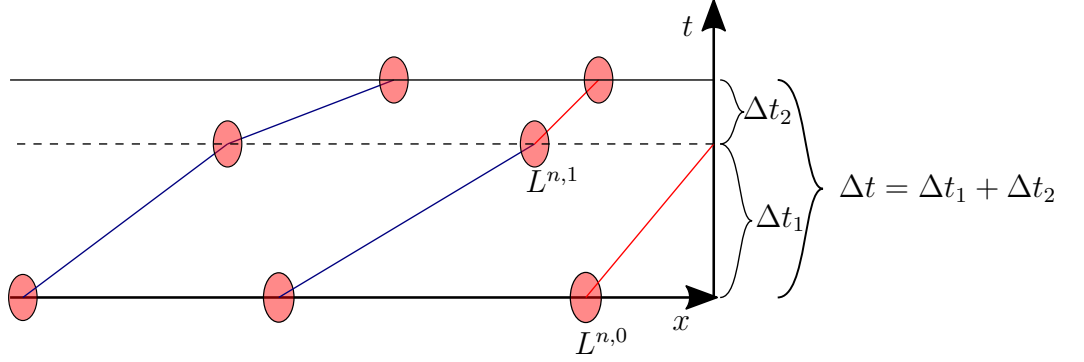


FIGURE 3.5 – Evolution entre deux pas de temps t et $t + \Delta t$

$$\rho^{n,k} = \frac{1}{|X_{L^{n,k}}^{n_0,k_0}|}, \quad n_0, k_0 = \{n, k / X_{L^{n,k}+1}^{n,k} = 0\}, \quad (3.37)$$

$$F_{-1/2}^n = \frac{1}{\Delta t} \sum_k \Delta t_k \rho^{n,k} v(\rho_0^n). \quad (3.38)$$

Notons en particulier que dans l'équation (3.72) k_0, n_0 est bien défini grâce au choix de Δt^k

Remarque 3.1. Le flux numérique $F_{-1/2}^n$ dans l'équation (3.38) peut aussi s'écrire :

$$F_{-1/2}^n = \rho_b^n v(\rho_0^n), \quad (3.39)$$

$$\rho_b^n = \frac{1}{\Delta t} \sum_k \Delta t^k \rho^{n,k}. \quad (3.40)$$

Ce flux peut donc être vu comme le flux numérique entre la première cellule et une cellule dont la densité serait une moyenne particulière des densités à l'intérieur. Cette formulation permet d'avoir un principe du maximum pour le flux à l'interface, qui n'était pas assuré dans le schéma numérique précédent.

Tests numériques. Nous présentons ici un cas test dont les paramètres sont donnés par le tableau 3.2. Deux niveaux de discrétisation différents sont donnés pour le modèle macroscopique avec $\Delta x = 5$ m et $\Delta x = 0.5$ m (figure 3.6). Une simulation du modèle FTL sans couplage avec les mêmes paramètres est aussi donnée comme référence.

Remarque 3.2. Il serait possible de prendre un maillage plus fin en espace pour le modèle macroscopique, mais l'on se retrouverait avec une taille de maille inférieure à la taille d'un individu. À ces échelles, le modèle macroscopique n'a plus de sens en termes de modélisation, nous serions en train de raisonner à une échelle plus petite que la taille de la particule élémentaire du « fluide » que nous modélisons.

Une comparaison entre les résultats des deux tailles de mailles montre que le modèle couplé converge qualitativement vers le modèle FTL seul.

paramètre	valeur
$v(\rho)$	$U(1 - \frac{\rho}{\rho_{max}})$
U	1.2 m/s
ρ_{max}	2.5 m^{-1}
ρ_0	$2.5 \times \mathbf{1}_{[0,40]} + 2.2 \times \mathbf{1}_{[40,50]} + 1.0 \times \mathbf{1}_{[50,100]}$
Δx	0.5 - 5 m
Δt	0.2 - 2 s

TABLE 3.2 – Paramètres pour le cas test considéré en figure 3.6 et figure 3.7.

On remarque ensuite la disparition des oscillations observées avec la précédente modélisation. On peut assez simplement comprendre d'où vient cette stabilité en s'intéressant au cas d'une densité uniforme à droite $\bar{\rho} > 0$ de l'interface et de piétons uniformément répartis avec la même densité à gauche de l'interface à un pas n, k quelconque. En supposant toujours que $\varphi(1/\rho) = v(\rho)$, la vitesse des agents microscopiques sera égale à la vitesse locale de la densité macroscopique. À l'interface le flux vaut $\rho_b v(\bar{\rho}) = \bar{\rho} v(\bar{\rho})$ si la densité ρ_b a été initialisée correctement, donc le leader avancera aussi à vitesse constante et le flux à l'interface sera constant lui aussi. Lorsque le leader rencontre l'interface, la nouvelle densité ρ_b devenant la condition de bord sera aussi $\bar{\rho}$ grâce à la stabilité du modèle FTL.

3.1.3 Couplage macro vers micro

Nous étudions maintenant la situation opposée aux couplages abordés dans ce manuscrit jusqu'ici, le problème de couplage entre le modèle LWR à gauche et le modèle FTL à droite. Pour le modèle microscopique on adapte nos notations dans cette partie, les piétons seront numérotés de façon décroissante de 0 à $L(t)$, de façon à ce que la notation de l'équation (3.1) soit toujours valable.

La difficulté de ce couplage va consister à estimer le flux à l'interface, cette fois en générant des individus dans le modèle microscopique. Pour résoudre ce problème nous allons nous appuyer sur une interprétation lagrangienne du modèle macroscopique afin de déterminer une densité vue avec le premier voisin.

Procédure d'émission d'individu. Nous commençons par déterminer pour une densité donnée dans le modèle macroscopique une position dans le modèle macroscopique à laquelle générer un individu. Cet individu servira à déterminer une vitesse qui va servir de condition au bord pour le modèle macroscopique.

On introduit la fonction $t \mapsto x_g(t)$ telle que :

$$x_g(t) = \inf \left\{ x^* < 0, \int_{x^*}^0 \rho(t, x) dx \leq 1 \right\}. \quad (3.41)$$

Pour une densité donnée, x_g est donc la position dans le modèle macroscopique telle que la densité entre x_g et l'interface se somme à 1.

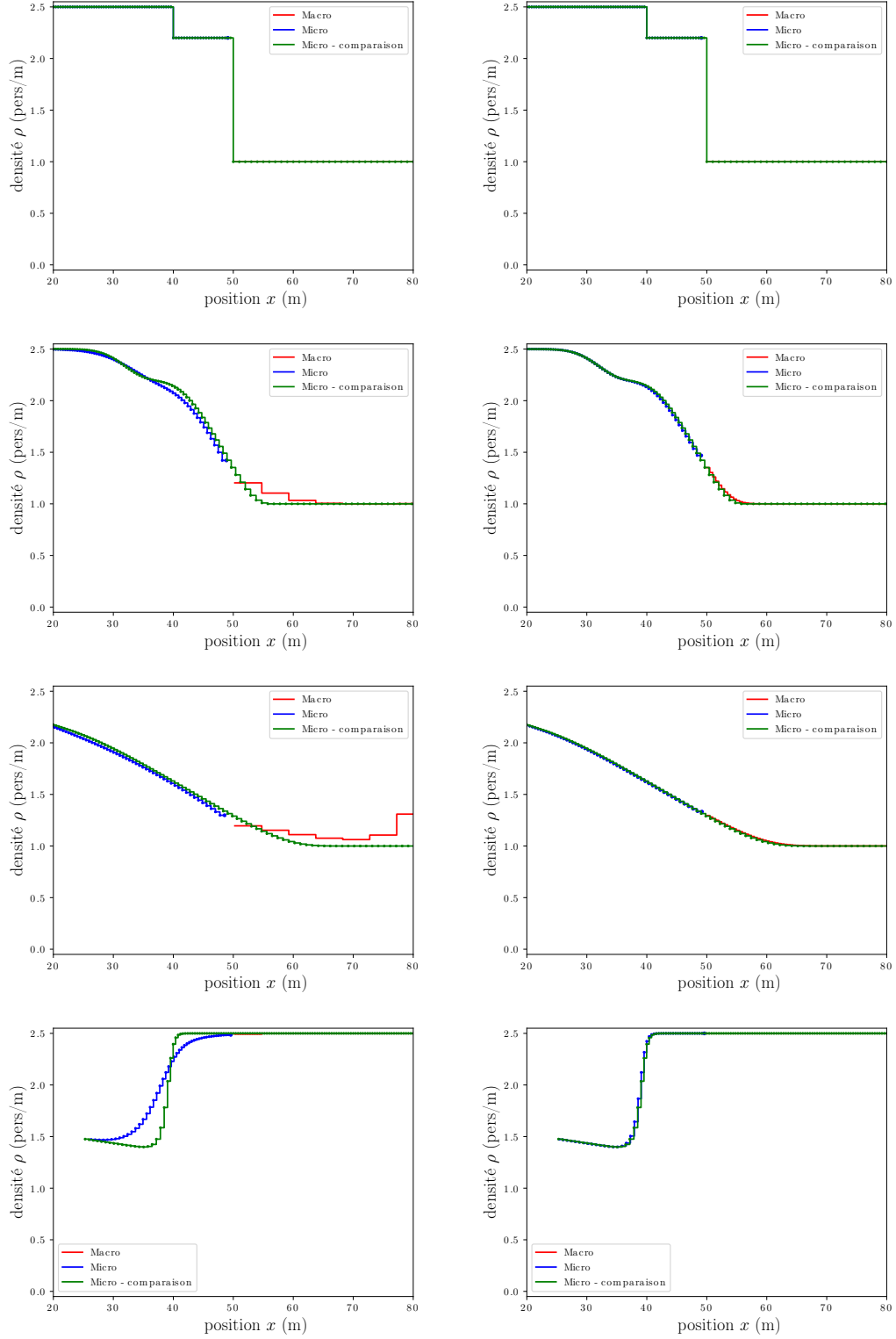


FIGURE 3.6 – Résultats de simulation pour le couplage micro-macro avec $\Delta x = 5$ m (gauche) et 0.5 m (droite). Les deux discrétisations donnent des résultats qualitativement proches. Le pas d'espace le plus fin produit les résultats les plus proches du modèle FTL seul, tandis que sur les mailles plus grandes des différences notables apparaissent.

On suppose qu'il existe un instant $t_0 \geq 0$ pour lequel on a $X_{L(t_0)}(t_0) = 0$. Si $x_g(t)$ est fini, on pose alors $X_{L(t_0)-1}(t_0) = x_g(t_0)$. On pose ensuite :

$$\dot{X}_{L(t_0)-1} = \varphi(X_{L(t_0)} - X_{L(t_0)-1}) \quad \forall t \in]t_0, t_1[, \quad (3.42)$$

$$f(t, 0^-) = \frac{1}{X_{L(t_0)} - X_{L(t_0)-1}} \dot{X}_{L(t_0)-1} \quad \forall t \in]t_0, t_1[, \quad (3.43)$$

où $t_1 = \sup\{t, X_{L(t_0)-1}(t) < 0\}$, soit le temps où le piéton généré atteint l'interface. Si t_1 est fini (donc si le nouvel individu atteint bien l'interface et n'est pas bloqué avant), on aura $L(t_1) = L(t_0) - 1$, et on se retrouvera dans une situation similaire à celle que l'on avait à t_0 , recommençant ainsi une autre procédure d'émission. Si $x_g(t_0)$ n'est pas fini, on remplace $\varphi(X_{L(t_0)} - X_{L(t_0)-1})$ par la vitesse maximale U , et on aura $t_1 = \infty$.

Pour compléter notre description on doit préciser ce qui est fait en $t = 0$. On choisit ici de prendre $X_{L(0)-1}(0) = x_g(0)$.

Schéma numérique. Nous reprenons la discrétisation des équations (3.5) et (3.6). L'écriture du schéma numérique nécessite de retranscrire la génération possible d'un individu lagrangien entre deux pas de temps.

Pour un temps t où on calcule le modèle microscopique hiérarchiquement en partant de la droite de l'interface, le calcul de $X_{L^n-1}^{n+1}$ se fait sans dépendre du domaine macroscopique :

$$X_{L^n-1}^{n+1} = X_{L^n-1}^n + \Delta t \varphi(X_{L^n}^{n+1} - X_{L^n-1}^{n+1}). \quad (3.44)$$

Il faut ensuite gérer le cas de dépassement possible de l'interface :

- si $X_{L^n-1}^{n+1} > 0$, cela signifie que t_1 est entre t^n et t^{n+1} , et la valeur de L change à t^{n+1} . On doit alors générer un individu $X_{L^{n+1}-1}^{n+1}$. Pour ce faire, $X_{L^{n+1}-1}^{n+1}$ est initialisé à $x_g(t_1)$ puis on applique le schéma numérique du modèle FTL entre t_1 et le prochain pas de temps.
- si ce n'est pas le cas, $X_{L^{n+1}-1}^{n+1} = X_{L(t^n)-1}^{n+1}$.

Le flux numérique est ensuite discrétisé tel que :

$$F_{-1/2}^n = \frac{1}{X_{L(t_0)}^n - X_{L(t_0)-1}^n} \frac{(X_{L^n-1}^n - X_{L^n-1}^{n+1})}{\Delta t} \quad \forall t \in]t_0, t_1[, \quad (3.45)$$

Tests numériques. Nous reprenons en figure 3.7 les mêmes cas tests et paramètres que dans tableau 3.2.

Le modèle couplé LWR vers FTL converge là aussi vers le modèle FTL seul. On observe ici aussi une absence d'oscillation, que l'on peut encore une fois comprendre en remarquant que lorsque la densité est constante, la procédure d'émission génère un flux stable.

La convergence montre la pertinence du schéma numérique couplé. En revanche dans un cadre applicatif on souhaiterait utiliser des mailles de plus grandes tailles, ainsi que des schémas numériques potentiellement de plus haut ordre. Le schéma que nous avons montré pourra être adapté, mais peut aussi donner des conditions au bord en microscopique et macroscopique tout en utilisant des schémas numériques plus adaptés dans le reste du domaine.

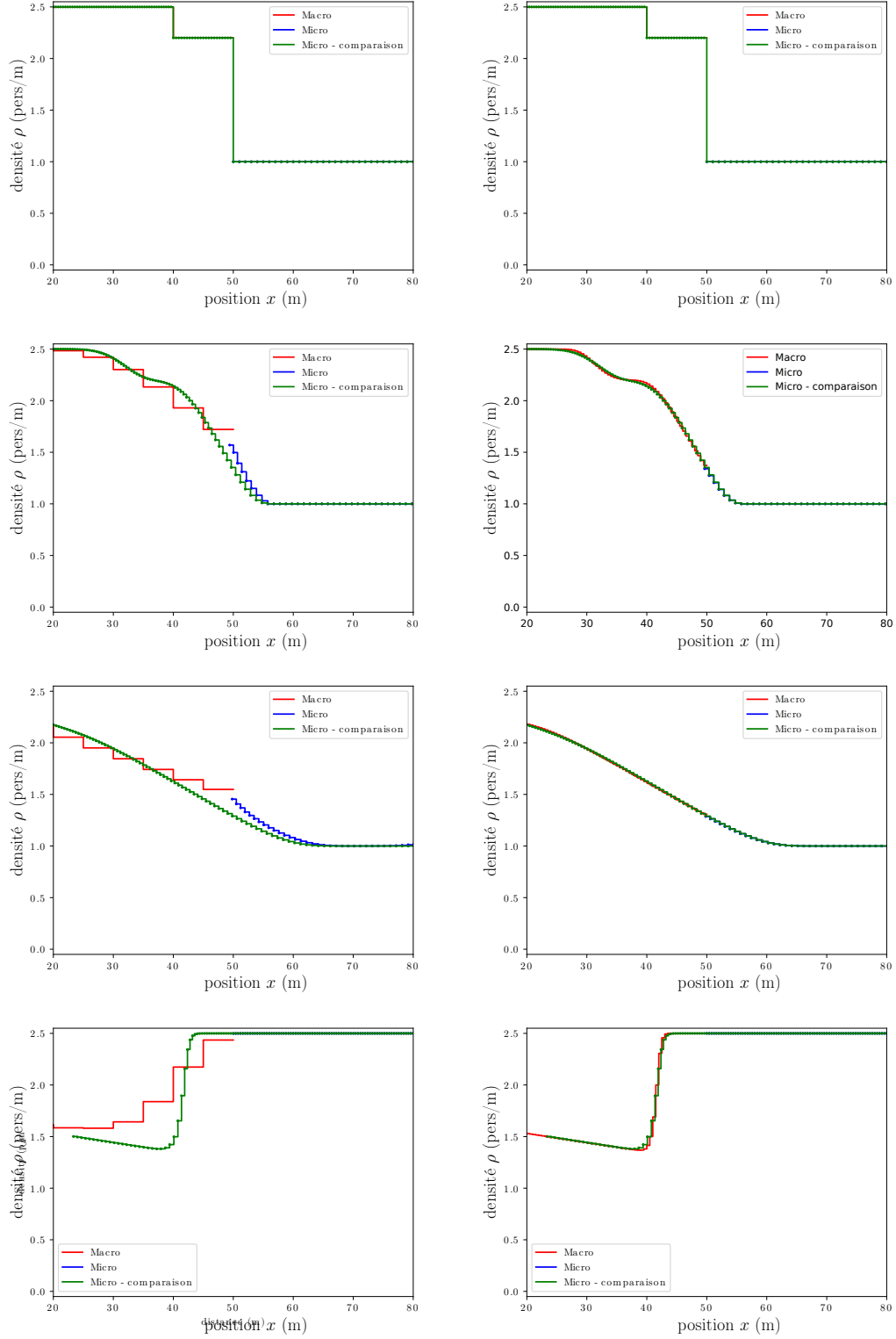


FIGURE 3.7 – Résultats de simulation pour le couplage macro-micro avec $\Delta x = 5$ m (gauche) et 0.5 m (droite). Les deux discrétisations donnent des résultats qualitativement proches. Le pas d'espace le plus fin produit les résultats les plus proches du modèle FTL seul, tandis que sur les mailles plus grandes des différences notables apparaissent.

3.2 Couplage de modèles d'inhibition

Dans le cadre des mouvements de foules à haute densité, nous avons montré dans le Chapitre 2 l'intérêt en modélisation mathématique et en simulation numérique d'utiliser des modèles granulaires.

Les modèles granulaires purs s'écrivent en projetant au sens des moindres carrés les vitesses souhaitées des individus sur un ensemble de vitesses qui n'entraînent pas de violation de la contrainte micro ou macro. Cette représentation du mouvement considère que les individus sont prêts à pousser pour atteindre leur vitesse souhaitée, un comportement peu observé dans des cas d'évacuations ordonnées. Une amélioration du modèle microscopique a été introduite ([Red17]) dans le but d'intégrer une notion d'inhibition individuelle, empêchant aux individus de se bousculer à tout prix. Une version macroscopique de ce modèle existe à une dimension seulement.

Nous étudions dans cette section le couplage pour ces modèles à une dimension, où les individus se déplacent et regardent vers les x positifs. Nous allons de plus montrer des liens que nous avons pu faire entre ces modèles, d'inspiration principalement granulaire, et les modèles basés sur l'anticipation introduits précédemment.

3.2.1 Présentation des modèles d'inhibition

Modèle microscopique. Dans le modèle microscopique d'inhibition, chaque individu décide de sa vitesse en fonction de la vitesse de son prédécesseur. La contrainte de non-recouvrement est résolue localement, et non globalement comme cela est le cas dans le modèle granulaire pur. Cette résolution locale se fait en partant du piéton le plus loin devant et en remontant jusqu'au dernier.

On se place dans une situation à N individus de positions $(X_i)_i$, de rayon r et de vitesses souhaitées $(U_i)_i$. Les individus doivent respecter une contrainte de non-recouvrement, c'est-à-dire qu'à tout instant les positions $X = (X_1, \dots, X_N)$ doivent appartenir à l'ensemble K tel que :

$$K = \{X \in \mathbb{R}^N, X_{i+1} - X_i \geq 2r, \quad 1 \leq i < N\} \quad (3.46)$$

On introduit pour chaque individu l'ensemble des vitesses qui ne vont pas conduire à une violation de cette contrainte de non-recouvrement $\mathcal{C}_i$:

$$\mathcal{C}_i = \{v \in \mathbb{R}, X_{i+1} - X_i = 2r \Rightarrow v - v_{i+1} \leq 0\}. \quad (3.47)$$

Il s'agit ici d'une version locale de l'ensemble de vitesses admissibles dans le cas du modèle granulaire (voir chapitre 2)

La vitesse effective u_i d'un individu sera la projection de sa vitesse souhaitée sur cet ensemble :

$$u_i = \operatorname{argmin}_{w \in \mathcal{C}_i} |w - U_i|^2 \quad \forall i < N \quad (3.48)$$

À tout instant l'ensemble des vitesses $(u_i)_{i < N}$ est bien défini, car il est possible d'effectuer la projection de l'équation (3.48) pour chaque individu hiérarchiquement en partant de $i = N$ et en redescendant les individus jusqu'au dernier $i = 1$.

Une différence notable avec le modèle FTL est que la vitesse d'un individu n'est pas déterminée uniquement par la distance avec son prédécesseur, mais par la vitesse de ce voisin. Lorsque plusieurs individus forment un cluster, ce cluster se déplacera à la vitesse de l'individu le plus en avant du cluster (si celui-ci a la vitesse la plus faible).

On peut discrétiser ces équations en remarquant que la contrainte de non-recouvrement s'écrit pour les positions discrétisées :

$$X_{i+1}^n - X_i^n = 2r \implies \frac{X_i^{n+1} - X_i^n}{\Delta t} \leq \frac{X_{i+1}^{n+1} - X_{i+1}^n}{\Delta t}. \quad (3.49)$$

Cela revient à prendre :

$$X_i^{n+1} \leq X_{i+1}^{n+1} - 2r, \quad (3.50)$$

et la projection des vitesses a une solution simple :

$$\frac{X_i^{n+1} - X_i^n}{\Delta t} = \min \left(U_i, \frac{X_{i+1}^{n+1} - 2r - X_i^n}{\Delta t} \right).$$

Modèle macroscopique. Le modèle granulaire avec inhibition précédent n'a pas d'équivalent macroscopique identifié très clairement pour le moment : il existe un modèle macroscopique basé sur une contrainte de densité maximale ([Rou11]), mais il n'y a pas de modèle analogue au modèle granulaire avec inhibition qui soit écrit et analysé mathématiquement. Des premières étapes pour l'écriture de ce modèle ont été proposées dans la thèse de F. Al Reda ([Red17]), et dans le chapitre 5 nous proposons des développements numériques sur ce modèle. En dimension 1, bien que le modèle ne soit pas identifié une expression du champ de vitesse peut être donnée.

On se place sur un intervalle I de $\mathbb{R}$. Soit $U(x)$ la vitesse souhaitée en tout point x de I , $\rho_{max} \in \mathbb{R}$ la densité maximale qui ne devra pas être dépassée. Pour respecter cette contrainte dans le temps, la vitesse effective v doit être à divergence positive sur les zones saturées (c'est-à-dire celles où $\rho = \rho_{max}$). À une dimension, cette contrainte sur v signifie simplement qu'elle doit être croissante sur les zones saturées.

Dans le modèle granulaire macroscopique sans inhibition, le champ de vitesse effectif est donné par une projection du champ U sur l'ensemble des vitesses respectant cette contrainte. Dans certains cas, cette projection peut cependant donner des vitesses localement plus élevées que la vitesse souhaitée U . Cela signifie alors qu'à ces endroits, les individus sont localement poussés par ceux présents derrière eux. Pour modéliser l'effet de la politesse, la vitesse effective du modèle macroscopique d'inhibition ne devrait à l'inverse pas dépasser la vitesse souhaitée, comme illustré en figure 3.8

À une dimension pour chaque zone saturée, on peut donner l'expression d'une vitesse respectant cette condition supplémentaire :

$$v(x) = \min_{y \geq x, \rho(y) = \rho_{max}} U(y).$$

Chaque zone saturée est traitée séparément pour le calcul de la vitesse.

Résolution numérique. Nous proposons une approche numérique du problème basée sur la version discrétisée des contraintes présentées ainsi qu'une résolution hiérarchique du problème, à

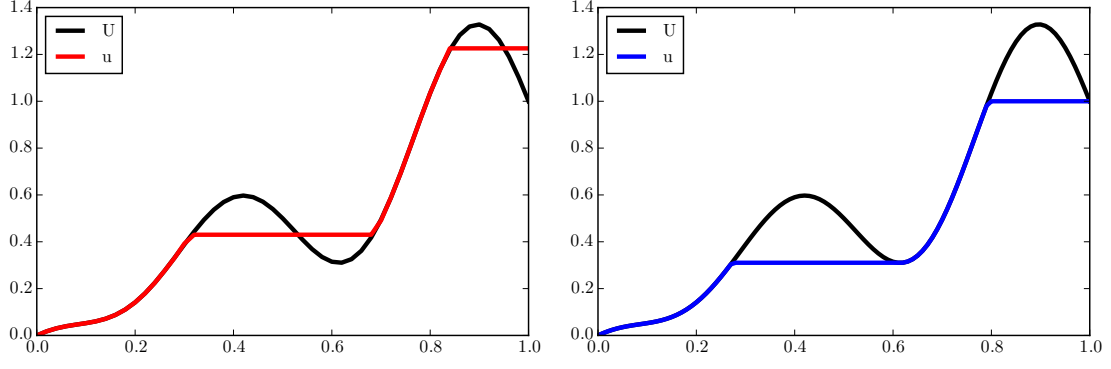


FIGURE 3.8 – Vitesse souhaitée et vitesse effective dans le cas du modèle granulaire et du modèle d’inhibition. Figure tirée de [Red17]

l’instar du modèle d’inhibition microscopique. Temps et espace sont discrétisés avec pas de temps Δt , un pas d’espace Δx et on adopte un schéma de type volumes finis :

$$\frac{\rho_j^{n+1} - \rho_j^n}{\Delta t} + \frac{F_{j+1/2}^n - F_{j-1/2}^n}{\Delta x} = 0, \quad (3.51)$$

où ρ_j^n est la valeur approchée de ρ sur la cellule j pendant entre $n\Delta t$ et $(n+1)\Delta t$ et $F_{j+1/2}^n$ le flux à l’interface entre les mailles j et $j+1$ dont nous allons préciser l’expression par la suite.

Les flux sont calculés successivement en itérant en j décroissant. La contrainte de non-dépassement de la densité maximale $\rho_i^n \leq \rho_{max}$ devient pour une itération entre n et $n+1$:

$$\rho_j^{n+1} \leq \rho_{max} \implies \rho_j^n - \frac{\Delta t}{\Delta x} (F_{j+1/2}^n - F_{j-1/2}^n) \leq \rho_{max}.$$

En exprimant $F_{j-1/2}^n$ avec un schéma décentré, on a :

$$F_{j-1/2}^n = \rho_{j-1}^n v_{j-1/2}^n,$$

où $v_{j-1/2}^n$ est la vitesse à l’interface entre deux mailles, que l’on exprime :

$$v_{j-1/2}^n = \underset{v \in \mathcal{C}_j^n}{\operatorname{argmin}} |v - U_{j-1/2}|^2, \quad (3.52)$$

où :

$$\mathcal{C}_j^n = \left\{ v \in \mathbb{R}, \rho_j^n - \frac{\Delta t}{\Delta x} (\rho_j^n v_{j+1/2}^n - \rho_{j-1}^n v) \leq \rho_{max} \right\}.$$

Notons que l’expression de la contrainte fait apparaître $v_{j+1/2}^n$ qui est connu au moment du calcul de $v_{j-1/2}^n$ lorsque l’on itère en j décroissant. Pour exprimer l’ensemble des vitesses, on résout frontalement chacun de ces problèmes de minimisation en partant de $j = N_x$ si N_x est le nombre de cellules en espace. Le terme $v_{N_x+1/2}^n$ est supposé connu (donné par une condition au bord), et l’équation (3.52) devient en $j = N_x - 1$ la minimisation d’une fonction convexe sous une contrainte linéaire. On peut exprimer par récurrence la valeur de $v_{j+1/2}^n$ pour tout j . On peut même exprimer la solution de cette projection pour chaque cellule :

$$v_{j-1/2}^n = \min \left(U_{j-1/2}^n, \frac{\Delta x}{\rho_{j-1}^n} \left(\frac{\rho_{max} - \rho_j^n}{\Delta t} - \frac{\rho_j^n v_{j+1/2}^n}{\Delta x} \right) \right). \quad (3.53)$$

Pour ce modèle, si un cluster de densité maximale se forme, la vitesse sera limitée par la vitesse de l'avant du cluster, comme pour le modèle microscopique. De la même façon, si la vitesse s'annule en un point les vitesses de tous les individus en amont dans le même cluster s'annuleront aussi instantanément. Dans le schéma numérique, si la vitesse pour une maille est nulle, la projection de toutes les cellules en amont en sera impactée. Si ces cellules sont toutes saturées, les vitesses s'annuleront aussi en amont en une seule itération.

3.2.2 Limite raide de modèles d'anticipation

Avant de s'intéresser plus en détail aux couplages des modèles de contacts en eux-mêmes, nous présentons ici quelques résultats faisant le lien entre des modèles d'anticipation et des modèles de contact. L'idée principale est de montrer comment à partir des modèles FTL et LWR on peut identifier des comportements de modèles de contacts en faisant tendre la raideur des fonctions vitesses aux abords de la densité maximale.

Nous reprenons les fonctions vitesse utilisées pour les modèles FTL et LWR, en prenant une autre forme faisant intervenir un paramètre de raideur ε :

$$\varphi_\varepsilon(w) = U \left(1 - \exp\left(-\frac{w - 2r}{\varepsilon}\right) \right), \text{ pour } w \geq 2r \quad (3.54)$$

$$v_\varepsilon(\rho) = U \left[1 - \exp \left(- \left(\frac{1}{\rho} - \frac{1}{\rho_{max}} \right) / \varepsilon \right) \right], \text{ pour } 0 < \rho \leq \rho_{max}. \quad (3.55)$$

On peut observer (figure 3.9) que lorsque $\varepsilon \rightarrow 0$ on obtient pour φ_ε une fonction vitesse dont la pente en $2r$ (respectivement ρ_{max} pour la fonction v_ε) tend vers ∞ .

FTL. On étudie le comportement du système pour deux agents. On s'intéresse au système d'équations (3.56) - (3.59) :

$$\dot{X}_2(t) = V(t), \quad (3.56)$$

$$\dot{X}_1^\varepsilon(t) = \varphi_\varepsilon(X_2(t) - X_1^\varepsilon(t)), \quad (3.57)$$

$$X_2(0) = \bar{X}_2 \in \mathbb{R}, \quad (3.58)$$

$$X_1^\varepsilon(0) = \bar{X}_1 \in \mathbb{R}, \quad (3.59)$$

$$\bar{X}_2 - \bar{X}_1 > 2r, \quad (3.60)$$

où V est une fonction à valeurs continues dans $[0, U]$. On a tout d'abord une première proposition :

Proposition 3.1. *Les équations (3.56) - (3.59) avec la relation d'ordre (3.60) ont une solution maximale unique sur $[0, T]$ pour tout $\varepsilon > 0$ et $T > 0$.*

Preuve. Quelque soit $\varepsilon > 0$, la fonction φ_ε étendue par 0 sur $] - \infty, 2r[$ est globalement Lipschitz. On peut appliquer le théorème de Cauchy-Lipschitz au système de Cauchy (3.56) - (3.59) sur $[0, +\infty[\times \mathbb{R}^2$.

On regarde ensuite la limite quand $\varepsilon \rightarrow 0$.

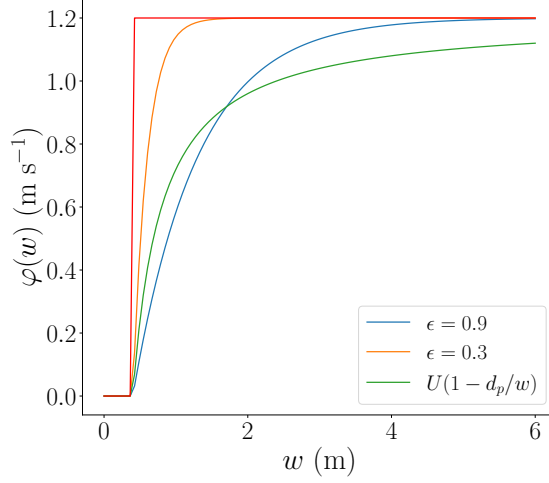


FIGURE 3.9 – Fonctions de vitesses de plus en plus raides

Proposition 3.2. *Il existe une limite X_1^* vers laquelle converge $\{X_1^\varepsilon\}_{\varepsilon>0}$ à une sous-suite près sur $[0, T]$ lorsque $\varepsilon \rightarrow 0$.*

Preuve. D’après le théorème d’Arzela-Ascoli, si une suite de fonctions à valeurs réelles f_ε définies sur un intervalle fermé borné $[a, b]$ est uniformément bornée et équicontinue, alors il existe une sous-suite uniformément convergente. On peut donc appliquer ce théorème à la séquence $\{X_1^\varepsilon\}_{\varepsilon>0}$ et obtenir la convergence.

Lorsque la distance entre les deux individus est supérieure à $2r$, la vitesse du second individu tend vers U lorsque ε tend vers 0. Si un contact arrive cependant le comportement de la solution limite est difficile à déterminer.

LWR. La limite raide pour ce type de modèles est bien plus complexe à obtenir. Nous décrivons à la place le comportement des solutions du problème de Riemann pour illustrer la difficulté du problème. Notons tout d’abord que si on reprend une équation du modèle LWR avec vitesse v_ε :

$$\partial_t \rho + \partial_x(\rho v_\varepsilon(\rho)) = 0, \quad (3.61)$$

et que l’on remplace directement v_ε par sa limite en $\varepsilon \rightarrow 0$ pour une densité $\rho < \rho_{max}$, on obtient une équation de transport à vitesse constante U . Des situations plus complexes apparaissent cependant lorsque l’on a une densité ρ_{max} .

On s’intéresse donc pour l’équation (3.61) au problème de Riemann :

$$\rho_0(x) = \begin{cases} \rho_L & \text{si } x < 0, \\ \rho_R & \text{si } x > 0, \end{cases}$$

où ρ_L ou ρ_R vaut ρ_{max} . Pour tout $\varepsilon > 0$, la solution entropique de ce problème est donnée par la théorie des systèmes hyperboliques, voir par exemple [Ser99].

Dans le cas où $\rho_R = \rho_{max}$, la solution est un choc de vitesse σ donnée par la relation de Rankine-Hugoniot, avec le flux $f_\varepsilon(\rho) = \rho v_\varepsilon(\rho)$:

$$\sigma = \frac{f_\epsilon(\rho_R) - f_\epsilon(\rho_L)}{\rho_R - \rho_L} = \frac{-f_\epsilon(\rho_L)}{\rho_{max} - \rho_L}.$$

On aura donc un choc dont la vitesse tend en ϵ vers une vitesse limite

$$\sigma^* = \frac{-\rho_L U}{\rho_{max} - \rho_L}$$

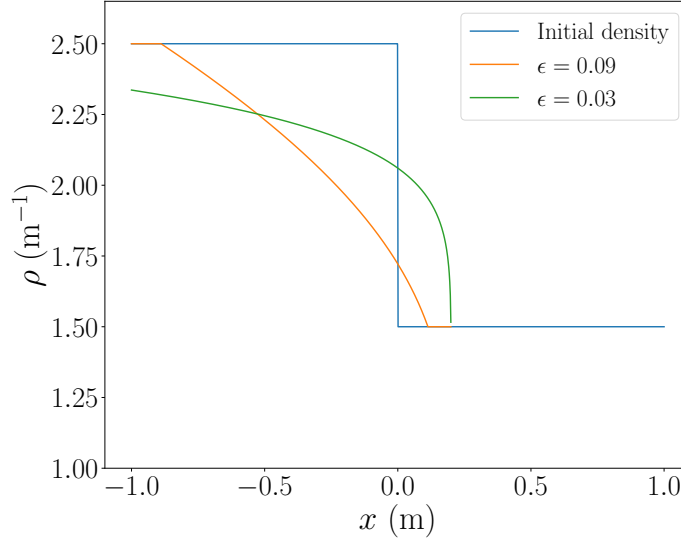


FIGURE 3.10 – Solution du problème de Riemann pour différents ϵ , avec $\rho_L = \rho_{max} = 2.5m^{-1}$ et $\rho_R = 1.5m^{-1}$, à $t = 0.2 s$ et $U = 1 m/s$: on observe la formation d'une onde de raréfaction de plus en plus raide lorsque $\epsilon \rightarrow 0$.

On s'intéresse maintenant au cas inverse, où $\rho_R < \rho_{max}$ et $\rho_L = \rho_{max}$. La solution pour un ϵ donné est une onde de raréfaction telle que :

$$\rho(x, t) = \begin{cases} \rho_{max} & \text{si } x/t \leq f'_\epsilon(\rho_{max}), \\ (f'_\epsilon)^{-1}(x/t) & \text{si } f'_\epsilon(\rho_{max}) \leq x/t \leq f'_\epsilon(\rho_R), \\ \rho_R & \text{si } f'_\epsilon(\rho_R) \leq x/t, \end{cases} \quad (3.62)$$

avec :

$$f'_\epsilon(\rho) = U \left[1 - \left(1 + \frac{1}{\epsilon \rho} \right) \exp \left(-\left(\frac{1}{\rho} - \frac{1}{\rho_{max}} \right) / \epsilon \right) \right].$$

On observe en particulier que pour $\epsilon \rightarrow 0$ on a $f'_\epsilon(\rho_{max}) \rightarrow -\infty$, c'est-à-dire que l'information se propage en remontant le courant à vitesse infinie. Ce phénomène empêche l'utilisation des techniques classiques reposant sur l'analyse des solutions de Cauchy du problème de Riemann.

3.2.3 Couplage des modèles de contact

Microscopique vers macroscopique. Revenons sur le problème du couplage microscopique vers macroscopique avec les modèles d'inhibition. Les piétons $1, \dots, L(t)$ à gauche de l'interface $x = 0$ sont décrits par le modèle d'inhibition, et à droite de l'interface la densité évolue selon la version macroscopique du modèle. Il faut alors donner une vitesse pour le piéton $L(t)$ en tout temps du modèle microscopique ainsi que des conditions au bord pour le modèle à droite.

Notre étude précédente sur l'importance du choix de la densité au bord ρ_b vue par le modèle macroscopique a montré l'intérêt de choisir :

$$\rho_b(t) = \frac{1}{|X_{L(\tau)}(\tau)|},$$

avec :

$$\tau(t) = \begin{cases} 0 & \text{si } L(t) = L(0) \\ \sup \{t, X_{L(t)+1} < 0\} & \text{sinon.} \end{cases}$$

Ce choix permet d'attribuer une densité diffuse à chaque individu devant traverser l'interface. Le choix de la vitesse du leader doit permettre de transmettre l'information de la vitesse lorsqu'un bouchon se propage et la vitesse du leader doit s'adapter au contact de l'interface. Pour prendre en compte ces effets, lorsqu'un bouchon est présent au niveau de l'interface la vitesse du leader doit vérifier :

$$V_L(t) = \min_{y \geq x, \rho(y) = \rho_{max}} U(y),$$

et $U(0)$ sinon. En cas de bouchon, le leader adopte la vitesse des individus du bouchon. Avec ce choix le flux au bord pour le modèle macroscopique vaut $f(t, 0^+) = \rho_b(t)V_L(t)$.

Discrétisation et tests numériques. De la même façon que pour le couplage des modèles FTL vers LWR, un sous-pas de temps Δt_k est nécessaire pour prendre en compte les changements de la densité $\rho_b(t)$ lors de passages de l'interface par un individu. Pour calculer la vitesse du leader, nous raisonnons de la même façon que pour obtenir le schéma numérique pour le modèle macroscopique, en étudiant la contrainte de densité maximale pour la première cellule macroscopique.

En notant $V_L^{n,k}$ la vitesse du leader et $\rho_b^{n,k}$ la valeur discrétisée de $\rho_b(t)$, la contrainte de densité maximale pour la première cellule macroscopique à l'interface lors du passage du leader pendant un temps Δt devient :

$$\rho_0^n - \frac{\Delta t}{\Delta x} (F_{1/2}^n - \rho_b^{n,k} V_L) \leq \rho_{max} \quad (3.63)$$

On peut donc écrire directement la vitesse souhaitée du leader en prenant en compte la contrainte :

$$V_{L^{n,k}}^{n,k} = \min \left(U_L, \left[\frac{\Delta x}{\Delta t} (\rho_{max} - \rho_0^n) + F_{1/2}^n \right] / \rho_b^n \right), \quad (3.64)$$

où U_L est la vitesse maximale du leader. Le reste du schéma s'écrit en reprenant le schéma numérique du couplage entre les modèles FTL et LWR. Le sous-pas de temps est donné par :

$$\Delta t_k = \min \left\{ \Delta t - \sum_{j < k} \Delta t_j, \frac{|X_{L^{n,k}}^{n,k}|}{V_L^{n,k}} \right\}, \quad (3.65)$$

et les positions des individus sont données par :

$$X^{n,0} = X^n, \quad (3.66)$$

$$X_i^{n,k+1} = X_i^{n,k} + \Delta t^k V_i^{n,k} \quad \forall i < L^{n,k}, \quad (3.67)$$

$$V_i^{n,k} = \min \left(U_i, \frac{X_{i+1}^{n,k+1} - 2r - X_i^{n,k}}{\Delta t} \right), \quad (3.68)$$

$$X_{L^{n,k}}^{n,k+1} = X_{L^{n,k}}^{n,k} + \Delta t^k V_{L^{n,k}}^{n,k}, \quad (3.69)$$

$$X^{n+1} = X^{n,K^n}, \quad (3.70)$$

et K^n est tel que :

$$\sum_{k \leq K^n} \Delta t_k = \Delta t. \quad (3.71)$$

Le flux et la densité à l'interface sont donnés par :

$$\rho^{n,k} = \frac{1}{|X_{L^{n,k}}^{n_0,k_0}|}, \quad n_0, k_0 = \{n, k / X_{L^{n,k}+1}^{n,k} = 0\}, \quad (3.72)$$

$$F_{-1/2}^n = \frac{1}{\Delta t} \sum_k \Delta t_k \rho^{n,k} V_{L^{n,k}}^{n,k}. \quad (3.73)$$

et le reste du modèle macroscopique est donné par (3.51).

Des résultats de simulation pour ce schéma sont présentés en figure 3.11, pour un cas de densité uniforme de part et d'autre de l'interface avec une remontée d'un bouchon. La remontée d'information se fait correctement lorsque la densité atteint sa valeur maximale.

Couplage macroscopique vers microscopique. Le couplage dans le sens inverse fonctionne de la même manière que dans le cas LWR vers FTL : le seul changement est dans le calcul du déplacement des individus microscopiques et les calculs en dehors de l'interface.

Le résultat d'une simulation de propagation de bouchon est représenté en figure 3.12.

3.3 Cas d'application et extensions

3.3.1 Cas d'application : fouille ponctuelle

Nous nous intéressons ici à un cas d'application potentiel de la procédure de couplage développée. Notre but est de montrer l'avantage d'un point de vue du modélisateur de pouvoir choisir localement le modèle à adopter en fonction de la situation. Nous avons choisi pour ça de modéliser le cas d'une fouille de sac ponctuelle. On se place sur un tronçon de chemin de 60m représentant l'entrée d'un stade par exemple. On suppose qu'en un point à 30m, un point de contrôle est installé et arrête 1 personne sur 5. Ce contrôle prend un temps fixe de 8 secondes.

On modélise cette condition en forçant 1 piéton sur 5 à adopter une vitesse nulle pendant 8 secondes en atteignant les 30m. Pour cela il est nécessaire d'avoir un modèle microscopique au niveau du point de contrôle, mais l'on peut garder un modèle macroscopique au-delà.

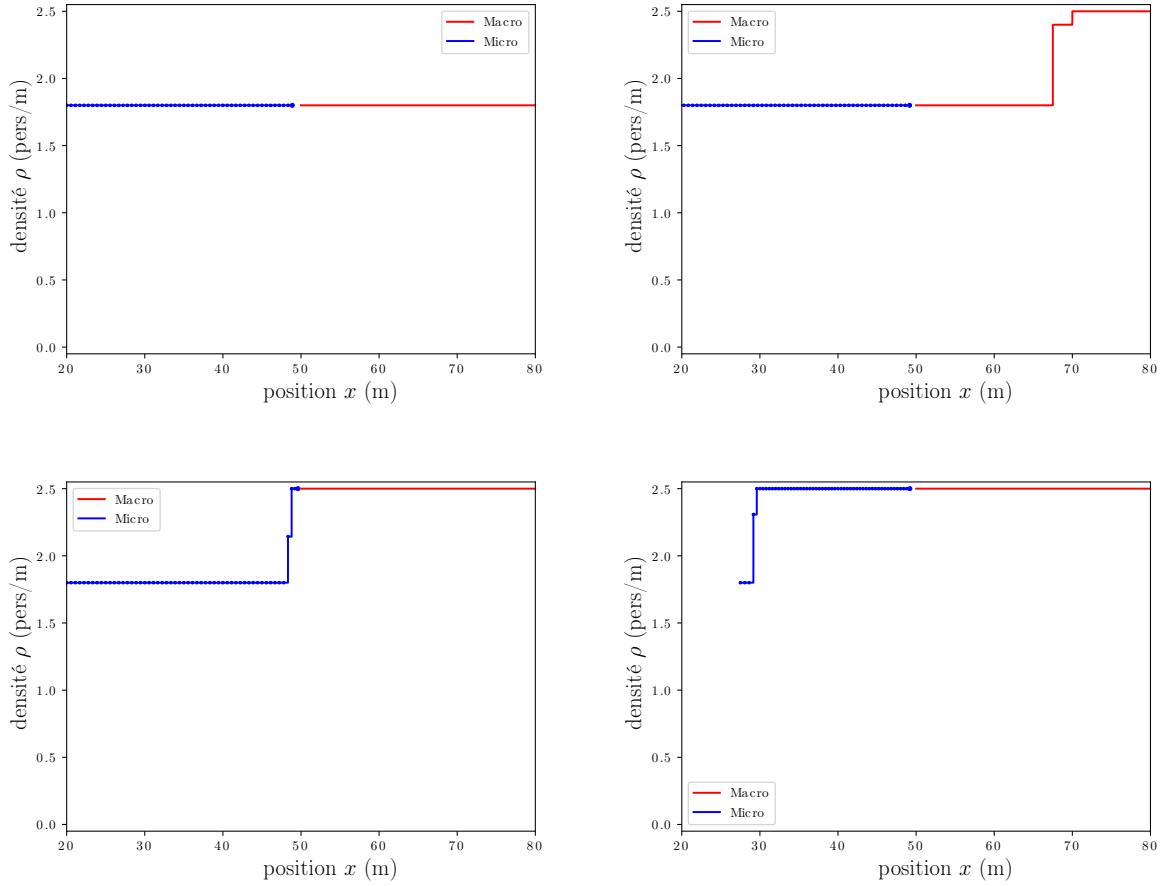


FIGURE 3.11 – Résultats de simulation du couplage microscopique vers macroscopique pour les modèles d’inhibition, à $t = 0, 10, 16, 23$ s et avec $\Delta t = 1$ s, $\Delta x = 2.5$ m et $U = 1.2$ m/s.

Remarque 3.3. *Il est possible de modéliser de façon macroscopique l’effet de cette fouille, en diminuant localement le flux admissible par exemple ([CG07]). L’approche présentée ici permet en revanche une description microscopique de l’interaction, qui peut s’avérer importante si l’on souhaite aller plus loin dans le modèle de temps d’attente, en incluant des comportements propres aux individus.*

Les autres paramètres propres à la modélisation et la discrétisation sont indiqués tableau 3.3.

On montre en figure 3.13 les résultats de la simulation avec les modèles LWR et FTL. On voit que la présence du poste de contrôle entraîne l’apparition de vagues de congestion (faisant penser à un phénomène de *stop-and-go*) qui remontent le courant des agents, et à l’inverse en aval des vagues de faibles densités. Notons que l’on est capable d’observer ces vagues grâce au maillage suffisamment fin pour les observer, avec un maillage moins fin on ne pourrait observer que les effets de congestion augmentée en amont et de décongestion en aval, sans capturer les oscillations.

Les résultats de figure 3.14 sont obtenus en prenant les modèles d’inhibition raide à la place. On considère ainsi la même situation en supposant que le comportement des individus est moins

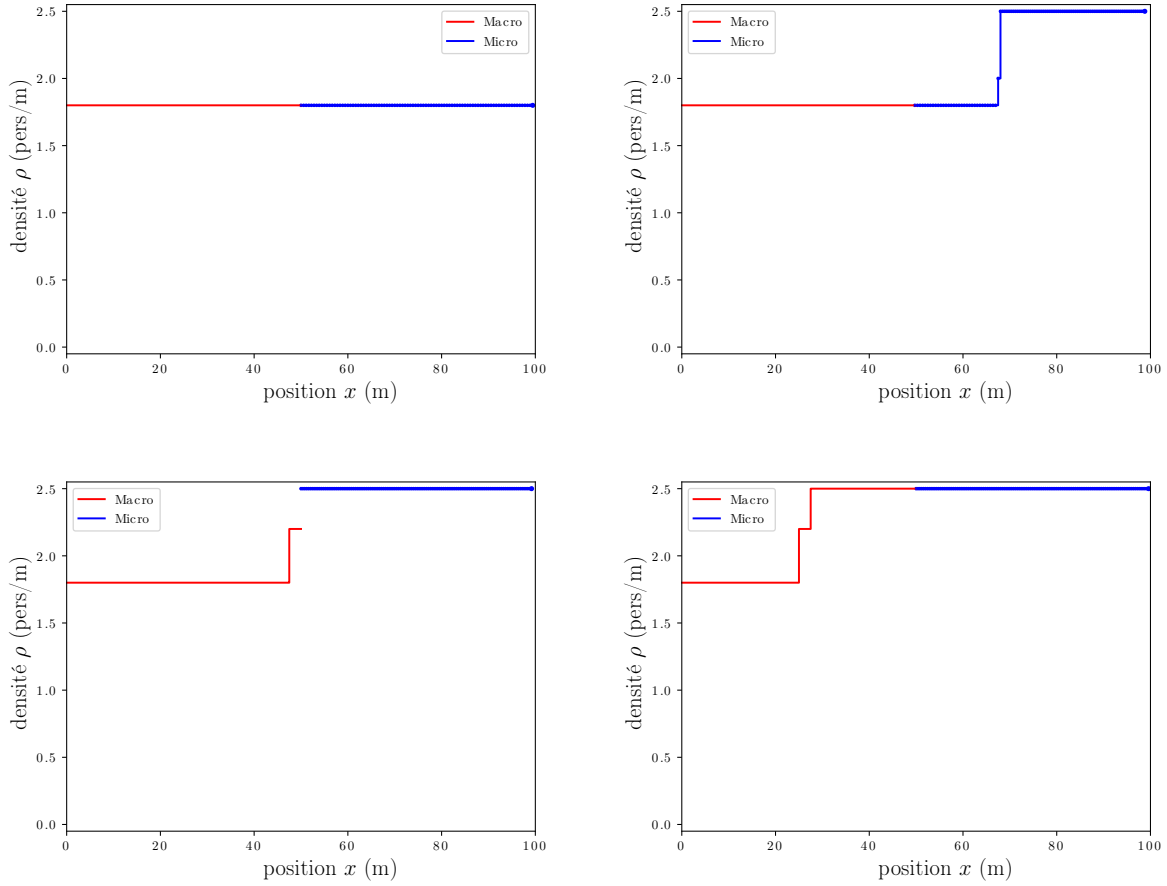


FIGURE 3.12 – Résultats de simulation couplage macroscopique vers microscopique pour les modèles d’inhibition, à $t = 0, 10, 16, 23$ s et avec $\Delta t = 1$ s, $\Delta x = 2.5$ m et $U = 1.2$ m/s.

anticipatif et plus pressé. On retrouve naturellement des vagues en aval, mais plus de vagues en amont du point de contrôle. À la place on observe un bloc de densité maximale qui se forme au fur et à mesure. Dans ce cas-là on s’attend donc à une congestion plus élevée et plus concentrée, tout en gardant par moment des vitesses élevées et donc une situation plus dangereuse.

L’utilisation du couplage au sein de chaque simulation permet ainsi de récupérer des informations locales sur une échelle globale. La comparaison de deux modèles différents nous permet d’obtenir une vision multiple des évolutions possibles : dans le cas d’une foule ordonnée, le poste de contrôle va induire une élévation de la densité en amont qui va se propager en remontant le courant, créant potentiellement des vagues de *stop-and-go*. Pour une foule plus pressée et nerveuse, on peut en revanche s’attendre à une congestion très importante juste avant le point de contrôle. Ces deux informations de natures différentes pourraient être utiles pour des organisateurs d’évènement, pouvant tester l’effet de stratégies différentes selon la configuration testée.

paramètre	valeur
$v(\rho)$	$U(1 - \frac{\rho}{\rho_{max}})$
U	$1.2m/s$
ρ_{max}	$2.5m^{-1}$
ρ_0	$2.5 \times \mathbf{1}_{[0,40]} + 2.2 \times \mathbf{1}_{[40,50]} + 1.0 \times \mathbf{1}_{[50,100]}$
Δx	$0.5 - 5.m$
Δt	$0.1 - 1.0s$

TABLE 3.3 – Paramètres pour les cas d’application considérés en Figures 3.13, 3.14

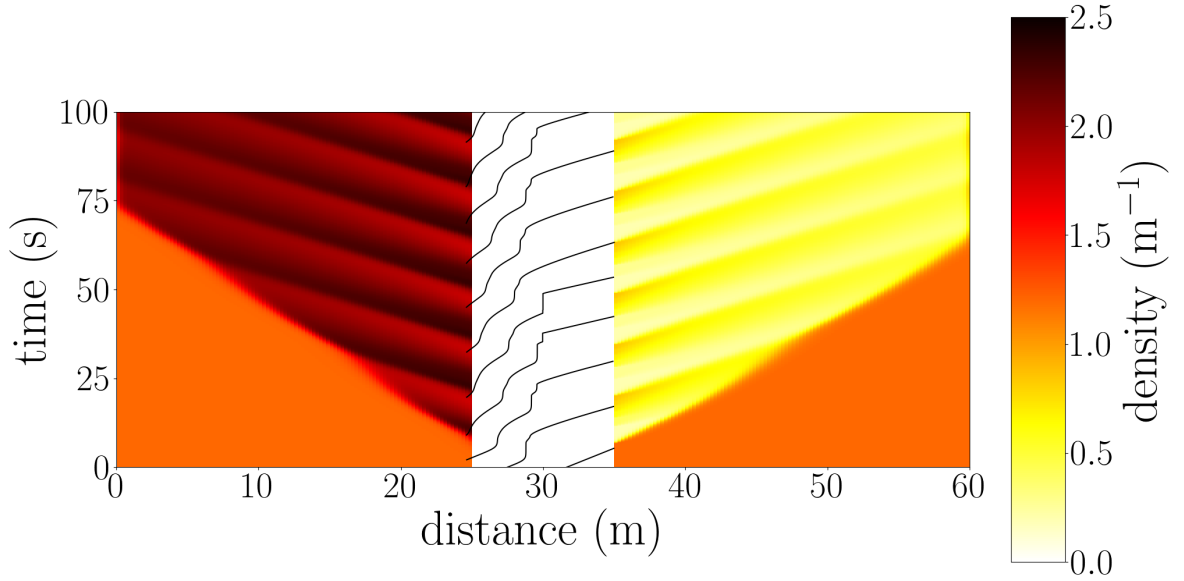


FIGURE 3.13 – Cas d’application - modèles FTL et LWR. 1 piéton sur 5 est tracé pour la partie microscopique.

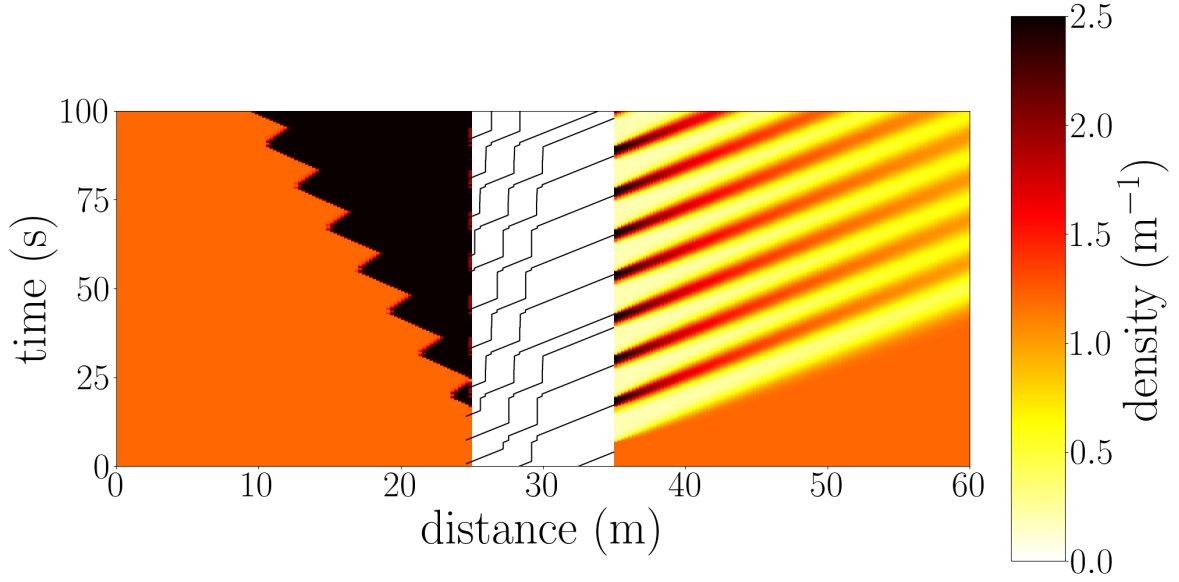


FIGURE 3.14 – Cas d’application - modèles d’inhibition micro et macro.
1 piéton sur 5 est tracé pour la partie microscopique.

3.3.2 Extensions

Choix de la vitesse pour le leader dans le couplage FTL-LWR. On a précédemment pris dans le cadre des modèles FTL et LWR pour le leader comme vitesse $v(\rho(t, 0^+))$, un choix qui correspond à une forme de vision lagrangienne du modèle macroscopique considérant localement le déplacement du flux à l’interface. Sur les cas tests montrés ce choix donne des résultats satisfaisant, mais il peut aussi entraîner un artéfact de modélisation : lorsque la densité à l’interface atteint ρ_{max} et que la vitesse s’annule, le leader va s’arrêter quelle que soit sa distance à l’interface. Afin de résoudre ce problème on peut essayer de remplacer la vitesse du leader par une expression prenant en compte d’une part la distance à l’interface et d’autre part la densité à l’interface. On peut prendre par exemple une combinaison pondérée des vitesses microscopiques et macroscopiques :

$$\dot{X}_{L(t)} = \frac{\rho(t, 0^+)(|X_L|)\varphi(|X_L|) + v(\rho(t, 0^+))}{1 + (|X_L|)\rho(t, 0^+)}. \quad (3.74)$$

Cette expression réintroduit cependant des effets d’oscillations au niveau de l’interface. On ne peut donc simplement corriger l’expression de la vitesse du leader à l’interface sans modifier aussi l’expression de la densité et du flux.

En se concentrant maintenant sur le modèle LWR, avec la reformulation de notre problème et en particulier la densité ρ_b donnée par l’équation (3.30) on peut s’intéresser à ce que donne le modèle LWR si on lui impose ρ_b comme condition au bord. Des travaux comme ceux de Lebacque ([Leb96]) suggèrent comme condition au bord d’introduire les fonctions :

$$f_D(\rho) = \begin{cases} f(\rho) & \text{si } \rho < \rho_c \\ f(\rho_c) & \text{si } \rho > \rho_c \end{cases}, \quad (3.75)$$

$$f_S(\rho) = \begin{cases} f(\rho_c) & \text{si } \rho < \rho_c \\ f(\rho) & \text{si } \rho > \rho_c \end{cases}, \quad (3.76)$$

où ρ_c est la densité pour laquelle le flux est maximal, et d'exprimer la condition au bord comme étant :

$$f(t, 0) = \min \{f_S(\rho(t, 0^+)), f_D(\rho_b(t))\}. \quad (3.77)$$

On peut ensuite définir la vitesse du leader comme étant $f(t, 0)/\rho_b(t)$, qui sera bien défini puisque $\rho_b(t)$ ne peut être nul.

Le schéma numérique peut être facilement adapté suite à ces considérations : dans l'équation (3.38) on calcule pour chaque instant un flux $\rho^{n,k}v(\rho_0^n)$ qui est la conséquence du déplacement du leader à vitesse $v(\rho_0^n)$. On peut cependant étendre le schéma pour prendre en compte d'autres flux numériques : si on suppose qu'on se donne un flux numérique $g(\rho_1, \rho_2)$ on peut définir à la place :

$$F_{-1/2}^n = \frac{1}{\Delta t} \sum_k \Delta t^k g(\rho^{n,k}, \rho_0^n), \quad (3.78)$$

$$X_{L^{n,k}}^{n,k+1} = X_{L^{n,k}}^{n,k} + \Delta t^k \frac{g(\rho^{n,k}, \rho_0^n)}{\rho^{n,k}}. \quad (3.79)$$

Couplage sur des jonctions. Une extension envisagée pour ce type de couplage consiste à regarder le passage micro-macro comme une jonction et non une interface. On considère le cas de plusieurs chemins se croisant à un point qui servira de jonction entre les différentes voies, chaque voie pouvant être modélisée par un modèle microscopique ou macroscopique (comme illustré en figure 3.15). Ce problème a été très étudié du point de vue des modèles macroscopiques pour le trafic routier ([Cos14] pour des équations de Hamilton-Jacobi sur des réseaux, ou [GP09] par exemple pour une approche basée sur un problème de Riemann pour une jonction).

Dans le cas purement macroscopique de ces modélisations, chaque voie est décrite par un modèle du type LWR. La jonction est vue pour chaque voie comme une condition de bord, dépendante des autres voies entrantes et sortantes. Le point central dans ce type de modélisation est ainsi le calcul du flux passant au travers de la jonction : la propagation des différents bouchons dans les voies sera fortement conditionnée par le calcul du flux pour chaque voie au niveau de cette jonction.

Ces modèles ont été brièvement étudiés durant cette thèse avec l'objectif de les appliquer pour des cages d'escalier. Dans le cadre d'essais d'évacuation réalisés sur un immeuble de bureaux, nous avons remarqué des différences dans les mouvements dans les cages d'escalier et dans les étages. Dans les cages d'escaliers, le mouvement des individus est contraint et les flux de chaque étage se rencontrent, tandis que dans les bureaux de chaque étage peu de congestion ont été observées. Pour une description de cet essai [JDW+20], la suite ne repose pas sur le détail de cette expérience. Différentes études expérimentales exhibent un diagramme fondamental pour le mouvement de piétons dans des escaliers ([HM12],[PHK12]). La modélisation d'une cage d'escalier comme une « voie » de

passage macroscopique pourrait ainsi rendre compte simplement de ces effets, tandis qu'en dehors des escaliers un modèle microscopique sera préférable. Le but de cette partie est ainsi d'illustrer quelques développements pour intégrer les problématiques de couplage à des modèles de calcul de flux pour une jonction.

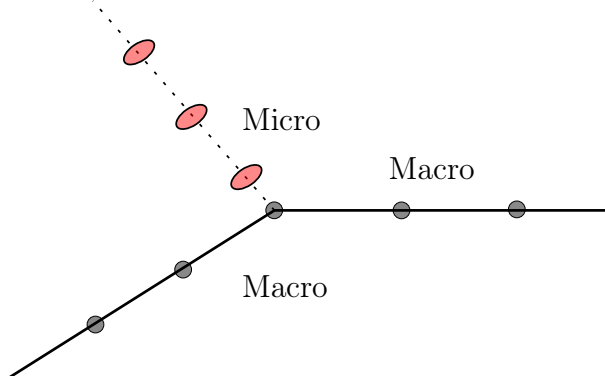


FIGURE 3.15 – Exemple de jonction considéré

Nous utilisons un modèle pour la jonction issu de [CLM15] : celle-ci est séparée en deux chemins entrants (un micro et un macro) et un chemin sortant, bien que cette approche ait été généralisée à plusieurs chemins entrants et sortants. Les chemins entrants et sortants macroscopiques vérifient une équation de conservation sur une demi-droite en espace :

$$\partial_t \rho^1 + \partial_x (f(\rho^1)) = 0 \quad \forall (t, x) \in [0, \infty] \times]-\infty, 0[, \quad (3.80)$$

$$\partial_t \rho^2 + \partial_x (f(\rho^2)) = 0 \quad \forall (t, x) \in [0, \infty] \times]0, +\infty[, \quad (3.81)$$

où $] -\infty, 0[$ est le domaine du chemin entrant, et $]0, +\infty[$ celui du chemin sortant. La jonction est ainsi située en 0 pour les deux modèles.

Le modèle microscopique évolue selon le modèle FTL :

$$\dot{X}_i = \varphi(X_{i+1} - X_i) \quad \forall 1 \leq i < L(t), \quad (3.82)$$

où $L(t)$ est le leader à l'instant t .

Le schéma de résolution de la jonction fait appel à deux fonctions définies à partir du flux et de la densité ρ_c où celui-ci atteint son maximum :

$$f_D(\rho) = \begin{cases} f(\rho) & \text{si } \rho < \rho_c \\ f(\rho_c) & \text{si } \rho > \rho_c \end{cases} \quad (3.83)$$

$$f_S(\rho) = \begin{cases} f(\rho_c) & \text{si } \rho < \rho_c \\ f(\rho) & \text{si } \rho > \rho_c \end{cases} \quad (3.84)$$

On se donne enfin un coefficient γ qui donne la proportion du flux dû au modèle micro à la jonction, et $1 - \gamma$ sera la proportion de flux macroscopique.

Le flux traversant l'interface sera donné par :

$$F_{-1/2}^n = \min \left(\frac{1}{\gamma} f_D(\rho_b^n), \frac{1}{1-\gamma} f_D(\rho_{-1}^{1,n}), f_S(\rho_0^{2,n}) \right), \quad (3.85)$$

où ρ_b est la densité à l'interface définie à la section précédente. Pour les deux jonctions entrantes on multiplie ensuite ce flux par le coefficient approprié γ ou $1 - \gamma$. Pour le modèle FTL couplé, la vitesse du leader est donnée en remplaçant ce flux dans l'équation (3.79).

Application : modélisation de cages d'escaliers. Nous utilisons le schéma numérique précédent pour modéliser une jonction à chaque étage d'un immeuble, comme illustré en figure 3.16. Chaque étage est modélisé en microscopique et les escaliers en macroscopique, avec une jonction au niveau de la porte menant pour chaque étage à la cage d'escalier.

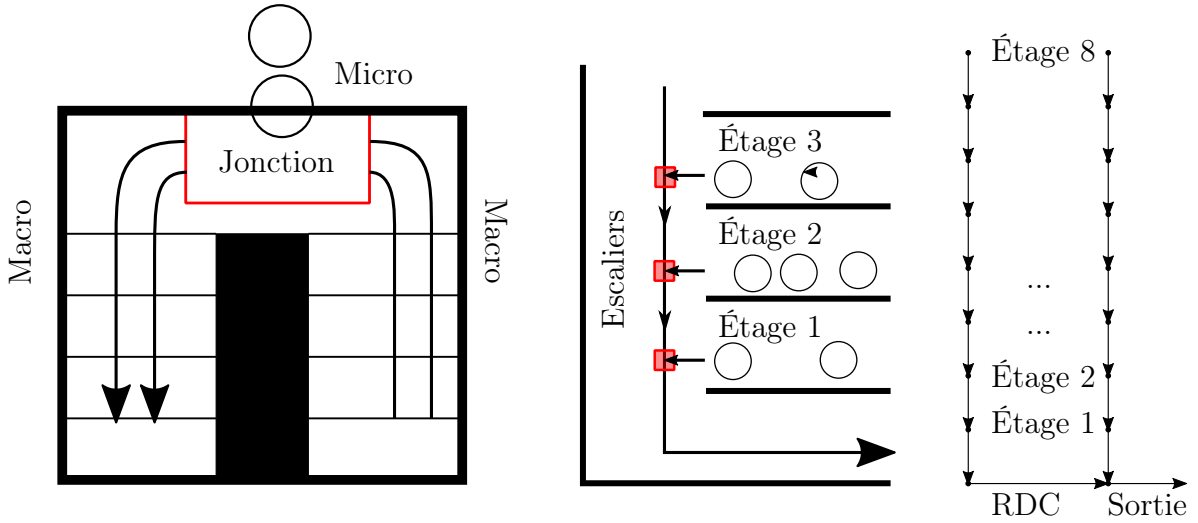


FIGURE 3.16 – Représentation d'une jonction entre une cage d'escalier et un étage vue de dessus et vue de côté, avec une représentation globale des chemins entre étages et vers la sortie. Le point d'interaction entre l'étage et la cage d'escalier est modélisé par une jonction, on peut ainsi modéliser plusieurs étages comme représenté à droite.

On représente de cette façon deux cages d'escalier se rejoignant au rez-de-chaussée pour un immeuble de 8 étages, comme illustré en figure 3.16. Les étages sont modélisés par des couloirs de 20 m de long, avec 40 personnes réparties uniformément à l'initialisation, avec une vitesse souhaitée de 1.2 m/s. Les escaliers sont modélisés en macroscopique par des segments de 6 m de long et 2 m de large, où la vitesse souhaitée est fixée à 0.9 m/s. La relation entre vitesse et densité choisie est celle de Weimann ([Wei92], ou voir section 2.3). Pour chaque jonction entre étage et escalier, le paramètre de priorité dans la jonction est pris à égalité entre toutes les jonctions entrantes.

Remarque 3.4. *Pour les escaliers, nous parlons ici de vitesse et de distance à l'horizontale. La composante verticale de la vitesse n'est pas prise en compte, bien qu'étant donné que les individus descendent cette dernière n'est pas nulle. Les représentations des résultats de simulation suivants affichent les densités dans les escaliers à la verticale pour une meilleure visibilité, mais il s'agit bien d'un calcul à l'horizontale.*

Cette situation est modélisée avec un pas de temps $\Delta t = 0.2$ s et d'espace $\Delta x = 0.8$ m, et les résultats sont représentés en figure 3.17. Durant une première phase les individus des étages

entrent dans la cage d’escalier avec peu de congestion, puis les flux arrivant des étages supérieurs se rencontrent au niveau de chaque jonction. La congestion se forme dans chaque cage d’escalier et se vide progressivement jusqu’à la fin de la simulation.

Lors d’essais d’évacuation réels ([JDW+20]) sur un bâtiment dont est inspiré cette géométrie, nous avons pu observer un phénomène semblable, où les cages d’escalier ont été congestionnées pendant une majeure partie de l’évacuation.

Ces premières étapes démontrent l’intérêt d’une utilisation de jonctions et de modèles 1D pour étudier les mouvements dans les escaliers. Le modèle microscopique 1D en revanche est ici trop restrictif pour une application à des cas réels et un modèle à deux dimensions serait plus adapté pour avoir des résultats pertinents.

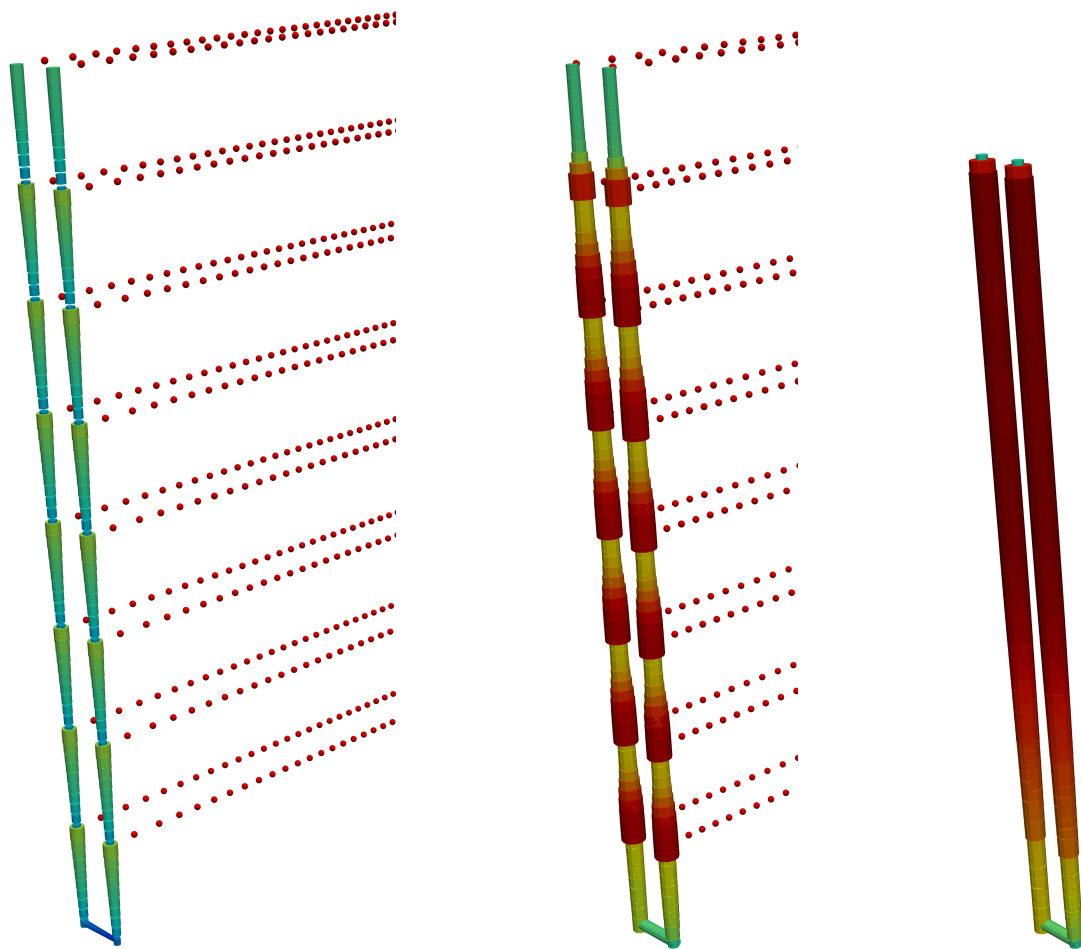


FIGURE 3.17 – Simulation de deux cages d’escalier d’un immeuble de 8 étages avec des jonctions à chaque étage, à $t = 12, 40, 80$ s. Pour le modèle macroscopique, la densité est représentée par des cylindres de rayon proportionnel à la densité locale. La congestion se forme dans les cages d’escalier au niveau des jonctions, puis s’étend à l’ensemble de la cage d’escalier avant qu’elle ne se vide lorsqu’elle n’est plus alimentée par les flux provenant des étages.

Chapitre 4

Couplage de modèles d'inhibition

4.1	Couplage microscopique vers macroscopique de modèles d'inhibition	65
4.2	Couplage macroscopique vers microscopique de modèles d'inhibition	70
4.3	Extensions, perspectives	75

À une dimension, nous avons montré une formulation du couplage de modèles granulaires avec inhibition microscopique et macroscopique à travers une interface. Ce couplage permet de capter la propagation de l'information dans les deux sens sans créer de perturbation au niveau de l'interface. Dans ce chapitre, nous présentons les premiers pas vers l'écriture d'un couplage pour un modèle de mouvement de foules à deux dimensions microscopique.

Le modèle qui nous intéresse principalement est le modèle granulaire avec inhibition issu de la thèse de F. Al Reda [Red17] et présenté dans le chapitre 2. Dans ce modèle, chaque individu est représenté par un disque possédant une vitesse souhaitée et soumis à une contrainte de non-recouvrement avec les autres agents. La vitesse effective des individus est calculée en deux phases :

- Une phase d'inhibition, où chaque agent corrige sa vitesse souhaitée : tout individu va choisir sa vitesse en fonction des voisins qu'il observe dans son champ de vision personnel. La vitesse après correction est la plus proche de la vitesse souhaitée initiale parmi celles qui ne vont pas entraîner de violation de la contrainte de non-recouvrement.
- Une phase de projection globale : la contrainte non-recouvrement est imposée en contraignant les vitesses que peuvent physiquement adopter les individus. La vitesse effective est la projection (au sens des moindres carrés) de l'ensemble des vitesses souhaitées corrigées sur l'ensemble des vitesses qui ne vont pas entraîner de violation de la contrainte de non-recouvrement.

Le modèle granulaire pur, sans inhibition, ne contient pas la première phase qui modélise un comportement humain. Seule la projection physique des individus est présente, qui modélise un comportement plus égoïste où les individus cherchent à atteindre leur vitesse souhaitée à tout prix et sont prêts à bousculer leurs voisins. Notons qu'il est possible dans certains cas particuliers que la phase d'inhibition ne modifie pas ou peu les vitesses individuelles : dans ce cas le modèle d'inhibition se rapproche du modèle granulaire pur.

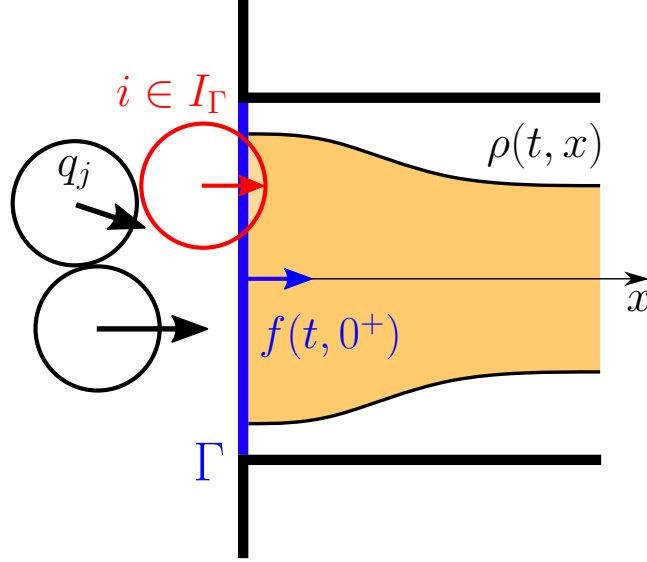


FIGURE 4.1 – Illustration de la représentation de la densité aux abords du couplage.

Ce choix du modèle d'inhibition s'est fait pour deux raisons. D'une part les situations de poussée extrêmes comme celles modélisées par le modèle granulaire pur sont peu souvent observées en cas d'évacuation([Kul16]), les comportements d'inhibition sont donc plus appropriés. La deuxième raison derrière ce choix est liée à la dissymétrie que présente le modèle d'inhibition : dans le cas du modèle granulaire, l'information d'un bouchon peut se transmettre en une seule itération de l'algorithme à une vitesse infinie vers l'amont ou l'aval. Chercher à coupler cette propagation avec un modèle macroscopique implique alors de devoir résoudre potentiellement les deux modèles conjointement. Dans le cadre du modèle d'inhibition, la phase de décision initiale a pour effet de limiter la propagation de l'information dans la direction du mouvement. Cet effet permet de séparer les contributions des modèles en amont et en aval.

Le modèle d'inhibition est couplé dans ce chapitre au travers d'une interface fixe, un segment dans le plan. Ce choix est fait pour se concentrer sur la propagation de l'information le long de la direction normale à l'interface. La situation typique modélisée dans ce cas est celle où la séparation entre deux modèles se fait au travers d'une porte. Le problème est de plus simplifié en considérant un modèle macroscopique à une dimension uniquement. La propagation de l'information étant étudiée uniquement selon la direction normale de l'interface, ce choix simplifie fortement l'écriture et l'étude du problème tout en préservant l'évolution du modèle dans la direction d'intérêt.

Nous présentons tout d'abord les modèles utilisés et les schémas numériques associés. Nous étudions ensuite le couplage entre le modèle granulaire avec inhibition et son analogue macroscopique à une dimension introduit dans le chapitre précédent, dans le sens microscopique vers macroscopique puis macroscopique vers micro. Une dernière partie montrera des extensions considérées pour ces couplages et leur cas d'application.

Modèles et schémas numériques utilisés

Les modèles utilisés ici ont déjà été présentés dans le chapitre 2, nous rappelons ici simplement l'expression des équations importantes ainsi que les schémas numériques.

Modèle granulaire avec inhibition. On note $q = (q_1, q_2, \dots, q_N)$ les positions des individus, identifiés à des disques de rayons $r_1, \dots, r_N$ et de vitesses souhaitées $U_1, \dots, U_N$. On note $D_{ij}(q) = |q_i - q_j| - r_i - r_j$ la distance entre les individus i et j , et $e_{ij} = \frac{q_j - q_i}{|q_i - q_j|}$ le vecteur pointant de i vers j . On note pour chaque individu i son ensemble d'influence I_i , c'est-à-dire les individus dans son cône de vision. Le calcul de la vitesse effective des individus se fait en deux étapes : une phase d'inhibition individuelle, où chaque individu modifie sa vitesse souhaitée pour ne pas pousser les agents dans son ensemble d'influence, puis une projection granulaire de l'ensemble des vitesses obtenues sur les vitesses qui ne vont pas entraîner de recouvrement.

La vitesse résultante de la première étape de correction est donnée par la solution $\bar{u} = (\bar{u}_1, \dots, \bar{u}_N)$ de N problèmes de minimisation résolus successivement :

$$\bar{u}_i = \underset{w \in \mathcal{C}_i(q, \bar{u}_{-i})}{\operatorname{argmin}} \frac{1}{2} |w - U_i|^2, \quad (4.1)$$

où $\bar{u}_{-i} = (\bar{u}_1, \dots, \bar{u}_{i-1}, \bar{u}_{i+1}, \dots, \bar{u}_N)$ et l'espace des contraintes $\mathcal{C}_i(q, \bar{u}_{-i})$ est défini par :

$$\mathcal{C}_i(q, \bar{u}_{-i}) = \{v \in \mathbb{R}^2, \forall j \in I_i, D_{ij}(q) = 0 \implies e_{ij} \cdot (v - \bar{u}_j) \leq 0\}, \quad (4.2)$$

Les problèmes de minimisation ont une solution unique lorsqu'ils sont résolus en commençant par les individus les plus en amont dans le mouvement, c'est-à-dire ceux qui ne sont influencés par personne. Cet ordre de résolution est déterminé en effectuant un tri topologique des individus à partir de leurs ensembles d'influence respectifs.

Ces nouvelles vitesses souhaitées sont ensuite elles-mêmes projetées globalement sur l'ensemble des contraintes données par les contacts entre individus et la vitesse effective est alors donnée par :

$$u = \underset{w \in \mathcal{C}(q)}{\operatorname{argmin}} \frac{1}{2} |u - \bar{u}|^2, \quad (4.3)$$

où l'ensemble des contraintes est :

$$\mathcal{C}(q) = \{v \in \mathbb{R}^{2N}, D_{ij} = 0 \implies e_{ij} \cdot (v_j - v_i) \geq 0\}. \quad (4.4)$$

Schéma numérique. On note Δt le pas de temps et les positions des individus sont calculées itérativement par $q^{n+1} = q^n + \Delta t u^{n+1}$, où u^{n+1} est la vitesse effective calculée en deux étapes. Ces deux étapes sont basées sur un développement au premier ordre des contraintes de non-recouvrement ([Red17; MF18]) :

$$\begin{aligned} D_{ij}(q^n + \Delta t v) &\simeq D_{ij}(q^n) + \Delta t \nabla D_{ij}(q^n) \cdot v \\ &= D_{ij}(q^n) + \Delta t e_{ij} \cdot (v_j - v_i) \end{aligned}$$

La première étape de correction de la vitesse souhaitée devient alors :

$$\bar{u}_i^n = \operatorname{argmin}_{w \in \mathcal{C}_i^n(q^n, \bar{u}_{-i}^n)} \frac{1}{2} |w - U_i|^2,$$

où les contraintes sont cette fois :

$$\mathcal{C}_i^n(q^n, \bar{u}_{-i}^n) = \{v \in \mathbb{R}^2, \forall j \in I_i, D_{ij}(q^n) + \Delta t e_{ij}(q^n) \cdot (\bar{u}_j^n - v) \geq 0\}.$$

Ces projections sont résolues en remontant frontalement le graphe d'influence défini par les ensembles I_i à chaque étape. De cette façon, toutes les vitesses $\bar{u}_j$ des individus dans l'ensemble d'influence de i ont déjà été calculées. La vitesse obtenue est ensuite projetée globalement :

$$u^n = \operatorname{argmin}_{w \in \mathcal{C}^n(q^n)} \frac{1}{2} |w - \bar{u}|^2,$$

où l'on prend en compte tous les contacts :

$$\mathcal{C}^n(q^n) = \{v \in \mathbb{R}^{2N}, \forall j \neq i, D_{ij}(q^n) + \Delta t e_{ij}(q^n) \cdot (v_j - v_i) \geq 0\}.$$

Modèle macroscopique d'inhibition. On considère une densité ρ dont l'évolution est donnée par une équation de conservation :

$$\partial_t \rho + \partial_x(\rho v) = 0,$$

où la vitesse vaut $U(x)$ lorsque $\rho(x) < \rho_{max}$, et pour chaque cluster où $\rho = \rho_{max}$:

$$v(x) = \min_{y \geq x, \rho(y) = \rho_{max}} U(y). \quad (4.5)$$

Ces équations sont approchées par une méthode de type volume fini. Pour un pas de temps Δt et un pas d'espace Δx , et on a :

$$\frac{\rho_i^{n+1} - \rho_i^n}{\Delta t} + \frac{F_{j+1/2}^n - F_{j-1/2}^n}{\Delta x} = 0, \quad (4.6)$$

où ρ_j^n est la valeur approchée de ρ sur la cellule j entre $n\Delta t$ et $(n+1)\Delta t$ et $F_{j+1/2}^n$ le flux à l'interface entre les mailles j et $j+1$ dont nous allons préciser l'expression par la suite.

La contrainte de non-dépassement de la densité maximale $\rho_i^n \leq \rho_{max}$ devient pour une itération entre n et $n+1$:

$$\rho_i^n - \frac{\Delta t}{\Delta x} (F_{j+1/2}^n - F_{j-1/2}^n) \leq \rho_{max}.$$

En exprimant $F_{j-1/2}^n$ avec un schéma décentré, on a :

$$F_{j-1/2}^n = \rho_{j-1}^n v_{j-1/2}^n,$$

où $v_{j-1/2}^n$ est la vitesse à l'interface entre deux mailles, que l'on exprime :

$$v_{j-1/2}^n = \operatorname{argmin}_{v \in \mathcal{C}_j^n} |v - U_{j-1/2}|^2, \quad (4.7)$$

où $U_{j-1/2}$ est la vitesse souhaitée à l'interface entre deux mailles, et :

$$\mathcal{C}_j^n = \left\{ v \in \mathbb{R}, \rho_i^n - \frac{\Delta t}{\Delta x} \left(F_{j+1/2}^n - \rho_{j-1}^n v \right) \leq \rho_{max} \right\}$$

Notons que la modélisation de l'effet de l'inhibition se situe dans ce calcul de la vitesse effective entre deux mailles.

La vitesse peut être donnée explicitement pour chaque maille :

$$v_{i-1/2}^n = \min \left(U_{i-1/2}^n, \frac{\Delta x}{\rho_{i-1}^n} \left(\frac{\rho_{max} - \rho_i^n}{\Delta t} - \frac{\rho_i^n v_{j+1/2}^n}{\Delta x} \right) \right). \quad (4.8)$$

4.1 Couplage microscopique vers macroscopique de modèles d'inhibition

Les deux modèles que nous couplons sont établis sur des contraintes de densité maximale ou de non-recouvrement et sont tous les deux basés sur une hiérarchie dans leur résolution numérique. Cette hiérarchie permet une stratégie de résolution du couplage : résoudre le problème en amont du couplage dans le sens du mouvement, traiter l'interface entre les deux modèles et enfin traiter le problème en aval.

Formalisme et schéma numérique

Problème instantané. On note $\Omega \subset \mathbb{R}^2$ le domaine microscopique borné, de bord $\partial\Omega$ et $I \subset [0, \infty[$ le domaine du modèle macroscopique. On notera Γ l'interface entre les modèles vue par le modèle microscopique. Par supposition, la normale en tout point à Γ est un vecteur constant que l'on note n .

On note $q = (q_1, \dots, q_N)$ les positions des individus dans le domaine Ω , et on introduit leur distance algébrique $D_i^\Gamma(q)$ à l'interface :

$$D_i^\Gamma(q) = (q_i - x_i) \cdot n,$$

où x_i est le projeté de q_i sur Γ . Il s'agit de la distance algébrique entre le centre de l'individu et Γ , chaque individu est en contact avec l'interface à partir du moment où $D_i^\Gamma \leq r_i$, et il aura complètement dépassé l'interface lorsque $D_i^\Gamma = -r_i$. L'ensemble des individus en train de traverser l'interface est noté I_Γ , et est défini par :

$$I_\Gamma = \{i \leq N, D_i^\Gamma(q) \leq r_i\}.$$

Il s'agit des individus dont la vitesse va être influencée par le modèle macroscopique.

Afin de traduire leur déplacement en flux, on va choisir d'exprimer le flux linéairement en fonction de la vitesse avec le même coefficient quelle que soit la position du piéton couplé. Un piéton $i \in I_\Gamma$ avec une vitesse u_i induit dans le modèle macroscopique un flux :

$$f_i = \frac{1}{2r_i} u_i \cdot n.$$

Pour le modèle macroscopique, ce flux entrant doit respecter la contrainte de non-dépassement de la densité maximale, sans dépasser la vitesse souhaitée. La seconde contrainte peut s'exprimer :

$$u_i \cdot n \leq U(0) \quad \forall i \in I_\Gamma,$$

où $U(x)$ est la vitesse souhaitée pour le modèle macroscopique en un point x . En ce qui concerne la contrainte de non-dépassement de la densité maximale, on doit prendre en compte le fait que plusieurs individus peuvent rentrer par la même interface. On exprime ainsi cette contrainte lorsque $\rho(t, 0) = \rho_{max}$ pour les flux comme :

$$\sum_{i \in I_\Gamma} f_i \leq f(t, 0^+) l, \quad (4.9)$$

où $f(t, 0^+) = \rho_{max} v(t, 0^+)$ est le flux au bord pour le modèle macroscopique, et l la largeur du couloir.

Ces contraintes permettent de définir dans la phase de correction de la vitesse souhaitée du modèle d'inhibition une étape pour prendre en compte le couplage avec le modèle macroscopique. On note $u_\Gamma = (u_i)_{i \in I_\Gamma}$ et on pose :

$$\bar{u}_\Gamma = \operatorname{argmin}_{w \in \mathcal{C}_\Gamma} \frac{1}{2} \sum_{i \in I_\Gamma} |w_i - U_i|^2, \quad (4.10)$$

avec :

$$\mathcal{C}_\Gamma(q) = \left\{ w \in \mathbb{R}^{2N}, \forall i \in I_\Gamma, w_i \cdot n \leq U(0), \sum_{i \in I_\Gamma} \frac{w_i}{2r_i} \leq f(0^+) \right\}.$$

Ces vitesses sont utilisées pour continuer l'étape de correction de la vitesse souhaitée du modèle granulaire avec inhibition. Les vitesses résultantes sont ensuite projetées globalement sur l'ensemble des vitesses admissibles comme décrit dans l'équation (4.3). Les vitesses résultantes de cette projection peuvent cependant ne plus respecter la contrainte à l'interface (4.10). On fait ici le choix de rajouter les contraintes liées à l'interface dans la projection granulaire. La projection de l'équation (4.3) est ainsi effectuée à la place sur l'espace :

$$\mathcal{C}_{tot} = \mathcal{C}(q) \cap \mathcal{C}_\Gamma(q).$$

Ce choix est un choix de modélisation de l'interface dans le couplage : on considère ainsi que la contrainte à l'interface de l'équation (4.9) est une contrainte physique imposée par le modèle macroscopique sur le modèle micro.

Schéma numérique. Les deux modèles seront discrétisés avec le même pas de temps Δt , et le modèle macroscopique sera de plus discrétisé avec un pas d'espace Δx . Ce choix est fait afin de simplifier l'expression des schémas numériques, nous discuterons brièvement des pistes possibles pour utiliser des pas de temps différents par la suite.

Dans le schéma numérique du modèle macroscopique, le flux à l'interface $F_{-1/2}^n$ doit respecter la contrainte :

$$F_{-1/2}^n \leq F_{1/2}^n + \frac{\Delta x}{\Delta t} (\rho_{max} - \rho_0^n). \quad (4.11)$$

On note $F_\Gamma^n = F_{1/2}^n + \frac{\Delta x}{\Delta t} (\rho_{max} - \rho_0^n)$.

On note $f_i^n(w)$ le flux correspondant à un individu i proche de l'interface Γ ayant une vitesse w . Entre deux itérations, on exprime ce flux :

$$f_i^n(w) = \frac{1}{2r_i} (w \cdot n). \quad (4.12)$$

Ces contraintes doivent s'appliquer sur les individus qui peuvent effectivement rejoindre l'interface durant l'intervalle de temps Δt . On note par la suite I_Γ^n l'ensemble des individus qui vérifient cette condition :

$$I_\Gamma^n = \{i, D_i^\Gamma(q) - \Delta t U_i \cdot n \leq r_i\}.$$

Le problème discrétisé de projection pour les vitesses des individus au niveau de l'interface peut ainsi s'écrire :

$$\bar{u}_\Gamma^n = \operatorname{argmin}_{w \in \mathcal{C}_\Gamma^n} \sum_{i \in I_\Gamma^n} \frac{1}{2} |w - U_i|^2, \quad (4.13)$$

où les contraintes sont :

$$\mathcal{C}_\Gamma^n = \left\{ v \in \mathbb{R}^{2N}, \forall j \in I_\Gamma^n, v_j \cdot n \leq U_j, \sum_{i \in I_\Gamma^n} f_i(v_i) \leq F_\Gamma^n \right\}. \quad (4.14)$$

Le calcul des autres individus microscopiques est ensuite inchangé par rapport au modèle non couplé.

Le flux macroscopique à l'interface est obtenu en sommant les flux créés par les individus ayant effectivement traversé l'interface :

$$F_{-1/2}^n = \sum_{i \in I_{\Gamma,2}^n} f_i^n(w),$$

où $I_{\Gamma,2}^n$ est l'ensemble des individus touchant effectivement l'interface :

$$I_{\Gamma,2}^n = \{i, D_i^\Gamma(q) - \Delta t v_i^n \cdot n \leq r_i\}.$$

Test numérique

Nous simulons une pièce rectangulaire simple, avec 80 personnes évacuant à travers une porte de 75 cm donnant sur un couloir de 2 m. Cette configuration simple correspond à ce que l'on peut retrouver dans des expériences d'évacuation contrôlées. En particulier cette situation est inspirée d'expériences menées dans [NBK17] qui étudient le flux de passage d'individus au comportement soit individualiste soit poli. Nous allons vérifier qu'en utilisant le modèle d'inhibition en amont de la porte et un modèle macroscopique en aval, la transmission d'information se fait correctement en prenant l'exemple d'un bouchon qui remonterait le courant et passerait la porte.

La simulation du couloir est basée sur le modèle macroscopique à partir de la porte et fait une longueur de 5 m. On impose une vitesse nulle à la sortie du couloir. Les résultats des simulations avec et sans couplages sont présentés en figure 4.2.

Dans les deux cas l'information du blocage remonte le long du couloir et un bouchon se forme et les individus s'arrêtent au fur et à mesure. Le modèle couplé arrive à transmettre cette information dans les contraintes au bord, alors qu'un piéton est partiellement entré dans le modèle macroscopique. En comparant de plus la densité entre le modèle microscopique et macroscopique on voit que le bouchon se propage à la même allure dans les deux cas.

Une différence notable apparaît au niveau de la porte entre les deux modélisations : on remarque que l'espace au-dessus et en dessous de la porte est libre dans le cas du modèle microscopique pur, tandis que dans le modèle couplé cet espace est rempli.

Le modèle macroscopique simplifié permet donc de capturer l'information de blocage après le passage de la porte et faire remonter cette information jusqu'à la porte. La modélisation simpliste de la géométrie autour de la porte ne reproduit pas l'organisation des piétons après celle-ci. Cet effet découle directement du choix d'un modèle macroscopique après la porte, si les effets que l'on souhaite modéliser finement se situent après la porte il suffira de choisir de continuer la modélisation microscopique après la porte et de passer au modèle macroscopique plus loin.

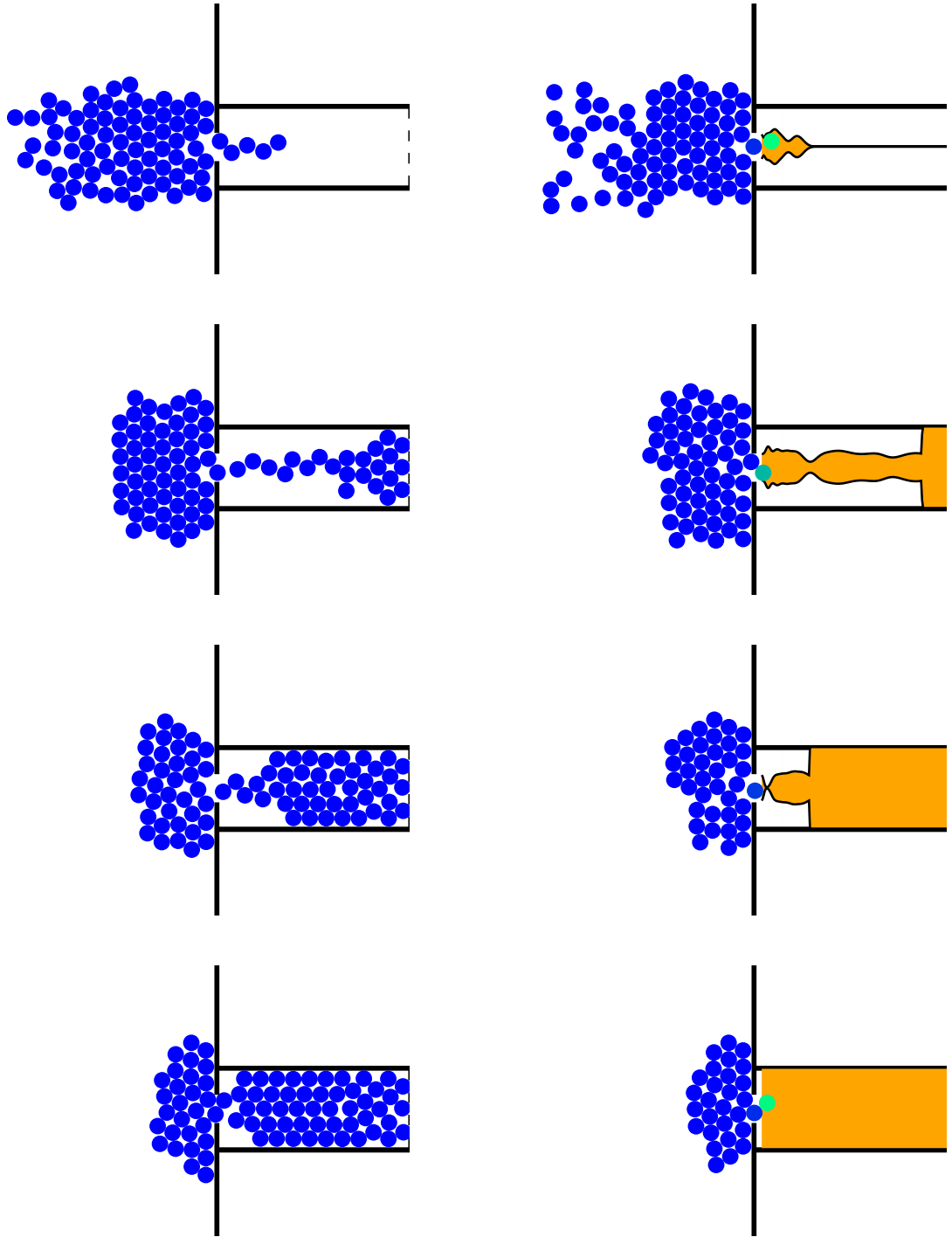


FIGURE 4.2 – Simulation de cas test de propagation de bouchon avec couplage, à $t = 2, 8, 14, 20$ s. Les individus sont colorés en vert en fonction de leur proportion de matière dans le modèle macroscopique. La densité est représentée dans le couloir en orange. Le bouchon se propage à une vitesse comparable entre les deux cas.

4.2 Couplage macroscopique vers microscopique de modèles d'inhibition

Pour le couplage macroscopique vers micro, une différence très nette apparaît par rapport au cas 1D pur, où la génération d'un individu pouvait se faire lorsque l'intégrale du flux traversant était égale à 1. À deux dimensions, pour éviter des recouvrements entre particules, il est nécessaire d'assurer que la génération d'une nouvelle particule est possible. Le calcul du flux ne peut se faire sans informations locales au niveau microscopique (voir figure 4.3). D'autres approches présentes dans la littérature ([CPT14; BCB21]) utilisent et calculent explicitement des positions lagrangiennes dans leur approche de la modélisation macroscopique. Les individus lagrangiens sont générés à l'initialisation et sont ensuite transportés par le mouvement eulérien jusqu'à une zone microscopique où leur mouvement est calculé avec un modèle microscopique. Le problème mentionné précédemment n'apparaît donc pas, cependant cela nécessite de s'assurer que les particules lagrangiennes soient correctement suivies dans le modèle macroscopique et implique des contraintes dans le choix du modèle macroscopique ou de sa discrétisation numérique.

Dans cette section, nous montrons une méthode ne supposant pas que le mouvement lagrangien soit connu en amont, et qui génère uniquement en amont à l'interface les individus. Chacun de ces individus est ensuite intégré au calcul microscopique pour déterminer sa vitesse. Le flux pour le modèle macroscopique en sera ensuite déduit des mouvements microscopiques. L'intérêt de cette approche est qu'elle peut ensuite être appliquée indifféremment pour coupler différents modèles microscopiques et macroscopiques.

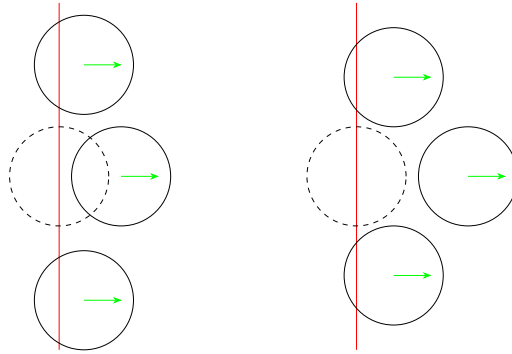


FIGURE 4.3 – Deux configurations différentes donnant un même flux pour le passage macroscopique vers micro. Dans les deux cas, la densité en aval est similaire, mais à gauche il est impossible de générer un individu sans créer de recouvrement, alors qu'à droite l'espacement des individus le permet.

Génération d'individus. Nous souhaitons générer des individus dans une maille selon une densité ρ donnée. Notons A l'aire de cette maille, le nombre N d'individus doit être tel que :

$$N = \lfloor \rho A \rfloor$$

où $\lfloor x \rfloor$ désigne la partie entière de x . La densité ρ évolue dans le temps et le nombre d'individus dans la cellule aussi.

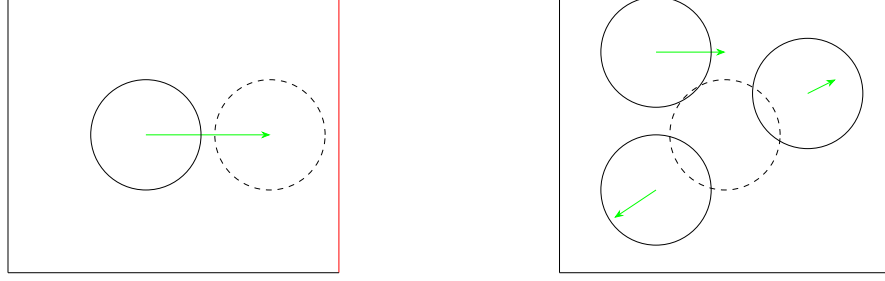


FIGURE 4.4 – Deux configurations indésirables possibles lors du tirage aléatoire de positions. À gauche, un agent est généré devant un individu macroscopique déjà présent dans la cellule. À droite, le placement d’un individu nécessite de pousser les autres pour se faire.

Une première solution pour générer ces individus est de tirer aléatoirement des positions initiales, puis de projeter ces positions sur l’ensemble des positions sans recouvrement entre particules. Numériquement cela revient à effectuer la projection du modèle de mouvement granulaire pur avec une vitesse nulle. Cette méthode permet de générer des configurations initiales pour une simulation microscopique même à forte densité. Dans le cas d’application présent, deux phénomènes peuvent se produire (voir figure 4.4) :

- Des individus peuvent être générés devant une autre individu microscopique dans la cellule. Du point de vue de la modélisation ce phénomène est contre-intuitif : le flux arrivant en amont de l’interface peut créer un individu après ceux déjà générés.
- La projection peut déplacer les individus microscopiques déjà présents. Pour des modèles autres que le modèle granulaire pur et basés sur des anticipations de collision, cet effet n’est pas souhaitable.

Nous proposons une autre méthode pour générer des individus qui n’ait pas ces problèmes. Cette méthode repose à la place sur une génération d’individus séparés par une distance $d(\rho)$ dépendante de la densité. $d(\rho)$ est la distance inter-particules moyenne, et on s’attend à ce qu’elle vérifie $d(\rho) \sim \rho^{-1/2}$.

Lorsque nous avons N individus dans la cellule et devons générer un individu $N + 1$, nous procédons en deux étapes :

- on tire aléatoirement la composante tangentielle à l’interface,
- on calcule l’intersection entre les cercles de rayon $d(\rho)$ et la hauteur perpendiculaire à l’interface issue de la coordonnée tangentielle tirée, comme montré en figure 4.5

La coordonnée normale retenue sera le maximum entre le point d’intersection le plus loin de l’interface et $d(\rho)$. L’individu généré est ainsi à une distance $d(\rho)$ de l’interface ou d’un autre individu déjà présent.

Cas test. Nous appliquons le couplage macroscopique vers microscopique au cas d’un couloir de 2 m de large pour étudier la propagation de l’information. La même situation est simulée à basse densité $\rho = 1 \text{ m}^{-2}$ (figure 4.6) et $\rho = 4 \text{ m}^{-2}$ figure 4.7). Dans le cas de haute densité, la sortie du couloir est bloquée pour former un bouchon qui se propage en remontant le courant.

À basse et haute densité, on observe une sur-congestion dans la dernière cellule macroscopique. La densité dans la dernière cellule oscille ensuite au cours de la simulation tant que le flux est stable.

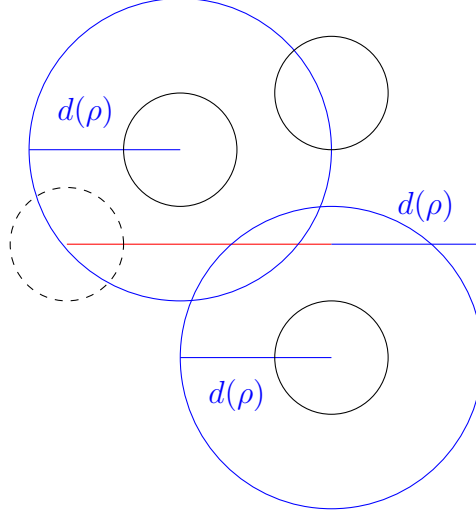


FIGURE 4.5 – Tirage de nouveaux piétons suivant l’algorithme proposé. La coordonnée tangentielle à l’interface est choisie aléatoirement, tandis que la coordonnée normale est choisie de manière à avoir une distance $d(\rho)$ au plus proche voisin (ou à l’interface s’il n’y a pas de plus proche voisin).

Pour la simulation à haute densité, le bouchon remonte ensuite le courant et traverse l’interface entre microscopique et macroscopique. Les oscillations dans la cellule s’arrêtent alors.

Discussion. Dans le chapitre 3 nous avons constaté que les oscillations étaient apparues à cause du choix inapproprié de représentation de la densité, qui entraînait de fortes variations dans le temps du flux au travers de l’interface. Il n’est donc pas surprenant d’observer de nouveau ces oscillations dans ces simulations. À cela s’ajoute le fait que nous avons ici choisi un pas de temps adapté au modèle microscopique, et qui aggrave donc le problème.

L’effet de sur-congestion peut être attribuée plus spécifiquement à la méthode de génération d’individus. Cet effet est présent à basse et haute densité, malgré le fait d’avoir pris en compte explicitement dans la méthode de tirage que la distance moyenne entre individus dépend de la densité. Une cause probable de ce phénomène est un défaut de la méthode en particulier au niveau de l’initialisation : les premiers piétons générés le sont à une distance trop grande de l’interface, ce qui entraîne une accumulation ponctuelle de masse. Une fois cette initialisation dépassée, la procédure de génération parvient à établir un régime stable sans pour autant parvenir à corriger l’accumulation initiale de densité.

Certaines pistes pourraient être explorées pour corriger les effets d’oscillation observés. Dans un premier temps adopter deux pas de temps différents pour les modèles permettrait de stabiliser le flux à l’interface, en calculant sur plusieurs pas de temps microscopiques. En revanche pour espérer faire disparaître les oscillations, il serait nécessaire de lisser le calcul du flux issu d’un piéton ou d’utiliser une meilleure représentation de la densité d’un individu.

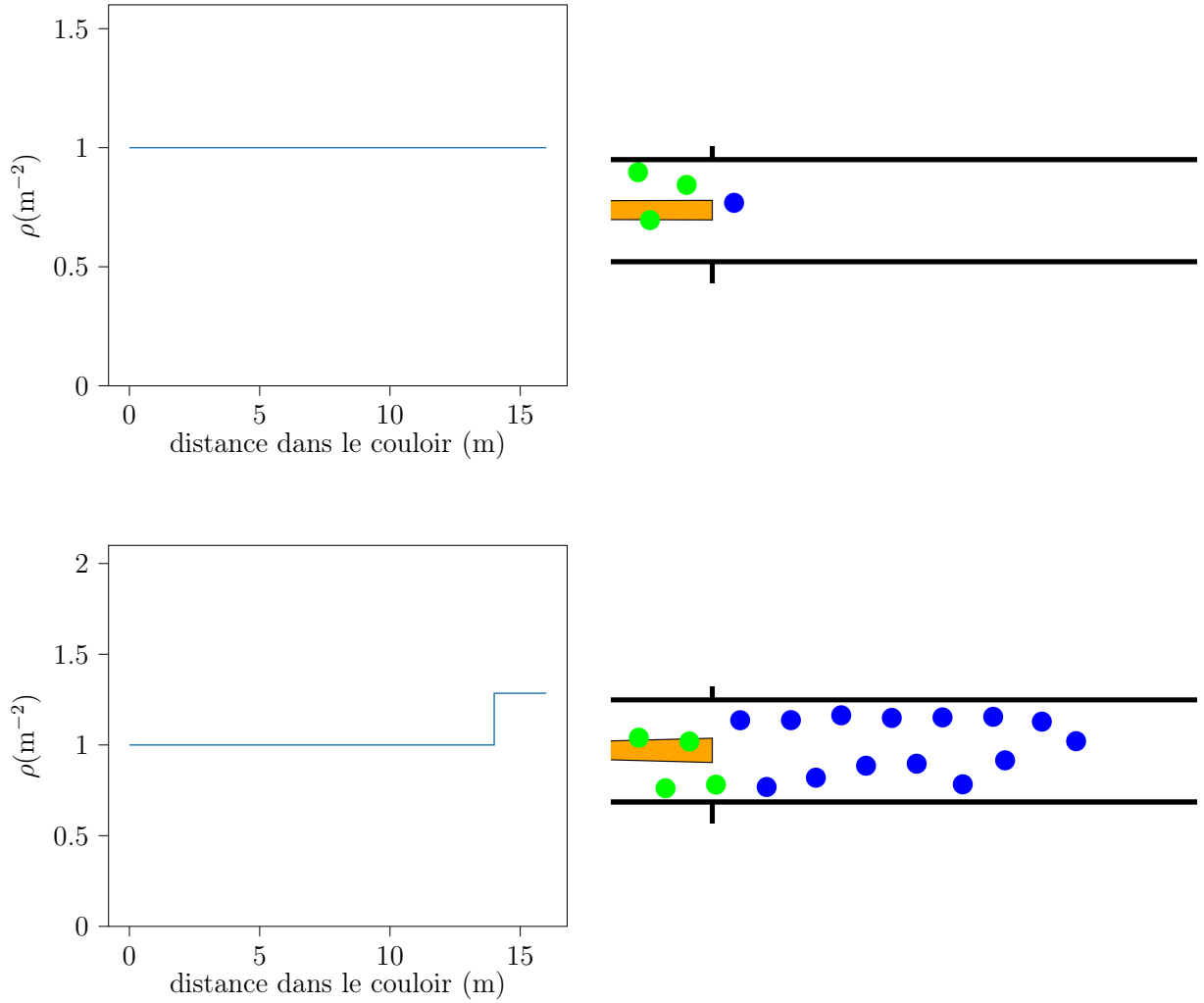


FIGURE 4.6 – Densité dans le modèle macroscopique pour la simulation à basse densité (1 m^{-2}) à $t = 0, 5 \text{ s}$. La densité dans la cellule couplée avec le modèle microscopique est surcongestionnée et oscille durant la simulation du fait des oscillations du flux de piétons générés en microscopique.

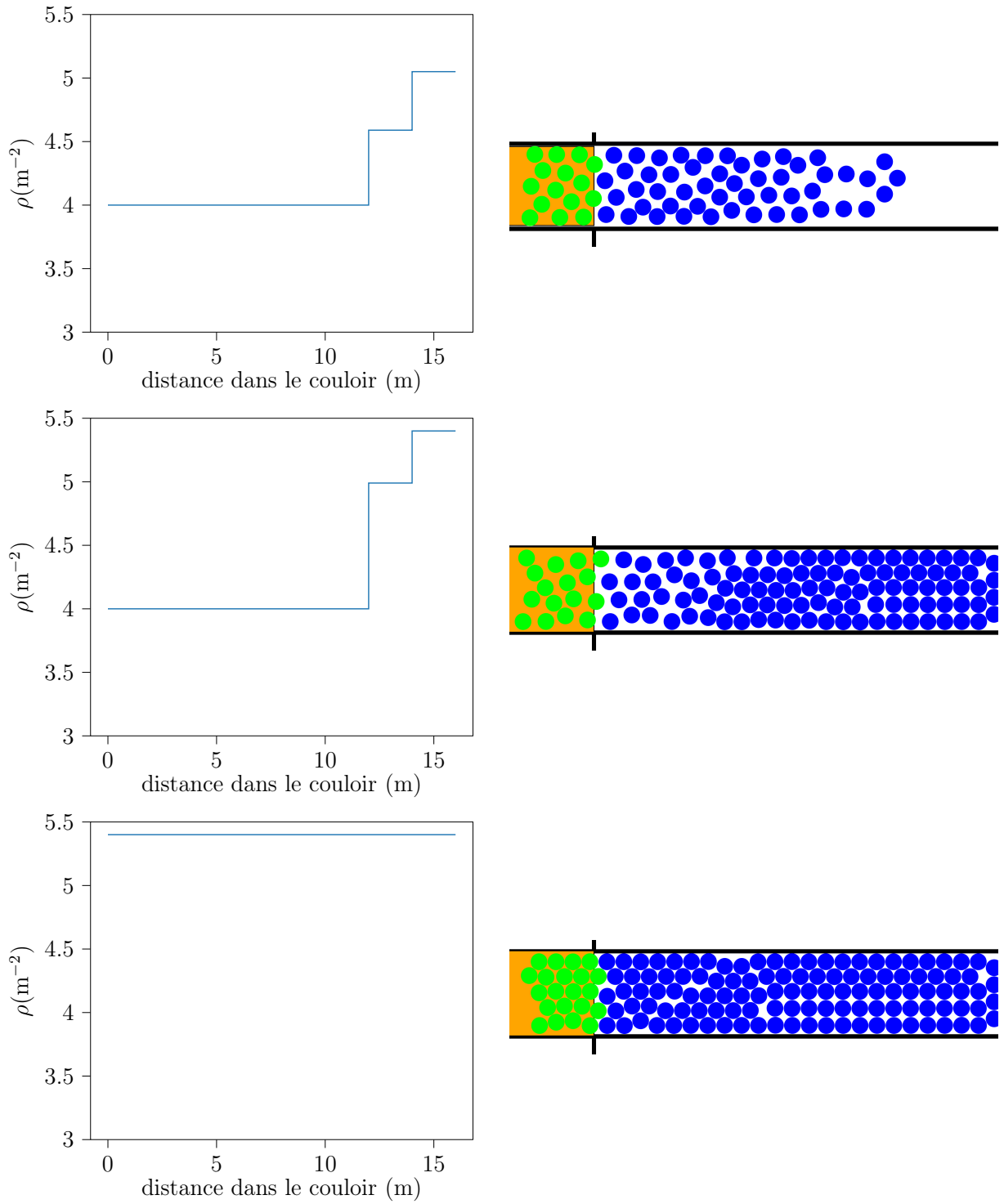


FIGURE 4.7 – Densité dans le modèle macroscopique pour la simulation à haute densité (4 m^{-2}) à $t = 5, 10, 15$ s. La densité dans la cellule couplée avec le modèle microscopique est surélevée et oscille durant la simulation du fait des oscillations du flux de piétons générés en microscopique.

4.3 Extensions, perspectives

Extension : croisement de flux

Une extension aux différents couplages proposés consiste à étudier des formulations où les flux microscopiques et macroscopiques de couplage se croisent. Un cas applicatif en particulier a motivé cette extension. Il s'agit d'une étude portant sur l'évacuation d'une rame de tramway dans un tunnel dans un cas d'incendie. Le tunnel n'a qu'un couloir très étroit comme issue de sortie et l'incendie contraint l'évacuation à ne se faire d'un seul côté, ce qui rend cette configuration potentiellement dangereuse. La situation est représentée en figure 4.8.

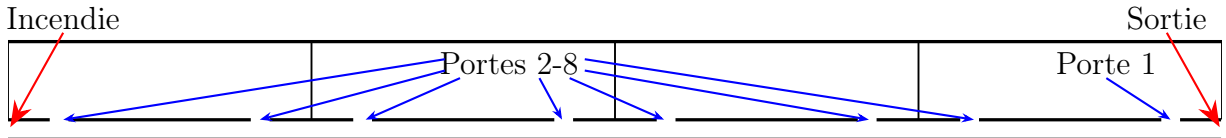


FIGURE 4.8 – Représentation du domaine modélisé pour la simulation d'évacuation de tram.

Des premiers essais de simulation de cette situation en utilisant le modèle granulaire avec inhibition seul donnent des résultats insatisfaisants : le couloir très resserré (60 cm de largeur) où peuvent se croiser des individus au niveau des portes entraîne la création de bouchons stables dans une simulation¹. Le modèle d'inhibition n'est pas adapté pour gérer ces croisements.

Nous allons présenter une solution alternative, basée sur un couplage avec un modèle uni-dimensionnel pour le couloir. Un mouvement uni-dimensionnel dans le couloir semble une supposition raisonnable, puisque deux personnes ne peuvent de toute façon passer de front en même temps. Les différences avec les cas étudiés auparavant sont que cette fois le mouvement uni-dimensionnel ne se fera pas dans la direction normale à l'interface et que les flux microscopiques et macroscopiques vont se croiser.

Afin d'appliquer le couplage à ce cas, nous devons donc préciser une modélisation pour le croisement de flux au niveau de chaque porte : on aura en même temps des piétons microscopiques et un flux macroscopique qui vont se rencontrer au sein d'une cellule, comme illustré en figure 4.9. Nous présentons un modèle heuristique minimaliste pour le croisement de flux dans le cas de modèles d'inhibition. Celui-ci se base sur la contrainte de densité maximale pour le schéma numérique des équations macroscopiques (4.7) et pour le couplage microscopique (4.13), (4.14). Nous allons considérer que pour une cellule où des flux microscopiques et macroscopiques se rencontrent, le flux résultant sera la projection des flux souhaités sur l'ensemble des flux ne dépassant pas la densité maximale autorisée.

1. L'apparition de bouchons stable est présente usuellement dans le modèle granulaire pur, et non dans le modèle d'inhibition. Ici cependant, la géométrie force les individus traversant une porte à essayer de passer à côté de ceux venant du couloir. Lorsque deux individus sont à côté, aucun n'est dans le champ de vision de l'autre et la phase d'inhibition du modèle ne modifie pas ou peu leur vitesse : la phase d'inhibition des individus n'a quasiment aucun effet et la phase de projection du modèle granulaire pur est prépondérante, donnant lieu à des bouchons stables.

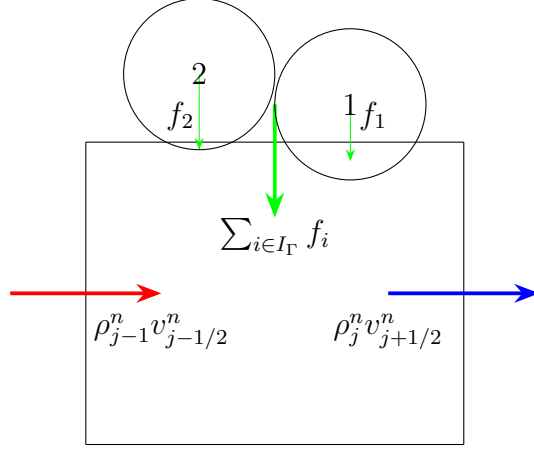


FIGURE 4.9 – Modélisation au niveau d’une sortie avec croisement de flux.

Plus précisément, en notant I_Γ les individus microscopiques couplés, pour la cellule macroscopique j nous avons :

- un flux maximal autorisé

$$F_\Gamma^n = \rho_j^n v_{j+1/2}^n + \frac{\Delta x}{\Delta t} (\rho_{max} - \rho_j^n),$$

- des individus I_Γ , dont le flux entrant dans la cellule vaut

$$f_i^n(v_i, q) = \frac{1}{2r_i} (v_i \cdot n),$$

- un flux macroscopique entrant

$$F_{j+1/2}^n = \rho_{j-1}^n v_{j-1/2}^n.$$

Les vitesses microscopiques et macroscopiques sont calculées par une projection des vitesses souhaitées sur l’espace des vitesses admissibles, tel que :

$$(\bar{u}_\Gamma^n, v_{j-1/2}^n) = \operatorname{argmin}_{(w,v) \in \mathcal{C}_\Gamma^n} \sum_{i \in I_\Gamma^n} \frac{1}{2} |w - U_i|^2 + \frac{1}{2} |v - U_{j-1/2}|^2,$$

où les contraintes sont :

$$\begin{aligned} \mathcal{C}_\Gamma^n = \{ & (w, v) \in \mathbb{R}^{2N} \times \mathbb{R}, \forall j \in I_\Gamma^n, w_i \cdot n \leq U_\Gamma, \\ & v \leq U_{j-1/2}, \rho_j^n v + \sum_{i \in I_\Gamma^n} f_i(v_i, q) \leq F_\Gamma^n \}. \end{aligned}$$

Nous utilisons cette modification locale du schéma de résolution numérique pour le cas d’évacuation du tram : chaque voiture de la rame est représentée par une boîte de 2.5 m par 16 m avec deux sorties de 1 m donnant sur un couloir de 60 m. Chaque rame est remplie avec 180 individus qui doivent sortir de leur rame puis traverser le couloir pour sortir. Dans cette modélisation, nous attribuons à chaque individu une sortie fixe à laquelle se diriger. La vitesse dans les rames est fixée à 1.4 m/s, tandis que dans le couloir où les mouvements sont plus difficiles nous la fixons à 1.0 m/s.

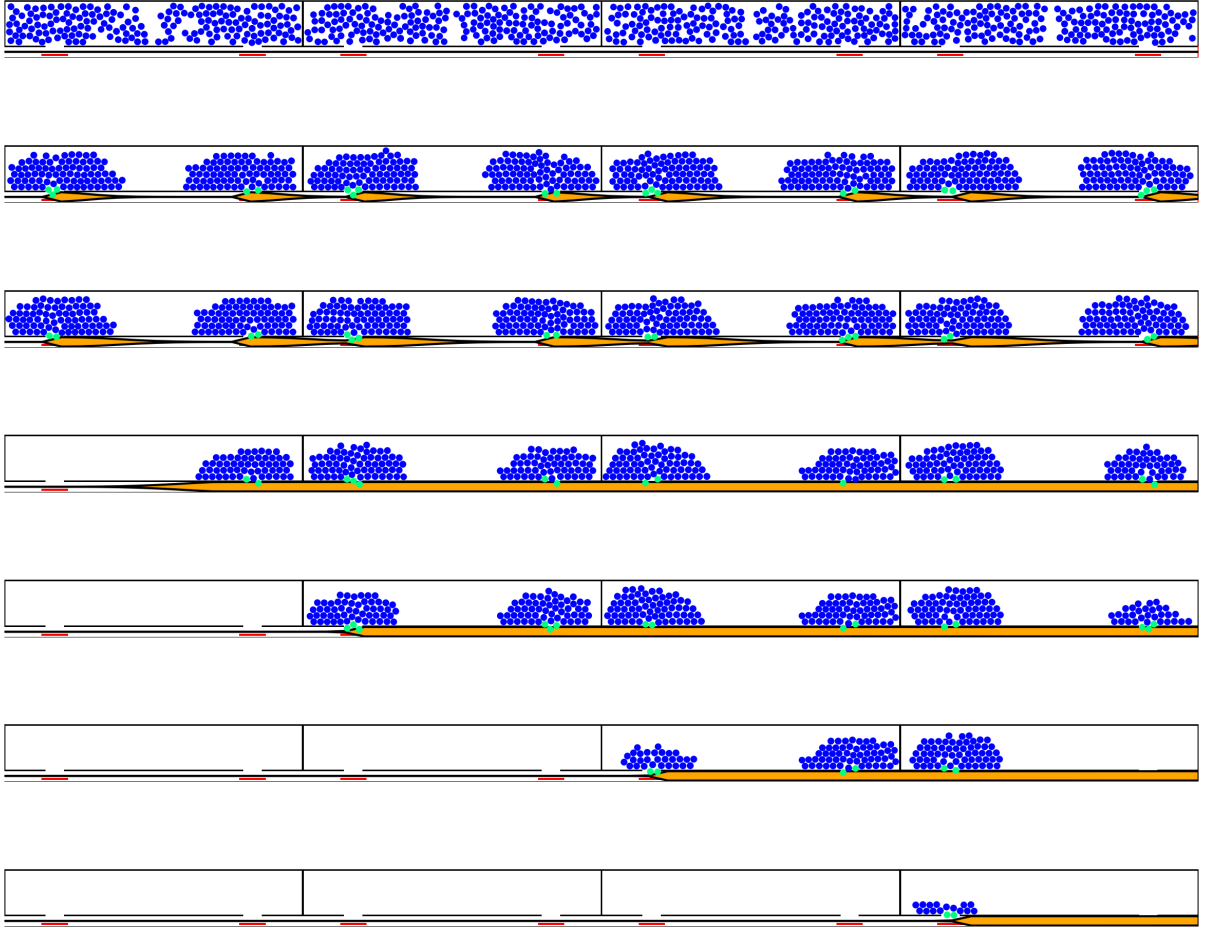


FIGURE 4.10 – Simulation de la rame de tram, aux instants $t = 0, 2.5, 5, 50, 100, 200, 280$ s. Les couleurs indiquent les individus en train d’être couplés dans le modèle macroscopique.

En figure 4.10 sont représentés des instantanés de la simulation, et en figure 4.11 sont représentés le cumul d’individus traversant chaque porte de la rame en fonction du temps.

Discussion. On remarque sur les instantanés que les premiers compartiments à se vider sont ceux du fond, ainsi que celui le plus proche de la sortie. Notons la présence de fortes congestions en amont de chaque porte. Sur la figure 4.11, on peut vérifier plus précisément l’évolution du débit de sortie des compartiments : on observe en premier une phase où les individus entrent dans le couloir sans encombre. Lorsque la congestion apparaît, les compartiments du fond de la rame se vident successivement très rapidement, et en parallèle la porte la plus proche de la sortie se vide à vitesse plus faible.

Le modèle de croisement donne un flux faible en sortie du compartiment le plus près de la porte, en revanche pour les autres un blocage important à la porte apparaît et les individus microscopiques laissent passer le flux devant eux sans oser pénétrer. Cet effet semble plausible dans des cas d’évacuation ordonnée, en faisant le rapprochement à une sortie de bus ou d’avion ordonnée où les individus doivent se déverser dans le couloir permettant de sortir.

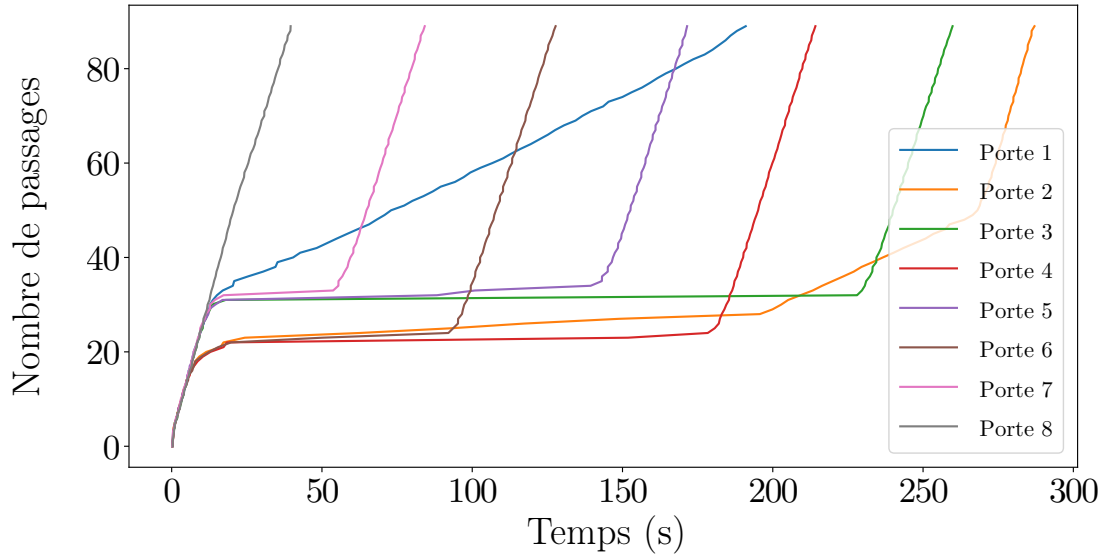


FIGURE 4.11 – Cumul du nombre de personnes sorties en fonction du temps pour chaque porte de sortie de la rame de tram.

Pistes d'améliorations futures

Un point possible d'amélioration et d'extension concerne la représentation de la densité des individus microscopiques dans le couplage microscopique vers macroscopique. Le choix actuel d'un flux linéaire à partir du moment où un individu est en contact avec l'interface conduit à des instabilités.

Pour des calculs de densité à partir d'entités microscopiques, un outil couramment utilisé est le diagramme de Voronoi (dans le cadre de mesures expérimentales par exemple [SS10]). Pour un individu i donné parmi N , sa cellule de Voronoi est l'ensemble des points de l'espace dont i est le plus proche parmi les N individus. Cette cellule permet d'attribuer une aire propre à chaque individu et donc une densité uniforme s'intégrant à 1. L'avantage d'une telle approche serait d'avoir une définition de la densité capable de gérer des situations de hautes et basses densités locales. Le problème est cependant que les cellules ainsi définies évoluent dans le temps d'une part, et d'autre part que la prise en compte de l'interface dans le calcul de la cellule n'est pas immédiat et doit être précisé. Ces améliorations doivent de plus être accompagnées de comparaisons à des données réelles sur des cas pertinents.

Chapitre 5

Modèle numérique macroscopique avec inhibition basé sur une résolution hiérarchique

5.1	Discrétisation, hiérarchie et positions du problème	82
5.2	Différentes approches	85
5.3	Discussion, perspectives	93

Durant cette thèse, nous avons à plusieurs reprises utilisé le modèle granulaire avec inhibition issu de la thèse de Fatima Al Reda ([Red17]). Ce modèle microscopique pour les évacuations de foules repose sur un principe de projection pour chaque piéton. Chacun choisit la vitesse la plus proche de sa vitesse souhaitée en s'adaptant aux individus dans son champ de vision sans les pousser. Ce choix est calculé par une projection locale de la vitesse souhaitée individuelle sur l'ensemble des vitesses qui ne vont pas entraîner de recouvrement avec d'autres individus ou des obstacles.

Ces projections individuelles s'effectuent dans un ordre particulier : Les individus sont connectés

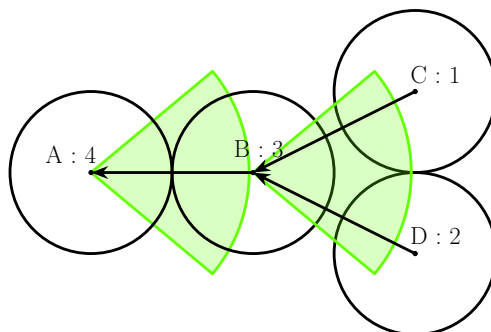


FIGURE 5.1 – Exemple de hiérarchie basée sur les cônes de vision : A voit B, qui voit lui-même C et D. Un tri topologique doit donner pour A un indice supérieur à B, et pour B supérieur à C et D.

via un graphe tel que j pointera vers i si j est dans le cône de vision de i . En partant des individus qui n'ont personne dans leur cône de vision, il est possible de trier les individus pour obtenir un ordre de parcours des nœuds du graphe qui remonte les cônes de vision individuels (voir figure 5.1 pour un exemple de résultat d'un tel tri, appelé tri topologique). Le choix des vitesses des individus se fait en projetant localement selon cette hiérarchie et crée une asymétrie dans les interactions entre piétons. Les agents les plus proches de la sortie décident de leur vitesse en premier et chaque agent va ralentir en fonction de la vitesse de ses voisins de devant.

Durant cette thèse, nous avons exploré une piste pour décrire un équivalent macroscopique à ce modèle d'inhibition. L'idée centrale de notre approche est d'écrire un modèle numérique basé sur cette notion de hiérarchie présente dans l'approche microscopique. Plutôt que de raisonner sur des individus, nous modélisons des phénomènes à l'échelle d'une cellule. La suite de projections individuelles deviendra une suite de projections pour chaque cellule. Le cœur de notre modélisation résidera dans l'écriture de ces projections et des contraintes locales sur les vitesses.

Historique du développement du modèle d'inhibition macroscopique. Ces travaux ne sont pas la première tentative d'écriture d'un tel modèle macroscopique. Dans sa thèse, F. Al Reda propose des approches différentes pour écrire ce modèle, en prenant comme point de départ le modèle granulaire (sans inhibition) macroscopique étudié par Aude Roudneff-Chupin dans sa thèse ([Rou11]) et évoqué en Section 2.2.

Le modèle granulaire macroscopique décrit l'évolution d'une densité ρ dans un domaine Ω borné de $\mathbb{R}^2$, contrainte à ne pas dépasser une valeur fixée ρ_{max} . Lorsque la densité est inférieure à cette valeur, elle se déplace à la vitesse souhaitée $U = U(x)$. Pour imposer la contrainte de densité maximale en tout temps, une condition sur la vitesse effective est formulée pour les points où la densité est égale ρ_{max} . Pour ces zones de l'espace, le champ de déplacement appliqué à ρ ne doit pas entraîner une violation de la contrainte de densité maximale et doit donc être non-concentrant. Cette condition peut être vérifiée en se restreignant aux champs de vitesse à divergence positive. La vitesse effective sera la projection de la vitesse souhaitée sur ces champs de vitesse non-concentrant. Le modèle s'écrit alors :

$$\begin{cases} \partial_t \rho + \nabla \cdot (\rho u) = 0, \\ u = P_{C_\rho} U, \end{cases} \quad (5.1)$$

où P_{C_ρ} désigne la projection pour la norme L^2 sur l'espace des vitesses admissibles C_ρ , qui sont les vitesses dont la divergence est positive sur les zones où $\rho = \rho_{max}$. L'espace C_ρ se définit de façon duale :

$$C_\rho = \left\{ v \in L^2(\Omega), \int_{\Omega} v \cdot \nabla p \leq 0, \forall p \in H_\rho^1(\Omega), p \geq 0 \text{ p.p.} \right\}, \quad (5.2)$$

avec

$$H_\rho^1(\Omega) = \{ p \in H^1(\Omega), p(\rho_{max} - \rho) = 0 \text{ p.p.} \}. \quad (5.3)$$

Le problème de projection de U sur l'ensemble C_ρ est bien défini, car cet ensemble est un convexe fermé de $L^2(\Omega)$. Le cadre théorique permettant de donner un sens à l'équation de transport fait appel à des techniques d'analyse poussées qui sont en dehors de la portée du présent manuscrit. Le lecteur intéressé pourra se reporter à [Rou11, Chapitre 2] pour trouver ces éléments.

Un premier candidat de modèle d'inhibition s'écrit en partant de ce modèle, et en rajoutant une contrainte sur les vitesses admissibles. On note $V(x, U)$ un champ décrivant les cônes de vision

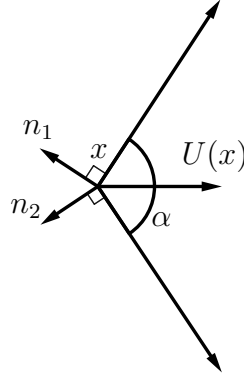


FIGURE 5.2 – Cône de vision $V(x, U)$, vecteurs normaux et cône polaire.

liés à la vitesse souhaitée. En tout point x , $V(x, U)$ est un cône centré sur U et d'angle α donné inférieur à π . En un point x le cône de vision est défini par :

$$V(x, U) = \{w \in \mathbb{R}^2, w \cdot n_i(x, U) \leq 0, i = 1, 2\},$$

où n_1 et n_2 sont les vecteurs normaux aux plans formant le cône de vision, comme illustré en figure 5.2. L'effet de l'inhibition est modélisé en imposant que la correction sur la vitesse effective ne se fasse que dans une direction opposée au cône de vision individuel. Cette formulation reprend l'aspect directionnel du modèle d'inhibition microscopique, où chaque individu va chercher à aller le plus vite possible sans pousser. Cette contrainte modélise que la vitesse ne peut que diminuer pour s'adapter à la congestion, les individus vont ralentir pour s'adapter à la congestion et non se pousser.

La modification principale que cela entraîne sur la formulation du modèle est de remplacer l'espace C_ρ par :

$$C_\rho(U) = \left\{ v \in L^2(\Omega), \int_{\Omega} v \cdot \nabla p \leq 0, p \geq 0 \text{ p.p.}, \forall p \in H_\rho^1(\Omega), (v - U) \in N(\cdot, U) \text{ p.p.} \right\},$$

où $N(x, U)$ est le cône polaire de U formé par les vecteurs n_1 et n_2 :

$$N(x, U) = \{w \in \mathbb{R}^2, w \cdot n_i(x, U) \geq 0, i = 1, 2\},$$

Ce nouvel ensemble admissible est convexe et fermé, la projection sur cet espace est donc bien définie de façon unique. Le modèle granulaire macroscopique possédait cependant une structure de flot-gradient dans l'espace de Wasserstein. Cette structure est perdue ici du fait de la dissymétrie des interactions.

Une seconde approche est basée sur la relaxation de la contrainte précédente. Nous reprenons les notations pour le modèle granulaire macroscopique et introduisons la fonction D donnant en tout point la distance à la sortie du domaine, notée Γ . La vitesse effective v dans ce nouveau modèle est obtenue en plusieurs étapes. Dans un premier temps, considérons le problème de minimisation suivant :

$$\min_{v \in C_\rho} J_\epsilon(v) = \min_{v \in C_\rho} \frac{1}{2} \int_{\Omega} a_\epsilon |v - U|^2,$$

où a_ϵ est une fonction dépendante d'un paramètre ϵ et telle que :

- $\forall x, y \in \Omega^2$, si $D(x) < D(y)$, alors $\lim_{\epsilon \rightarrow 0} a_\epsilon(x)/a_\epsilon(y) = +\infty$
- a_ϵ est positive et maximale en Γ .

La solution v_ϵ d'un tel problème pour un ϵ donné peut être interprétée comme la projection de U sur l'ensemble des champs à divergences positives par une distance L^2 pondérée par cette fonction a_ϵ , qui attribue un poids plus important aux vitesses proches de la sortie Γ . La vitesse effective pour cette version du modèle d'inhibition est donnée par la limite lorsque $\epsilon \rightarrow 0$ des fonctions v_ϵ . À une dimension, cette limite existe et coïncide avec la vitesse effective donnée par la première approche présentée ([Red17, Chapitre 5]). L'inhibition est modélisée en attribuant aux individus les plus proches de la porte un poids plus grand que tous ceux derrière eux. Cette formulation met en avant le caractère hiérarchique du modèle, en l'encodant dans cette fonction de poids a_ϵ . L'étude de ce passage à la limite est à l'étude dans une thèse entamée en 2022.

Une approche numérique. Nous présentons une approche heuristique pour écrire un modèle, couramment utilisée en mouvements de foules, en raisonnant sur un modèle discret en espace et en temps. Afin de nous rapprocher de la résolution du modèle microscopique, nous supposons explicitement que le champ de vitesse souhaitée U est défini à partir du gradient du temps de trajet à une sortie du domaine et calculons localement la vitesse en itérant hiérarchiquement suivant le temps de trajet. Plus précisément, nous suivons ces étapes :

1. calcul du temps de trajet à la sortie et définition d'une hiérarchie associée,
2. discrétisation de l'équation de transport pour la vitesse souhaitée,
3. calcul des vitesses effectives par projections hiérarchiques.

Le troisième point est celui où nous montrerons l'effet des différentes hypothèses de modélisation. L'avantage de cette approche est que les schémas obtenus sont assurés d'avoir un temps de calcul évoluant linéairement avec le nombre de cellules et ne nécessitent pas d'effectuer une projection globale sur toutes les cellules, très coûteuse en temps de calcul.

5.1 Discrétisation, hiérarchie et positions du problème

On se place sur un domaine $\Omega = [0, 1]^2$, discrétisé avec un maillage cartésien de pas d'espace h , comme illustré en figure 5.3. On note Γ la sortie du domaine, c'est-à-dire une partie connexe du bord de Ω .

Temps de trajet à la sortie, hiérarchie. On note $T(x)$ le temps de trajet le plus court à Γ en tout point de Ω . Ce temps est donné par la solution d'une équation eikonale :

$$|\nabla T(x)| = \frac{1}{c} \quad \forall x \in \Omega \quad (5.4)$$

$$T(x) = 0 \quad \forall x \in \Gamma \quad (5.5)$$

La célérité $c \in \mathbb{R}$ est un paramètre donnant la vitesse de déplacement.

Remarque 5.1. *Il est possible de choisir à la place un champ scalaire dépendant de la position pour prendre en compte des zones de vitesse différentes (par exemple en prenant en compte la densité locale de piétons comme dans le modèle de Hughes [Hug02]), ou de développer des cadres plus généraux permettant de prendre en compte des cas de célérités anisotropes ([Mir]).*

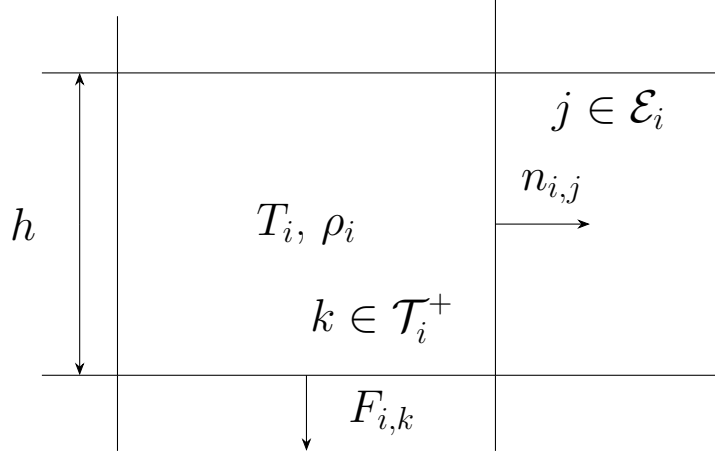


FIGURE 5.3 – Représentation d'une maille et des quantités associées

Cette équation peut se discrétiser sur un maillage cartésien et être résolue avec une méthode de Fast Marching, que nous décrivons plus en détail dans l'annexe A. Nous supposons ici que pour chaque maille i la donnée de T_i , l'approximation numérique du temps de trajet le plus court entre cette maille et la sortie, est connue.

La vitesse souhaitée est définie à partir de ce temps :

$$U = -c^2 \nabla T. \quad (5.6)$$

Dans la suite, nous prenons $c = 1$ afin d'alléger les notations. Pour toute cellule du maillage, on note $\mathcal{E}_i$ l'ensemble de ses cellules voisines. Soit alors $n_{i,j}$ la normale à l'interface entre les deux cellules de i vers $j \in \mathcal{E}_i$. La composante de la vitesse selon cette normale est approchée par :

$$U_{i,j} = -\frac{T_j - T_i}{h} \quad (5.7)$$

Remarque 5.2. *Les conventions de notation suivies dans ce chapitre ne sont pas usuellement prises pour des maillages cartésiens, puisque les indices i et j ne désignent pas les lignes ou colonnes de la matrice de positions. Ce choix est fait pour simplifier les notations par la suite lorsque la notion de résolution frontale apparaîtra.*

Nous trions ensuite les cellules du maillage selon T_i , qui vérifient pour la suite :

$$\forall i, j \ / \ 1 \leq i < j \leq N, \quad T_i \leq T_j. \quad (5.8)$$

Précisons que ce tri n'est pas unique, du fait de la présence possible de cas d'égalité. Nous discuterons au cas par cas l'influence du tri sur les résultats obtenus.

Volumes finis. Nous choisissons d'adopter une discrétisation de type volumes adaptée à l'équation de transport :

$$\partial_t \rho + \nabla \cdot (U \rho) = 0. \quad (5.9)$$

On note ρ_i la variable discrète associée à la densité au sein d'une maille i . La version discrétisée de (5.9) au sens des volumes finis est donnée de façon générale par :

$$\frac{\rho_i^{n+1} - \rho_i^n}{\Delta t} + \frac{1}{h} \left(\sum_{j \in \mathcal{E}_i} F_{i,j}^n \right) = 0 \quad \forall i, \ 1 \leq i \leq N, \quad (5.10)$$

$j \in \mathcal{T}_i^+$ $T_l = 3$	i $T_i = 2$	$l \in \mathcal{T}_i^0$ $T_j = 2$
	$k \in \mathcal{T}_i^-$ $T_k = 1$	
	Γ	Γ

FIGURE 5.4 – Représentations des différents types de voisinages pour une même cellule.

avec un choix décentré amont pour le flux entre deux interfaces :

$$F_{i,j} = F(\rho_i, \rho_j, U_{i,j}) = \begin{cases} \rho_i U_{i,j} & \text{si } U_{i,j} \geq 0 \\ U_{i,j} \rho_j & \text{si } U_{i,j} < 0. \end{cases}$$

En tenant compte de la structure du problème, nous pouvons reformuler cette écriture. Nous introduisons pour chaque cellule i plusieurs types de cellules voisines :

$$\mathcal{T}_i^- = \{k \in \mathcal{E}_i, T_k < T_i\}, \quad (5.11)$$

$$\mathcal{T}_i^+ = \{j \in \mathcal{E}_i, T_j > T_i\}, \quad (5.12)$$

$$\mathcal{T}_i^0 = \{l \in \mathcal{E}_i, T_l = T_i\}, \quad (5.13)$$

Ces différents types de cellules permettent de distinguer les cellules selon leur contribution par rapport à la cellule i , comme illustré en figure 5.4. Les cellules dans $\mathcal{T}_i^-$ ont une valeur de temps de trajet à la sortie plus faible, donc la vitesse souhaitée est dirigée pour ces mailles vers l'extérieur de la cellule i . Ces cellules précéderont de plus toujours la cellule i dans la hiérarchie. Les cellules appartenant à $\mathcal{T}_i^+$ seront à l'inverse en amont : la vitesse souhaitée à l'interface avec ces mailles pointe vers la cellule i , et ces cellules seront toujours après i dans la hiérarchie.

Les cellules $\mathcal{T}_i^0$ sont un cas particulier : le tri hiérarchique ne permet pas de dire si elles seront avant ou après i dans la hiérarchie et la vitesse souhaitée sera nulle à l'interface.

En utilisant ces notations, la formulation volumes finis s'écrit sous une forme montrant plus explicitement les contributions de chaque type de cellules :

$$\frac{\rho_i^{n+1} - \rho_i^n}{\Delta t} = \frac{1}{h} \left(- \sum_{k \in \mathcal{T}_i^-} F_{i,k}^n + \sum_{j \in \mathcal{T}_i^+} F_{j,i}^n - \sum_{l \in \mathcal{T}_i^0} F_{i,l}^n \right) \quad \forall i, 1 \leq i \leq N. \quad (5.14)$$

Avec cette formulation, chacun des termes de flux est positif ou nul.

Afin de calculer la vitesse effective, nous allons calculer les vitesses aux interfaces en fonction de la contrainte de densité maximale. La vitesse effective à l'interface entre les cellules i et j sera notée $v_{i,j}$, et le flux induit $f_{i,j}$ tel que :

$$f_{i,j} = F(\rho_i, \rho_j, v_{i,j}) = \begin{cases} \rho_i v_{i,j} & \text{si } v_{i,j} \geq 0 \\ v_{i,j} \rho_j & \text{si } v_{i,j} < 0. \end{cases}$$

Nous montrerons dans la prochaine section différentes méthodes pour calculer cette vitesse effective. En notant ρ_{max} la densité maximale, la contrainte de non-dépassement lors d'une itération en temps pour la densité ρ_i^{n+1} revient pour tout n à :

$$\rho_i^n + \frac{\Delta t}{h} \left(- \sum_{k \in \mathcal{T}_i^-} f_{i,k}^n + \sum_{j \in \mathcal{T}_i^+} f_{j,i}^n - \sum_{l \in \mathcal{T}_i^0} f_{i,l}^n \right) \leq \rho_{max} \quad \forall i, 1 \leq i \leq N. \quad (5.15)$$

Conditions au bord et cellules voisines à T_i égaux. Nous imposons pour une interface avec un mur une vitesse nulle : $v_{i,j} = 0$, et pour la sortie du domaine $v_{i,j} = U_0$ fixé.

Nous imposons de plus une condition de flux nul entre deux cellules pour lesquelles T_i est identique. Cela signifie que par la suite nous aurons toujours :

$$f_{i,l}^n = 0, \quad \forall l \in \mathcal{T}_i^0, \quad \forall i \leq N$$

5.2 Différentes approches

Première approche : schéma *push*

Un premier essai pour ce modèle consiste à calculer pour chaque cellule les vitesses à l'interface avec les cellules en aval. Ce modèle imite la progression de la phase d'inhibition du modèle microscopique, où chacun cherche à avancer le plus possible en fonction de la place disponible. Nous traduisons cela d'un point de vue macroscopique par l'idée que les individus dans chaque cellule cherchent à pousser autant que possible les individus dans la cellule de devant (d'où le nom du schéma *push*).

Plus précisément, pour chaque cellule i remontée dans le sens de la hiérarchie :

- les cellules $k \in \mathcal{T}_j^-$ sont en aval et ont déjà été traitées précédemment dans la hiérarchie,
- pour chaque k , la vitesse effective $v_{i,k}^n$ est la plus proche de $U_{i,k}$ et pour laquelle le flux induit $f_{i,k}$ n'entraîne pas une violation de la contrainte de densité maximale,
- une fois $v_{i,k}$ déterminé, le flux $f_{i,k}$ est advecté entre les deux cellules.

Un exemple de la progression obtenue est donnée en figure 5.5.

Notons que le troisième point implique que les densités des cellules évoluent au cours d'un parcours complet de l'ensemble des cellules. Ce choix est fait pour qu'à chaque étape, la contrainte de densité maximale prenne en compte tous les flux déjà calculés dans la hiérarchie.

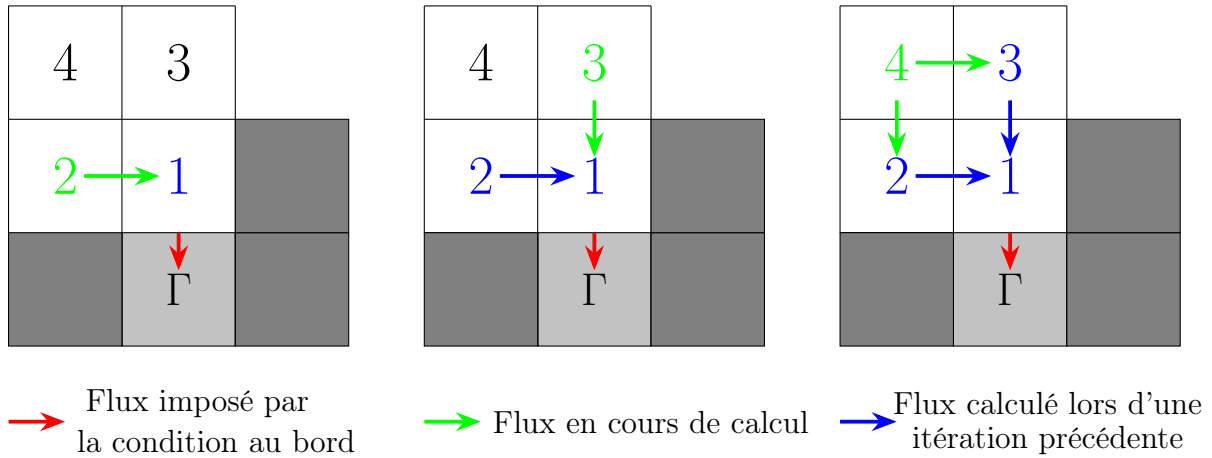


FIGURE 5.5 – Calcul des projections dans la hiérarchie pour le schéma *push*. Pour la cellule 1, la vitesse est donnée par la condition au bord. Pour les cellules 2 et 3, la vitesse est calculée pour ne pas donner lieu à une sur-densité dans la cellule 1. Enfin la cellule 4 calcule pour chaque face la vitesse qui ne va pas faire dépasser la densité maximale pour les cellules 2 et 3.

La densité de la cellule j à l'étape i de parcours de la hiérarchie sera notée $\rho_j^{n,i}$ par la suite. Avec cette notation, lors du parcours hiérarchique pour chaque étape i , nous calculons la valeur de $\rho_i^{n,i+1}$ ainsi que celles de $\rho_k^{n,i+1} \forall k \in \mathcal{T}_i^-$ en allant de $i = 0$ à $i = N$.

L'expression de $v_{i,k}^n$ peut être donnée exactement. Reprenons l'inégalité (5.15) exprimée du point de vue d'une cellule $k \in \mathcal{T}_i^-$ à l'étape i de l'itération dans la hiérarchie, en ne conservant que le terme de flux à l'interface entre les deux :

$$\rho_j^{n,i} + \frac{\Delta t}{h} (f_{i,k}^n) \leq \rho_{max}.$$

Cette équation peut se développer et se mettre sous la forme :

$$v_{i,k}^n \leq \frac{h}{\Delta t} \frac{\rho_{max} - \rho_k^{n,i}}{\rho_i^{n,i}}.$$

La vitesse souhaitée peut ainsi s'écrire pour chaque étape dans le parcours hiérarchique :

$$v_{i,k}^n = \min \left(U_{i,k}, \frac{h}{\Delta t} \frac{\rho_{max} - \rho_k^{n,i}}{\rho_i^{n,i}} \right).$$

L'algorithme complet de calcul des vitesses admissibles est donné dans l'algorithme 1.

Nous pouvons montrer que la contrainte de densité maximale est respectée lors du calcul de ρ_i^{n+1} pour tout i .

Proposition 5.1. *Pour tout n, i , sous la condition CFL $\max_{i,k} |U_{i,k}| \leq \frac{\Delta x}{\Delta t}$:*

$$0 \leq \rho_i^n \leq \rho_{max} \implies 0 \leq \rho_i^{n+1} \leq \rho_{max}$$

Preuve. La preuve consiste en une récurrence durant les étapes de parcours hiérarchique de l'algorithme. Initialement nous avons $\rho_j^{n,0} = \rho_j^{n,0} \forall j \leq N$. Pour la première étape, la cellule $i = 0$ est nécessairement en bordure du domaine¹. Toutes les cellules voisines de la cellule $i = 0$ sont dans Γ et la valeur du flux est fixée par la condition au bord. On a :

$$\rho_0^{n,1} = \rho_0^{n,0} \left(1 - \frac{\Delta t}{h} \sum_{k \in \mathcal{T}_0^-} v_{i,k}^n \right),$$

qui assure $\rho_0^{n,1} \geq 0$ grâce à la condition CFL. Supposons qu'à l'étape i maintenant la condition $0 \leq \rho_j^{n,i} \leq \rho_{max}$ soit vérifiée pour tout j . L'itération à l'étape i donne :

$$v_{i,k}^n = \min \left(U_{i,k}, \frac{h}{\Delta t} \frac{\rho_{max} - \rho_k^{n,i}}{\rho_i^{n,i}} \right) \quad \forall k \in \mathcal{T}_i^- \quad (5.16)$$

$$\rho_i^{n,i+1} = \rho_i^{n,i} \left(1 - \frac{\Delta t}{h} \sum_{k \in \mathcal{T}_i^-} v_{i,k}^n \right) \quad (5.17)$$

$$\rho_k^{n,i+1} = \rho_k^{n,i} + \left(\frac{\Delta t}{h} \rho_i^{n,i} v_{i,k}^n \right) \quad \forall k \in \mathcal{T}_i^- \quad (5.18)$$

Le calcul de $v_{i,k}$ assure $\rho_k^{n,i+1} \leq \rho_{max}$, et comme $\rho_k^{n,i} \leq \rho_{max}$, l'équation (5.16) donne aussi $v_{i,k} \geq 0$ et assure la positivité de $\rho_i^{n,i+1}$. La condition CFL assure $\rho_i^{n,i+1} \geq 0$, et $v_{i,k} \geq 0$ implique $\rho_i^{n,i+1} \leq \rho_i^{n,i} \leq \rho_{max}$ aussi, permettant de conclure la récurrence.

Données : ρ_i^n pour tout i , $U_{i,k}$ pour tout i, k

Résultat : ρ_i^{n+1} pour tout i , $v_{i,k}^n$ pour tout i, k

```

1  $\tilde{\rho}_i \leftarrow \rho_i^n$  pour tout  $i$ 
2 pour  $1 \leq i \leq N$  faire
3   pour  $k \in \mathcal{T}_i^-$  faire
4      $v_{i,k}^n \leftarrow \min \left( U_{i,k}, \frac{h}{\Delta t} \frac{\rho_{max} - \tilde{\rho}_k}{\tilde{\rho}_i} \right)$ 
5      $\tilde{\rho}_k \leftarrow \tilde{\rho}_k + \frac{\Delta t}{h} \tilde{\rho}_i v_{i,k}^n$ 
6      $\tilde{\rho}_i \leftarrow \tilde{\rho}_i - \frac{\Delta t}{h} \tilde{\rho}_i v_{i,k}^n$ 
7   fin
8 fin
9  $\rho_i^{n+1} \leftarrow \tilde{\rho}_i$  pour tout  $i$ 
```

Algorithme 1 : Algorithme d'itération dans la hiérarchie pour le schéma *push*.

1. Il peut cependant y avoir d'autres cellules au bord du domaine après la cellule 0.

Limites de l’algorithme. Un résultat de simulation est présenté en figure 5.6, dans le cas d’une salle simple avec une seule sortie. Dans un premier temps la densité est advectée vers la porte et la vitesse est diminuée dans les zones de congestion. Une fois le bouchon formé, celui-ci se vide au fur et à mesure par la porte de sortie. Le bouchon se vide en biaisant fortement certaines zones de l’espace, ce qui donne une forme atypique au bouchon en amont de la porte.

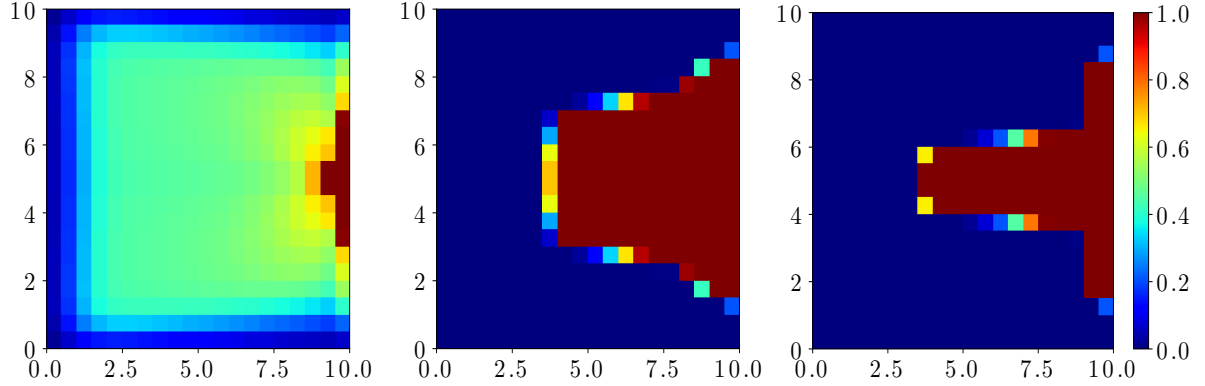


FIGURE 5.6 – Simulation avec le schéma *push*, à $t = 1, 10, 25$ s. La salle est un carré de 10 m de côté, avec une porte de 1 m de côté. La densité initiale est fixée à 0.4 m^{-2} , avec $\Delta t = 0.1$ s et $h = 0.5$ m.

Revenons sur la figure 5.5 pour comprendre ce phénomène. Dans le cas présenté, la cellule 2 est la première pour laquelle un calcul de vitesse est effectué. En cas de situation fortement congestionnée, la vitesse $v_{2|1}$ sera réduite pour ne pas violer la contrainte de densité maximale en cellule 1. En passant ensuite à la cellule 3, la densité en cellule 1 est déjà à sa valeur maximale : la vitesse $v_{3|1}$ est donc nulle. Cet effet se retrouve ensuite à chaque cellule en remontant, favorisant en fonction de la hiérarchie certaines cellules par rapport à d’autres.

Cet effet compromet l’utilisation de cet algorithme dans cet état, du fait de sa dépendance très forte au tri choisi.

Deuxième approche : schéma *pull*

Pour chaque cellule, nous calculons non plus les vitesses aux interfaces avec les cellules en aval, mais celles en amont. Les vitesses sont projetées sur l’ensemble des vitesses n’entraînant pas de violation de la contrainte de densité maximale, en prenant en compte toutes les vitesses susceptibles de violer cette contrainte. Le premier schéma calculait les vitesses en « poussant » autant que possible, tandis que celui-ci « tire » les flux des cellules en amont. Ce nouveau parcours est illustré en figure 5.7.

Projections. Reprenons l’équation générale de la contrainte de densité (5.15) :

$$\rho_i^n + \frac{\Delta t}{h} \left(- \sum_{k \in \mathcal{T}_i^-} f_{i,k}^n + \sum_{j \in \mathcal{T}_i^+} f_{j,i}^n \right) \leq \rho_{max} \quad \forall i, 1 \leq i \leq N.$$

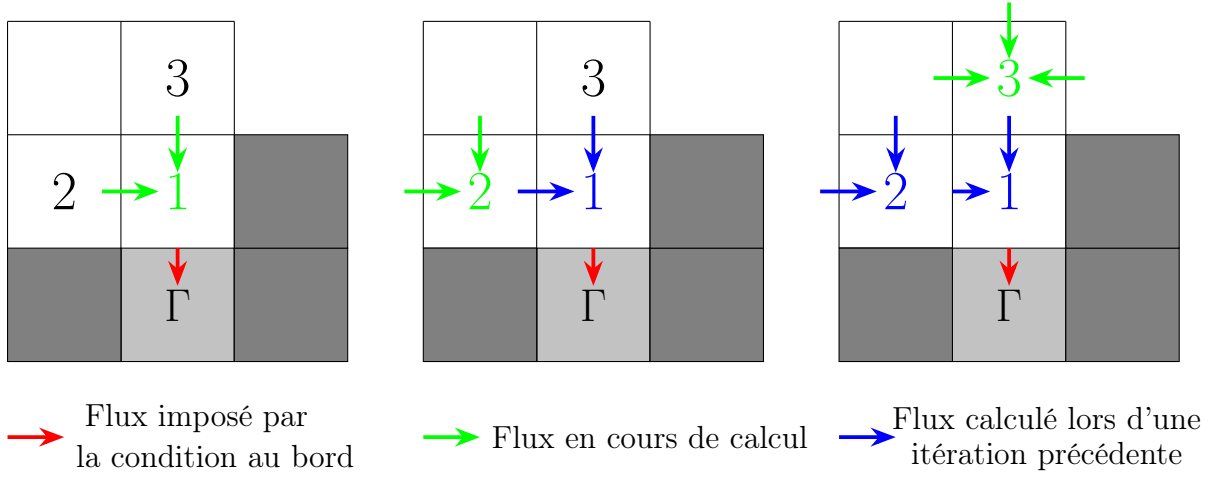


FIGURE 5.7 – Calcul des projections dans la hiérarchie pour le schéma *pull*. Pour la cellule 1, les flux des cellules 2 et 3 sont pris en compte et les vitesses projetées sur la contrainte en densité. Pour les cellules 2 et 3, les flux d'échange avec la cellule 1 sont déjà calculés et pris en compte dans la contrainte.

En séparant les contributions dues à $\mathcal{T}_i^-$ et $\mathcal{T}_i^+$, cette inégalité devient :

$$\sum_{j \in \mathcal{T}_i^+} f_{j,i}^n \leq \frac{h}{\Delta t} (\rho_{max} - \rho_i^n) + \sum_{k \in \mathcal{T}_i^-} f_{i,k}^n \quad \forall i, 1 \leq i \leq N. \quad (5.19)$$

Lorsque les flux $f_{i,k}^n$, $k \in \mathcal{T}_i^-$ en aval sont connus, l'inégalité (5.19) est une contrainte affine sur les flux $f_{j,i}^n$, $j \in \mathcal{T}_i^+$ en amont.

Nous pouvons écrire pour chaque cellule la projection au sens des moindres carrés des flux amont sur ces contraintes pour toute cellule i :

$$(f_{i,j})_{j \in \mathcal{T}_i^+} = \operatorname{argmin}_{f \in \mathcal{C}_i} \frac{1}{2} \sum_j |f_j - F_{j,i}|^2, \quad (5.20)$$

où :

$$\mathcal{C}_i = \mathcal{C}(\rho_i, (f_{i,j}^n)_{j \in \mathcal{T}_i^-}) = \left\{ f \in \mathbb{R}^{|\mathcal{T}_i^+|}, \sum_j f_j \leq \frac{h}{\Delta t} (\rho_{max} - \rho_i^n) + \sum_{k \in \mathcal{T}_i^-} f_{i,k}^n \right\}.$$

Le problème de minimisation sous contraintes (5.20) a une solution unique : à chaque étape, tous les flux $(f_{i,j}^n)_{j \in \mathcal{T}_i^-}$ ont été calculés par une condition au bord ou une itération précédente.

Simulations, limites du modèle. Nous reprenons la même configuration que précédemment en figure 5.8. La densité se concentre selon deux diagonales partant de la porte. Le bouchon se vide ensuite sans les effets de zones vidées différemment par la hiérarchie observées pour le modèle précédent.

L'algorithme développé crée des flux de matière vers l'amont lors des projections successives. Avec la résolution hiérarchique cependant, toute matière remontant vers l'amont est prise en compte plus tard dans la hiérarchie et participera à saturer la condition de densité maximale sur d'autres cellules en amont. Cet effet implique que chaque flux vers l'amont participe potentiellement à l'apparition d'autres flux vers l'amont, comme illustré en figure 5.9. Les cas facilitant l'apparition de flux vers l'amont sont ceux où les vitesses en x et en y ont des valeurs très différentes. La matière sera donc concentrée aux points où la vitesse selon x et y est sensiblement équivalente, ce qui s'avère être le cas pour les diagonales partant des portes pour le champ de vitesse choisi.

Une des conséquences de ce phénomène est de perdre au cours des itérations le principe du maximum qui était observé pour le schéma *push* : prenons deux cellules voisines ρ_i et ρ_k telles que $k \in \mathcal{T}_i^-$ et toutes les deux saturées. La cellule k sera traitée avant la cellule i , donc si $f_{i,k}$ est négatif, cela signifie que :

$$\rho_i - \frac{\Delta t}{h} f_{i,k} \geq \rho_{max}. \quad (5.21)$$

Lorsque la cellule i sera à son tour traitée, le calcul des flux en amont devra répartir cet excès de matière sur d'autres cellules en amont. Ce phénomène peut donner des résultats incohérents dans certains cas près des murs, comme illustré en figure 5.10. Dans le cas extrême où la cellule i n'a pas de cellule en amont, aucun flux en amont ne peut être créé et la densité dépasse la contrainte de densité maximale.

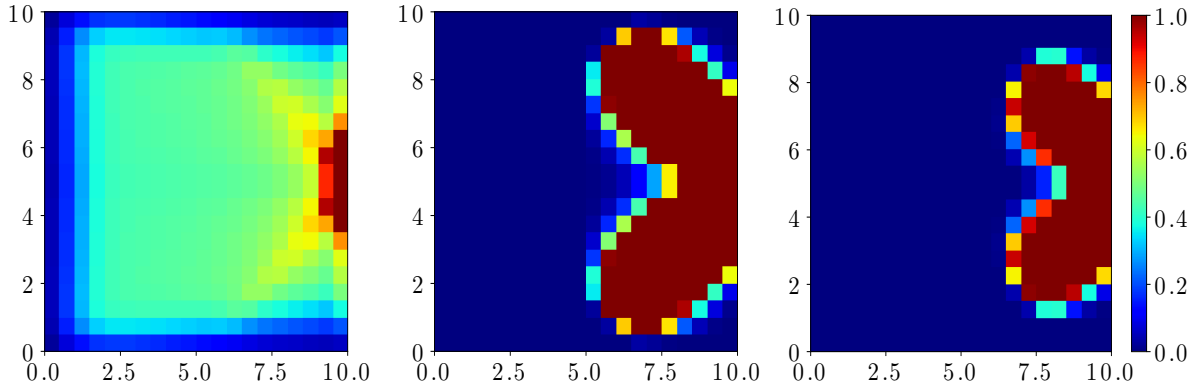


FIGURE 5.8 – Simulation avec le schéma *pull*, à $t = 1, 10, 25$ s. La salle est un carré de 10 m de côté, avec une porte de 1 m de côté. La densité initiale est fixée à 0.4 m^{-2} , avec $\Delta t = 0.1$ s et $h = 0.5$ m.

Correction au schéma *pull* : le schéma *pull* contraint

Nous avons montré que le schéma *pull* crée des flux remontant en amont, vers des cellules qui n'ont pas été calculées dans la hiérarchie. Pour éviter ce problème, nous pouvons alors contraindre les flux obtenus lors de la projection à rester positifs.

Signification en modélisation. Avant d'expliquer comment rajouter cette contrainte dans le modèle, interrogeons-nous sur sa signification du point de vue de la modélisation. Pour rappel, nous projetons les flux en amont d'une cellule pour que ceux-ci n'entraînent pas le dépassement

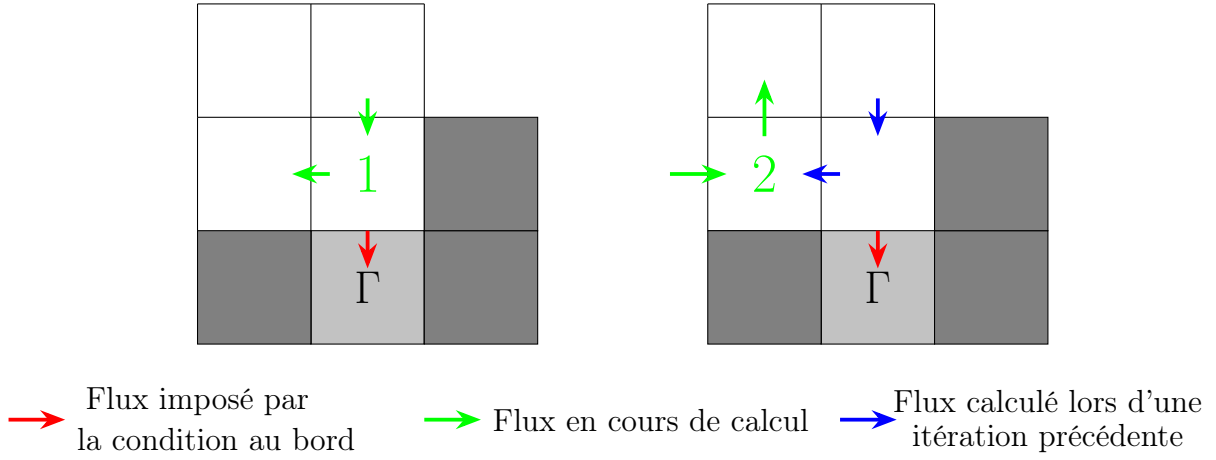


FIGURE 5.9 – Avec l'algorithme *pull*, la première projection pour la cellule 1 peut donner un flux négatif, qui entraîne un mouvement de densité vers la cellule 2. Lors de la projection pour la cellule 2, le flux provenant de la cellule 1 sature la contrainte en densité et cette projection crée d'autres flux négatifs en amont.

de la densité maximale. Si un de ces flux devient négatif, cela signifie que localement une vitesse à l'interface devient négative : les individus sont contraints à reculer par rapport à leur direction souhaitée.

Une explication de ce phénomène est la suivante : les entités modélisées sont susceptibles de se pousser localement en passant d'une cellule à l'autre. Les individus ayant une vitesse souhaitée plus élevée localement poussent la matière à l'extérieur de la cellule, par les autres interfaces où la vitesse souhaitée est plus faible. La positivité d'un flux en amont à une interface force donc les autres flux à ne pas pousser la matière hors de la cellule pour respecter la contrainte de densité maximale.

Écriture du schéma En ajoutant la positivité des flux amont, l'espace des contraintes devient pour toute cellule i :

$$C_i = C(\rho_i, (f_{i,j}^n)_{j \in \mathcal{T}_i^-}) = \left\{ f \in \mathbb{R}^{|\mathcal{T}_i^+|}, f_j \geq 0, \sum_j f_j \leq \frac{h}{\Delta t} (\rho_{max} - \rho_i^n) + \sum_{k \in \mathcal{T}_i^-} f_{i,k}^n \right\},$$

et les flux sont projetés sur cette contrainte :

$$(f_{i,j})_{j \in \mathcal{T}_i^+} = \operatorname{argmin}_{f \in C_i} \frac{1}{2} \sum_j |f_j - F_{j,i}|^2. \quad (5.22)$$

Il s'agit pour chaque cellule i d'un système de contraintes affines de difficulté comparable au schéma *pull* précédent.

Simulations. Nous reprenons en figure 5.11 les résultats de ce modèle pour le même cas test. La congestion apparaît toujours devant la porte de sortie, mais les densités ne sont plus regroupées uniquement le long des diagonales de la porte.

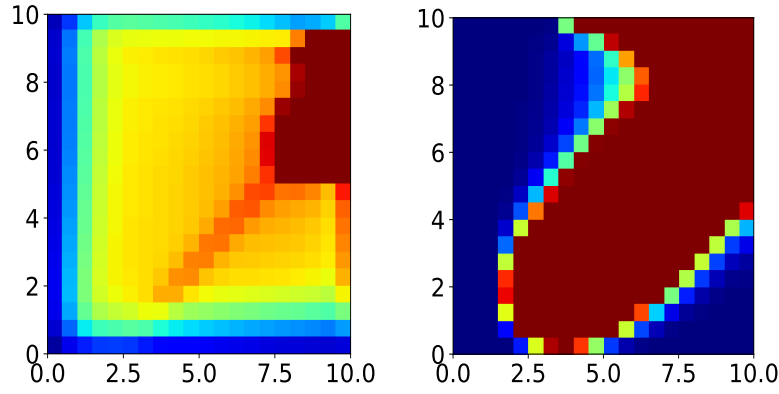


FIGURE 5.10 – Simulation avec le schéma *pull*, à $t = 1, 6$ s, avec les mêmes paramètres que pour la figure 5.8, et une densité initiale à 0.6 m^{-2} et une porte de sortie centrée à $y = 8 \text{ m}$.

Ce modèle est préférable au modèle *pull* précédent : la remontée de matière au cours des itérations hiérarchique n'est plus possible et l'on assure en tout temps que la densité obtenue respecte effectivement la contrainte de densité maximale. Comparons au modèle *push* pour comprendre la différence entre les deux : dans le modèle *push*, les flux entre une cellule i et ses cellules en amont sont calculés individuellement au cours du parcours de la hiérarchie. La première cellule traitée en amont aura la possibilité de transmettre plus de matière et les cellules suivantes ont une quantité de matière disponible plus petite. Dans le modèle *pull* contraint, toutes les cellules en amont sont traitées au cours de la même itération.

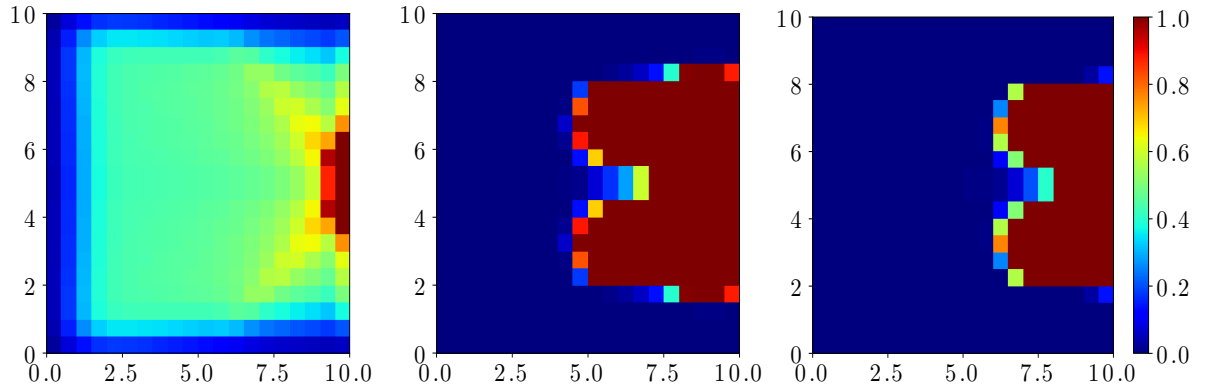


FIGURE 5.11 – Simulation avec le schéma *pull* contraint, à $t = 1, 10, 25$ s. La salle est un carré de 10 m de côté, avec une porte de 1 m de côté. La densité initiale est fixée à 0.4 m^{-2} , avec $\Delta t = 0.1 \text{ s}$ et $h = 0.5 \text{ m}$.

5.3 Discussion, perspectives

Comparaison des approches. Les deux approches *push* et *pull* contraint conduisent à des comportements différents, bien que leur construction repose sur des arguments similaires. En l’absence d’un modèle de référence, nous les comparons ici à un autre modèle issu de la littérature pour comprendre leur différence. En section 3.2 nous avons comparé le modèle d’inhibition macroscopique au modèle LWR, plus connu dans la littérature. Pour rappel, ce dernier modélise le mouvement de la densité à l’aide d’une relation entre densité ρ et vitesse $v(\rho)$. Nous avons montré qualitativement comment ce modèle d’inhibition s’apparente au modèle LWR pris lorsque la relation entre vitesse et densité devient de plus en plus raide.

En dimension 2, le modèle LWR peut s’écrire :

$$\partial_t \rho + \nabla \cdot (\rho v(\rho)),$$

où U est la même vitesse souhaitée, et $v(\rho)$ peut par exemple être donnée par :

$$v_\varepsilon(\rho) = U \left[1 - \exp \left(- \left(\frac{1}{\rho} - \frac{1}{\rho_{max}} \right) / \varepsilon \right) \right]. \quad (5.23)$$

Cette équation peut être discrétisée par exemple sur un maillage cartésien et résolue avec un schéma explicite, dont nous donnons les résultats en figure 5.12 pour $\varepsilon = 0.01$. Avec une vitesse aussi raide, les conditions de stabilité imposent l’utilisation de pas de temps très petits.

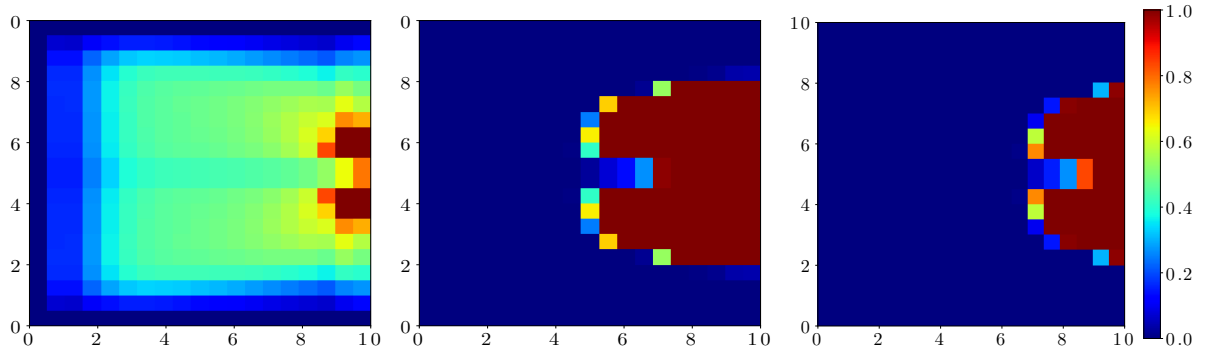


FIGURE 5.12 – Simulation avec le modèle LWR à 2 dimensions, à $t = 1, 10, 25$ s. La salle est un carré de 10 m de côté, avec une porte de 1 m de côté. La densité initiale est fixée à 0.4 m^{-2} , avec $\Delta t = 0.1$ s et $h = 0.5$ m.

Les résultats du schéma *pull* contraint sont qualitativement proches de ceux obtenus avec ces simulations. Cette analogie entre les deux modèles permet de plus d’avoir un premier point de validation qualitative : l’algorithme ainsi développé modélise des situations où les individus dans la foule cherchent à aller le plus vite possible que l’espace disponible leur permet. Cette comparaison ne remet cependant pas en cause l’utilité du modèle *push* en soi.

Perspectives

Les modèles décrits ici ne tiennent compte que de la projection de la vitesse souhaitée sur les vitesses admissibles, avec une résolution hiérarchique. Plusieurs extensions sont envisageables afin de continuer à développer ces modèles.

Nous avons toujours utilisé un champ de vitesse souhaitée fixe en temps jusqu'ici. La principale contrainte imposée par le schéma numérique est de pouvoir définir un parcours hiérarchique des cellules du maillage. Tant que la direction du champ de vitesse souhaitée est définie à partir du gradient d'un temps de trajet à la sortie. Dans le modèle de Hughes ([Hug02] et voir section 2.2) une modification proposée est de prendre en compte dans l'estimation du temps de trajet à la sortie l'influence de la densité sur la vitesse locale. La direction est ensuite donnée par le gradient normalisé de ce temps de trajet.

Appliquer cette modification avec l'un des schémas décrits plus haut changerait le calcul des valeurs de T_i . Cependant, tant qu'il est possible de définir une hiérarchie pour les cellules le processus de projection hiérarchique reste inchangé. D'autres méthodes de modification du champ de vitesse souhaitée (par des modèles de chimiotaxie par exemple) ne sont pas immédiatement applicables à ce modèle.

D'autres modifications possibles peuvent porter sur la variation de la norme du champ de vitesses souhaitée. Le schéma numérique global se décomposerait par exemple en deux étapes :

- le calcul du champ de vitesse souhaitée en fonction de la densité,
- le calcul du champ de vitesse effective avec le schéma numérique de projection hiérarchique.

Dans le modèle de Hughes, cette norme est donnée par une relation à la densité locale. Il serait possible de modifier de la même façon le champ de vitesse souhaitée à chaque étape, pour prendre en compte un effet de ralentissement dû à l'augmentation de la densité. Dans cette optique, aux basses densités les effets sociaux dominant et la vitesse est donnée par une relation $U(\rho)$, tandis que pour les hautes densités les effets physiques de la congestion sont calculés par l'étape de projection. Les vitesses utilisées usuellement pour le modèle de Hughes s'annulent cependant pour une certaine valeur ρ_{max} de la densité, pour que la projection ait un intérêt, il est nécessaire de choisir une fonction de vitesse appropriée et adaptée à la situation.

Une autre possibilité serait de modifier le parcours hiérarchique. Nous avons seulement considéré le cas où la hiérarchie est fixée en amont et donnée par la valeur du temps de trajet à la sortie. Serait-il possible de relâcher cette contrainte dans la formulation d'un modèle? Une piste serait de prendre en compte un ensemble de cellules se propageant de proche en proche, à l'instar de l'algorithme de Fast Marching.

Chapitre 6

Conclusion générale et perspectives

6.1 Conclusion sur les études menées

Différents aspects de la modélisation et la simulation de mouvements de foules ont été abordés dans cette thèse. Le couplage de modèles microscopiques et macroscopiques au passage d'une interface a été étudié dans un premier temps. Dans un sens, ce passage nécessite de décrire une procédure d'absorption d'entités lagrangienne dans un milieu eulérien. Les individus ont une vitesse attribuée en fonction du modèle macroscopique, tandis qu'un flux au travers de l'interface est exprimé en cohérence. À l'inverse, lorsque le flux de matière traverse l'interface en provenant du continuum, les entités lagrangiennes sont générées à la volée au niveau de l'interface.

Le chapitre 3 porte sur l'étude des couplages à une dimension, dans un premier lieu pour modèles d'anticipation de collision (LWR et FTL) issus de la littérature. Nous proposons un schéma numérique couplé pour les deux modèles et ne créant pas d'oscillations au niveau de l'interface entre les modèles. Dans un second temps des modèles granulaires avec inhibition sont couplés aussi sans apparition d'artéfact à l'interface. Ces deux types de modèles sont comparés sur un cas applicatif, puis des extensions possibles sur des jonctions entre chemins sont discutées.

Dans le chapitre 4 le modèle granulaire avec inhibition à deux dimensions est couplé à une interface avec son analogue macroscopique à une dimension. Dans le cas du passage en eulérien, la condition de couplage à l'interface est basée sur la contrainte de non-dépassement de la densité maximale pour le modèle macroscopique. Cette contrainte s'exprime pour les entités lagrangiennes sur la composante normale à l'interface de leur vitesse. Dans le sens inverse, nous proposons une méthode pour générer à la volée des individus sans recouvrement. Cette génération crée des oscillations dans la densité, mais l'information d'un bouchon est correctement propagée.

Le chapitre 5 montre différentes approches pour développer un modèle numérique inspiré du modèle macroscopique avec inhibition. L'espace est divisé en cellules pour lesquelles sont déterminée une hiérarchie basée sur le temps de trajet entre chaque cellule et la sortie. Les flux entre cellules sont calculés en itérant selon l'ordre donné par la hiérarchie. Pour chaque cellule, les flux en amont ou en aval peuvent être calculés en forçant un respect local de la contrainte de densité maximale autorisée.

6.2 Perspectives - couplages et modélisation de foules

Développements pour les couplages. Plusieurs pistes restent à explorer sur le développement de couplages pour le futur, dans la suite des travaux du chapitre 4. Une première direction est de coupler le modèle d’inhibition avec les modèles macroscopiques numériques développés dans le chapitre 5. Ce couplage est cependant complexe à réaliser tant que le modèle macroscopique d’inhibition n’a pas été clairement identifié (comme expliqué en début de chapitre 5). Ce passage présente des difficultés techniques de mise en place, mais une question importante est dans la modélisation de l’interaction entre les flux d’origine microscopique et macroscopique. L’ordre dans lequel effectuer la remontée hiérarchique ainsi que la place que prend la projection granulaire du modèle microscopique et l’écriture des contraintes de contact ont un effet sur les flux de passage. Le cas d’application d’évacuation de tram qui est présenté montre une première approche de ces questions, dans laquelle nous avons observé notamment qu’avec le modèle de croisement de flux proposé le modèle microscopique semblait laisser le passage aux individus provenant du modèle macroscopique. L’origine de ce phénomène dans le modèle de croisement n’est pas clairement identifiée à l’heure actuelle et un travail sur la modélisation de ces croisements de flux microscopique et macroscopiques est encore à faire.

Un autre point important pour la suite est dans le couplage de modèle de nature différentes, comme le modèle d’inhibition en microscopique et le modèle de Hughes en macroscopique par exemple. L’écriture de ces couplages en modélisation n’est pas unique et l’écriture du couplage à l’interface peut être déterminante sur le comportement des solutions. Des travaux sont en cours pour écrire ces couplages, mais aussi bien dans du microscopique vers le macroscopique que dans l’autre sens, le calcul du flux instantané mène à des effets d’oscillations. Pour des modèles macroscopiques avec des phénomènes d’anticipation, ces oscillations ont des conséquences importantes. Des solutions pour parvenir à lisser numériquement l’expression de ces flux sont à l’étude.

Étude du modèle macroscopique d’inhibition. Plusieurs approches numériques ont été proposées dans le chapitre 5 pour écrire un équivalent macroscopique au modèle d’inhibition. Parmi les deux approches retenues, le modèle *pull* contraint est celui qui semble le plus abouti et prêt à l’utilisation dans sa forme actuelle.

Le modèle continu reste cependant toujours à écrire et à étudier théoriquement. Cela permettrait d’identifier si les algorithmes proposés sont effectivement des discrétisations pertinentes de ce modèle. Selon les équations utilisées pour écrire le modèle, l’apparition de la hiérarchie notamment n’est pas nécessairement évidente dans la construction du modèle et serait une propriété à vérifier. Ces questions sont en cours d’étude dans une thèse commencée cette année.

Mise en application des modèles. Une des motivations à l’origine de cette thèse fût l’arrivée prochaine des Jeux Olympiques à Paris en 2024. À cette occasion, de nombreuses *fan zones* sont en cours de planification et les besoins de réponse sur les problématiques de sécurité est grandissant. De nombreux modèles microscopiques sont disponibles et sont couramment utilisés à l’heure actuelle et plusieurs projets permettent même d’utiliser plusieurs modèles conjointement ([KZGK19; vGG+20; FM] entre autres). Certains calculs de grande ampleur sont à l’heure actuelle en cours, mais nous n’avons pas l’autorisation de les montrer dans ce manuscrit. Ces calculs nous ont en revanche permis d’identifier plusieurs problématiques à étudier dans les prochaines années pour pouvoir intégrer la simulation de foule aux démarches des responsables de la sécurité des événements.

Une première piste importante d'étude est la détermination des sorties empruntées par les individus lors d'une évacuation. Cette question est encore délicate, puisque le choix de la sortie à emprunter pour chaque individu peut être influencé par un grand nombre de facteurs environnementaux et psychologiques encore très durs à modéliser de façon systématiques. La majorité des modèles mis en avant s'intéressent aux interactions avec les individus proches et aux changements de trajectoire locaux, mais le choix de l'objectif global vers lequel se diriger reste encore peu abordé par les modèles microscopiques. Durant cette thèse, nous avons en parallèle travaillé sur des nouveaux modèles pour déterminer ce choix de la sortie à emprunter en cas d'évacuation avec un stagiaire (Adrien Gautier).

Dans ce modèle, chaque individu est capable d'attribuer à chaque sortie un temps estimé d'évacuation, pondéré par le temps de trajet pour aller à cette sortie et le temps estimé de congestion dû aux autres individus se dirigeant vers cette même sortie et plus proche de cette sortie. En formalisant ce problème dans un cadre de théorie des jeux, chaque sortie devient une stratégie dépendante des autres joueurs. Sous quelques hypothèses simples sur les coûts et placements des joueurs, il est possible de montrer une solution unique à ce jeu dans la plupart des situations¹. Des premiers résultats appliqués sont encourageants (voir figure 6.1) et d'autres approches et perfectionnement dans cette direction sont envisagés pour l'avenir.

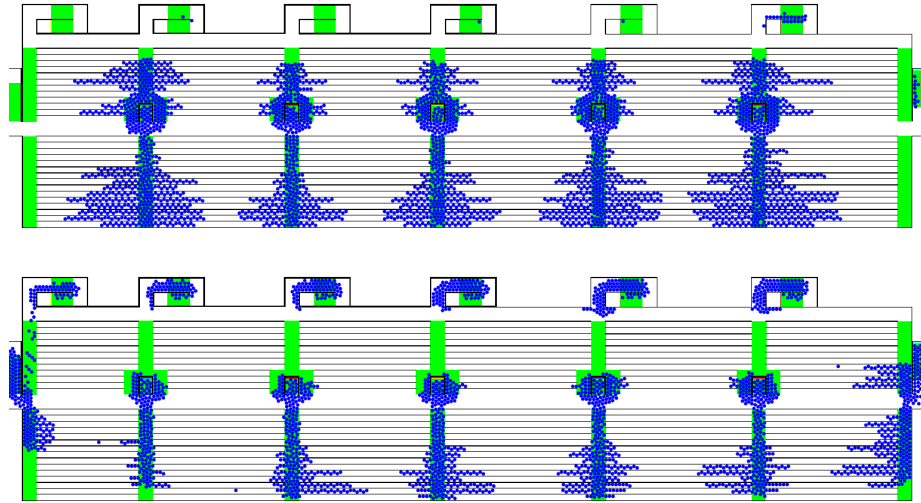


FIGURE 6.1 – Résultats de simulation de gradins avec le modèle d'inhibition, où les individus se dirigent vers la sortie la plus proche (au-dessus) ou en choisissant leur sortie avec l'algorithme de théorie des jeux (en-dessous). Les zones vertes indiquent des escaliers, où la vitesse de marche est réduite. Dans le premier cas les individus se concentrent vers les vomitoires au milieu des gradins, et n'empruntent pas les sorties du haut. Sur la figure du bas les individus sont capables de savoir qu'une sortie plus éloignée, mais peu utilisée, est favorable.

Un second point bloquant pour les études de mouvements de foules est l'accès à des données

1. La preuve repose sur le fait qu'il est possible d'éliminer itérativement des stratégies pénalisantes pour chaque joueur en partant de ceux les plus près de la porte de sortie. Cette preuve donne de ce fait aussi une méthode effective de calcul de cette stratégie en parcourant frontalement les individus.

pertinentes et des expériences répétables : si aujourd’hui certains sujets comme le flux de passage par une porte ont fait l’objet de beaucoup d’expériences, beaucoup de configurations envisagées pour des fan zones à venir ne sont pas couvertes par des expériences de la littérature. La mise en place de ces expériences peut cependant être trop complexes à réaliser dans les temps ou trop coûteuse en moyen. La réalité virtuelle permet de contourner ces difficultés de mise en place en plongeant numériquement des individus dans une situation donnée. Une thèse dans la suite des travaux de [Ber20] a été entamée au LCPP pour étudier ces pistes en détail.

Une troisième piste de développement concerne l’utilisation plus fréquente de modèles macroscopiques. Peu de logiciels de simulation macroscopiques sont proposés dans le cadre de la simulation de mouvements de foules. Ces modèles apportent une solution complémentaire des simulations microscopiques pour étudier la viabilité de configurations données, mais sont à notre connaissance peu utilisés par les industriels du métier. À ce titre une librairie proposant quelques modèles provenant de la littérature et des outils facilitant la mise en place d’une simulation à partir d’un plan sont à développer. Nous envisageons à terme un apport à la librairie python *cromosim* ([FM]) qui propose déjà l’implémentation de plusieurs modèles microscopiques.

Une dernière piste plus ouverte repose sur une réflexion plus générale de l’utilisation des outils de calculs : à l’heure actuelle, beaucoup d’outils se concentrent sur la simulation du mouvement à partir d’une configuration donnée établie à l’avance. À notre connaissance, peu d’outils proposent des méthodes afin d’optimiser la configuration en elle-même, comme cela peut être fait dans la thèse de M. Mimault ([Mim, Chapitre 3]), où une optimisation sur la répartition de la densité initiale est proposée. Ce type d’outils, plus complexe à mettre en place, serait cependant plus proches des problématiques des différents services responsables de la prévention qui cherchent des méthodes de guidage efficaces ou souhaitent comment améliorer les configurations qui leur sont proposées pour améliorer la sécurité des événements.

Annexe A

Algorithme de Fast Marching

Pour la modélisation d'évacuation de foules, le calcul de la vitesse souhaitée peut se faire à partir du temps de trajet à la plus proche sortie $T(x)$:

$$U = -c^2 \nabla T(x),$$

où c est la célérité des piétons et $T(x)$ est donné par l'équation eikonale :

$$|T(x)| = 1/c. \quad (\text{A.1})$$

Pour simplifier nous prenons $c = 1$ dans la suite.

Les solutions de cette équation peuvent être approchées sur un maillage cartésien avec l'algorithme de Fast Marching ([Set99]). Cet algorithme consiste à calculer au fur et à mesure le temps de trajet $T(x)$, en propageant depuis la sortie de proche en proche. Les points du maillage sont séparés en trois zones évoluant au cours des itérations :

- une zone d'ombre, où la valeur de T n'est pas encore calculée et est fixée à $+\infty$,
- une zone éclairée, où T est calculé et fixé,
- une zone de pénombre où T est calculé, mais peut encore évoluer.

À l'initialisation, la zone éclairée est la sortie, les cellules voisines des cellules à la sortie sont dans la zone de pénombre et le reste des cellules sont dans la zone d'ombre.

Calcul de T . En notant $T_{i,j}$ la valeur approchée de T au centre de la cellule (i, j) et h le pas du maillage, la version discrète de (A.1) est donnée par :

$$\max \left(\tau_{i,j}^{-x}, -\tau_{i,j}^{+x}, 0 \right)^2 + \max \left(\tau_{i,j}^{-y}, -\tau_{i,j}^{+y}, 0 \right)^2 = 1, \quad (\text{A.2})$$

avec :

$$\tau_{i,j}^{+x} = \frac{T_{i+1,j} - T_{i,j}}{h} \quad \text{et} \quad \tau_{i,j}^{-x} = \frac{T_{i,j} - T_{i-1,j}}{h}, \quad (\text{A.3})$$

$$\tau_{i,j}^{+y} = \frac{T_{i,j+1} - T_{i,j}}{h} \quad \text{et} \quad \tau_{i,j}^{-y} = \frac{T_{i,j} - T_{i,j-1}}{h}. \quad (\text{A.4})$$

Ici $\tau_{i,j}^{+x}$ est la dérivée approchée de T en (i, j) utilisant les valeurs de $i+1$ et i et donc adaptée pour une propagation de l'information de droite à gauche. Inversement, le terme $\tau_{i,j}^{+x}$ propage l'information de

gauche à droite et de façon similaire les termes $\tau_{i,j}^{+y}$ et $\tau_{i,j}^{-y}$ propagent l'information dans la direction verticale.

Remarque A.1. *L'expression (A.2) provient de l'étude d'un schéma hyperbolique pour l'équation de propagation des lignes de niveau de T (voir [Set99, Chapitre 5] pour une dérivation détaillée de ce schéma).*

Pour chaque cellule (i, j) dans la zone de pénombre, la valeur de $T_{i,j}$ est calculé à partir de l'équation (A.2) en n'utilisant que les cellules dans la zone éclairée. Tant que la zone éclairée ne recouvre pas tout le domaine de calcul, l'algorithme de Fast Marching consiste à répéter les étapes :

1. déterminer le $T_{i,j}$ minimum parmi la zone de pénombre,
2. ajouter ce $T_{i,j}$ à la zone éclairée,
3. ajouter ses voisins à la zone de pénombre et calculer leur valeur de T à l'aide de (A.2).

Références

- [ADR14] B. ANDREIANOV, C. DONADELLO et M. D. ROSINI, “Crowd dynamics and conservation laws with nonlocal constraints and capacity drop,” *Mathematical Models and Methods in Applied Sciences*, t. 24, n° 13, p. 2685-2722, déc. 2014. DOI : 10.1142/S0218202514500341.
- [AGS10] B. ANDREIANOV, P. GOATIN et N. SEGUIN, “Finite volume schemes for locally constrained conservation laws,” *Numerische Mathematik*, t. 115, n° 4, p. 609-645, juin 2010. DOI : 10.1007/s00211-009-0286-7.
- [AL20] Y. ACHDOU et M. LAURIÈRE, “Mean Field Games and Applications : Numerical Aspects,” in *Mean Field Games*. Cham : Springer International Publishing, 2020, t. 2281, p. 249-307. DOI : 10.1007/978-3-030-59837-2_4.
- [AR00] A. AW et M. RASCLE, “Resurrection of "Second Order" Models of Traffic Flow,” *SIAM Journal on Applied Mathematics*, t. 60, n° 3, p. 916-938, jan. 2000. DOI : 10.1137/S0036139997332099.
- [AS20] B. ANDREIANOV et A. SYLLA, “A Macroscopic Model to Reproduce Self-organization at Bottlenecks,” in *Finite Volumes for Complex Applications IX - Methods, Theoretical Aspects, Examples*, R. KLÖFKORN, E. KEILEGAVLEN, F. A. RADU et J. FUHRMANN, éd., t. 323, Cham : Springer International Publishing, 2020, p. 243-254. DOI : 10.1007/978-3-030-43651-3_21.
- [BCB21] D. H. BIEDERMANN, J. CLEVER et A. BORRMANN, “A generic and density-sensitive method for multi-scale pedestrian dynamics,” *Automation in Construction*, t. 122, p. 103 489, fév. 2021. DOI : 10.1016/j.autcon.2020.103489.
- [BČG+14] A. BRESSAN, S. ČANIĆ, M. GARAVELLO, M. HERTY et B. PICCOLI, “Flows on networks : recent results and perspectives,” *EMS Surveys in Mathematical Sciences*, t. 1, n° 1, p. 47-111, 2014. DOI : 10.4171/EMSS/2.
- [BCL+18] G. BRETTI, E. CRISTIANI, C. LATTANZIO, A. MAURIZI et B. PICCOLI, “Two algorithms for a fully coupled and consistently macroscopic PDE-ODE system modeling a moving bottleneck on a road,” *Mathematics in Engineering*, t. 1, n° 1, p. 55-83, 2018. DOI : 10.3934/Mine.2018.1.55. arXiv : 1807.07461.
- [Ber20] F. BERTON, “Immersive Virtual Crowds : Evaluation of Pedestrian Behaviours in Virtual Reality,” These de Doctorat, Rennes 1, déc. 2020.

- [BH02] E. BOURREL et V. HENN, “Mixing Micro and Macro Representations of Traffic Flow : A First Theoretical Step,” in *Proceedings of the 9th Meeting of the Euro Working Group on Transportation*, 2002, p. 610-616.
- [BL03] E. BOURREL et J.-B. LESORT, “Mixing Microscopic and Macroscopic Representations of Traffic Flow : Hybrid Model Based on Lighthill–Whitham–Richards Theory,” *Transportation Research Record : Journal of the Transportation Research Board*, t. 1852, n° 1, p. 193-200, jan. 2003. DOI : 10.3141/1852-24.
- [BML92] O. BIHAM, A. A. MIDDLETON et D. LEVINE, “Self-organization and a dynamical transition in traffic-flow models,” *Physical Review A*, t. 46, n° 10, R6124-R6127, nov. 1992. DOI : 10.1103/PhysRevA.46.R6124.
- [Cep09] E. M. CEPOLINA, “Phased evacuation : An optimisation model which takes into account the capacity drop phenomenon in pedestrian flows,” *Fire Safety Journal*, t. 44, n° 4, p. 532-544, mai 2009. DOI : 10.1016/j.firesaf.2008.11.002.
- [CG07] R. M. COLOMBO et P. GOATIN, “A well posed conservation law with a variable unilateral constraint,” *Journal of Differential Equations*, t. 234, n° 2, p. 654-675, mars 2007. DOI : 10.1016/j.jde.2006.10.014.
- [CGGR08] H. CHATÉ, F. GINELLI, G. GRÉGOIRE et F. RAYNAUD, “Collective motion of self-propelled particles interacting without cohesion,” *Physical Review E*, t. 77, n° 4, p. 046113, avr. 2008. DOI : 10.1103/PhysRevE.77.046113.
- [CLM15] G. COSTESEQUE, J.-P. LEBACQUE et R. MONNEAU, “A convergent scheme for Hamilton–Jacobi equations on a junction : application to traffic,” *Numerische Mathematik*, t. 129, n° 3, p. 405-447, mars 2015. DOI : 10.1007/s00211-014-0643-z.
- [CM15] R. M. COLOMBO et F. MARCELLINI, “A mixed ODE-PDE model for vehicular traffic,” *Mathematical Methods in the Applied Sciences*, t. 38, n° 7, p. 1292-1302, mai 2015. DOI : 10.1002/mma.3146.
- [Cos14] G. COSTESEQUE, “Contribution à l’étude du trafic routier sur réseaux à l’aide des équations d’Hamilton–Jacobi,” p. 297, 2014.
- [CP02] G. M. COCLITE et B. PICCOLI, *Traffic Flow on a Road Network*, fév. 2002. arXiv : math/0202146.
- [CPT14] E. CRISTIANI, B. PICCOLI et A. TOSIN, *Multiscale Modeling of Pedestrian Dynamics* (MS&A). Cham : Springer International Publishing, 2014, t. 12. DOI : 10.1007/978-3-319-06620-2.
- [CSV97] A. CZIRÓK, H. E. STANLEY et T. VICSEK, “Spontaneously ordered motion of self-propelled particles,” *Journal of Physics A : Mathematical and General*, t. 30, n° 5, p. 1375-1385, mars 1997. DOI : 10.1088/0305-4470/30/5/009.
- [Daa04] W. DAAMEN, “Modelling passenger flows in public transport facilities,” thèse de doct., DUP Science / Delft University of Technology, 2004.
- [DC97] M. DOMEIER et P. COLIN, “Tropical Reef Fish Spawning Aggregations : Defined and Reviewed,” *Bulletin of Marine Science*, t. 60, p. 698-726, juill. 1997.

- [Deg19] P. DEGOND, “Mathematical models of collective dynamics and self-organisation,” in *Proceedings of the International Congress of Mathematicians (ICM 2018)*, Rio de Janeiro, Brazil : WORLD SCIENTIFIC, mai 2019, p. 3925-3946. DOI : 10.1142/9789813272880_0206.
- [DG14a] M. L. DELLE MONACHE et P. GOATIN, “Scalar conservation laws with moving constraints arising in traffic flow modeling : An existence result,” *Journal of Differential Equations*, t. 257, n° 11, p. 4015-4029, déc. 2014. DOI : 10.1016/j.jde.2014.07.014.
- [DG14b] M. L. DELLE MONACHE et P. GOATIN, “A front tracking method for a strongly coupled PDE-ODE system with moving density constraints in traffic flow,” *Discrete & Continuous Dynamical Systems - S*, t. 7, n° 3, p. 435-447, 2014. DOI : 10.3934/dcdss.2014.7.435.
- [DH13] P. DEGOND et J. HUA, “Self-organized hydrodynamics with congestion and path formation in crowds,” *Journal of Computational Physics*, t. 237, p. 299-319, mars 2013. DOI : 10.1016/j.jcp.2012.11.033.
- [DM08] P. DEGOND et S. MOTSCH, “Continuum limit of self-driven particles with orientation interaction,” *Mathematical Models and Methods in Applied Sciences*, t. 18, n° supp01, p. 1193-1215, août 2008. DOI : 10.1142/S0218202508003005. arXiv : 0710.0293 [math-ph].
- [DR15] M. DI FRANCESCO et M. D. ROSINI, “Rigorous derivation of nonlinear scalar conservation laws from follow-the-leader type models via many particle limit,” *Archive for Rational Mechanics and Analysis*, t. 217, n° 3, p. 831-871, sept. 2015. DOI : 10.1007/s00205-015-0843-4. arXiv : 1404.7062.
- [FFL+16] M. FABRE, S. FAURE, M. LAURIÈRE, B. MAURY et C. PERRIN, “Non classical solution of a conservation law arising in vehicular traffic modelling,” *ESAIM : Proceedings and Surveys*, t. 55, E. FRÉNOT, E. MAITRE, A. ROUSSEAU, S. SALMON et M. SZOPOS, éd., p. 131-147, déc. 2016. DOI : 10.1051/proc/201655131.
- [FM] S. FAURE et B. MAURY, *Cromosim*.
- [Fru71] J. J. FRUIN, “Pedestrian Planning and Design,” rapp. tech., 1971.
- [FZG+20] C. FELICIANI, I. ZURIGUEL, A. GARCIMARTÍN, D. MAZA et K. NISHINARI, “Systematic experimental investigation of the obstacle effect during non-competitive and extremely competitive evacuations,” *Scientific Reports*, t. 10, n° 1, p. 15 947, déc. 2020. DOI : 10.1038/s41598-020-72733-w.
- [GB16] S. M. GWYNNE et K. BOYCE, “Engineering Data,” in *SFPE Handbook of Fire Protection Engineering*, New York, NY : Springer New York, 2016, p. 2070-2114. DOI : 10.1007/978-1-4939-2565-0_64.
- [GH15] C. GERSHENSON et D. HELBING, “When slower is faster,” *Complexity*, t. 21, n° 2, p. 9-15, nov. 2015. DOI : 10.1002/cplx.21736. arXiv : 1506.06796.

- [GHR61] D. C. GAZIS, R. HERMAN et R. W. ROTHERY, "Nonlinear Follow-the-Leader Models of Traffic Flow," *Operations Research*, t. 9, n° 4, p. 545-567, août 1961. DOI : 10.1287/opre.9.4.545.
- [GLG+15] E. GALEA, P. LAWRENCE, S. GWYNNE, L. FILIPPIDIS, D. BLACKSHIELDS et D. COONEY, *buildingEXODUS V6.2 Technical Manual and User Guide*, Manual, University of Greenwich, Greenwich, London, UK, déc. 2015.
- [GMP+18] A. GARCIMARTÍN, D. MAZA, J. M. PASTOR, D. R. PARISI, C. MARTÍN-GÓMEZ et I. ZURIGUEL, "Redefining the role of obstacles in pedestrian evacuation," *New Journal of Physics*, t. 20, n° 12, p. 123 025, déc. 2018. DOI : 10.1088/1367-2630/aaf4ca.
- [GMQD20] F. GNOGBO, L. MAILLO, T. QUEVEAU et C. DAUNIS, "La Sécurité Des "Fan-zones" et Autres Grands Rassemblements," ENSOP, Paris, rapp. tech., 2020.
- [GP09] M. GARAVELLO et B. PICCOLI, "Time-varying Riemann solvers for conservation laws on networks," *Journal of Differential Equations*, t. 247, n° 2, p. 447-464, juill. 2009. DOI : 10.1016/j.jde.2008.12.017.
- [GP17] M. GARAVELLO et B. PICCOLI, "Boundary coupling of microscopic and first order macroscopic traffic models," *Nonlinear Differential Equations and Applications NoDEA*, t. 24, n° 4, p. 43, août 2017. DOI : 10.1007/s00030-017-0467-5.
- [GPP+16] A. GARCIMARTÍN, D. R. PARISI, J. M. PASTOR, C. MARTÍN-GÓMEZ et I. ZURIGUEL, "Flow of pedestrians through narrow doors with different competitiveness," *Journal of Statistical Mechanics : Theory and Experiment*, t. 2016, n° 4, p. 043 402, avr. 2016. DOI : 10.1088/1742-5468/2016/04/043402.
- [GZP+14] A. GARCIMARTÍN, I. ZURIGUEL, J. PASTOR, C. MARTÍN-GÓMEZ et D. PARISI, "Experimental Evidence of the "Faster Is Slower" Effect," *Transportation Research Procedia*, t. 2, p. 760-767, 2014. DOI : 10.1016/j.trpro.2014.09.085.
- [Hag20] M. HAGHANI, "Empirical methods in pedestrian, crowd and evacuation dynamics : Part II. Field methods and controversial topics," *Safety Science*, t. 129, p. 104 760, sept. 2020. DOI : 10.1016/j.ssci.2020.104760.
- [HBJW05] D. HELBING, L. BUZNA, A. JOHANSSON et T. WERNER, "Self-Organized Pedestrian Crowd Dynamics : Experiments, Simulations, and Design Solutions," *Transportation Science*, t. 39, n° 1, p. 1-24, fév. 2005. DOI : 10.1287/trsc.1040.0108.
- [HD05] S. P. HOOGENDOORN et W. DAAMEN, "Pedestrian Behavior at Bottlenecks," *Transportation Science*, t. 39, n° 2, p. 147-159, mai 2005. DOI : 10.1287/trsc.1040.0102.
- [HFV00] D. HELBING, I. FARKAS et T. VICSEK, "Simulating Dynamical Features of Escape Panic," *Nature*, t. 407, n° 6803, p. 487-490, sept. 2000. DOI : 10.1038/35035023. arXiv : cond-mat/0009448.

- [HFV02] D. HELBING, I. J. FARKAS et T. VICSEK, “Crowd Disasters and Simulation of Panic Situations,” in *The Science of Disasters : Climate Disruptions, Heart Attacks, and Market Crashes*, Berlin, Heidelberg : Springer Berlin Heidelberg, 2002, p. 330-350. DOI : 10.1007/978-3-642-56257-0â&A1.
- [HJ13] D. HELBING et A. JOHANSSON, “Pedestrian, Crowd, and Evacuation Dynamics,” *arXiv :1309.1609 [physics]*, sept. 2013. arXiv : 1309.1609 [physics].
- [HJA07] D. HELBING, A. JOHANSSON et H. Z. AL-ABIDEEN, “The Dynamics of Crowd Disasters : An Empirical Study,” *Physical Review E*, t. 75, n° 4, p. 046 109, avr. 2007. DOI : 10.1103/PhysRevE.75.046109. arXiv : physics/0701203.
- [HKHE12] S. HELIÖVAARA, T. KORHONEN, S. HOSTIKKA et H. EHTAMO, “Counterflow model for agent-based simulation of crowd dynamics,” *Building and Environment*, t. 48, p. 89-100, fév. 2012. DOI : 10.1016/j.buildenv.2011.08.020.
- [HM12] B. L. HOSKINS et J. A. MILKE, “Differences in measurement methods for travel distance and area for estimates of occupant speed on stairs,” *Fire Safety Journal*, t. 48, p. 49-57, fév. 2012. DOI : 10.1016/j.firesaf.2011.12.009.
- [HM95] D. HELBING et P. MOLNÁR, “Social force model for pedestrian dynamics,” *Physical Review E*, t. 51, n° 5, p. 4282-4286, mai 1995. DOI : 10.1103/PhysRevE.51.4282.
- [HMC06] M. HUANG, R. P. MALHAME et P. E. CAINES, “Nash Certainty Equivalence in Large Population Stochastic Dynamic Games : Connections with the Physics of Interacting Particle Systems,” in *Proceedings of the 45th IEEE Conference on Decision and Control*, San Diego, CA, USA : IEEE, 2006, p. 4921-4926. DOI : 10.1109/CDC.2006.377683.
- [HR17] H. HOLDEN et N. H. RISEBRO, “Follow-the-Leader models can be viewed as a numerical approximation to the Lighthill-Whitham-Richards model for traffic flow,” *arXiv :1702.01718 [math]*, sept. 2017. arXiv : 1702.01718 [math].
- [Hug02] R. L. HUGHES, “A continuum theory for the flow of pedestrians,” *Transportation Research Part B : Methodological*, t. 36, n° 6, p. 507-535, juill. 2002. DOI : 10.1016/S0191-2615(01)00015-7.
- [Hur97] J. HURET, *La Catastrophe Du Bazar de La Charité (4 Mai 1897)*. 1897.
- [IMO07] IMO, “Guidelines for Evacuation Analysis for New and Existing Passenger Ships,” 2007.
- [JALP12] A. JELIĆ, C. APPERT-ROLLAND, S. LEMERCIER et J. PETTRÉ, “Properties of pedestrians walking in line : Fundamental diagrams,” *Physical Review E*, t. 85, n° 3, p. 036 111, mars 2012. DOI : 10.1103/PhysRevE.85.036111.
- [JDW+20] Q. JULLIEN et al., “Evacuation modelling – application to an office building using several simulation tools,” *Fire and Evacuation Modeling Technical Conference (FEMTC) 2020*, p. 15, 2020.
- [KH09] T. KORHONEN et S. HOSTIKKA, “Fire Dynamics Simulator with Evacuation : FDS+ Evac,” *Technical Reference and User’s Guide. VTT Technical Research Centre of Finland*, 2009.

- [Kul16] E. D. KULIGOWSKI, “Human Behavior in Fire,” in *SFPE Handbook of Fire Protection Engineering*, M. J. HURLEY et al., éd., New York, NY : Springer New York, 2016, p. 2070-2114. DOI : 10.1007/978-1-4939-2565-0_58.
- [KZGK19] B. KLEINMEIER, B. ZÖNNCHEN, M. GÖDEL et G. KÖSTER, “Vadere : An Open-Source Simulation Framework to Promote Interdisciplinary Understanding,” *Collective Dynamics*, t. 4, 2019-01-01, 2019. DOI : 10.17815/CD.2019.21, published.
- [Leb96] J.-P. LEBACQUE, “The Godunov Scheme and What It Means for First Order Traffic Flow Models,” in *International Symposium on Transportation and Traffic Theory*, 1996, p. 647-677.
- [LJK+12] S. LEMERCIER et al., “Realistic following behaviors for crowd simulation,” *Computer Graphics Forum*, t. 31, n° 2pt2, p. 489-498, mai 2012. DOI : 10.1111/j.1467-8659.2012.03028.x.
- [LL06a] J.-M. LASRY et P.-L. LIONS, “Jeux à champ moyen. I – Le cas stationnaire,” *Comptes Rendus Mathématique*, t. 343, n° 9, p. 619-625, nov. 2006. DOI : 10.1016/j.crma.2006.09.019.
- [LL06b] J.-M. LASRY et P.-L. LIONS, “Jeux à champ moyen. II – Horizon fini et contrôle optimal,” *Comptes Rendus Mathématique*, t. 343, n° 10, p. 679-684, nov. 2006. DOI : 10.1016/j.crma.2006.09.018.
- [LL07] J.-M. LASRY et P.-L. LIONS, “Mean field games,” *Japanese Journal of Mathematics*, t. 2, n° 1, p. 229-260, mars 2007. DOI : 10.1007/s11537-007-0657-8.
- [LLE10] R. LUKEMAN, Y.-X. LI et L. EDELSTEIN-KESHET, “Inferring individual rules from collective behavior,” *Proceedings of the National Academy of Sciences*, t. 107, n° 28, p. 12 576-12 580, juill. 2010. DOI : 10.1073/pnas.1001763107.
- [LP10] C. LATTANZIO et B. PICCOLI, “Coupling of microscopic and macroscopic traffic models at boundaries,” *Mathematical Models and Methods in Applied Sciences*, t. 20, n° 12, p. 2349-2370, déc. 2010. DOI : 10.1142/S0218202510004945.
- [LRK20] R. LOVREGGIO, E. RONCHI et M. J. KINSEY, “An Online Survey of Pedestrian Evacuation Model Usage and Users,” *Fire Technology*, t. 56, n° 3, p. 1133-1153, mai 2020. DOI : 10.1007/s10694-019-00923-8.
- [LW55] M. J. LIGHTHILL et F. R. S. WHITHAM, “On kinematic waves II. A theory of traffic flow on long crowded roads,” *Proceedings of the Royal Society of London. Series A. Mathematical and Physical Sciences*, t. 229, n° 1178, p. 317-345, mai 1955. DOI : 10.1098/rspa.1955.0089.
- [MF18] B. MAURY et S. FAURE, *Crowds in Equations : An Introduction to the Microscopic Modeling of Crowds*. World Scientific, 2018.
- [MFAB18] B. MAURY, S. FAURE, J. ANGELÉ et R. BACHIMONT, “A Time-continuous Compartment Model for Building Evacuation,” *Journal of Physics : Conference Series*, t. 1107, p. 072 003, nov. 2018. DOI : 10.1088/1742-6596/1107/7/072003.

- [MGM+12] M. MOUSSAÏD et al., “Traffic Instabilities in Self-Organized Pedestrian Crowds,” *PLoS Computational Biology*, t. 8, n° 3, E. BEN-JACOB, éd., e1002442, mars 2012. DOI : 10.1371/journal.pcbi.1002442.
- [MHT11] M. MOUSSAÏD, D. HELBING et G. THERAULAZ, “How simple rules determine pedestrian behavior and crowd disasters,” *Proceedings of the National Academy of Sciences*, t. 108, n° 17, p. 6884-6888, avr. 2011. DOI : 10.1073/pnas.1016507108.
- [Mim] M. MIMAULT, “Crowd Motion Modeling by Conservation Laws,” p. 145,
- [Mir] J.-M. MIREBEAU, “Numerical schemes for anisotropic PDEs on cartesian grid domains,” p. 153,
- [Mor77] J. J. MOREAU, “Evolution problem associated with a moving convex set in a Hilbert space,” *Journal of Differential Equations*, t. 26, n° 3, p. 347-374, déc. 1977. DOI : 10.1016/0022-0396(77)90085-7.
- [MRSV11] B. MAURY, A. ROUDNEFF-CHUPIN, F. SANTAMBROGIO et J. VENEL, “Handling congestion in crowd motion modeling,” *arXiv :1101.4102 [math]*, jan. 2011. arXiv : 1101.4102 [math].
- [MV11] B. MAURY et J. VENEL, “A discrete contact model for crowd motion,” *ESAIM : Mathematical Modelling and Numerical Analysis*, t. 45, n° 1, p. 145-168, jan. 2011. DOI : 10.1051/m2an/2010035.
- [NBK17] A. NICOLAS, S. BOUZAT et M. N. KUPERMAN, “Pedestrian flows through a narrow doorway : Effect of individual behaviours on the global flow and microscopic dynamics,” *Transportation Research Part B : Methodological*, t. 99, p. 30-43, mai 2017. DOI : 10.1016/j.trb.2017.01.008.
- [PHK12] R. PEACOCK, B. HOSKINS et E. KULIGOWSKI, “Overall and local movement speeds during fire drill evacuations in buildings up to 31 stories,” *Safety Science*, t. 50, n° 8, p. 1655-1664, oct. 2012. DOI : 10.1016/j.ssci.2012.01.003.
- [PM78] V. M. PREDTECHENSKII et A. I. MILINSKII, *Planning for Foot Traffic Flow in Buildings*. New Delhi : Amerind Pub. Co., 1978.
- [PMLW14] J. PORZYCKI, M. MYCEK, R. LUBAŚ et J. WAS, “Pedestrian Spatial Self-organization According to its Nearest Neighbor Position,” *Transportation Research Procedia*, t. 2, p. 201-206, 2014. DOI : 10.1016/j.trpro.2014.09.033.
- [PPD07] S. PARIS, J. PETTRÉ et S. DONIKIAN, “Pedestrian Reactive Navigation for Crowd Simulation : a Predictive Approach,” *Computer Graphics Forum*, t. 26, n° 3, p. 665-674, sept. 2007. DOI : 10.1111/j.1467-8659.2007.01090.x.
- [Red17] F. A. REDA, “Crowd motion modelisation under some constraints,” thèse de doct., Université Paris-Saclay, 2017.
- [Rey+99] C. W. REYNOLDS et al., “Steering Behaviors for Autonomous Characters,” in *Game Developers Conference*, Citeseer, t. 1999, 1999, p. 763-782.
- [Ric56] P. I. RICHARDS, “Shock Waves on the Highway,” *Operations Research*, t. 4, n° 1, p. 42-51, 1956. DOI : 10.1287/opre.4.1.42. eprint : <https://doi.org/10.1287/opre.4.1.42>.

- [Rou11] A. ROUDNEFF, “Modelisation macroscopique de mouvements de foule,” thèse de doct., 2011.
- [Ser99] D. SERRE, *Systems of Conservation Laws 1 : Hyperbolicity, Entropies, Shock Waves*, First, trad. par I. N. SNEDDON. Cambridge University Press, mai 1999. DOI : 10.1017/CB09780511612374.
- [Set99] J. A. SETHIAN, *Level Set Methods and Fast Marching Methods : Evolving Interfaces in Computational Geometry, Fluid Mechanics, Computer Vision, and Materials Science*. Cambridge university press, 1999, t. 3.
- [SS10] B. STEFFEN et A. SEYFRIED, “Methods for measuring pedestrian density, flow, speed and direction with minimal scatter,” *Physica A : Statistical Mechanics and its Applications*, t. 389, n° 9, p. 1902-1910, mai 2010. DOI : 10.1016/j.physa.2009.12.015.
- [SSG+06] B. SZABO, G. J. SZOLLOSI, B. GONCI, Z. JURANYI, D. SELMECZI et T. VICSEK, “Phase transition in the collective migration of tissue cells : experiment and model,” *Physical Review E*, t. 74, n° 6, p. 061 908, déc. 2006. DOI : 10.1103/PhysRevE.74.061908. arXiv : q-bio/0611045.
- [SSNI95] M. SCHRECKENBERG, A. SCHADSCHNEIDER, K. NAGEL et N. ITO, “Discrete stochastic models for traffic flow,” *Physical Review E*, t. 51, n° 4, p. 2939-2949, avr. 1995. DOI : 10.1103/PhysRevE.51.2939.
- [TGD14] M. TWAROGOWSKA, P. GOATIN et R. DUVIGNEAU, “Macroscopic modeling and simulations of room evacuation,” *Applied Mathematical Modelling*, t. 38, n° 24, p. 5781-5795, déc. 2014. DOI : 10.1016/j.apm.2014.03.027.
- [Thi17] A. THIRY-MULLER, “Guide de Bonnes Pratiques Pour Les Études d’ingénierie Du Désenfumage Dans Les Établissements Recevant Du Public,” Laboratoire Central de la Préfecture de Police, rapp. tech., 2017.
- [Thu] THUNDERHEAD ENGINEERING, *Pathfinder*.
- [TP21] W. TOLL et J. PETTRÉ, “Algorithms for Microscopic Crowd Simulation : Advancements in the 2010s,” *Computer Graphics Forum*, t. 40, n° 2, p. 731-754, mai 2021. DOI : 10.1111/cgf.142664.
- [VCB+95] T. VICSEK, A. CZIRÓK, E. BEN-JACOB, I. COHEN et O. SHOCHET, “Novel Type of Phase Transition in a System of Self-Driven Particles,” *Physical Review Letters*, t. 75, n° 6, p. 1226-1229, août 1995. DOI : 10.1103/PhysRevLett.75.1226.
- [vGG+20] W. VAN TOLL et al., “Generalized Microscopic Crowd Simulation Using Costs in Velocity Space,” in *Symposium on Interactive 3D Graphics and Games*, sér. I3D ’20, New York, NY, USA : Association for Computing Machinery, 2020. DOI : 10.1145/3384382.3384532.
- [Vio13] P. VIOT, “Le Territoire Sécurisé Des Grandes Manifestations Contemporaines,” jan. 2013.

- [vMM08] J. VAN DEN BERG, MING LIN et D. MANOCHA, “Reciprocal Velocity Obstacles for real-time multi-agent navigation,” in *2008 IEEE International Conference on Robotics and Automation*, Pasadena, CA, USA : IEEE, mai 2008, p. 1928-1935. DOI : 10.1109/ROBOT.2008.4543489.
- [Wei92] U. WEIDMANN, “Transporttechnik der Fussgänger : Transporttechnische Eigenschaften des Fussgängerverkehrs, Literaturlauswertung,” ETH Zurich, rapp. tech., 1992, 84 S. DOI : 10.3929/ETHZ-A-000687810.
- [XLC+10] M. XIONG, M. LEES, W. CAI, S. ZHOU et M. Y. H. LOW, “Hybrid modelling of crowd simulation,” *Procedia Computer Science*, t. 1, n° 1, p. 57-65, mai 2010. DOI : 10.1016/j.procs.2010.04.008.
- [Zha02] H. M. ZHANG, “A non-equilibrium traffic model devoid of gas-like behavior,” p. 16, 2002.