



Modélisation mathématique et simulation de TCP par des méthodes de champ moyen.

Julien Reynier

► To cite this version:

Julien Reynier. Modélisation mathématique et simulation de TCP par des méthodes de champ moyen..
Modélisation et simulation. Ecole Polytechnique X, 2006. Français. NNT : . pastel-00002032

HAL Id: pastel-00002032

<https://pastel.hal.science/pastel-00002032>

Submitted on 28 Jul 2010

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.



École Doctorale de l'École Polytechnique

Modélisation mathématique et simulation de TCP par des méthodes de champ moyen

THÈSE

présentée et soutenue publiquement le 15 septembre 2006

pour l'obtention du

Doctorat de l'école Polytechnique – Palaiseau
(Informatique)

par

Julien Reynier

Composition du jury

<i>Directeur :</i>	François Baccelli	Directeur de recherche INRIA à l'ENS
<i>Co-directeur :</i>	David McDonald	Professeur à l'Université d'Ottawa
<i>Rapporteur et Président :</i>	Armand Makowski	Professeur à l'Université du Maryland
<i>Rapporteurs :</i>	Frank den Hollander Fabrice Guillemin	Professeur à l'Université de Leiden Chercheur à France Télécom R&D
<i>Membre :</i>	Carl Graham	Chercheur du CNRS à l'École Polytechnique

Remerciements

Je remercie tous les gens qui m'ont soutenu, j'espère ne pas oublier trop de monde. Je commence par mon jury et en particulier mes rapporteurs dont la patience et les remarques pertinentes m'ont permis d'améliorer le manuscrit. Je continue avec un grand merci à François Baccelli et David McDonald sans qui cette thèse n'aurait pas existé ; leur soutien et leur encadrement de tous les instants depuis le début de mon stage de DEA en 2001 ne sont pas étrangers à la réussite du projet. En particulier David a su utiliser les ressources de toute sa famille pour améliorer mes condition de travail : Nasrin, sa femme, a bien voulu m'héberger et me nourrir en Ottawa. Omid et ses meilleurs amis d'Ottawa Gad Bentolila, Philip Painter et Chris Lee ont su me divertir pendant tout un été.

Poursuivons en remerciant Augustin Chaintreau, Ki-Baek Kim, Fabien Mathieu et James Roberts pour leur aide précieuse dans la rédaction et la préparation de la soutenance. Ainsi que mes collègues de bureau pendant ces années, Florent Tournois, Neil O'Connell, Dohy Hong, Bartek Blaszczyzyn, Giovanna Carofiglio, Bruno Kauffmann, Marc Lelarge, Charles Bordenave, Minh-Anh Tran, Emmanuel Roy, Prasanna Chaporkar, Pierre Brémaud, Alexandre Proutière, Thomas Bonald, Nidhi Hegde et Bozidar Radunovic. Je remercie Laurence Lenormand, Florence Barbara, Joëlle Isnard, Valerie Mongiat, Jacques Beigbeder, Gabrielle Baron et Léonor Lemée pour leur aide logistique. Mes co-auteurs d'articles et de programmes que je n'ai pas encore cité, Anh-Tuan Gai, Fabien de Montgolfier, Fabrice Poppe, Guido Petit et Laurent Fournier.

Je remercie ensuite Anamaria, ma femme, Ghighi, mon beau-frère, Chantal et Philippe, mes parents et Adriana et Traian, mes beaux-parents pour leur bon encadrement familial, c'est important. Velvet, mon chien noir, pour me réchauffer les pieds en toutes saisons en particulier pendant l'été.

Je n'oublie pas bien sur de remercier chaleureusement Nikos et Alina Paragios, Xavier Payet, François Alter, Laurent Viennot, Erik Bergelin et le professeur Patrick Juillard qui m'ont fourni leurs conseils avisés concernant l'après thèse.

Table des matières

Résumé	1
Abstract	3
Introduction	5
Chapitre 1 Situation du problème	7
1.1 Le défi relevé par TCP	7
1.1.1 Principe général de fonctionnement	8
1.1.2 Le cahier des charges	8
1.1.3 La congestion dans Internet	9
1.1.4 Position du problème d’optimisation	10
1.2 Comment TCP Reno résout le problème	11
1.2.1 Détection des pertes	11
1.2.2 Fenêtre et contrôle du flux de données	12
1.2.3 Fenêtre de congestion et détection de pertes	13
1.3 Les algorithmes de gestion de taille des files d’attente des routeurs . . .	14
1.3.1 La «Drop Tail»	14
1.3.2 Destructures de paquets anticipées	14
1.3.3 RED	15
1.3.4 Gentle RED	16
1.3.5 AQM	16
1.4 Les modèles mathématiques de TCP	16
1.4.1 Quelques types de modèles	17
1.4.2 Pourquoi s’intéresser au champ moyen	18
Chapitre 2 Modèles de TCP menant à des champs moyens	21
2.1 Généralités	21

2.1.1	Autour du champ moyen	21
2.1.2	Le problème étudié, premières notations	22
2.1.3	Les équations fluides	23
2.2	Modèles par champ moyen	25
2.2.1	Champ moyen Baccelli-Hong	25
2.2.2	Champ moyen Bain-Kelly-Kelly	28
2.2.3	Champ moyen Deb-Shakkottai-Srikant	30
2.2.4	Champ moyen Makowski-Tinnakornsriruphap	31
2.3	Les modèles étudiés dans la thèse	33
2.3.1	Le modèle TCP-RED avec sources persistantes de la thèse	33
2.3.2	Adaptation HTTP-RED avec sources intermittentes	35
2.3.3	Type de résultats que nous démontrons	36
2.3.4	Positionnement par rapport aux autres modèles	37

Partie I Éléments mathématiques du champ moyen

Chapitre 3	Champ moyen et tirages d'une urne	41
3.1	Rappels mathématiques et motivations générales	42
3.1.1	Systèmes de particules	42
3.1.2	Convergence faible sur les espaces Polonais	42
3.1.3	Métriques de la convergence faible	43
3.1.4	Motivations mathématiques	45
3.2	Spécification du problème particulier étudié dans ce chapitre	46
3.2.1	Notations et définitions	46
3.2.2	Résultats	47
3.2.3	Échangeabilité	47
3.2.4	Discussion	47
3.3	Champ moyen par la méthode classique : tension, limites, unicité	48
3.3.1	Explications générales sur la méthode classique	48

3.3.2	Preuve du théorème par la méthode classique du champ moyen	49
3.3.3	Une approche alternative	49
3.4	Champ moyen par la méthode de couplage trajectoriel à un système auxiliaire	50
3.4.1	Explications générales	50
3.4.2	Couplage entre les trajectoires et convergence	51
3.4.3	Définition du système auxiliaire : nouvelle version	52
3.4.4	Énoncés et démonstrations des théorèmes	52
3.4.5	Définition du système auxiliaire : version originale	53
3.4.6	Comparaison des méthodes	54
3.5	Apports mathématiques de cette thèse	54
Chapitre 4	Problème du champ moyen dans $\mathbb{D}$	57
4.1	Rappels de mathématiques générales	58
4.1.1	Prévisibilité	58
4.1.2	Topologie de Skorokhod	61
4.1.3	Module de continuité et compacts de $\mathbb{D}$	62
4.1.4	Tension d'une suite de processus càdlàg au sens de Skorokhod	63
4.2	Première présentation de l'équation différentielle stochastique que nous allons étudier	63
4.3	Éléments mathématiques spécifiques	64
4.3.1	Un mode de définition d'un processus de Poisson doublement stochastique	64
4.3.2	Marginales	65
4.3.3	Crochet marginal double	66
4.3.4	Fonctions prévisibles de E	67
4.4	Formulation du problème	67
4.4.1	Notations et définitions	67
4.4.2	Une première formulation du problème	68
4.4.3	Hypothèses techniques	69
4.5	Couplage et norme	69
4.5.1	Norme de convergence d'un système infini de particules	69
4.5.2	Couplage du système initial	70
4.5.3	Reformulation du problème	71
4.6	Discussion autour du problème étudié	71

Chapitre 5	Système auxiliaire dans $\mathbb{D}$	73
5.1	Existence et étude de la particule du champ moyen	74
5.1.1	Définition de la particule du champ moyen	74
5.1.2	Un lemme technique	74
5.1.3	Existence et unicité forte des solutions de l'EDS du champ moyen	76
5.1.4	Compatibilité par rapport aux représentations*	78
5.2	Système auxiliaire : définition et propriétés	79
5.2.1	Équations distributionnelles et continuité du champ moyen* . .	79
5.2.2	Flot de solutions et définition du système auxiliaire	81
5.2.3	Quelques caractéristiques du flot	81
5.2.4	Loi des grands nombres pour le système auxiliaire	82
5.3	Convergence du système original vers le système auxiliaire	85
5.3.1	Un lemme technique	85
5.3.2	Le théorème de convergence	86
5.3.3	Conséquences de la convergence	86
5.4	Discussion et prolongements*	86
5.4.1	Particules couplées k à $k^{(*)}$	87
5.4.2	Dépendance dans le passé*	87
Chapitre 6	Problème aux frontières dans $\mathbb{D}$	89
6.1	Spécification du problème	90
6.1.1	Discussion	90
6.1.2	Premières définitions	90
6.1.3	Couplage des sauts	92
6.1.4	Problème avec frontière	93
6.2	Particule du champ moyen et convergence	93
6.2.1	Définition de la particule du champ moyen	93
6.2.2	Loi des grands nombres pour le champ moyen	94
6.2.3	Passage de la convergence après un saut	95
6.2.4	Convergence par l'étude du candidat limite	97
6.3	Un exemple d'utilisation : convergence du modèle Baccelli-Hong	98
6.3.1	Adaptation du modèle à N particules	98
6.3.2	Preuve de convergence en utilisant les résultats des chapitres précédents	100

Chapitre 7 Un premier exemple : champ moyen pour le modèle HTTP-AIMD de Baccelli-Hong **103**

7.1	Rappel des modèles et problèmes à résoudre pour prouver la convergence du champ moyen	103
7.1.1	Le modèle AIMD original	104
7.1.2	Les équations du champ moyen intuitives pour HTTP sur AIMD	105
7.2	Adaptation du modèle à N particules	106
7.2.1	Les débuts et fins de transmission	106
7.2.2	Le modèle de pertes et couplage des sauts synchrones	107
7.2.3	Précisions sur les processus de Poisson	107
7.3	Preuve de convergence en utilisant les résultats des chapitres précédents	108
7.3.1	Définition de la limite	108
7.3.2	Convergence hors des temps d'arrêt	108
7.3.3	Traversée des frontières	108

Partie II Application du champ moyen à la modélisation de flots TCP

Chapitre 8 Une modélisation de TCP **115**

8.1	Le modèle et ses équations d'évolution : processus de Markov à N particules	116
8.1.1	Modélisation de TCP : hypothèses et équations d'évolution . . .	116
8.1.2	Modélisation des débits à partir de la valeur des fenêtres	117
8.1.3	Génération et prise en compte des pertes de paquets	117
8.1.4	Équation différentielle sur la taille de la file d'attente :	118
8.2	Discussion et explication du modèle	120
8.2.1	Modélisation de RED et réductions de fenêtres	120
8.2.2	Équation différentielle pour les fenêtres :	121
8.2.3	Limites de la modélisation	122

8.3	Hypothèses techniques	122
8.3.1	Couplage des processus de Poisson	122
8.3.2	Hypothèses portant sur l'état initial	122
8.3.3	Borne pour la taille de la fenêtre au temps t	123
8.3.4	Relation entre le RTT et la taille de la file d'attente	123
8.4	Mathématisation du modèle	124
8.4.1	Énoncé du problème	124
8.4.2	Reformulation en termes de processus prenant pour valeur des mesures	124
8.4.3	Énoncé du théorème	126
8.4.4	Discussion	127
8.5	Vers une preuve du champ moyen	128
8.5.1	Plan de la preuve	128
8.5.2	Théorèmes sur l'équation de la limite et convergence	130
8.5.3	Difficulté pour mener à bien la démarche classique par tension*	131
Chapitre 9 Démonstration du champ moyen		135
9.1	Existence d'une limite	135
9.1.1	Système infini de particules auxiliaire	136
9.1.2	Premières propriétés du système auxiliaire	137
9.1.3	Existence d'une limite pour le système auxiliaire	139
9.1.4	Équations de la limite $(\mathcal{W}, Q)$:	142
9.2	Unicité de la solution forte	143
9.2.1	Convergence loin de la frontière	147
9.2.2	Quand on traverse ou qu'on frôle le bord	153
9.2.3	Convergence en champ moyen sur la frontière	156
9.3	RED est la limite faible de Gentle RED*	159
Chapitre 10 EDP du champ moyen : équations et points fixes		163
10.1	Équations aux dérivées partielles du champ moyen	164
10.1.1	La martingale qui tend vers 0	164
10.1.2	L'équation aux dérivées partielles	166
10.1.3	Simplifications de l'EDP	167
10.2	Étude de l'équation simplifiée	168
10.2.1	Point fixe des équations simplifiées du champ moyen	168

10.2.2	Une discussion autour des timeouts	171
10.2.3	Preuve du théorème sur les timeouts	173
10.3	Contrôle et stabilité des points fixes de l'équation approximée	175
10.3.1	Étude des moments de la distribution des fenêtres	175
10.3.2	Hypothèse pour l'étude de la stabilité	176
10.3.3	Stabilité sous les contrôles classiques	177
10.3.4	Un contrôle stable sur les débits	183
10.3.5	Un contrôle stable sur les débits avec utilisation maximale du lien	185
Chapitre 11 Adaptation du champ moyen à HTTP		189
11.1	Modélisation des particules	189
11.1.1	Environnement	189
11.1.2	Équations d'une particule	190
11.1.3	Les équations du champ moyen	191
11.1.4	Les extensions possibles	191
11.2	Point fixe	191
11.2.1	Le point fixe exact pour les inactifs	192
11.2.2	Point fixe approximé pour les utilisateurs actifs	192
11.2.3	Formules autour du point fixe et normalisation	193
11.2.4	Formule de la racine carrée	195
11.3	Exemple de stabilité dans un cas simple	196
11.3.1	Hypothèse simplificatrice	196
11.3.2	Équation pour M_1	196
11.3.3	Équation de contrôle du débit	197
11.3.4	Petites perturbations	197

Partie III Simulations

Chapitre 12 Le simulateur	203
12.1 Conception du simulateur	204

12.1.1	Analyse numérique	204
12.1.2	Dans tous les cas : problème des retards et des pertes	206
12.1.3	Pour HTTP : tailles de fichiers et temps de sommeil	208
12.1.4	Difficultés techniques	213
12.2	Validations du simulateur pour des sources persistantes	213
12.2.1	Comparaison entre le simulateur et OPNET	213
12.2.2	Analyse de la qualité de service	218
12.2.3	Comparaison entre le simulateur et des réalisations du système à N particules	220
12.3	Amélioration du modèle de débit	222
12.3.1	Diagnostic d'un problème dans la prise en compte des débits	222
12.3.2	Changement du modèle de lien et de file d'attente	224
12.3.3	Implémentation du modèle de débits réel	226
12.4	Intégration des Timeouts	228
Chapitre 13 Exemples d'utilisation du simulateur		231
13.1	Effet de la congestion pour un routeur sous ses réglages habituels	232
13.1.1	Données numériques	232
13.1.2	Premiers résultats	232
13.1.3	Détails sur la congestion	232
13.1.4	Conclusion sur la congestion avec HTTP	233
13.2	Turbulences	233
13.2.1	Première approche	234
13.2.2	Facteurs de turbulence	235
13.2.3	Des cas réalistes de régime transitoire long et de turbulences pour TCP-AIMD (valable pour SACK, Reno et New-Reno)	236
13.2.4	Conclusions sur le phénomène de turbulence	236
Annexes		
Annexe A		
Des files d'attente de petite taille		
A.1	Pourquoi une file d'attente?	241
A.2	Limite de fonctionnement des files d'attente face à TCP	243
A.3	Dimensionnement minimal du buffer d'un routeur	244

Annexe B**Quelques théorèmes de mathématiques générales**

B.1	Lemmes de Gronwall	249
B.2	Autour des lemmes de Dini et son corollaire probabiliste	250
B.2.1	Deuxième généralisation : sur les fonctions	252
Bibliographie		255

Résumé

Avec l'avènement d'Internet, le monde est entré dans une ère nouvelle. Celle d'un partage quasiment sans limite de l'information. Ce n'est pas seulement un medium de communication nouveau qui est apparu, mais bien aussi un changement dans la façon de vivre de chacun. L'un des aspects de cette révolution est le transport effectué par le réseau utilisant souvent le protocole TCP. Aujourd'hui encore les usages sont parfois limités par les débits disponibles. La nouvelle génération d'Internet, peut-être issue des recherches sur Internet 2, offrira des possibilités sans doute encore plus étonnantes.

Ce mémoire porte sur l'étude du partage de bande passante par TCP. Elle commence par deux chapitre introductifs sur TCP (chapitre 1) et ses modèles mathématiques (chapitre 2) qui conduisent à trois parties. La première partie est très mathématique, partant d'un exemple très simple au chapitre 3 sur des tirages avec ou sans remise d'une urne, elle montre ensuite comment généraliser la méthode de champ moyen de «lookdown process» à des trajectoires avec des dynamiques complexes à travers trois chapitres ; le chapitre 4 pose le problème d'utilisateurs qui interagissent à travers une ressource partagée par l'intermédiaire de sauts poissonniens ; le chapitre 5 développe une méthode inspirée du «lookdown process» pour résoudre ce premier problème ; enfin, le chapitre 6 propose une généralisation où l'interaction se fait par des sauts à intensités mais aussi par des sauts synchrones.

La deuxième partie présente un travail joint avec David McDonald et François Baccelli sur la modélisation de TCP dans le cas de téléchargements de longue durée et une généralisation pour tenir compte d'utilisateurs non persistants. Le chapitre 8 présente le modèle utilisé, il s'agit d'étudier un grand nombre d'utilisateurs qui partagent une ressource commune qui est la bande passante. Ces utilisateurs se servent du protocole TCP sous un certain nombre d'hypothèses simplificatrices (modèle fluide, nombre d'utilisateurs constant avec le temps). Le chapitre 9 poursuit en proposant la démonstration de la convergence en champ moyen : lorsque le nombre d'utilisateurs qui partagent une ressource rare devient grand, les équations des files d'attente/débits au routeur deviennent des équations aux dérivées partielles déterministes. L'étude mathématique de ces équations de transport est l'objet du chapitre 10. Elles y sont établies, simplifiées et étudiées. En particulier, on y voit une étude de la stabilité d'une file d'attente sous certains contrôles centralisés comme RED. Cette partie s'achève avec le chapitre 11 qui propose une généralisation du modèle pour prendre en compte les utilisateurs du protocole HTTP sur TCP.

La troisième partie est dédiée aux simulations. Le chapitre 12 montre certains aspects de la conception du simulateur des équations aux dérivées partielles des utilisateurs de TCP persistants et de leur généralisation à HTTP. Le simulateur est ensuite validé expérimentalement puis nous discutons des principales améliorations qu'il faudrait apporter au simulateur et au modèle. Enfin le chapitre 13 propose deux exemples d'utilisations du simulateur, le premier illustre le problème de la congestion pour les utilisateurs de HTTP en insistant sur le fait que le problème est plus grave qu'on pourrait s'y attendre ; le second

présente un effet de turbulence mathématique pour HTTP/TCP, c'est à dire que dans certaines conditions, nous trouvons deux états limites, un où le débit est constant et un état oscillatoire où le débit moyen est moindre.

Mots-clés: HTTP, TCP, RED, AQM, Champ moyen.

Abstract

The dramatic spread of the Internet all around us marked the world forever. It has now become possible to send and receive an almost infinite amount of data to and from anywhere in the world. Not only does this change the way we communicate, it also has a considerable impact on our perception of the world we live in. The Internet revolution is partly due to data transport, often via TCP. The optimization of the information highway is one of the most exciting challenges for the Internet architects.

The two introduction chapters of this PhD thesis straightforwardly explain TCP and the related works in the mean field approach. We focus on TCP New-Reno and some particular problems TCP-networks face. Our angle on the issue will mainly be bottleneck congestion.

The first part deals with the mathematics of the mean field we develop. First we show a very simple example of what a mean field is, then we develop some theoretical enhancements to finish with the proof of Baccelli-Hong's model for TCP.

The second part endeavors to explain our mathematical mean field model to study large TCP controlled population sharing a common resource, namely the bandwidth of some internet access router. We make some simplifying assumptions such as fluid behavior or Little's formula for the window-bandwidth relationship. The evolution of the histogram of the connection window sizes is approximated by a deterministic partial differential equation. From a technical point of view, we can furnish QOS estimations. From a mathematical point of view, compared to the traditional tightness-limit-unicity mean field method, the auxiliary system-coupling methodology that we use offers certain advantages. It is more intuitive and allows us to solve the delay problem which otherwise would be difficult to handle. We then adapt the model to the HTTP case with on-off sources.

The third part of the thesis focuses on showing practical results we obtain by simulating mean field equations. We first validate the mean field model against discrete event simulators and the N -particle system from which it is extracted for large N , then show how such a model is different from (ie. better than) other simpler models. We conclude by simulating the equations which arise from the models which exhibit some interesting congestion behavior.

Keywords: HTTP, TCP, RED, AQM, Mean field.

Introduction

Cette thèse examine un problème pratique directement issu de l'étude du réseau de télécommunications Internet en utilisant des outils mathématiques issus en particulier de la théorie des probabilités. Aux frontières entre deux communautés scientifiques, cette thèse peut être lue autant par l'ingénieur que par le mathématicien.

Ce travail a débuté durant mon stage de DEA à l'Université d'Ottawa avec David McDonald. L'objectif que nous nous étions fixé était de comprendre comment fonctionnait le partage macroscopique de bande passante effectué par un routeur d'Internet confronté à de nombreux utilisateurs. Ceci dans le but de savoir ce qui se passait d'une part, et de tenter de remédier aux problèmes de congestions constatés d'autre part. Certes les algorithmes employés sont bien connus et assez simples, mais l'effet de leur utilisation par un très grand nombre d'utilisateurs et l'interaction qui en résulte sont plus mystérieux.

Pour effectuer cette modélisation, l'idée d'utiliser des méthodes de champ moyen nous a guidé depuis le début. Cependant ce qui devait être une simple application des connaissances classiques issues des méthodes utilisées par les mathématiciens et les physiciens pour la physique statistique s'est révélé plus compliqué que prévu à cause des délais et des effets de bords présents dans l'objet que nous étudions. Aussi nous avons suivi une approche dite du «lookdown process». Cette approche consiste à substituer à la dynamique des particules étudiées les évolutions soupçonnées lorsque le système a une taille infinie. Ceci permet d'obtenir un système «auxiliaire» couplé avec le premier où les particules évoluent indépendamment les unes des autres et interagissent uniquement avec le champ créé par le groupe. Par rapport à la manière classique de faire, cette approche présente plusieurs avantages. Avant tout, elle est plus simple à comprendre que les méthodes classiques; en particulier elle est accessible sans trop d'efforts à un non spécialiste car tous les résultats ont une interprétation trajectorielle facile à visualiser. De plus, elle est plus générale par certains aspects et ne nous a pas obligé pas à remanier notre modèle pour tenir compte des difficultés mathématiques, ce qui a permis de rajouter des caractéristiques aux modèles de manière incrémentale.

Le spectre d'application de la méthode développée dans cette thèse est bien plus vaste que le problème initial comme l'illustrent les généralisations obtenues aux chapitres 5 et 6. On imagine sans mal qu'elle pourrait servir dans d'autres domaines.

Du point de vue de l'étude de TCP, nous donnons certains critères analytiques de stabilité d'un routeur, et étudions numériquement quelques problèmes liés à l'usage de HTTP.

Plan du mémoire

Les deux chapitres introductifs présentent les aspects techniques du problème étudié (Chap. 1) et les principaux modèles associés qui utilisent une approche de champ moyen. Ceci permet de comprendre comment se place cette thèse dans l'étude mathématique de TCP par champ moyen.

La première partie concerne les mathématiques du «lookdown process», approche que nous désignons aussi par méthode «couplage-système auxiliaire». Nous montrons quelques unes des méthodes existantes sur un exemple simple (Chap. 3) avant d'expliquer comment appliquer la méthode à des évolutions de systèmes de particules interdépendants parce qu'ils utilisent une ressource commune (Chap. 4). Ceci passe par l'introduction d'un couplage et d'un système auxiliaire (Chap. 5). Dans un premier temps nous envisageons le cas où la dépendance par rapport à la ressource commune agrégée est continue ou avec des sauts à intensité bornée, puis nous ajoutons des problèmes de frontières où des sauts synchronisés peuvent avoir lieu (Chap. 6), ce qui nous amène en application directe à prouver la convergence du champ moyen pour le modèle Baccelli-Hong exposé plus tôt.

La deuxième partie explique le modèle de TCP que nous utilisons (Chap. 8), adapte les démonstrations de la première partie dans le cas particulier de structure de délais et de ressource partagée qui nous intéresse (Chap. 9), puis étudie de manière simplifiée les équations obtenues (Chap. 10). Enfin le chapitre 11 généralise l'approche pour prendre en compte des sources intermittentes ce qui modélise le comportement d'utilisateurs du protocole HTTP.

La troisième montre les résultats de simulations du champ moyen en comparant cela d'une part à des simulations d'autres simulateurs très répandus dans la communauté des réseaux, et d'autre part à des réalisations du système fini pour différents nombres d'utilisateurs d'où nous avons dérivé le modèle en champ moyen (Chap. 12). Le but étant dans un premier temps d'illustrer les étapes de la conception du simulateur pour ensuite l'utiliser dans des scénarii proches de la réalité (Chap. 13) ; dans le cas d'utilisateurs de HTTP, nous montrons ainsi par simulation l'effet de la congestion sur les performances ressenties par les utilisateurs puis un problème congestion dans des cas de sous-utilisation de la bande passante.

Principaux apports de cette thèse

Un premier aspect intéressant de cette thèse est de mener de bout en bout la démarche de la modélisation du système réel jusqu'à l'application numérique en passant par l'étude mathématique. Les parties mathématiques présentent une méthode à la marge du travail habituel sur les champs moyens en utilisant la technique de couplage associée à un système auxiliaire plutôt que la compacité. Ceci nous permet de traiter des problèmes avec des délais ainsi que des effets de bords avec des changements de dynamiques. Du point de vue pratique nous étudions les populations d'utilisateurs de TCP en montrant des faits bien plus forts que ceux qui existaient puisque nous obtenons des résultats trajectoriels avec la possibilité de traiter des effets de bords et des délais. De plus nous adaptons nos résultats à l'étude de sources intermittentes ce qui constitue une avancée intéressante. Enfin nous mettons en pratique tout ceci en présentant des résultats de simulations issues du modèle.

Chapitre 1

Situation du problème

Pour commencer, nous introduisons le problème qui a motivé cette thèse, c'est-à-dire l'étude de TCP (Transmission Control Protocol), le principal protocole d'échange de données d'Internet. TCP rentre en particulier en jeu dans le téléchargement de page web ou dans l'échange de donnée «peer to peer» qui représentent à eux deux une grande partie du volume des transferts effectués sur Internet.

Sommaire

1.1	Le défi relevé par TCP	7
1.1.1	Principe général de fonctionnement	8
1.1.2	Le cahier des charges	8
1.1.3	La congestion dans Internet	9
1.1.4	Position du problème d'optimisation	10
1.2	Comment TCP Reno résout le problème	11
1.2.1	Détection des pertes	11
1.2.2	Fenêtre et contrôle du flux de données	12
1.2.3	Fenêtre de congestion et détection de pertes	13
1.3	Les algorithmes de gestion de taille des files d'attente des routeurs	14
1.3.1	La «Drop Tail»	14
1.3.2	Destructions de paquets anticipées	14
1.3.3	RED	15
1.3.4	Gentle RED	16
1.3.5	AQM	16
1.4	Les modèles mathématiques de TCP	16
1.4.1	Quelques types de modèles	17
1.4.2	Pourquoi s'intéresser au champ moyen	18

1.1 Le défi relevé par TCP

Les informations techniques sur les protocoles d'Internet sont librement consultables, ce sont les RFC (Request For Comments). Par ailleurs cette section s'aide de la présen-

tation très claire de TCP Reno qui est faite dans la thèse de Thomas Kelly [84].

Le protocole d'échange d'Internet que nous étudions se nomme TCP Reno. TCP Reno fait partie de la famille des protocoles TCP où on trouve aussi TCP Vegas ou TCP Tahoe. D'autres familles de protocoles d'envoi de données sur Internet existent comme UDP ou RTP par exemple, ces derniers étant plus adaptés à la diffusion de données («streaming»).

1.1.1 Principe général de fonctionnement

Internet permet de véhiculer des paquets de données, c'est-à-dire une séquence de données d'une taille fixée. Lorsqu'un paquet est envoyé sur le réseau, tout est fait pour qu'il arrive au mieux (c'est le «best effort»), mais aucune assurance n'est donnée que ce soit en termes d'intégrité, de délai ou d'ordonnancement des paquets. L'envoi de paquets sur Internet est donc analogue à l'envoi du courrier par la poste.

Taille des paquets Nous ne rentrerons pas dans les détails techniques, mais il existe plusieurs notions selon qu'on compte ou pas les données utilisées pour permettre le transport. Par analogie la question est de savoir si on compte ou pas le poids de l'enveloppe quand on envoie une lettre. Retenons qu'il existe une taille minimale de paquet que les media d'Internet peuvent nécessairement transporter, elle est de l'ordre de 500 octets ; il y a aussi un maximum protocolaire de 64 kilooctets. La taille typique normalement utilisée est de l'ordre de 1500 octets. Un paquet n'a pas de taille minimale pour les données transportées : il est possible d'envoyer un paquet contenant 1 bit de donnée, cependant une fois « encapsulé » par les protocoles de transport la taille varie selon les protocoles et les media utilisés.

Le choix de la taille des paquets dépend de considérations techniques sur la qualité du medium de transport des données. Pour nous, la taille du paquet sera considérée comme une donnée figée.

1.1.2 Le cahier des charges

Le rôle de TCP est de fournir un service de transfert fiable d'un fichier de taille quelconque sur un medium qui ne transporte que des paquets de taille limitée sans aucune assurance d'intégrité, de délai et d'ordonnancement. Ceci tant qu'à faire en utilisant «au mieux» les ressources disponibles.

Intégrité des données TCP contrôle l'intégrité des données et s'assure que la totalité des paquets envoyés arrivent dans de bonnes conditions. Ceci se fait en ajoutant des «sommes de contrôle», c'est-à-dire une clé calculée à partir données envoyées. Une clé est beaucoup plus courte que le message original et permet de contrôler à la réception que le message reçu correspond bien à la clé qui lui est attaché. Le mécanisme de calcul fait qu'il est très peu probable d'avoir des collisions (c'est-à-dire deux clés identiques pour deux paquets proches).

De plus chaque paquets porte un NUMÉRO DE SÉQUENCE et un NUMÉRO D'ACK (ACK pour acknowledgement). Le numéro de séquence est celui du premier octet de

donnée du paquet considéré, le numéro d'ACK, celui du prochain octet à envoyer. Pour s'assurer que les données arrivent, le programme qui implémente TCP à la destination renvoie des accusés de réception pour tous les paquets reçus avec succès. Ces ACK sont eux-même des paquets TCP qui contiennent certaines données de contrôle dont le numéro d'ACK qui correspond au numéro du premier octet dont la destination ne dispose pas encore. Si la source des données ne reçoit pas d'ACK avant un temps d'expiration nommé le TIMEOUT, la donnée est alors considérée comme perdue et retransmise.

Bonne utilisation du débit disponible Pour utiliser au mieux les ressources disponibles, TCP utilise un protocole de fenêtre. Ceci signifie que la source s'autorise à envoyer à un destinataire un nombre de paquets correspondant à la taille de la fenêtre sans attendre la réception d'ACK. La taille des données correspondant aux paquets ainsi envoyés et dont l'ACK n'est pas encore parvenu à la source du flot de données se nomme la FENÊTRE (dont l'unité standard est l'octet ou le bit). La valeur de la fenêtre sera dans la suite de cette thèse assez souvent donnée en paquets, ce qui en toute rigueur est impropre puisqu'il n'existe pas de correspondance fixe entre nombre de paquet et nombre de bits de données.

Le problème de TCP Si le problème de l'intégrité des données dépend des contraintes physiques, l'optimisation de l'usage est un problème difficile. En effet très souvent la ressource utilisée, ici, la BANDE PASSANTE (c'est-à-dire la capacité à transporter de l'information) vient à manquer. Une adaptation adéquate de la taille de la fenêtre devient alors un problème particulièrement délicat.

1.1.3 La congestion dans Internet

Aux débuts d'Internet, le problème de la congestion ne se posait pas car les liens n'étaient pas très fiables et les machines trop peu puissantes pour échanger de grandes quantités de données. Une fois ces deux limites levées, la capacité disponible est devenue une ressource limitante dans les réseaux filaires. Les premiers problèmes se posèrent aux nœuds d'Internet, les routeurs et les switches. La limite fut levée (provisoirement) en installant des dispositifs de stockage temporaire des paquets, ou FILES D'ATTENTES, dans ces équipements. Cependant ceci ne fut pas suffisant car la demande de bande passante croissait toujours beaucoup plus vite que l'offre disponible.

Comment régler ce problème dans le cadre d'Internet ?

Par l'instauration de règles de fonctionnement : Au niveau du protocole IP, Internet n'est pas muni d'un dispositif de facturation de l'usage¹, ce qui fait qu'il n'est pas possible d'augmenter le prix de la bande passante pour limiter mécaniquement son usage. La seule solution possible - sans changer la physionomie d'Internet - est donc la régulation du trafic. Cependant le réseau étant décentralisé, cette régulation ne peut être qu'une autorégulation des utilisateurs ou un contrôle imposé par les ressources. C'est-à-dire que

¹bien sûr les opérateurs peuvent mesurer les débits. Par exemple dans le WAP ou les anciens accès par câble on paye en fonction de l'utilisation.

dans leur propre intérêt les utilisateurs et les fournisseurs de service décident tous de respecter des règles d'usage pour permettre un fonctionnement agréable du réseau.

Par l'utilisation de TCP, un contrôle décentralisé couplé avec RED, un contrôle au niveau des ressources : l'autorégulation passe dans TCP par un contrôle de la taille des fenêtre, c'est-à-dire de l'agressivité du protocole. Ce contrôle est fait par les utilisateurs de manière décentralisée en utilisant les seules informations dont ils disposent à savoir si les paquets envoyés arrivent ou pas et si oui en combien de temps.

De plus, compte tenu du mécanisme d'autorégulation des utilisateurs, l'idée que le réseau pouvait envoyer des informations sur son état en détruisant à l'avance des paquets a été introduite en 1993 dans [51], c'est l'algorithme RED (Random Early Detection). Cette idée rentre quelque peu en contradiction avec le principe de couches dans Internet, puisque le protocole de transport des paquets se met à tenir compte de la manière dont certains utilisateurs utilisent ces paquets limitant ainsi les évolutions possibles des protocoles. Nous répondrons que RED commence à fonctionner quand la ressource est déjà en état de congestion. Que des paquets soient détruits un petit peu en avance ne change globalement pas le problème et permet aux utilisateurs d'anticiper s'ils le souhaitent le phénomène de débordement de la file d'attente qui va se produire.

Le problème dont il est ici question est l'optimisation décentralisée par un grand nombre d'utilisateurs indépendants qui partagent une ressource munie d'un contrôle. C'est ce qui a motivé cette thèse dont le but est de comprendre, mathématiques à l'appui, une partie des phénomènes qui rentrent en jeu (l'objectif dans un premier temps étant plus de comprendre le phénomène que de résoudre le problème).

1.1.4 Position du problème d'optimisation

Internet est un réseau physique où de nombreux utilisateurs peuvent échanger des paquets. La traversée des liens et des nœuds occasionne des délais de transmission, des pertes et plus rarement des altérations des données. Nous supposons que le réseau n'a accès qu'à des données agrégées sur les utilisateurs telles que le débit en un point du réseau. Les utilisateurs ne reçoivent d'informations du réseau que par les délais de transmission éventuellement infinis lorsque le paquet s'est perdu ou par l'arrivée de données erronées.

Compte tenu de ces informations, les utilisateurs qui ont des données à s'échanger doivent mettre en place un mécanisme de transmission qui permet l'envoi de leurs données en assurant l'intégrité et un temps de transfert minimal. Certains points du réseau sont limitants en terme de débit disponible, et le réseau peut essayer d'indiquer aux utilisateurs son état en détruisant à l'avance des paquets.

Le caractère optimal des solutions recherchées est un sujet de débat. Nous retiendrons que le but est de maximiser une notion dénommée l'UTILITÉ. L'utilité collective est la somme des utilités individuelles. Nous allons voir comment ce problème apparemment insoluble a été résolu de manière approchée par l'utilisation de TCP Reno. nous n'essaierons pas de résoudre le problème d'optimisation que nous venons de poser mais plutôt de comprendre le fonctionnement de la solution mise en place par TCP Reno.

1.2 Comment TCP Reno résout le problème

TCP et le réseau Internet furent créés entre autres par les équipes de Vint Cerf pour la défense américaine entre la fin des années 70 et le début des années 80. Le seul mécanisme de contrôle était le timeout et une fenêtre à croissance linéaire. Les premières améliorations introduisant la notion de contrôle de congestion dans TCP furent introduites à la fin des années 80 par Van Jacobson ([76] où les algorithmes SLOW START et CONGESTION AVOIDANCE furent aussi introduits). L'algorithme FAST RETRANSMIT a été introduit par Mark Allman, Vern Paxson, et Richard Stevens (Fast retransmit RFC 2581 [2]) et SACK que nous évoquerons parfois par Matthew Mathis, Jamshid Mahdavi, Sally Floyd, et Allyn Romanow (Selective ACK RFC2018, [104]).

1.2.1 Détection des pertes

Les problèmes de transmission viennent des paquets qui sont soit perdus soit altérés (le problème de la duplication de paquet est secondaire). Détecter l'altération d'un paquet est facile car il est peu probable qu'il ait la bonne somme de hachage dans ce cas. La détection des pertes est plus délicate parce que les paquets arrivent dans le désordre. Un paquet perdu est simplement un paquet qui n'arrive jamais. Il n'est donc pas possible de savoir avec certitude si un paquet a bel et bien été perdu.

Convention sur les ACK La destination renvoie un ACK à chaque paquet reçu avec succès. Cet ACK contient le numéro du prochain octet attendu par la destination. Lorsque des paquets arrivent désordonnés à destination ou sont perdus, il se crée pour la destination un trou dans le fichier. L'ACK envoyé indique alors le même numéro que l'ACK précédent jusqu'à ce que le trou soit comblé. Les ACK qui arrivent en plusieurs exemplaires sont nommés ACK DUPLIQUÉS (Duplicate ACK).

Par cette convention un ACK peut accuser réception de plusieurs paquets à la fois. Aussi la source selon sa fenêtre agit alors comme si elle avait reçu un ACK par paquet. Si la fenêtre est constante par exemple la source renvoie autant de paquets que le nombre de paquets dont le dernier ACK accuse réception.

Le timeout La perte d'un paquet peut être détectée de plusieurs manières. La méthode standard est un mécanisme de timeout basé sur la connaissance approximative du temps pour recevoir un ACK une fois un paquet envoyé. Ce temps de retour se nomme ROUND TRIP TIME (RTT). Pour mémoire, le RTO, temps avant un timeout vaut la somme d'un estimé lissé du RTT (nommé SRTT) et de quatre fois sa variance estimée lissée (VRTT). Lorsque le temps d'attente entre l'envoi d'un paquet et la réception d'un ACK indiquant que ce paquet a été reçu dépasse la valeur RTO, on dit qu'il y a un timeout. La source renvoie alors le paquet suspect.

Fast retransmit Les premières améliorations de ce mécanisme ont été fast retransmit et SACK. Comme nous l'avons déjà évoqué, quand un paquet arrive qui n'est pas celui qu'on attendait immédiatement à destination, un ACK identique au précédent est renvoyé.

La source recevra donc des ACK dupliqués. Le mécanisme fast retransmit déduit de la réception d'un nombre donné d'ACK dupliqués qu'une perte a eu lieu (3 identiques dans les implémentations habituelles, c'est-à-dire en tout 4 ACK identiques à la suite²). TCP réagit normalement à la source à la détection d'une perte, il envoie le paquet le plus ancien qui n'a pas eu d'ACK. Ceci est une amélioration par rapport au mécanisme de timeout, si la fenêtre est grande, une perte est détectée par fast retransmit immédiatement alors qu'il faudrait le temps RTO sinon.

SACK SACK utilise des bits de l'entête TCP laissés au départ pour un usage futur. SACK permet d'envoyer à la source dans une ACK dupliqué l'information des paquets qui ont été reçus au delà du paquet manquant. Il est évident que SACK présente une amélioration puisque la source dispose d'une information supplémentaire qu'elle est libre d'utiliser ou pas. La quantification précise de l'amélioration apportée n'est cependant pas évidente et les paramètres sont encore un sujet d'étude.

1.2.2 Fenêtre et contrôle du flux de données

Nous allons détailler ce qu'est exactement le mécanisme de contrôle de la valeur de la fenêtre. Il y a à tout instant un certain nombre de paquets en vol, c'est-à-dire envoyés et sans ACK de retour à la source. Leur nombre est contrôlé par deux valeurs nommées FENÊTRE DE RÉCEPTION et FENÊTRE DE CONGESTION. À tout instant le nombre de paquets en vol est inférieur au minimum des deux fenêtres.

Fenêtre de réception La fenêtre de réception (receiver window *rwnd*) est écrite par la destination dans les ACK. Cela indique la place disponible dans la file d'attente réservée à cette connexion chez l'utilisateur qui reçoit les paquets. Cette file d'attente est entre TCP et l'application qui utilise le flux de données. Dans le cas où la liaison est plus rapide que la capacité de traitement à destination, cela évite de devoir jeter inutilement des paquets faute de pouvoir s'en servir.

Fenêtre de congestion La fenêtre de congestion (congestion window, *cwnd*) a pour but de résoudre le problème du partage des ressources d'Internet. L'hypothèse employée par ce mécanisme est que la fiabilité des liens est devenue très importante. Ainsi, une perte est nécessairement due à la congestion dans le réseau³. *cwnd* a plusieurs modes d'évolution, chacun dépendant du fait que tout se passe bien ou qu'il y ait une perte détectée.

La première règle possible est SLOW-START. Le but de slow-start est de trouver rapidement une fenêtre opérationnelle. Chaque ACK reçu fait augmenter *cwnd* d'un paquet. La source envoie donc deux paquets chaque fois qu'un paquet a bien été reçu et son ACK renvoyé.

²On considère qu'un ACK plus ancien qui arriverait après un autre ACK, est un ACK dupliqué, même si dans la pratique ils ne sont pas identiques.

³Ceci n'est pas du tout vrai dans les réseaux sans fil, ce qui explique en partie pourquoi les versions actuelle de TCP fonctionnent difficilement sur les réseaux WiFi par exemple.

La deuxième est l'ÉVITEMENT DE CONGESTION (congestion avoidance). Le but de congestion avoidance est de gérer *cwnd* pour rester au plus élevé possible sans provoquer de congestion. *cwnd* augmente en gros de $1/cwnd$ à chaque ACK (exactement de $\text{paquet}/\lfloor \frac{cwnd}{\text{paquet}} \rfloor$). Comme on n'envoie pas de portion de paquets, cela signifie que pour tous les *cwnd* paquets reçus on renvoie un paquet de plus dans le réseau.

Pour choisir dans quel mode faire évoluer *cwnd*, on utilise la variable *ssthresh* (Slow Start THRESHold). Si *cwnd* est en dessous de *ssthresh* on utilise slow-start, sinon on utilise l'évitement de congestion. C'est une convention, mais FAST RETRANSMIT/FAST RECOVERY sont des sous-modes de slow-start ou de l'évitement de congestion et timeout n'est pas véritablement un mode de fonctionnement. Tant que le timer n'a pas expiré, on est encore dans le mode précédent de mise à jour de *cwnd*. Ensuite, on rentre instantanément en slow-start. Cependant, ce timer peut mettre beaucoup plus qu'un RTT typique (timeouts répétés) pour se décider à expirer.

1.2.3 Fenêtre de congestion et détection de pertes

Nous rentrons ici dans des détails un peu techniques. *cwnd* réagit de la même manière face aux pertes qu'il soit en slow-start ou bien en congestion avoidance. Pour bien comprendre le mécanisme il faut noter *W* la fenêtre effective, c'est-à-dire le nombre de paquets en vol.

Si la perte est détectée par timeout, on a instantanément :

$$cwnd = 1 \text{ paquet} \quad (1.1)$$

$$RTO = 2RTO \quad (1.2)$$

$$ssthresh = \max(2 \text{ paquets}, \frac{W}{2}). \quad (1.3)$$

Ceci signifie que la connexion est réinitialisée, avec deux conséquences "négatives", un délai de timeout doublé (si on enchaîne les timeouts cela devient exponentiel, sinon, RTO reviendra rapidement à des valeurs proches du RTT) et un seuil pour passer à congestion avoidance plus bas (c'est négatif parce que congestion avoidance est beaucoup moins agressif que slow start). Globalement faire un timeout n'est vraiment pas une bonne chose, mais cela arrive dans le cas des petites fenêtres. Par exemple si $W = 4$, la perte réelle d'un paquet ne fera pas suffisamment d'ACK dupliqués (il faudrait en tout 4 ACK identiques), et on aura automatiquement un timeout. Dans le cas de grandes fenêtres, par contre, il est très peu probable de faire un timeout.

Si la perte est détectée par fast retransmit, les conséquences sont bien moindres : on fait fast retransmit/fast recovery. Ces deux algorithmes fonctionnent en tandem. Le présumé est que le réseau IP ne duplique pas les paquets (un petit malin qui voudrait faire planter tout un réseau enverrait en triple des ACK qu'il intercepterait). Voici ce qui est fait quand 4 ACK identiques sont reçus :

1. comme dans le cas d'un timeout : $ssthresh = \max(2\text{paquets}, \frac{W}{2})$;
2. on retransmet le paquet suspect ;

3. on choisit la nouvelle fenêtre en tenant compte des 3 ACK reçus en avance,

$$cwnd = ssthresh + 3;$$

4. à chaque nouvel ACK dupliqué,

$$cwnd = cwnd + 1;$$

5. si $W < \min(rwnd, cwnd)$, on envoie les paquets qui suivent les derniers envoyés (dont on ne suspecte pas encore la perte) ;

6. dès qu'un ACK différent est reçu on revient en mode normal avec :

$$cwnd = ssthresh.$$

Notons que le fait d'augmenter $cwnd$ d'un paquet à chaque ACK reçu correspond à une compensation du fait que cet ACK n'a pas diminué la valeur du nombre de paquets en vol, ce qui aurait dû être le cas si on avait reçu ce qui précède. *L'effet global de tout ceci est de diviser la fenêtre par deux dans l'espace d'un RTT lorsqu'une perte est constatée.*

1.3 Les algorithmes de gestion de taille des files d'attente des routeurs

Nous n'expliquerons pas en détails comment placer les files d'attente à un routeur en entrée ou en sortie. On suppose qu'il est muni d'une unique file FIFO (First In First Out ou premier arrivé, premier servi) à l'entrée et qu'il a un unique lien de sortie, comme c'est le cas pour un routeur d'accès.

1.3.1 La «Drop Tail»

À priori une file d'attente peut stocker une taille fixée de bits de données. Le fonctionnement normal d'une file d'attente est d'une part que le processeur du routeur envoie à la vitesse du lien sortant des données en provenance de la file ; et d'autre part que la file se remplit - tant qu'elle a de la place disponible - avec des paquets qui arrivent au routeur. Quand un paquet arrive et qu'il n'y a pas assez de place disponible pour le stocker, il est détruit. On nomme ce mode normal de fonctionnement «DROP TAIL».

1.3.2 Destructures de paquets anticipées

Comme nous venons de l'évoquer (page 10), Sally Floyd et Van Jacobson (dans [51]) proposèrent ceci : compte tenu du fait qu'une grande partie des utilisateurs d'Internet utilisent TCP Reno (sensible aux pertes de paquets), une file d'attente congestionnée peut envoyer un signal de congestion soit par une notification explicite⁴ (ce qui suppose d'améliorer l'algorithme installé sur toutes les machines), soit en détruisant des paquets. Dans la foulée l'article proposait l'algorithme nommé RED (Random Early Detection) qui mettait en pratique leur idée.

⁴ La norme de notification explicite de congestion se nomme ECN (Explicit Congestion Notification) dans la littérature.

1.3.3 RED

Nous reprenons ici les notations de [51] et mettons entre parenthèses les notations que nous utiliserons pour désigner ces quantités dans la suite de la thèse.

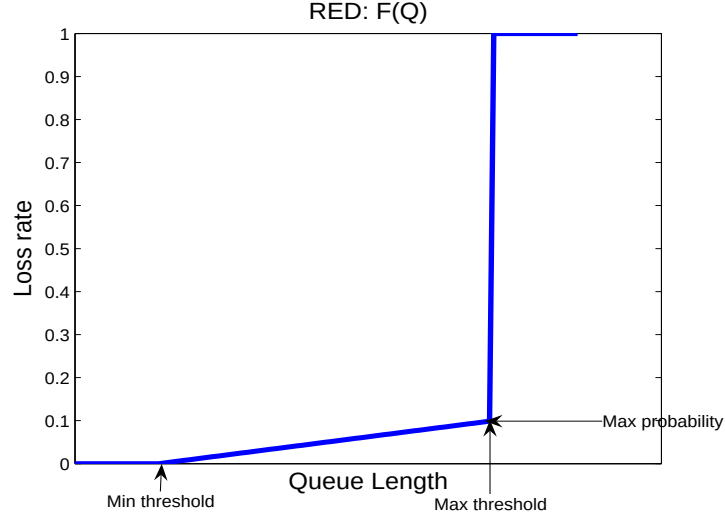


FIG. 1.1 – Fonction génération des pertes de RED

On se donne un seuil $\text{Min}_{\text{threshold}}$ sur la file d'attente (nous l'appellerons $Q_{\min}$ dans la suite), et au delà de ce seuil, la file d'attente provoque des destructions de paquets anticipées (le paquet est détruit au lieu de rentrer dans la file d'attente). L'idée de ce seuil est qu'en deçà on considère que le remplissage de la file est dû aux fluctuations stochastiques d'une file sous-utilisée ; quand on atteint ce seuil l'utilisation sur les derniers instants a dépassé la capacité, ce qui signifie que la file d'attente est congestionnée.

En réalité on ne regarde pas la taille réelle de la file d'attente qui varie de manière erratique, mais un estimé $\tilde{q}(t)$ lissé dans par une moyenne temporelle. Le principe général de l'application des pertes est de tirer une probabilité de perte $p(t) = F(\tilde{q}(t))$ qui croît linéairement de 0 à une valeur Max_p (que nous notons $p_{\max}$ dans la suite) lorsque $\tilde{q}(t)$ va de $\text{Min}_{\text{threshold}}$ à une valeur $\text{Max}_{\text{threshold}}$ (que nous noterons $Q_{\min}$ et $Q_{\max}$ dans la suite) qui peut valoir jusqu'à la taille B de la file d'attente. Une fois ce seuil dépassé, la file d'attente détruit tous les paquets entrants, comme si on avait dépassé la capacité réelle de la file.

Il existe de nombreuses manières de détruire des paquets avec probabilité $p(t)$, par exemple on peut détruire le premier paquet après l'initialisation puis laisser passer les $\lceil 1/p(t) \rceil - 1$ paquets suivants et ainsi de suite (ce qui est une méthode déterministe de destruction de paquets). RED doit son nom au fait que cet algorithme est randomisé dans les spécifications originales. On détruit chaque paquet avec probabilité $\tilde{p}(t)$, où $\tilde{p}(t)$ est choisi pour égaliser les pertes, c'est-à-dire que $\tilde{p}(t)$ vaut plus que $p(t)$ si aucun paquet n'a été détruit dans le passé proche et vaut moins dans le cas contraire.

1.3.4 Gentle RED

Dans la suite du mémoire nous parlerons parfois de *Gentle RED*. Introduit par Vincent Rosolen, Olivier Bonaventure et Guy Leduc ([137]), cet algorithme est en tout point identique à RED en prenant la fonction une fonction de pertes qui ne fait pas de saut à la valeur Q_{max} , mais qui croît linéairement de Max_p en $Max_{threshold}$ à 1 à Q_{max} , la taille de la file d'attente. Cette version de RED présente l'avantage de ne pas introduire de discontinuité dans la dynamique d'évolution des pertes.

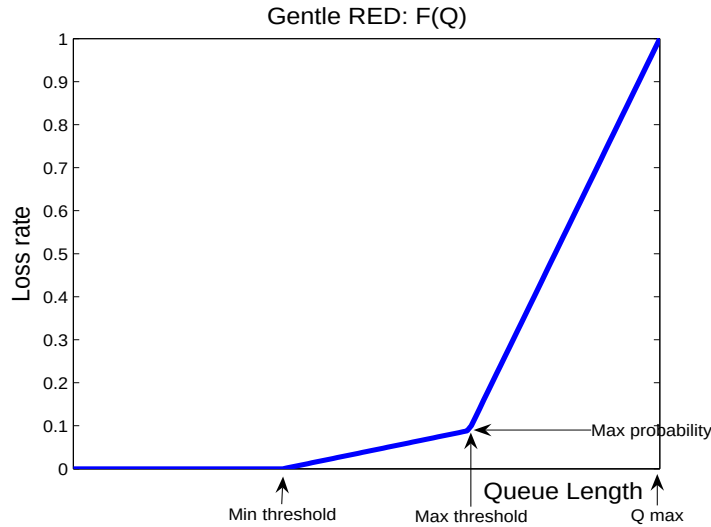


FIG. 1.2 – Fonction génération des pertes de Gentle RED

1.3.5 AQM

Les algorithmes de type AQM (Active Queue Management) sont inspirés de RED, leur point commun est de provoquer des pertes anticipées pour contrôler les utilisateurs de TCP Reno qui forment le flux élastique des routeurs. Tous ces algorithmes ont un problème en commun, c'est qu'ils sont très difficiles à régler au vu de la grande variété des situations possibles. De sorte que si RED est aujourd'hui implémenté dans tous les routeurs, il n'est presque jamais activé. Rappelons de plus que sur le plan de la conception des réseaux, l'utilisation de RED a de nombreux opposants pour qui il n'est pas sain que le réseau pénalise indifféremment tous les utilisateurs en détruisant des paquets pour contrôler uniquement certains d'entre eux qui utilisent un protocole qui pourrait bien disparaître un jour. Il convient donc de réserver les usages des AQM aux cas où la congestion est avérée ou en différenciant le trafic élastique des autres types de données échangées.

1.4 Les modèles mathématiques de TCP

Précisons que nous nous intéressons ici aux modèles dynamiques de TCP. En effet s'il est intéressant d'obtenir des distributions stationnaires présumées, les besoins pratiques

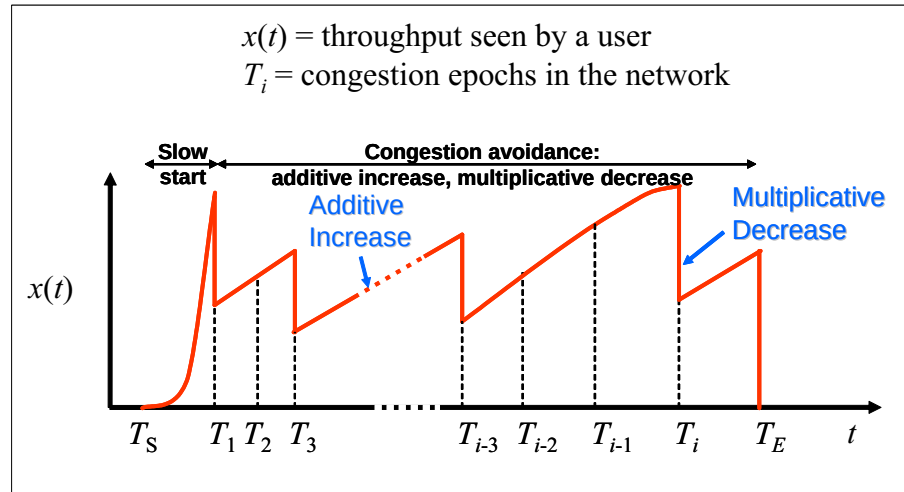


FIG. 1.3 – Modélisation mathématique de TCP. Point de vue de l'utilisateur. L'algorithme d'évitement de congestion correspond à peu près à des accroissements linéaires des débits et une décroissance multiplicative quand une perte est constatée, d'où le terme d'AIMD (Additive Increase Multiplicative Decrease).

concernent plus la stabilisation des débits pour minimiser les délais de transmission et leurs oscillations et utiliser au mieux les ressources de bande passante disponibles.

1.4.1 Quelques types de modèles

Le but de cette section est de situer l'utilisation que nous faisons de méthodes de champ moyen dans le contexte de l'étude de TCP. Dans la littérature, on retrouve plusieurs grandes familles de modèles :

- les modèles au niveau des paquet,
- les modèles fluides associés à la théorie du contrôle,
- les modèles économiques,
- les modèles champ moyen.

Les modèles «niveau paquet» Ils sont très proches de la réalité, mais relativement compliqués et ils passent en général mal à l'échelle lorsqu'il faut mettre en place des simulations ou essayer de tirer des résultats asymptotiques. Globalement, ils sont donc précis mais souvent difficiles à utiliser pour de grands réseaux. Ces modèles permettent avec un traitement dans l'algèbre $(max, +)$ de tirer des bornes dans le pire des cas (en fonction des paramètres) qui malheureusement sont souvent éloignées du cas moyen. Ce type de modèles se retrouvent par exemple dans [10, 7].

Les modèles fluides et théorie du contrôle L'idée de ces modèles consiste à gommer une partie des discontinuités du discret en rendant les évolutions continues comme l'écoulement d'un fluide. En s'inspirant de la dynamique observée, le principe est de déduire

des équations intuitives. Une fois les équations écrites, on étudie souvent leur dynamique à l'aide de la théorie du contrôle comme dans [90, 89, 111, 81, 101, 67].

Les modèles économiques Ces modèles plus macroscopiques sont basés sur la théorie économique. Ils ne modélisent pas TCP à proprement parler, mais un partage de bande passante en modélisant l'utilisation de TCP par de nombreux utilisateurs comme l'optimisation d'une fonction d'utilité convexe par chacun des utilisateurs. L'utilité générale est alors la somme des utilités particulières, une question qui se pose est de savoir si l'optimisation donne des équilibres de Pareto (c'est-à-dire qu'on est sur le bord de l'espace convexe des contraintes) ou mieux, des équilibres de Nash (la fonction d'utilité individuelle choisie conduit le groupe à optimiser l'utilité collective alors que chacun ne s'occupe que de son utilité personnelle). Ces modèles sont extrêmement macroscopiques et décrivent bien les incitations à mettre en place sur le long terme pour faire évoluer les protocoles, mais ils ne sont pas extrêmement précis et la phase de choix de la fonction d'utilité nécessite de trouver les liens entre l'approche microscopique au niveau paquet et l'approche macroscopique. Ceci se retrouve dans [80, 94].

Un autre type assez proche est la recherche d'invariants macroscopiques que l'on retrouve dans des travaux sur l'insensibilité [22, 24, 126]. On ne s'intéresse plus au fait que TCP est utilisé mais simplement aux propriétés macroscopiques du trafic.

Les modèles de Champ moyen Ce type de modèles propose de faire le lien entre le petit et le grand. Originnaire de la physique statistique, cette méthode consiste à dériver de comportements microscopiques connus des individus la loi du comportement global macroscopique d'un grand nombre d'individus et l'interaction macroscopique entre un individu et le groupe. L'utilisation d'une technique de champ moyen permet ainsi de partir d'un modèle d'utilisateur et pas directement d'une fonction d'utilité supposée. Ce modèle peut être au niveau paquet ou être fluide.

La preuve du champ moyen est un sujet qui est rarement abordé dans le domaine de l'étude de TCP. Les auteurs essaient plutôt de donner *ab nihilo* des équations fluides qui sont censées provenir d'un champ moyen sur les utilisateurs. L'intérêt de ces équations est qu'elles permettent d'étudier facilement les questions de stabilité. De plus elles sont faciles à simuler en pratique. En particulier elles n'ont aucun mal à prédire ce qui se passe avec beaucoup d'utilisateurs qui sont de difficulté croissante dans des simulations à événements discrets. Cependant, sans méthode mathématique il peut être très long de trouver ces équations. Il est de plus totalement impossible de prévoir les effets de petites variations sur les comportements individuels.

1.4.2 Pourquoi s'intéresser au champ moyen

Il est vrai que l'étude de l'interaction entre des utilisateurs par des méthodes de champ moyen permet plus d'accéder à un comportement asymptotique qu'à la «réalité» du réseau. De plus les techniques de champ moyen font intervenir des mathématiques d'une grande complexité, et mieux vaut être certain que leur usage est strictement nécessaire :

l'étude d'Internet par des simulations a longtemps été largement suffisante pour que les architectes du réseau puisse concevoir les protocoles et dimensionner les infrastructures.

Cependant voici pourquoi nous pensons que les méthodes de champ moyen constituent un outil adapté pour étudier Internet : Internet est un système chaotique régi par des règles d'interaction simples et parfaitement connues⁵. Le mot chaotique vient du fait que l'interaction qui résulte de ces règles est en apparence imprévisible par beaucoup d'aspects. Nous verrons justement dans cette thèse comment certains aspects du chaos disparaissent quand le nombre d'utilisateurs devient grand pour laisser la place à des évolutions prévisible.

L'utilisation du champ moyen permet donc d'aller du petit vers le grand. C'est à dire permet d'accéder précisément aux phénomènes qui interviennent au niveau des ressources partagées. Dans l'optique d'optimisation de l'usage de ces ressources, cette approche permet d'inverser le problème du grand vers le petit en expliquant comment le petit doit se comporter pour améliorer le grand. On a donc espoir de répondre à la question que se pose tout ingénieur face à un problème d'optimisation : «Comment peut-on choisir les paramètres d'Internet pour réaliser un optimum en termes d'utilisation du réseau?».

⁵Ces toutes les lois d'évolution, mis à part les délais sont entièrement choisies par les concepteurs.

Chapitre 2

Modèles de TCP menant à des champs moyens

Ce deuxième chapitre compare les principaux modèles de TCP qui ont été développés et qui ont mené à une étude par des méthodes de champ moyen. Nous commençons par expliquer le problème spécifique qui a motivé ces développements, puis nous présentons rapidement les différents modèles dans le but de pouvoir comparer leurs portées relatives.

Sommaire

2.1	Généralités	21
2.1.1	Autour du champ moyen	21
2.1.2	Le problème étudié, premières notations	22
2.1.3	Les équations fluides	23
2.2	Modèles par champ moyen	25
2.2.1	Champ moyen Baccelli-Hong	25
2.2.2	Champ moyen Bain-Kelly-Kelly	28
2.2.3	Champ moyen Deb-Shakkottai-Srikant	30
2.2.4	Champ moyen Makowski-Tinnakornsriruphap	31
2.3	Les modèles étudiés dans la thèse	33
2.3.1	Le modèle TCP-RED avec sources persistantes de la thèse	33
2.3.2	Adaptation HTTP-RED avec sources intermittentes	35
2.3.3	Type de résultats que nous démontrons	36
2.3.4	Positionnement par rapport aux autres modèles	37

2.1 Généralités

2.1.1 Autour du champ moyen

La méthode dite du champ moyen est originaire de la physique. Elle y apparaît dans des contextes très variés : depuis la formule de Clausius-Mossotti donnant la constante d'un milieu diélectrique polarisable, le champ moléculaire de Weiss dans la théorie du

magnétisme, la méthode de Landau en physique statistique, la notion de milieu effectif dans des milieux désordonnés, la méthode de Hartree-Fock en physique atomique ou dans le problème à N corps, jusqu'à l'analyse semi-classique des systèmes quantiques. L'idée est de «substituer à la dynamique ou à la statistique d'un grand nombre de variables équivalentes couplées, celle de l'interaction de l'une d'entre elles avec un champ effectif, produit par les variables restantes et déterminé de manière auto-cohérente» ([75]).

Cependant il n'existe pas une définition précise de ce qu'est le champ moyen : de nombreuses techniques se côtoient, leur point commun étant sans doute qu'il s'agit d'un type de loi des grands nombres fonctionnelle plus ou moins formellement traité. En mathématiques, une des études les plus anciennes concerne les équations de Boltzman pour obtenir l'équation des gaz parfaits.

2.1.2 Le problème étudié, premières notations

N stations de travail d'un département universitaire sont connectées à travers un routeur d'accès à Internet. Chacune des stations considérées télécharge des fichiers d'une machine distante qui se trouve au-delà du routeur. Dans ce cas le routeur risque d'être un goulot d'étranglement. Les stations de travail seront les destinations, et les serveurs distants les sources. Dans tous les modèles, la taille des paquets est une constante universelle.

Problème 2.1. TCP À TRAVERS UN GOULOT D'ÉTRANGLEMENT

DYNAMIQUE DISTRIBUÉE : *Les flots utilisent TCP Reno (dont les caractéristiques principales sont données dans 1.1) et peuvent être PERSISTANTS, c'est à dire que tous les N utilisateurs téléchargent en permanence un fichier dont la taille est infinie ou bien INTERMITTENTS c'est à dire que les sources téléchargent des successions de fichiers séparés par des temps de sommeil. Chaque utilisateur est caractérisé par une variable d'état dépendante du temps $X_i^N(t)$, en particulier cette variable peut être :*

- le débit, que l'on notera alors $x_i^N(t)$,
- la fenêtre (de congestion), que l'on notera $W_i^N(t)$ ⁶.

Par ailleurs chaque utilisateur a un délai de propagation T_i qui lui est associé (indépendant du nombre de ses concurrents). Ce délai correspond au délai optimal pour que le signal fasse le trajet aller-retour entre source et destination. Le temps de retour (RTT) de l'utilisateur i au temps t sera lui noté $R_i^N(t) \geq T_i$. Dans les cas où ce temps est considéré comme une constante on notera r la valeur du RTT.

CONTRÔLE CENTRALISÉ : *De plus le routeur d'accès implémente le contrôle de débit RED avec les notations suivantes :*

- une capacité du lien sortant $C^N = NL$,
- $Q^N(t)$ désigne la taille de la file d'attente toujours supposée FIFO (cad que le temps d'attente d'un paquet qui arrive au temps t à la file est $\frac{Q^N(t)}{NL}$) pour N utilisateurs ; on note $q^N(t) = \frac{Q^N(t)}{N}$, la taille normalisée de la file d'attente,
- En général $K^N(t)$ désignera la probabilité de marquage au temps t ⁷. Dans les cas

⁶Dans toute cette thèse le mot fenêtre désigne par défaut la fenêtre de congestion.

⁷ K^N peut tout aussi bien dépendre du débit ou de toute autre information dont le contrôleur (placé au niveau du routeur) peut disposer.

étudiés on aura $K^N(t) = F^N(Q^N)(t)$, et on notera $F^N(Q^N(t))$ pour le cas particulier où K^N ne dépend que de la valeur instantanée de Q^N .

- $Q_{min}^N < Q_{max}^N$ désignent les seuils de RED, et $B^N \geq Q_{max}^N$ est la capacité physique de la file d'attente. On note de plus $q_{min} := \frac{Q_{min}^N}{N}$, $q_{max} := \frac{Q_{max}^N}{N}$ et $b := \frac{B^N}{N}$. L'hypothèse de passage à l'échelle est que ces quantités ne dépendent pas de N . On pourra par exemple écrire $K^N(t) = F(q^N(t))$.

ÉQUATIONS D'ÉVOLUTION : Spécifier un modèle revient à donner les équations d'évolution et à préciser, si cela est nécessaire, ce qu'on appelle valeur d'une variable au temps t . Ces équations seront de 3 sortes :

- une sur le calcul du taux de pertes,
- une sur l'évolution de la taille de la file d'attente,
- N sur les évolutions des variables caractéristiques des utilisateurs (une par utilisateur).

QUESTION À RÉSOUDRE : Un modèle étant spécifié, on cherche à trouver une limite $q(t)$ pour la taille de la file d'attente normalisée $q^N(t) := \frac{Q^N(t)}{N}$ pour tout t ; une limite en loi $W(t)$ (resp $x(t)$) de $W_i^N(t)$ (resp $x_i^N(t)$) ; une limite $K(t)$ de la fonction de destruction de paquets $K^N(t)$; et enfin des équations d'évolution des limites $q(t)$ et $W(t)$ (resp $x(t)$) et $K(t)$.

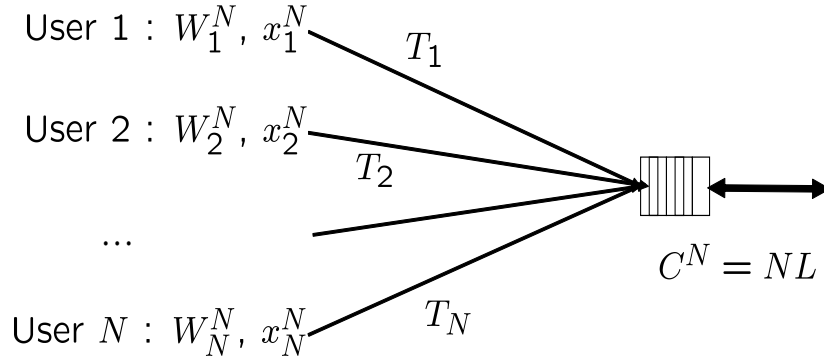


FIG. 2.1 – N utilisateurs partagent un routeur

Selon les modèles le temps est continu ou discret : pour une fonction f donnée on note $f[k]$ la valeur de f au k -ième slot de temps et $f(t)$ la valeur au temps t ; si le pas de temps est Δt dans un cas discret, $f[k] = f(k\Delta t)$.

Pour aider la lecture nous avons représenté à la figure (2.2) le débit crée par plusieurs utilisateurs de TCP.

2.1.3 Les équations fluides

Les équations fluides en moyenne TCP qui ont guidé notre recherche viennent des articles de 2000 et 2001 [113] et [67] de Vishal Misra, Don Towsley, Wei-Bo Gong et Chris Hollot.

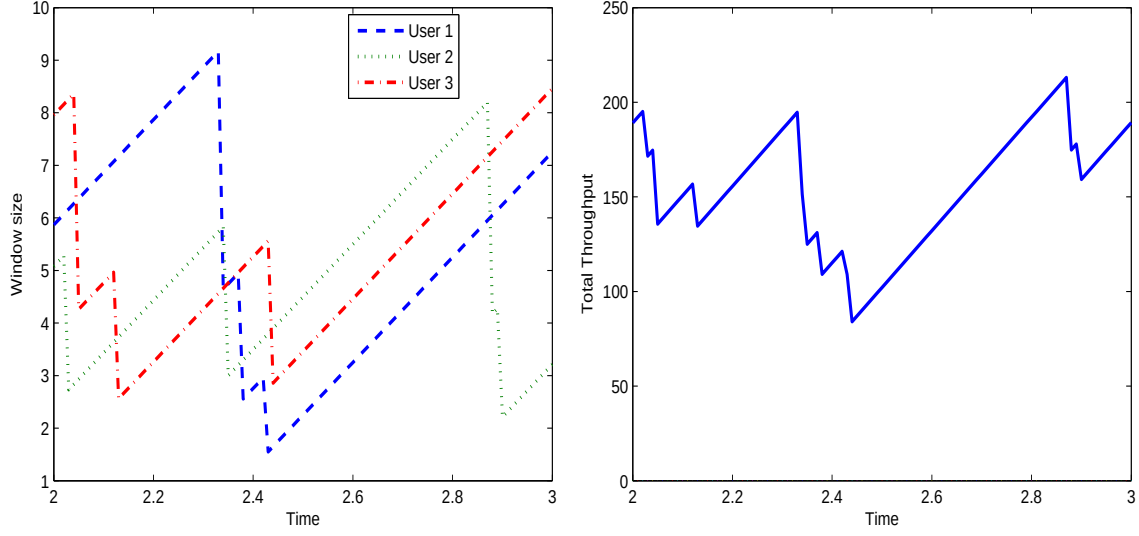


FIG. 2.2 – Le débit à droite est la somme des débits à gauche de trois utilisateurs de TCP. On peut voir le mécanisme de division par deux du débit de chaque utilisateur lorsqu’une perte est constatée [Pas de file d’attente, $T = 100ms$, taux de pertes $k = 1\%$. Cette figure correspond au modèle à N utilisateurs de TCP persistant que nous utilisons dans la thèse].

Les équations

1. $\frac{dw^N}{dt}(t) = \frac{1}{R^N(t)}dt - \frac{w^N(t)}{2}I^N(t),$
2. $\frac{dQ^N}{dt}(t) = N\frac{w^N(t)}{R^N(t)} - C^N$ dans la limite des bords,
3. $I^N(t) = \frac{w^N(t-R^N(t))}{R^N(t-R^N(t))}K^N(t-R^N(t)),$
4. $R^N(t) = T + \frac{Q^N(t)}{C^N},$
5. $K^N(t) = F^N(Q^N(t)).$

Explication des termes $w^N(t)$ est la valeur d’une fenêtre qui pourrait par exemple être la moyenne $\frac{1}{N} \sum_{i=1}^N W_i^N(t)$, et T est la valeur commune de tous les temps de transmission T_i . Le débit moyen supposé est $x^N(t) = \frac{w^N(t)}{R(t)}$, et la file d’attente évolue donc naturellement au gré du débit entrant et du débit sortant constant C^N ; $I^N(t)$ est le taux de marquage ou de destruction de paquets du point de vue de la source. $R^N(t)$ est défini du point de vue de la file d’attente, c’est le RTT d’un paquet qui rentre dans la file d’attente au temps t (et qui n’est pas détruit); ce temps résulte uniquement du temps de transmission et de la taille de la file d’attente.

Discussion Ces équations définissent le modèle proposé. On peut admirer l’effort d’ingénierie pour arriver à ceci sans passer par des étapes mathématiques. De plus ce modèle

fonctionne relativement bien. S'il permet d'étudier le routeur, il ne donne pas cependant d'information intéressante sur les utilisateurs, à part un débit moyen, que peut-être personne ne constate.

Nous verrons de plus que ces équations ne sont pas les meilleures possibles. Mais pour cela il faut une analyse mathématique plus précise. Un de nos objectifs est, en partant d'un modèle intuitif d'utilisateur fluide, de trouver des équations du même genre que celles que nous venons de présenter tout en remplissant le vide mathématique laissé entre la modélisation des individus et le comportement du groupe. Il ne faut pas attendre d'améliorations spectaculaires des études existantes basées sur ce type de modèles dans le cas simple étudié. L'idée dans un premier temps sera surtout de tester le domaine de validité de ce type d'équations qui ne sont confortées que par un nombre limité d'expériences et beaucoup d'intuition. Dans un second temps notre approche pourra être plus efficace pour attaquer les nouveaux problèmes plus complexes.

2.2 Modèles par champ moyen

Nous rappelons ici brièvement les principaux modèles de champ moyen développés pour l'étude de TCP. Leur point commun est d'être relativement récents (2001 pour le plus ancien). Ils se placent tous sur le même cas d'école qui est l'étude d'un routeur partagé par N utilisateurs. On étudie ensuite les propriétés asymptotiques du trafic lorsque le nombre d'utilisateurs devient grand sous des hypothèses appropriées de mise à l'échelle. Les modèles sont donnés par ordre alphabétique du premier auteur.

2.2.1 Champ moyen Baccelli-Hong

Modèle fluide avec pas de temps stochastique, sources intermittentes et sans file d'attente. Le modèle original pour des sources persistantes a été développé et utilisé par François Baccelli et Dohy Hong ([11, 13, 12, 29]). Ce modèle a été ensuite adapté à une limite en champ moyen puis prolongé au cas de sources non persistantes ([9, 17]) par le travail conjoint de François Baccelli, David McDonald, Augustin Chaintreau et Danny De Vleeschauwer. Il s'agit d'un modèle où le temps est découpé par slots aléatoires (qui tendent vers des slots déterministes à la limite). La congestion est modélisée par un taux aléatoire de coupure des fenêtres (taux de synchronisation). La modélisation des sources intermittentes (dénommées sources HTTP) utilise le modèle en boucle fermée de Arzad Kherani and Anurag Kumar [88].

Ce modèle est décrit en détail puis prouvé dans le chapitre 7 comme application des mathématiques développées au cours des chapitres 3, 4, 5 et 6. Voici les spécifications de ce modèle dans sa version avec sources persistantes ; la plupart des notations sont illustrées sur les figures (7.1) et (7.2).

Le modèle à N utilisateurs Le RTT r est fixe et identique pour tous les utilisateurs. Par défaut on utilise TCP Reno ; ainsi le taux de transmission augmente à la vitesse $1/r^2$ (en effet la fenêtre augmente de 1 tous les r et le débit vaut W/r en moyenne).

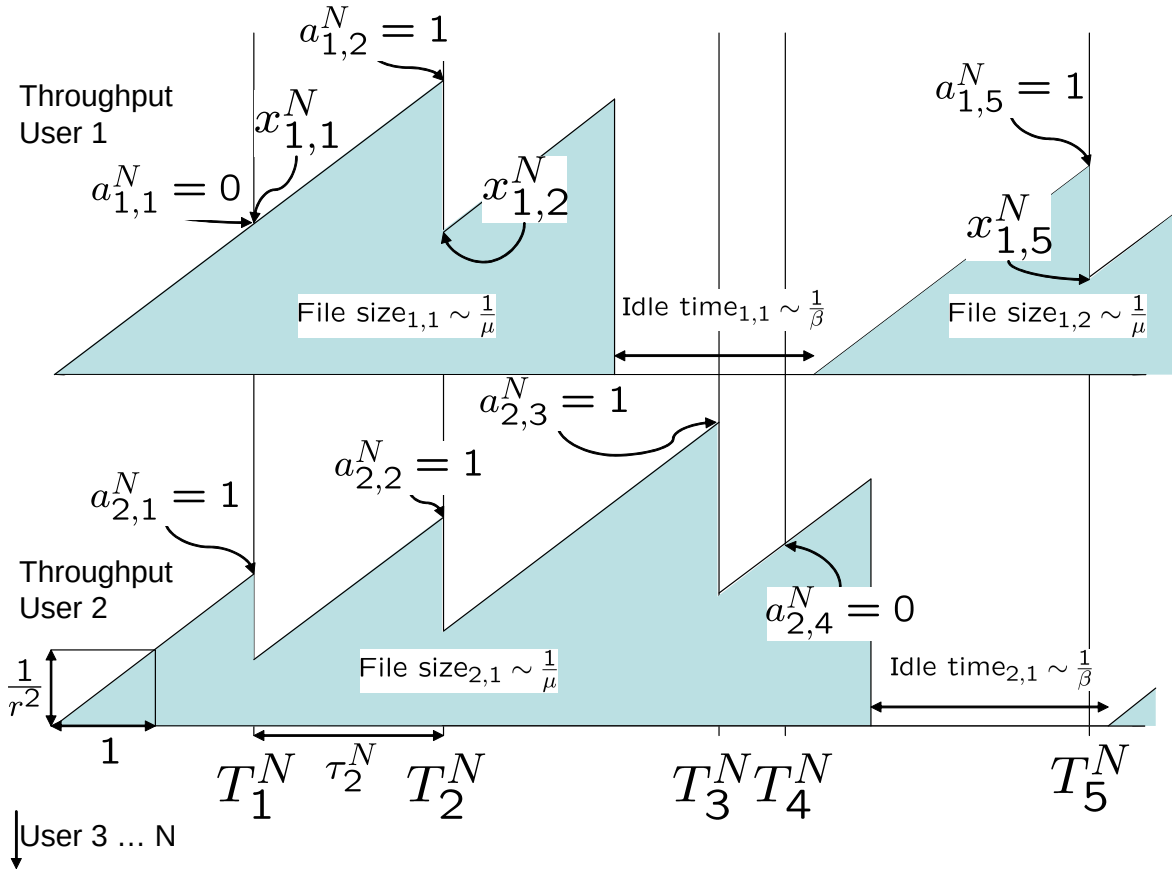


FIG. 2.3 – Illustration du modèle d'utilisateurs de HTTP, les utilisateurs adaptent leur débit en fonction des indications de congestion enchaînant période d'activité et périodes de sommeil.

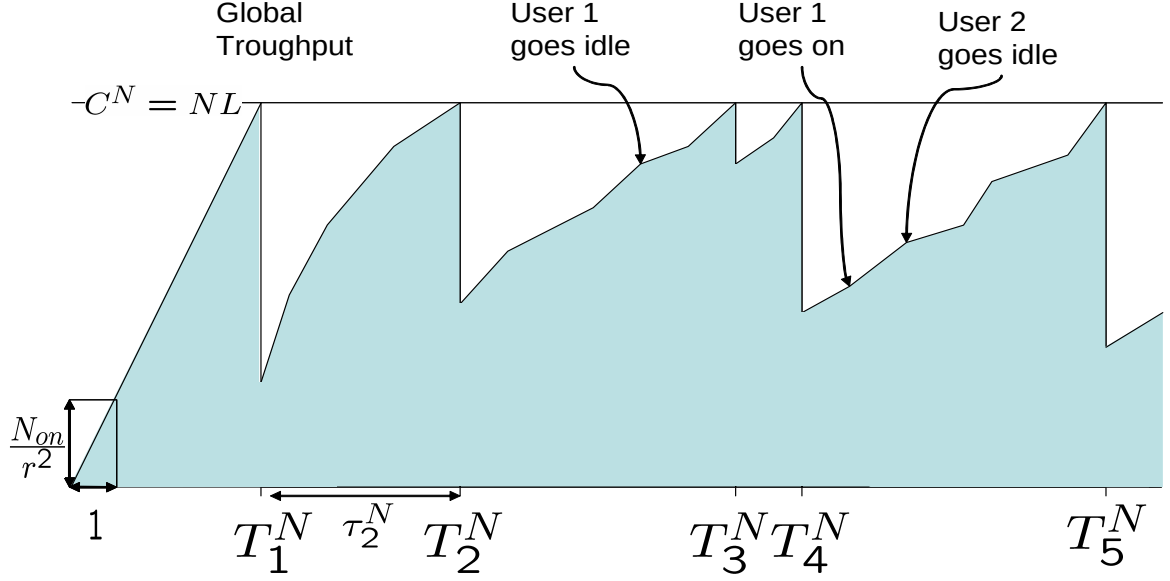


FIG. 2.4 – Illustration du modèle d'utilisateurs de HTTP, les temps d'arrêts T_j^N interviennent lorsque la condition de saturation $\sum_{i=1}^N x_i^N(t) = C^N = NL$ où $x_i^N(t)$ est le débit de la source i au temps t est vérifiée.

Les débits des flots évoluent au gré des pertes imposées par le routeur de la manière suivante. Par définition le j -ième temps de congestion, T_j^N correspond à la j -ième époque de congestion au cours de laquelle des pertes de paquets quasiment simultanées vont se produire. On notera $\tau_j = T_j - T_{j-1}$. Les autres notations utilisées sont les suivantes :

- $C^N = NL$ est la capacité du routeur et L la capacité ramenée au nombre d'utilisateurs,
- $x_{i,j}^N$ et $W_{i,j}^N$ sont le débit et la fenêtre de l'utilisateur i après la j -ième époque de congestion,
- $a_{i,j}^N$ sont des variables aléatoires qui ne dépendent pas de N , valent 0 ou 1 selon que l'utilisateur i perd un paquet au temps de congestion T_j ou pas,
- les $a_{i,j}^N$ étant supposés identiquement distribués et indépendants par rapport au numéro j du saut ; p , le taux de synchronisation vaut $p := \mathbb{E}(a_{i,j}^N)$,
- le taux de croissance des fenêtres vaut $1/r$ (donc celui du débit est $\frac{1}{r^2}$) et le taux de réduction des fenêtres par saut vaut $\frac{1}{2}$,

Les $a_{i,j}^N$ ne sont pas indépendants en i , mais issus d'un tirage de variables indépendantes conditionné par le fait que la somme des $a_{i,j}^N$ qui correspondent à des utilisateurs en cours de transmission vaut au moins 1. En effet une période de congestion est caractérisée par le fait qu'au moins une perte est constatée.

Ainsi les équations pendant la période de transmission sont :

$$x_{i,j+1}^N = (1 - \frac{1}{2}a_{i,j+1}^N)(x_{i,j}^N + \frac{1}{r^2}\tau_{j+1}^N), \quad (2.1)$$

$$\frac{N}{r^2}\tau_{j+1}^N = C^N - \sum_{i=1}^N x_{i,j}^N. \quad (2.2)$$

Explications intuitives Le débit x d'un utilisateur progresse à la vitesse constante pendant les périodes entre les congestions τ_j^N et est divisé par deux lors des congestions, ceci est à l'origine du nom AIMD (Additive Increase Multiplicative Decrease) pour ce type d'évolution. Les temps entre les époques de congestion sont calculées par la deuxième équation qui dit que la congestion intervient quand les débits réduits par la congestion précédente atteignent de nouveau la valeur limite C^N .

Les équations du champ moyen Le résultat principal est que lorsque le nombre d'utilisateurs devient grand, le débit constaté se comporte comme le résultat de l'évolution d'une unique particule $\bar{x}(t)$ dont on peut calculer l'évolution sans connaître la distribution exacte des débits.

Théorème 2.1. HYPOTHÈSES : *Le théorème est prouvé dans la section 7 et les hypothèses techniques sont rassemblées au théorème 6.3. Rappelons que la distance que nous nommons d^* engendre la topologie de la convergence faible sur les mesures (elle est définie formellement au chapitre suivant).*

CONCLUSION 1, CONVERGENCE : *il existe une particule du champ moyen $\bar{x}(t)$ et une distribution de débits M_t telle que pour tout t :*

$$M_t^N \xrightarrow[N \rightarrow \infty]{d^*} M_t \text{ en probabilité,} \quad (2.3)$$

$$\bar{x}(t) = \mathbb{E}(M_t), \quad (2.4)$$

$$\frac{1}{N} \sum_{i=1}^N x_i^N(t) \xrightarrow[N \rightarrow \infty]{} \bar{x}(t) \text{ en probabilité,} \quad (2.5)$$

où M_t^N est la mesure empirique des débits.

CONCLUSION 2, ÉQUATIONS DE LA LIMITE : *de plus $(M_t, \bar{x})$ croissent à la vitesse $1/r^2$. Lorsque $\bar{x}(t^-) = L$, une proportion p des utilisateurs (indépendamment de la valeur de leur fenêtre) divisent leur fenêtre par deux⁸. Ainsi $\bar{x}(t) = (1 - \frac{p}{2})L$. Ceci signifie que l'on peut étudier la particule du débit champ moyen $\bar{x}$ directement sans s'intéresser à l'évolution de la distribution des débits.*

2.2.2 Champ moyen Bain-Kelly-Kelly

Modèle avec paquets en temps continu pour une source persistante. Dans sa thèse de 2004, Alan Bain [20] montre une convergence en champ moyen par la méthode

⁸Le fait que les pertes constatées soient indépendantes de la fenêtre résulte directement d'une hypothèse de modélisation sur le système à N utilisateurs et pas d'un fait mathématique.

classique tension limites unicité, avec le modèle proposé par Frank Kelly de TCP pour étudier principalement *Scalable TCP*, l'évolution de TCP de la thèse de Tom Kelly [83, 84]. Nous en donnons les spécifications dans le cas particulier appelé MulTCP qui est assez proche de TCP-Reno.

Définition 2.1. *MulTCP est un protocole de gestion de la fenêtre de congestion $w(t)$ qui dépend d'un paramètre m . Lorsqu'un ACK arrive, $w(t)$ augmente de $\frac{m}{w(t)}$, lorsqu'un «ACK négatif» arrive (paquet marqué), $w(t)$ est décrémenté de $\frac{w(t)}{2m}$.*

Pour $m = 1$ MulTCP réagit comme TCP-Reno dans le mode d'évitement de congestion que nous étudions. Pour un coefficient m fixé, MulTCP réagit comme un utilisateur utilisant m connexions TCP en même temps avec la possibilité d'équilibrer la fenêtre de congestion entre les différentes connexions dont la valeur commune est $\frac{w(t)}{m}$. On voit en particulier que MulTCP est plus uniforme que m connexions TCP en parallèle dans sa réaction aux pertes (dans le cas de m connexions, un perte peut toucher une petite ou une très grande fenêtre, dans MulTCP, on touche toujours une valeur moyenne). D'autre part l'agressivité du protocole augmente lorsque m augmente.

Ce modèle n'étudie pas le système quand le nombre d'utilisateurs tend vers l'infini mais un unique utilisateur utilisant MulTCP pour un m fixé quand le débit tend vers l'infini. Le but est d'étudier les propriétés asymptotiques des protocoles pour dire lesquels s'adapteront le mieux à l'augmentation des débits disponibles des réseaux du futur. Répétons donc qu'il y a un seul utilisateur qui utilise une version fixée de MulTCP qui équivaut à m connexions TCP en parallèle, c'est bien sur le débit disponible et pas sur m ni sur $N = 1$ que va porter la mise à l'échelle.

Le modèle d'ordre N est composé de deux objets «physiques», un CONTRÔLEUR DE TAILLE DE FENÊTRE et une LIGNE DE RETARD⁹. Se rajoutent un processus de file d'attente $Q^N(t)$ et un processus de marquage $M^N(t)$.

Ligne de retard On ignore la file d'attente pour le calcul des retards et on s'intéresse uniquement à une transmission d'information. La source envoie des paquets dans une ligne de retard de durée T . Chaque paquet entrant est marqué à l'instant de son entrée dans la ligne de retard par un processus $M^N(t)$ à valeurs dans $\{0, 1\}$. Quand le temps résiduel dans la ligne de retard atteint la valeur 0, le paquet et sa marque peuvent être utilisés par le contrôleur en plus d'autres informations qu'il connaît déjà.

Action du contrôleur sur la taille de la fenêtre Quand un paquet non marqué arrive, on augmente la taille $w^N(t)$ de $\frac{m}{w^N(t)}$. Quand un paquet marqué arrive, cette taille est amputée de $w^N(t)/2m$. Contrairement aux autres modèles, le contrôleur a un comportement stochastique dans l'envoi des paquets. En effet notant $v^N(t)$ le nombre de paquets en vol et $w^N(t)$ la fenêtre de congestion, les paquets sont envoyés selon un processus de Poisson doublement stochastique d'intensité $I^N(t) := \gamma(w^N(t) - v^N(t))^+$ où $\gamma > 0$ est le coefficient de réponse. De sorte qu'en notant J un processus de Poisson, on considère le processus d'envoi de paquets $J(\int_0^t I^N(s)ds)$.

⁹«Delay line» en anglais : l'information ressort un temps T après être rentrée

C'est à dire que quand la fenêtre est réduite en dessous du nombre de paquets en vol, on n'envoie plus de paquets. Au contraire quand la fenêtre $w^N(t)$ dépasse le nombre de paquets en vol $v^N(t)$, on envoie un paquet selon un processus de Poisson au bout du temps moyen $\frac{1}{\gamma(w^N(t)-v^N(t))}$.

Processus de Marquage et pseudo-file d'attente On se donne un processus de marquage $M^N(t)$ à valeurs dans $\{0, 1\}$ et un processus de file d'attente $Q^N(t)$ à valeurs dans $\mathbb{R}^+$. $M^N(t)$ n'est pas Markovien tel quel, par contre $M^N(t)$ conditionné par $Q^N(t)$ est un processus de Markov. De plus $Q^N(t)$ désigne le nombre de paquets dont le temps résiduel dans la ligne à retard est dans l'intervalle $[T - \tau, T]$ pour un certain paramètres du modèle $0 \leq \tau < T$.

Résultat On définit alors les débits $x^N(t) := \frac{w^N(t)}{T}$.

Théorème 2.2. ([20] p. 174-175) *Sous quelques hypothèses et après certaines simplifications, $\frac{x^N(t)}{N}$ et $\frac{K^N(t)}{N}$ ont des limites déterministes presque sûres $x(t)$ et $K(t)$ lorsque $N \rightarrow \infty$. L'équation vérifiée par ces limites est :*

$$\frac{1}{T} \int_{t-T}^t x(s) ds = \kappa x(t-T) \left\{ (1 - K(t-T)) \frac{1}{\frac{1}{T} \int_{t-T}^t x(s) ds} + K(t-T) \frac{1}{2T} \int_{t-T}^t x(s) ds \right\}.$$

2.2.3 Champ moyen Deb-Shakkottai-Srikant

Modèle sur les débits en temps discret pour sources persistantes avec file d'attente sous perturbation stochastique stationnaire.

Dans [37] Supratim Deb et Rayadurgam Srikant prolongent le travail de R. Srikant et de Srinivas Shakkottai [94]. Ils analysent en temps discret au niveau faussement-paquets (en réalité c'est une analyse fluide qui ne dit pas son nom) des modèles basés sur le contrôle du débit ou de la file d'attente quand le nombre de sources tend vers l'infini. Ceci sous un bruit stochastique créé par des utilisateurs «non contrôlés» qui représentent les fameuses souris qui sont sensées compléter les éléphants d'Internet (certains auteurs pensent que le trafic sur Internet se compose exclusivement de très petits fichiers, les souris, et de très gros fichiers, les éléphants). Pour les modèles basés sur les files d'attentes ou les débits, ces articles contiennent des intuitions de preuve de la convergence de la file d'attente vers une limite déterministe et la propagation du chaos. Le grand mérite de ce modèle est de vérifier la stabilité des résultats sur les flots longs soumis à un bruit stochastique. C'est un point important, car il laisse espérer que les raisonnements qui ne prennent en compte que les utilisateurs de TCP peuvent avoir un intérêt quand on les met en compétition avec d'autre type d'utilisateurs.

Le système à N éléments se compose de N utilisateurs «contrôlés» et de N utilisateurs «non contrôlés». On commence par définir une queue $Q^N[k]$ associée à un lien de capacité $C^N = NL$. Le pas de temps Δt vaut le temps de servir un paquet, c'est à dire $\Delta t := \frac{1}{NL}$.

$$Q^N[k+1] = (Q^N[k] - 1)^+ + \Delta t \sum_{i=1}^N x_i^N[k] + \sum_{i=1}^N e_i^N \left(\frac{1}{NL} \right)$$

où les e_i^N forment le bruit de fond. Ce sont des processus stationnaires indépendants de moyenne a . De plus on note $e^N = \frac{1}{N} \sum e_i^N$. Le débit est modifié approximativement tous les N slots de temps. Aussi pour TCP, le modèle est que :

$$x_i^N[N(k+1)] = x_i^N[Nk] + \left(\frac{1}{r^2 L} - \frac{1}{2} x_i^N[Nk] K^N[N(k-rL+1)] x_i^N[N(k-rL+1)] \right).$$

K^N est la probabilité de marquage et de délai rL est tel que $rLN\Delta t$ corresponde à la valeur du RTT supposée fixe dans ce modèle. Dans le cas de RED, la probabilité de marquage $K^N[k] = F^N(Q^N[k])$ où $F^N(Q^N) = \min(1, p_{\max}(Q^N - Q_{\min}^N)^+)$ où $Q_{\min}^N = Nq_{\min}$.

Théorème 2.3. ([37] p. 13) *Alors sous certaines hypothèses, le bruit stochastique vérifie la loi forte des grands nombres :*

$$\frac{1}{N} \sum e_i^N \rightarrow a \text{ ps},$$

la file d'attente $q^N[Nk] = \frac{Q^N}{N}[Nk]$ et le débit total normalisé $\bar{x}^N[Nk] = \frac{1}{N} \sum_i x_i^N[Nk]$ ont des limites déterministes presque sûres $q[k]$ et $\bar{x}[k]$. De plus il existe une probabilité $x[k]$ telle que $x_i^N[Nk] \rightarrow^{\mathcal{L}} x[k]$ ¹⁰ (dans ce modèles, la condition initiale est la même pour tous). Enfin ces variables vérifient les équations (avec rL pas de temps par RTT qui correspondent à l'arrivée des ACK) :

$$x[k+1] - x[k] = \frac{1}{r^2 L} - \frac{1}{2} x[k+1-rL] M[k+1-rL]$$

$$M[k] =_{\mathcal{L}} \text{Poisson}(2x[k - \frac{rL}{2}] K[k])$$

$$\bar{x}[k+1] - \bar{x}[k] = \frac{1}{r^2 L} - \frac{1}{2} \mathbb{E}(x[k] x[k+1-rL] K[k+1-rL])$$

$$q[k+1] = \left\{ q[k] + \bar{x}[k - \frac{rL}{2}] + a - L \right\}^+$$

$$K[k] = F(q[k]) \text{ avec } F(q) = \min(1, p_{\max}(q - q_{\min})^+).$$

2.2.4 Champ moyen Makowski-Tinnakornsrisuphap

Modèle niveau paquets simplifié en temps discret pour sources persistantes avec file d'attente.

Armand Makowski et Peerapol Tinnakornsrisuphap dans [119, 118, 120], font une analyse parfaitement rigoureuse de la convergence en champ moyen de TCP associé à un théorème de la limite centrale. Le modèle utilisé est un modèle en temps discret au niveau des paquets, il s'agit sans doute du modèle le plus simple que l'on pouvait imaginer, ce

¹⁰le symbole $\mathcal{L}$ désigne la loi d'une variable aléatoire, ici $\rightarrow^{\mathcal{L}}$ signifie tendre en loi. Nous utiliserons aussi $=^{\mathcal{L}}$ pour parler de l'égalité en loi entre variables aléatoires et lois.

qui est spectaculaire est qu'il réussit à capturer les aspects principaux de la dynamique de TCP liée aux retards d'un RTT dans les décisions.

Voici brièvement quelques points du modèle. Le pas de temps $\Delta t = r$, un RTT. $W_i^N[k]$ est un entier. La file d'attente suit l'équation :

$$Q^N[k+1] = \left(Q^N[k] - NL + \sum_{i=1}^N W_i^N[k] \right)^+,$$

et

$$W_i^N[k+1] = W_{max} \wedge \begin{cases} W_i^N[k] + 1 & \text{si } M_i^N[k] = 1 \\ \lceil \frac{W_i^N[k]}{2} \rceil & \text{sinon} \end{cases},$$

où

$$M_i^N[k+1] = \prod_{j=1}^{W_i^N[k]} M_{i,j}[k+1].$$

Les $M_{i,j}[k]$ sont des variables aléatoires indépendantes (et indépendantes des conditions initiales $Q^N[0]$ et des $(W_i^N[0])_i$) de *non*-marquage des paquets à valeurs dans $\{0, 1\}$. $M_{i,j}[k+1]$ vaut 0 avec probabilité $K^N[k] := F^N(Q^N[k])$ où F^N est la fonction de marquage.

Théorème 2.4. ([120] p. 8) *Sous certaines hypothèses techniques, le modèle précédent vérifie la loi faible des grands nombres. Il existe une fonction $q[k]$ déterministe et une probabilité $W[k]$ sur $[0, W_{max}]$ tels que :*

$$\frac{Q^N[k]}{N} \rightarrow q[k] \text{ en probabilité}$$

et

$$W_i^N[k] \rightarrow^{\mathcal{L}} W[k].$$

De plus les W_i^N sont asymptotiquement indépendants et on a les équations simples :

$$q[k+1] = (q[k] - L + \mathbb{E}(W[k]))^+,$$

$$W[k+1] = W_{max} \wedge \begin{cases} W[k] + 1 & \text{si } M[k] = 1 \\ \lceil \frac{W[k]}{2} \rceil & \text{sinon} \end{cases},$$

où $M[k+1] = \chi_{U[k+1] \leq (1-K[k])W[k]}$ avec des variables $U[k]$ iid (indépendantes identiquement distribuées) uniformes sur $[0, 1]$ avec $K[k] = F(q[k])$

De plus on a un théorème fondamental de la statistique (théorème central limite) qu'il est bon de signaler car c'est le seul modèle où ceci ait été démontré. Notant

$$L_0^N[k] := \frac{Q^N[k]}{N} - q[k]$$

et

$$L_{avg}^N := \frac{1}{N} \sum_{i=1}^N W_i^N[k] - \mathbb{E}W[k].$$

Théorème 2.5. [120] *Sous certaines hypothèses techniques un peu plus fortes que précédemment il existe des variables aléatoires L_0 et L_{avg} telles que :*

$$\sqrt{N}L_0^N[k] \xrightarrow[N \rightarrow \infty]{\mathcal{L}} L_0[k] \text{ et } \sqrt{N}L_{avg}^N[k] \xrightarrow[N \rightarrow \infty]{\mathcal{L}} L_{avg}[k].$$

De plus $L_0[k]$ et $L_{avg}[k]$ vérifient une équation de récurrence assez simple.

2.3 Les modèles étudiés dans la thèse

Nous mettons en gras les principales hypothèses de modélisation.

2.3.1 Le modèle TCP-RED avec sources persistantes de la thèse

Nous utilisons une description fluide de la file d'attente et nous faisons l'approximation de continuité du taux de pertes de chaque connexion qui nous permet de construire un modèle d'évolution d'une file d'attente fluide couplée à l'histogramme des fenêtres. Ce modèle introduit dans [18] généralise celui de [67] qui donne la distribution stationnaire des fenêtres et comporte les éléments suivants :

Un processus de fenêtre qui contrôle les débits Ce modèle étudie la valeur de la fenêtre de congestion $W_i^N(t)$ des utilisateurs. Ces fenêtres sont des processus stochastiques qui évoluent selon l'équation fluide intuitive :

$$dW_i^N(t) = \frac{1}{R_i^N(t)} dt - \frac{W_i^N(t-)}{2} dJ_i^N(t), \quad (2.6)$$

qui représente une évolution de la fenêtre sous TCP-Reno en évitement de congestion. J_i^N est un processus de Poisson doublement stochastique que nous décrivons ensuite.

La fenêtre est reliée au débit en utilisant la valeur du RTT et **une formule type Little**¹¹ :

$$x_i^N(t) = \frac{W_i^N(t)}{R_i^N(t)}.$$

Une file d'attente qui détruit des paquets La file d'attente détruit des paquets soit par le mécanisme «drop tail», soit par RED. Nous conservons les notations du chapitre pour RED : F^N est la fonction de pertes définie dans le problème 2.1. La taille de la file d'attente est donc si $0 < Q^N(t) < Q_{max}^N$:

$$\frac{dQ^N}{dt}(t) = \left(\sum_{i=1}^N \frac{W_i^N(t)}{R_i^N(t)} \right) (1 - K^N(t)) - NL, \quad (2.7)$$

¹¹Pour respecter autant que possible les notations initiales des différents articles et avoir une présentation unifiée, nous avons noté $x(t)$ un débit, $W(t)$ une fenêtre. Plus généralement nous noterons $X(t)$ un vecteur de variables d'états d'un utilisateur.

où

$$K^N = F^N(Q^N(t)). \quad (2.8)$$

Ceci représente une simplification par rapport à RED puisque le taux de pertes est pris sur la valeur instantanée et non sur la moyenne temporelle. Ceci n'est pas essentiel dans les preuves ni dans les simulations mais simplifie un peu la présentation.

Un retard d'action Toutes les valeurs au temps t sont prises du point de vue de la file d'attente. Si nous notons τ le temps fixe entre l'utilisateur et la file d'attente, ceci signifie que si un paquet part au temps t_0 , passe à la file au temps $t_1 = t_0 + \tau$ et revient sous forme d'ACK au temps t_2 , on note :

- $W_i^N(t_1)$ la fenêtre au temps t_0 pour cet utilisateur,
- $R_i^N(t_2 + \tau) = t_2 - t_0$,
- $Q^N(t_1)$ la taille de la file au temps t_1 .
- $K^N(t_1)$ le taux de paquets entrants détruits au temps t_1 .

De sorte que :

$$R_i^N(t) = T_i + \frac{Q^N(t - R_i^N(t))}{L}. \quad (2.9)$$

La raison de prendre cette valeur du RTT est que le débit que nous calculons dépend du débit d'ACK à la source, c'est à dire du temps de parcours il y a un RTT. De plus cette définition est exactement le retard dont nous avons besoin pour le calcul des pertes ressenties.

Une relation entre les pertes K^N imposées et les sauts J_i^N ressentis $J_i^N(t)$ est un processus de Poisson doublement stochastique d'intensité au temps t :

$$I_i^N(t) := K^N(t - R_i^N(t)) \frac{W_i^N(t - R_i^N(t))}{R_i^N(t - R_i^N(t))}. \quad (2.10)$$

Des effets de bords : problème de la gigue Dans notre modèle, la dynamique des fenêtres reste invariable. Ce n'est pas le cas de celle de la file d'attente et du processus de pertes $K^N(t)$ qui sont soumises à des effets de bords.

Quand $Q^N(t) = 0$, alors on ne tient compte de la dérivée que si elle est positive, et $K^N(t) = 0$. Ceci ne pose pas de grand problèmes, puisque la dynamique peut toujours s'écrire avec une unique équation avec une indicatrice.

Un cas plus délicat se présente lorsque $Q^N(t) = Q_{max}^N$. Contrairement à ce que laisserait penser l'intuition du modèle paquet, *le problème avec une fonction de pertes qui fait un saut n'a pas de solution avec les équations que nous avons proposées*. Imaginons par exemple que le débit entrant dans la file d'attente est constant et vaut $2C^N$, et que la file s'est entièrement remplie. Alors l'intuition dit qu'environ la moitié des paquets sont détruits. Considérons de plus la fonction F^N de pertes fait un saut de $p_{max} = 5\%$ à 100% à la valeur Q_{max} de la file d'attente. Dans notre modèle fluide, nous voyons que le problème est mal défini, et n'admet aucune solution : Q^N étant continue, car définie en intégrant des quantités bornées, les deux conditions « $K(t) = 1 \Rightarrow Q^N(t)$ strictement décroissant» et « $K(t) = p_{max} \Rightarrow Q^N(t)$ strictement croissant» sont incompatibles.

Que se passe-t-il ? En fait dans un modèle paquet, on aurait un phénomène oscillatoire où les paquets arrivent une fois sur deux environ avec assez d'espace dans la file d'attente pour être candidat à l'entrée¹². Dans le modèle fluide, la file d'attente est censée être toujours pleine, et le taux de pertes résultant ne doit pas permettre à la file d'attente de dépasser la valeur Q_{max} .

Que fait-on ? Pour circonvenir à cette difficulté, nous avons décidé que lorsque $Q^N(t)$ touche la valeur Q_{max}^N , la dynamique devait changer pour tolérer les valeurs de $K(t)$ entre p_{max} et 1. Nous nommons ce deuxième état du système la GIGUE en référence au phénomène oscillatoire qui lui donne naissance dans un modèle paquet. La dynamique de la gigue maintient la file d'attente à sa valeur maximale tant que le débit entrant ne s'est pas réduit suffisamment pour pouvoir appliquer la dynamique originale.

Définition de l'état de gigue Quand $Q^N(t) = Q_{max}^N$, nous disons la file d'attente fait la gigue à Q_{max} et l'équation

$$\frac{dQ^N}{dt}(t) = 0, \quad (2.11)$$

permet de définir la valeur de $K^N(t)$:

$$K^N(t) := \frac{\sum_{i=1}^N \frac{W_i^N(t)}{R_i^N(t)} - NL}{\sum_{i=1}^N \frac{W_i^N(t)}{R_i^N(t)}}. \quad (2.12)$$

Ainsi lorsque la file d'attente rentre dans la gigue, le débit entrant à priori est supérieur strictement à la capacité de la file, ce qui signifie que $K^N(t)$ fait un saut vers le haut. La sortie de l'état de gigue s'effectue dès l'instant où $K^N(t) = p_{max}$ (la valeur maximale du taux de pertes imposé par le mécanisme RED). Cet instant correspond aussi au moment où le débit entrant vaut $\frac{L}{1-p_{max}}$, c'est à dire que la dynamique originale (c'est à dire sans la gigue) donne de nouveau des solutions au système. À partir de cet instant de sortie de gigue, si le débit continue à décroître, la file d'attente va se vider.

2.3.2 Adaptation HTTP-RED avec sources intermittentes

Environnement N utilisateurs utilisent HTTP. Nous modélisons HTTP comme une suite de téléchargements de fichiers en utilisant TCP séparés par des temps de réflexion comme dans [88]. La modélisation fine de TCP et l'évolution de la taille de la file d'attente se fait comme ci-dessus.

Une des nouveautés est l'arrivée de nouveaux utilisateurs, ils commencent leur téléchargement à la fenêtre 0. D'autre part chaque utilisateur peut être dans l'état actif où il se comporte comme précédemment en gardant de plus en mémoire la quantité de données qui a été reçue (c'est à dire dont la destination a bien accusé réception). Ou bien dans l'état d'attente où il ne se passe presque rien : on garde juste en mémoire depuis combien de temps on est dans cet état. Le passage d'un état à l'autre se fait ainsi : on se donne

¹²un peu plus d'une fois sur deux en fait puisque des paquets candidats admissibles sont parfois détruits par le mécanisme RED

$\mathfrak{S}(s)$, la probabilité qu'une source à d'être en attente pendant la durée s . $\mathfrak{F}(f)$ est la probabilité de devoir envoyer fichier de taille f (l'unité de mesure est le paquet). Chaque utilisateur va enchaîner des téléchargements de fichiers de tailles f iid et des temps d'attente de durée s iid. On se donne un état initial où pour chacun des N utilisateur on connaît :

- son état $E_i^N(t)$ (actif ou en attente),
- la taille de fichier qu'il a envoyé et sa fenêtre ($W_i^N(t), F_i^N(t)$) ou bien depuis combien de temps il attend ($S_i^N(t)$).

Les sauts de changement d'état quand le fichier est fini, se produisent donc avec la probabilité :

$$\mathbb{P}(F_n^N \text{ est la taille finale sachant que le fichier fait au moins } F_n^N).$$

Les sauts inverses pour revenir à l'état actif se font avec la même probabilité sur le temps d'attente et la distribution des temps d'attente.

Équations d'une particule L'espace total a 4 états $(W, F, S, \{0, 1\})$, la taille de la fenêtre, le nombre de paquets transmis, le temps d'inactivité et l'état de transmission ou d'inactivité.

Comme d'habitude, on modélise TCP. Ici on rajoute la donnée de la taille du fichier. Donc dans l'espace (W, F) :

$$dW_i^N(t) = \frac{1}{R_i^N(t)}dt - \frac{W_i^N(t)}{2R_i^N(t) - R_i^N(t)}dJ_i^N(t) \quad (2.13)$$

$$dF_i^N(t) = \frac{W_i^N(t)}{R_i^N(t)}(1 - K^N(t))dt \quad (2.14)$$

où J_i^N est un processus de Poisson d'intensité $I_i^N(t)$ qui a même valeur que précédemment.

De plus on a un changement d'état puisque l'utilisateur devient inactif après avoir fini son téléchargement. Ceci se produit après la réception du paquet numéro F_i^N de la transmission avec la probabilité

$$P_F(F_i^N) := \mathbb{P}(F = F_i^N | F \geq F_i^N) = \frac{\mathfrak{F}(F_i^N)}{\int_{F_i^N}^{\infty} \mathfrak{F}(u)du}.$$

Pendant la période d'inactivité, le temps I_i^N progresse (par l'équation $dI_i^N(t) = dt$), et on reprend la transmission avec la probabilité :

$$P_S(S_i^N) = \mathbb{P}(S = S_i^N | S \geq S_i^N) = \frac{\mathfrak{S}(S_i^N)}{\int_{S_i^N}^{\infty} \mathfrak{S}(u)du}.$$

2.3.3 Type de résultats que nous démontrons

Nous montrons des résultats suivants en termes de convergence.

Théorème 2.6. *Sous certaines hypothèses techniques détaillées au chapitre 8,*

Conclusion 1 : Il existe une suite $(\mathcal{X}_1, \dots, \mathcal{X}_i, \dots)$ telle que : $(X_1^N, \dots, X_N^N)$ converge vers $(\mathcal{X}_1, \dots, \mathcal{X}_i, \dots)$ dans un sens trajectoirel qui implique en particulier une convergence en probabilité des mesures empiriques lorsque $N \rightarrow \infty$.

Conclusion 2 : les $(\mathcal{X}_i)$ vérifient une loi forte des grands nombres pour tout t et $g \in \mathcal{C}_b$, il existe une particule $\overline{\mathcal{X}}$ telle que :

$$\frac{1}{N} \sum_{i=1}^N g(\mathcal{X}_i(t)) \xrightarrow[N \rightarrow \infty]{} \mathbb{E}(g(\overline{\mathcal{X}}(t))) \text{ ps.}$$

Conclusion 3 : $X_i^N \rightarrow \mathcal{X}_i$ en probabilité uniformément sur tout compact du temps.

Conclusion 4 : Les $(\mathcal{X}_1, \dots, \mathcal{X}_i, \dots)$ et $\overline{\mathcal{X}}$ vérifient certaines équations que nous pourrions simplifier et simuler (cf : chapitres 10 et 12).

Les résultats de convergence ci-dessus donnent plusieurs types d'informations exploitables :

- $\overline{\mathcal{X}}$ donne la dynamique de la ressource partagée. La distribution à tout temps prend une valeur déterministe et est solution d'équations aux dérivées partielles qui se prêtent à la simulation.
- $\mathcal{X}_i$ donne la dynamique que l'on peut espérer pour un utilisateur donné et ceci est vrai uniformément en probabilité, c'est à dire qu'on a de bonnes chances de l'observer sur des réalisations.

2.3.4 Positionnement par rapport aux autres modèles

Le travail effectué dans cette thèse se différencie des précédents par l'utilisation d'un modèle à *dynamique en temps continu* avec file d'attente. De plus nous prenons en compte des *mécanismes de rétroaction (feedback) avec retard* et traitons des *effets de bords avec des changements de dynamique*. Enfin nous montrons une *convergence trajectoirelle* sur les utilisateurs, et une convergence des distributions *en probabilité*.

L'étude ces modèles fera l'objet de la deuxième partie. Nous affinerons le travail de la première partie mathématique pour tenir compte des délais dans le cas particulier de l'étude de TCP (le mécanisme de délai de l'équation (2.9) est si particulier qu'il n'était pas possible de l'intégrer au cas général). Enfin la troisième partie présente les prolongements à la *modélisation de sources intermittentes*, ce qui constitue une avancée importante par rapport aux autres modèles par champ moyen qui s'arrêtent le plus souvent au cas de sources persistantes.

Première partie

Éléments mathématiques du champ
moyen

Chapitre 3

Champ moyen et tirages d'une urne

Le but de ce chapitre est de présenter le concept du champ moyen sur un exemple simple. Ceci a deux fins essentielles :

- poser des notations dont nous allons nous inspirer dans toute la suite
- permettre de présenter les contributions liées à cette thèse.

L'exemple fondamental sera le classique tirage avec ou sans remise de boules d'une urne. Ceci nous permettra d'expliquer ce que nous entendons par champ moyen, les problèmes et théorèmes que nous souhaitons résoudre et d'introduire les techniques essentielles que nous allons utiliser.

Sommaire

3.1	Rappels mathématiques et motivations générales	42
3.1.1	Systèmes de particules	42
3.1.2	Convergence faible sur les espaces Polonais	42
3.1.3	Métriques de la convergence faible	43
3.1.4	Motivations mathématiques	45
3.2	Spécification du problème particulier étudié dans ce chapitre	46
3.2.1	Notations et définitions	46
3.2.2	Résultats	47
3.2.3	Échangeabilité	47
3.2.4	Discussion	47
3.3	Champ moyen par la méthode classique : tension, limites, unicité	48
3.3.1	Explications générales sur la méthode classique	48
3.3.2	Preuve du théorème par la méthode classique du champ moyen	49
3.3.3	Une approche alternative	49
3.4	Champ moyen par la méthode de couplage trajectoriel à un système auxiliaire	50
3.4.1	Explications générales	50
3.4.2	Couplage entre les trajectoires et convergence	51
3.4.3	Définition du système auxiliaire : nouvelle version	52

3.4.4	Énoncés et démonstrations des théorèmes	52
3.4.5	Définition du système auxiliaire : version originale	53
3.4.6	Comparaison des méthodes	54
3.5	Apports mathématiques de cette thèse	54

3.1 Rappels mathématiques et motivations générales

Nous rappelons brièvement quelques résultats mathématiques pour fixer le vocabulaire et les notations. Le but ici n'est aucunement de refaire un cours sur la convergence des mesures, aussi les résultats sont cités sans démonstrations (se reporter à [23]). Tout au long de la thèse E désignera l'espace d'états que l'on considère.

3.1.1 Systèmes de particules

Définition 3.1. *Un SYSTÈME DE PARTICULES $X_1, \dots, X_N$ DE E est une suite finie de variables aléatoires à valeurs dans E .*

Définition 3.2. *Un SYSTÈME INFINI DE PARTICULES, X , de E est une suite infinie de systèmes de particules sur E de la forme suivante :*

$$X = (X^N)_{N \in \mathbb{N}} = (X_i^N)_{(i,N) \in \mathbb{N}^2, i \leq N}.$$

Les X_i^N sont donc des variables aléatoires à valeurs dans E ; Dans toute la thèse nous retiendrons la notation d'exposant pour le nombre de particules et d'indice pour le numéro de la particule.

3.1.2 Convergence faible sur les espaces Polonais

Espaces Polonais Dans cette thèse tous les espaces que nous étudierons seront Polonais.

Définition 3.3. *Un espace métrique (E, d) est dit ESPACE POLONAIS si :*

- E est complet pour la topologie induite par d ,
- (E, d) est séparable, ce qui signifie (dans le cas métrique) qu'il existe une famille dénombrable de points de E dense dans E pour la topologie induite par d .

Définition 3.4. *Un espace topologique $(E, \mathcal{O})$ étant donné, la TRIBU DES BORÉLIENS $\text{Bor}(\mathcal{O})$ de E pour la topologie $\mathcal{O}$ est la tribu engendrée par l'ensemble des ouverts $\mathcal{O}$.*

En particulier un espace Polonais est canoniquement mesurable pour la topologie induite par sa métrique. Ceci permet de parler de $\mathcal{P}(E)$, l'ensemble des mesures de probabilités sur E .

D'autre part les Boréliens d'un espace métrique contiennent toujours les singletons. Ainsi $\mathcal{P}(E)$ contient les mesures ponctuelles.

mesure empirique

Définition 3.5. La MESURE EMPIRIQUE associée à N variables aléatoires $X_1, \dots, X_N$ à valeurs dans un espace Polonais E est la variable aléatoire à valeurs dans $\mathcal{P}(E)$:

$$M^N = \frac{1}{N} \sum_{i=1}^N \delta_{X_i},$$

où $(\Omega, \mathcal{A}, \mathbb{P})$ est un espace de probabilité supportant toutes les variables aléatoires.

$M^N(\omega)$ est une mesure ponctuelle, et M^N est un élément de $(\mathcal{P}(E))^\Omega$. La convergence de mesures se fera pour la convergence faible dont nous rappelons la définition :

Définition 3.6. Une suite de probabilités M^N sur $(E, \mathcal{B})$ CONVERGE FAIBLEMENT vers une probabilité μ si et seulement si :

$$\forall f \in \mathcal{C}_b(E, \mathbb{R}), \langle f, M^N \rangle \xrightarrow{N \rightarrow \infty} \langle f, \mu \rangle$$

où l'opérateur de dualité $\langle \cdot, \cdot \rangle$ est défini par :

$$\langle f, \nu \rangle := \int_E f d\nu.$$

On note alors $M^N \rightarrow^* \mu$.

3.1.3 Métriques de la convergence faible

Une métrique complète Dans certains cas la topologie de la convergence faible est métrique, c'est à dire qu'il existe une distance d^* sur $\mathcal{P}(E)$ pour laquelle la topologie induite est égale à la topologie faible. C'est ce qui se passe dans les espaces Polonais.

Proposition 3.1. (la définition de la distance est dans [42] p. 306, la propriété p. 310, entre temps il est expliqué pourquoi cette définition plus facile équivaut à la définition classique de la distance de Prokhorov). Soit E , un espace Polonais. Alors $\mathcal{P}(E)$ est Polonais pour la distance d_{Lip}^* :

$$d_{Lip}^*(\nu, \mu) = \sup_{f \in Lip_{1,b}(E, \mathbb{R})} |\langle f, \nu \rangle - \langle f, \mu \rangle|$$

où $Lip_{1,b}$ désigne les fonctions 1-Lipschitz bornées.

Une métrique facile à manipuler (les résultats qui suivent sur la séparabilité large ou stricte et la détermination de convergence peuvent se trouver dans [44] p111-117)

E est séparable, c'est à dire qu'il existe une partie dénombrable $F \subset \mathcal{C}_b(E, \mathbb{R})$ telle que

$$(\forall f \in F, \langle f, \mu \rangle = \langle f, \nu \rangle) \Rightarrow \mu = \nu.$$

Une autre notion intéressante pour nous est :

Définition 3.7. Une partie $F \subset \mathcal{C}_b(E, \mathbb{R})$ DÉTERMINE LA CONVERGENCE si

$$(\forall f \in F, \langle f, \mu^N \rangle \rightarrow \langle f, \mu \rangle) \Rightarrow \mu^N \rightarrow^* \mu.$$

Proposition 3.2. Nous avons les faits suivants :

- F détermine la convergence implique F séparant,
- Si E est polonais et compact, déterminer la convergence équivaut à être séparant.

Théorème 3.1. STONE-WEIERSTRASS FORTEMENT SÉPARABLE Il existe une famille dénombrable de fonctions bornées positives continues qui déterminent la convergence sur $\mathbb{R}^k$.

Une famille dénombrable de fonctions $(f^k)_{k \in \mathbb{N}}$ qui détermine la convergence étant donnée, nous avons la définition :

Définition 3.8.

$$\|\mu\|_w = \sum_{k=1}^{\infty} (1 \wedge \langle f^k, \mu \rangle) \frac{1}{2^k}$$

et la distance $d^*(\mu, \nu) = \|\mu - \nu\|_w$ (en rajoutant une valeur absolue à l'intérieur).

Proposition 3.3. Soit μ^N une suite de variables aléatoires à valeurs dans $\mathcal{P}(E)$. E admet une famille dénombrable qui détermine la convergence avec une distance d^* associée. Si $\forall g \in \mathcal{C}_b(E, \mathbb{R}), \mathbb{E}|\langle g, \mu^N \rangle - \langle g, \mu \rangle| \rightarrow 0$, alors $d^*(\mu^N, \mu) \rightarrow 0$ en probabilité.

Cette proposition remarque essentiellement que la métrique de Tychonov d'un produit dénombrable d'espaces rend équivalentes la convergence d'une suite et la convergence de chacune des projections de cette suite. Comme nous pensons qu'il est important de comprendre que ce résultat est à peu près évident nous reprenons la démonstration de [44] p. 107. Par ailleurs nous reverrons plus loin l'idée que la convergence fini-dimensionnelle implique sous certaines conditions de compacité une convergence uniforme.

Preuve : Partons de $\forall g \in \mathcal{C}_b(E, \mathbb{R}), \mathbb{E}|\langle g, \mu^N \rangle - \langle g, \mu \rangle| \rightarrow 0$, en particulier pour la famille f^k qui détermine la convergence.

Fixons un $\epsilon > 0$, comme $\sum \frac{1}{2^k}$ est une série sommable, il existe k_ϵ tel que $\sum_{k_\epsilon}^{\infty} \frac{1}{2^k} < \frac{\epsilon}{2}$, en particulier $\sum_{k_\epsilon}^{\infty} \frac{1}{2^k} |1 \wedge \langle f^k, \mu^N \rangle| < \frac{\epsilon}{2}$. De plus la convergence en probabilité sur un espace est équivalente avec celle sur les composantes pour un produit fini. Ainsi pour tout $\eta > 0$, il existe $N_0 > 0$ tel que pour tout $N \geq N_0$:

$$\mathbb{P} \left(\sum_{k=1}^{k_\epsilon-1} \frac{1}{2^k} |\langle f^k, \mu^N \rangle - \langle f^k, \mu \rangle| > \eta \right) \leq \frac{\epsilon}{2},$$

ce qui achève la preuve. ■

Une manière plus moderne et rapide de faire cette preuve serait d'appliquer le théorème de convergence dominée.

3.1.4 Motivations mathématiques

Un système de particules (X_i^N) à valeur dans un espace Polonais E étant donné, nous souhaitons définir une notion de convergence et de limite pour ce système. Le but étant d'étudier le comportement du groupe autant que l'interaction d'un individu avec ce groupe lorsque la taille du groupe devient trop grande pour une attaque combinatoire.

Une première étape est de regarder la suite de mesures empiriques :

$$M^N = \frac{1}{N} \sum_{i=1}^N \delta_{X_i^N}. \quad (3.1)$$

La convergence de cette suite de mesures est une exigence minimale de la convergence que nous souhaiterions pour ce système. Cependant, ceci ne permet pas d'étude individuelle des particules.

Convergence des mesures empiriques Une méthode générique que nous pourrions employer si nous ne souhaitions pas nous intéresser aux trajectoires mais simplement au comportement du groupe serait la suivante :

Notant que (M^N) est une suite de processus prenant leurs valeurs dans un espace de mesures de probabilités, la convergence en loi de M^N amène à étudier une convergence faible dans l'espace $\mathcal{P}(\mathcal{P}(E))$.

Pour ce faire, on pourra s'aider du théorème de Prokhorov qui est donné de manière moderne dans [42] pp. 311–315.

Définition 3.9. Une famille de mesures $(\mu^N)_N$ d'un espace Polonais E est dite **TENDUE** si :

$$\forall \epsilon > 0, \exists K_\epsilon, \text{ compact}, \forall N \mu^N(E \setminus K_\epsilon) < \epsilon.$$

Théorème 3.2. PROKHOROV Une famille de mesures d'un espace Polonais est relativement séquentiellement compacte si et seulement si elle est tendue.

Notons de plus le corollaire suivant du théorème de Prokhorov :

Corollaire 3.1. Soit E , Polonais. Si E est compact, alors $\mathcal{P}(E)$ est compact.

Ainsi $\mathcal{P}(E)$ hérite de toutes les propriétés de E puisqu'on sait déjà que le caractère Polonais se transmet. De plus remarquons ici qu'une des raisons qui fait introduire l'hypothèse de séparabilité des espaces Polonais est que cette hypothèse permet de vérifier que toute loi μ est tendue ([42] p. 311).

Ainsi on peut extraire une limite. Les problèmes d'études et d'unicité de la limite ne sont pas génériques et doivent être étudiés sur les cas particuliers. Cependant la convergence en loi de la mesure empirique ne nous suffit pas, il faut réussir à obtenir un résultat au moins en probabilité.

3.2 Spécification du problème particulier étudié dans ce chapitre

Nous prendrons le cas le plus simple possible, c'est à dire $E = \{0, 1\}$ pour la distance induite par $(\mathbb{R}, ||)$. Ainsi E , $\mathcal{P}(E)$ et $\mathcal{P}(\mathcal{P}(E))$ sont tous trois Polonais.

3.2.1 Notations et définitions

Une urne, contient un certain nombre de boules soit rouges de valeur unitaire, soit bleues de valeur nulle. Mathématiquement on a la définition suivante.

Définition 3.10. Une URNE U est un couple $(U_1, U_2) \in \mathbb{N}^2$ avec $U_1 \geq U_2 \geq 1$. U_1 est le nombre de boules de l'urne, et U_2 est le nombre de boules rouges.

Une urne étant donnée, on définit un tirage et deux notions d'itération de tirages, avec ou sans remise.

Définition 3.11. Un TIRAGE d'une boule d'une urne U est une variable aléatoire X sur l'espace de probabilité $(\Omega, \mathbb{P})$ et à valeurs dans $\{0, 1\}$. Ce tirage est dit UNIFORME PAR RAPPORT À U s'il vérifie de plus $\mathbb{P}(X = 1) = \frac{U_2}{U_1}$.

Définition 3.12. Une urne U étant donnée, une suite de tirages $(X_i)_{i=1..k}$ est dite UNIFORME AVEC REMISE si les X_i sont des tirages uniformes d'une boule de U et qu'ils sont indépendants dans leur ensemble.

Définition 3.13. Une urne U étant donnée, une suite de tirages $(X_i)_{i=1..k}$ est dite UNIFORME SANS REMISE si les X_i vérifient :

- X_1 est un tirage uniforme de U ,
- $\mathbb{E}_{\mathbb{P}}(X_i | X_1, \dots, X_{i-1}) = \frac{U_2 - \sum_{j=1}^{i-1} X_j}{U_1 - i}$.

Problème 3.1. LE PROBLÈME ÉTUDIÉ

Nous faisons les hypothèses et adoptons les notations suivantes :

- notons $(U^N)_{N \in \mathbb{N}}$, la suite d'urnes $U^N = (U_1^N, U_2^N)$. Nous prendrons dans la suite le cas particulier où le nombre de boules $U_1^N = N^2$ (le carré de N),
- notons $p^N := \frac{U_2^N}{N^2}$, la probabilité de tirage d'une boule rouge dans U^N ,
- supposons que p^N converge vers une limite $p \in [0, 1]$,
- notons $X = (X^N)_{N \in \mathbb{N}} = ((X_i^N)_{i=1..N})_{n \in \mathbb{N}}$, la suite doublement indexée par $N \in \mathbb{N}$ et $i \in [1, N]$ de N tirages uniformes sans remise de l'urne U^N .

Nous souhaitons étudier le comportement asymptotique lorsque $N \rightarrow \infty$ de X^N .

Nous allons donner un sens et démontrer l'assertion suivante : « Quand on réalise une suite de N tirages sans remise dans l'urne à N^2 boules U^N , la proportion de boules rouges tirées est asymptotiquement probablement p ». Dans le cas d'espèce l'étude d'un tirage donné ne présente que peu d'intérêt pratique.

3.2.2 Résultats

Notons $(M^N)_{N \in \mathbb{N}}$ la suite des mesures empiriques sur E des (X_i^N) .

Le but est de prouver le théorème suivant :

Théorème 3.3. HYPOTHÈSES : *On suppose les notations et hypothèses du problème 3.1 vérifiées. Notons M^N la mesure empirique de X^N .*

CONCLUSION : *Alors il existe une mesure déterministe $\mu \in (\mathcal{P}(\{0,1\}))^\Omega$ telle que :*

$$M^N(\omega) \xrightarrow[N \rightarrow \infty]{*} \mu(\omega) = \mu \quad (3.2)$$

en probabilité.

De plus μ vérifie :

$$\langle Id_{\{0,1\}}, \mu \rangle = p. \quad (3.3)$$

3.2.3 Échangeabilité

Définition 3.14. *Un système de particules $(X_1, \dots, X_N)$ est dit ÉCHANGEABLE si :*

$$\forall \sigma \in \mathfrak{S}_N, \mathcal{L}(X_1, \dots, X_N) = \mathcal{L}(X_{\sigma(1)}, \dots, X_{\sigma(N)}),$$

où $\mathfrak{S}_N$ est l'ensemble des permutations de $[1, N]$ et $\mathcal{L}$ l'application loi.

Dans la suite de cette thèse, les particules ne seront jamais échangeables, en effet nous prendrons des conditions initiales différentes pour chaque particule. Cependant cette notion est bien pratique pour simplifier la présentation de ce chapitre. Nous avons la propriété suivante :

Proposition 3.4. *Soit X un système infini de particules du problème 3.1. Alors pour tout N , le système X^N est échangeable.*

Cette proposition dit en particulier que la première particule tirée dans un tirage sans remise a la même loi individuelle que la quatrième par exemple.

Preuve : Considérons le tirage uniforme d'une partie à N éléments de $[1, N^2]$ que l'on permute de l'ordre croissant en tirant de manière uniforme un $\sigma \in \mathfrak{S}_{N^2}$. Nommons alors $X_i^N = \chi_{\sigma(i) \leq p^N N^2}$. Nous avons créé le tirage sans remise (En effet X_1 est bien un tirage uniforme, et X_i conditionné par $X_1, \dots, X_{i-1}$ est bien un tirage avec la probabilité voulue), il est alors trivial via ce procédé de construction de constater que le système est échangeable. ■

3.2.4 Discussion

Ce qui est écrit au dessus est une manière alambiquée de se donner deux variables aléatoires $M^N(0) = \frac{\text{Card}\{i \in [1, N], X_i^N = 0\}}{N}$ et $M^N(1) = 1 - M^N(0)$.

Le théorème 3.3 est une loi faible des grands nombres puisqu'on peut l'écrire en termes moins savants sous la forme de sens identique :

Théorème 3.4. HYPOTHÈSES : On suppose les notations et hypothèses du problème 3.1 vérifiées.

CONCLUSION : Alors

$$\frac{1}{N} \sum_{i=1..N} X_i^N \xrightarrow[N \rightarrow \infty]{} p \quad (3.4)$$

en probabilité.

Preuve : Soit $\epsilon > 0$. Notons $Y^N = \frac{X_1^N + \dots + X_N^N}{N}$. Remarquons que $\mathbb{E}(Y^N) = \mathbb{E}(X_i^N) = p^N$, en utilisant l'inégalité de Tchebychev, il vient :

$$\mathbb{P}(|Y^N - p^N| \geq \epsilon) \leq \frac{\text{Var}(Y)}{\epsilon^2}.$$

De plus

$$\begin{aligned} \text{Var}(Y^N) &= \frac{1}{N^2} \sum_{i,j} \mathbb{E}(X_i^N X_j^N) - (p^N)^2 \\ &= \frac{\text{Var}(X_1^N)}{N} + \frac{N^2 - N}{N^2} (\mathbb{E}(\mathbb{E}(X_1^N X_2^N | X_1^N)) - (p^N)^2) \\ &= \frac{\text{Var}(X_1^N)}{N} + \frac{N-1}{N} (p^N \frac{p^N N^2 - 1}{N^2 - 1} - (p^N)^2) \\ &= \frac{\text{Var}(X_1^N)}{N} - \frac{N-1}{N} \frac{p^N(1 - p^N)}{N^2 - 1}. \end{aligned}$$

Ainsi

$$\mathbb{P}(|Y^N - p| \geq 2\epsilon) \leq \chi(|p^N - p| > \epsilon) + \frac{\text{Var}(Y^N)}{\epsilon^2} \xrightarrow[N \rightarrow \infty]{} 0,$$

ce qui prouve le théorème. ■

3.3 Champ moyen par la méthode classique : tension, limites, unicité

3.3.1 Explications générales sur la méthode classique

Il s'agit de montrer la convergence de M^N en loi. Pour cela on se place dans $\mathcal{P}(\mathcal{P}(\{0, 1\}))$. Cet espace est compact car $\{0, 1\}$ l'est. Ensuite on doit étudier les limites pour montrer l'unicité de celle-ci en étudiant la convergence vers le candidat limite.

La méthode dont il est question a été introduite par Prokhorov dans les années 1950 (1956 pour le théorème) mais la version actuelle de la présentation est proche de celle de Billingsley (1968) on trouvera des détails sur les affiliations dans [44] p. 154. un cas bien traité est l'équation de Boltzman qu'on pourra voir dans [58], on retrouve aussi la méthode dans l'étude des populations ([115, 41]). Dans le contexte de l'étude de TCP, la méthode

classique est aussi employée dans [118] sur un modèle à pas de temps ; dans [20] et [82] sur des modèles où TCP est adapté pour respecter une hypothèse de partage équitable de bande passante. On pourra aussi citer dans des contextes de dimensionnement de switchs [64].

3.3.2 Preuve du théorème par la méthode classique du champ moyen

La tension Il faut montrer que les mesures M_N sont tendues. C'est le cas ici. En fait ici on peut utiliser directement la compacité faible des mesures de probabilité sur l'ensemble compact $\{0, 1\}$.

Étude des limites On peut donc extraire une limite μ de M_N , dont on ne sait pas qu'elle est déterministe. Il va falloir caractériser cette limite. Ici par exemple on a l'équation,

$$\lim_N \mathbb{E}(M_N) = p.$$

qui a une unique limite déterministe. C'est notre candidat naturel pour la convergence.

Unicité Prouvons directement la convergence vers $\mu = p\delta_1 + (1-p)\delta_0$. Pour toute fonction g , continue bornée sur E , c'est à dire ici toute fonction de E :

$$\langle g, M^N \rangle = \left(\frac{1}{N} \sum_{i=1}^N X_i^N \right) g(1) + \left(1 - \frac{1}{N} \sum_{i=1}^N X_i^N \right) g(0)$$

et

$$\langle g, \mu \rangle = pg(1) + (1-p)g(0).$$

Ainsi

$$\mathbb{E} \left| \langle g, M^N \rangle - \langle g, \mu \rangle \right| \leq (|g(1)| + |g(0)|) \mathbb{E} \left| \frac{1}{N} \sum_{i=1}^N X_i^N - p \right|.$$

Ce qui nous ramène plus ou moins à refaire le même calcul que précédemment.

3.3.3 Une approche alternative

Description L'approche classique ne recherche que des résultats sur les marginales dans le cas où les particules dépendent du temps (en général on se restreint aux processus Markoviens par rapport à un espace d'état naturel). Pour faire mieux il faut avoir recours par exemple au caractère chaotique du système. Ce prolongement présenté de manière pédagogique par Carl Graham dans [59, 61]. On y utilise la démarche de [141] pour obtenir une convergence en probabilité. Ce cas qui traite aussi le problème des fluctuations sur le cas de files d'attente avec choix de la plus courte parmi un nombre donné d'entre elles.

Mise en oeuvre

Définition 3.15. Soit $\nu \in \mathcal{P}(E)$, le système de particules X est dit ν -CHAOTIQUE si et seulement si :

$$\forall i \in \mathbb{N}^*, \mathcal{L}(X_1^N, \dots, X_i^N) \xrightarrow[N \rightarrow \infty]{*} \nu^{\otimes i} \text{ dans } \mathcal{P}(E^i).$$

Le but est d'utiliser le résultat suivant rappelé par Sznitman ([141], p.177) :

Proposition 3.5. HYPOTHÈSES : X est un système de particules infini ν -chaotique
CONCLUSION : la mesure empirique $M^N \rightarrow^* \delta_\nu$ dans $\mathcal{P}(\mathcal{P}(E))$.

Preuve : Ici notons $\mu = p\delta_1 + (1-p)\delta_0$ notre candidat limite, et prouvons que (X_i^N) est μ -chaotique. La condition de chaoticité équivaut par définition à :

$$\forall i, \forall g_1, \dots, g_i \in \mathcal{C}_b, \mathbb{E}\left(\prod_{k=1}^i g_k(X_k^N)\right) \xrightarrow[N \rightarrow \infty]{} \prod_{k=1}^i \mathbb{E}(g_k(\mu)).$$

Dans notre cas une fonction de $\{0, 1\}$ revient à la donnée de deux réels.

Mais dans notre cas une condition suffisante est que $\mathbb{E}(X_i^N | X_1^N, \dots, X_{i-1}^N) \rightarrow p$ ps (en raisonnant par récurrence) puisque :

$$\mathbb{E}(g_i(X_i^N) \prod_{k=1}^{i-1} g_k(X_k^N)) \xrightarrow[N \rightarrow \infty]{} (pg_i(1) + (1-p)g_i(0)) \mathbb{E}\left(\prod_{k=1}^{i-1} g_k(X_k^N)\right).$$

Or

$$\mathbb{E}(X_i^N | X_1^N, \dots, X_{i-1}^N) = \frac{p^N N^2 - \sum_{k=1}^{i-1} X_k^N}{N^2 - i + 1}.$$

Ainsi

$$|\mathbb{E}(X_i^N | X_1^N, \dots, X_{i-1}^N) - p| \leq \left| \frac{i-1 - \sum_{k=1}^{i-1} X_k^N}{N^2 - i + 1} \right| + |p^N - p| \leq \left| \frac{1}{N-1} \right| + |p^N - p|.$$

Ce qui achève la démonstration. ■

3.4 Champ moyen par la méthode de couplage trajectoriel à un système auxiliaire

3.4.1 Explications générales

La méthode de *couplage* consiste à prendre deux variables aléatoires d'espaces de bases différents et de trouver une représentation sur un espace commun où les deux variables soient facilement comparables. Ici nous allons coupler les variables du système des tirages *sans* remise à un système auxiliaire qui est un système de tirages *avec* remise. Mais nous allons aussi coupler les tirages pour des N différents entre eux. L'intérêt de faire ceci est

de se ramener à l'étude d'un système plus simple avec des tirages indépendants où les versions les plus simples de la loi des grands nombres sont applicables.

Si le principe du couplage est un grand classique de l'étude des probabilités, l'utilisation du couplage dans un contexte de champ moyen se retrouve dans [63] mais sans utiliser exactement la même idée de système auxiliaire. La méthode exacte que nous utiliserons ici s'appelle «lookdown construction» et a été introduite en 1999 par Peter Donnelly et Thomas Kurtz dans [41]. Le système «auxiliaire» que nous introduirons se nomme «lookdown process» selon la terminologie de Donnelly-Kurtz.

L'idée derrière la terminologie étant qu'il faut introduire une dynamique auxiliaire «à l'infini», c'est à dire aller voir ce qui se passe en prenant une limite sur la dynamique sans augmenter le nombre de particules.

3.4.2 Couplage entre les trajectoires et convergence

Couplage

Définition 3.16. Notons $\Omega = [0, 1]$ l'espace support des probabilités que nous munissons de la mesure de Lebesgue. On définit une suite de variables aléatoires iid uniformes $(\tilde{\mathcal{X}}_i)_{i \in \mathbb{N}}$ de $\Omega \rightarrow [0, 1]$.

Définition 3.17. Notons $\pi_{p^N} := \chi_{[0, p^N]} : [0, 1] \rightarrow \{0, 1\}$, la fonction caractéristique du segment $[0, p^N]$.

Proposition 3.6. Les TIRAGES COUPLÉS SANS REMISE, X_i^N , associés à la suite $(\tilde{\mathcal{X}}_i)_{i \in \mathbb{N}}$ et au problème 3.1 peuvent être définis de la manière itérative suivante :

- $X_1^N(\tilde{\omega}) = \pi_{p^N}(\tilde{\mathcal{X}}_1(\tilde{\omega}))$,
- si $X_1^N(\tilde{\omega}), \dots, X_{i-1}^N(\tilde{\omega})$ définis :
 - i) posons la variable aléatoire $p_i^N = \frac{1}{N^2 - i + 1}(p^N N^2 - \sum_{n=1}^{i-1} X_n^N)$, la proportion des boules rouges qui restent encore dans l'urne après les $i - 1$ premiers tirages.
 - ii) $X_i^N(\tilde{\omega}) = \pi_{p_i^N}(\tilde{\mathcal{X}}_i(\tilde{\omega}_i))$

La convergence Le cahier des charges de la norme que nous souhaitons définir est :

- permettre de comparer X^N et un système limite de particules de taille infinie $(X_1, \dots)$,
- la convergence pour la norme doit impliquer la convergence en probabilité des mesures empiriques vers une limite déduite de $(X_1, \dots)$,
- la convergence doit fournir une information sur la convergence trajectoire par trajectoire.

Une norme réalisant cela est délicate à trouver. Voici celle que nous allons utiliser, nous avons essayé de lui donner un nom intuitif.

Définition 3.18. Soit E , Polonais, X un système infini de particules et $(X_1, \dots)$ une suite de variables aléatoires de E . On dit que X converge vers $(X_1, \dots)$ pour la CONVERGENCE CESARÒ-AUXILIAIRE si et seulement si :

$$\|X - (X_1, \dots)\| := \mathbb{E} \left(\frac{1}{N} \sum_{i=1}^N |X_i^N - X_i| \right) \xrightarrow{N \rightarrow \infty} 0.$$

Nous généraliserons cette notion de convergence en utilisant la norme sup sur $[0, T]$ dans le cas où $E = \mathbb{D}$ et pas une distance déduite de la topologie J_1 de Skorokhod. Par ailleurs le fait que le choix de la convergence est judicieux s'explique dans ce qui suit.

3.4.3 Définition du système auxiliaire : nouvelle version

Dans cette thèse nous avons préféré nous écarter un petit peu de la méthode «look-down process». La motivation principale est que cette dernière approche se désintéresse complètement d'un objet intuitivement important qu'est le représentant de la particule limite. Nous reviendrons ensuite sur ce qu'est la version originale.

Définition 3.19. *La PARTICULE DU CHAMP MOYEN $\overline{\mathcal{X}}$ du problème 3.1 est le tirage d'un 1 avec probabilité p et 0 le reste du temps.*

Remarque 3.1. *Un tirage avec probabilité p n'est pas en général le tirage d'une urne finie (ce n'est même jamais le cas si p est irrationnel).*

On redéfinit alors :

Définition 3.20. *Le SYSTÈME LIMITE AUXILIAIRE associé au système couplé X est : $\mathcal{X}_i = \pi_p(\tilde{\mathcal{X}}_i)$.*

3.4.4 Énoncés et démonstrations des théorèmes

Le symbole d'espérance est employé dans ce qui suit pour désigner l'espérance sur $(\Omega^{\otimes \mathbb{N}}, \lambda^{\otimes \mathbb{N}})$.

Ce premier théorème illustre le rapport entre les équations des systèmes X , $\mathcal{X}$, et la limite $\mathcal{X}_i$.

Théorème 3.5. *ÉQUATIONS SUR LE SYSTÈME ORIGINAL ET LE SYSTÈME AUXILIAIRE*
Si p^N tend vers un réel p alors $\mathcal{X}_i$ et X_i^N vérifient les formules :

$$X_i^N = \pi_{p_i^N}(\tilde{\mathcal{X}}_i), \quad (3.5)$$

$$\mathcal{X}_i = \pi_p(\tilde{\mathcal{X}}_i), \quad (3.6)$$

$$p = \mathbb{E}(p_i^N). \quad (3.7)$$

Rappelons que $p_i^N = \frac{1}{N^2 - i + 1}(p^N N^2 - \sum_{n=1}^{i-1} X_n^N)$ est une variable aléatoire.

Théorème 3.6. *LOI DES GRANDS NOMBRES POUR LE SYSTÈME AUXILIAIRE*
De plus $\mathcal{X}_i$ vérifie la loi des grands nombres presque sûre :

$$\frac{1}{N} \sum_{i=1..N} \mathcal{X}_i \rightarrow \mathbb{E}(\overline{\mathcal{X}}) = p. \quad (3.8)$$

Théorème 3.7. *CONVERGENCE TRAJECTORIELLE DES TIRAGES AVEC REMISE VERS LES TIRAGES SANS REMISE*

Enfin quand N tend vers l'infini on a la convergence trajectorielle :

$$\mathbb{E}\left(\frac{1}{N} \sum_{i=1..N} |\mathcal{X}_i - X_i^N|\right) \rightarrow 0 \quad (3.9)$$

Preuve : Le premier théorème est une simple conséquence des bonnes définitions et du fait que $\mathbb{E}(p_i^N) = p^N$. La loi des grands nombres est une forme classique avec des variables indépendantes. Le troisième provient du calcul suivant :

$$\begin{aligned} \mathbb{E}\left(\frac{1}{N} \sum_{i=1..N} |\mathcal{X}_i - X_i^N|\right) &\leq \frac{1}{N} \sum_{i=1..N} \mathbb{E}(|p - p^N| + |p^N - p_i^N|) \\ &\leq |p - p^N| + \frac{1}{N} \sum_{i=1..N} \frac{1}{N^2 - i + 1} \mathbb{E}\left|\left(\sum_{n=1}^{i-1} X_n^N\right) - (i-1)p^N\right| \\ &\leq |p - p^N| + \frac{1}{N} \sum_{i=1..N} \frac{i-1}{N^2 - i + 1} \xrightarrow{N \rightarrow \infty} 0 \end{aligned}$$

■

Preuve : du théorème original 3.3 : la convergence des mesures aléatoires est une conséquence du théorème de convergence trajectorielle comme l'illustre le raisonnement volontairement généraliste suivant :

Pour toute fonction g (de classe $\mathcal{C}_b$) sur $\{0, 1\}$,

$$\begin{aligned} &\limsup_{N \rightarrow \infty} |\mathbb{E}\langle g, M^N \rangle - \mathbb{E}\langle g, \mu \rangle| \\ &\leq \limsup_{N \rightarrow \infty} \left| \frac{1}{N} \sum_{i=1}^N \mathbb{E}[g(X_i^N) - g(\mathcal{X}_i)] \right| + \limsup_{N \rightarrow \infty} \left| \frac{1}{N} \sum_{i=1}^N g(\mathcal{X}_i) - \mathbb{E}\langle g, \mu \rangle \right| \\ &\leq \limsup_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N \mathbb{E}|g(X_i^N) - g(\mathcal{X}_i)| + 0 \\ &\leq C_g \limsup_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N \mathbb{E}|X_i^N - \mathcal{X}_i| \\ &= 0. \end{aligned}$$

D'où le résultat (la deuxième ligne utilise la loi des grands nombres du système auxiliaire, et la dernière la convergence trajectorielle).

■

3.4.5 Définition du système auxiliaire : version originale

Nous reprenons ici une manière alternative de définir la limite du champ moyen. Cette approche a plusieurs intérêts, c'est la manière originale de définir le lockdown process dans [41], et c'est la méthode que nous allons suivre dans la deuxième partie où nous reprenons l'article [109].

L'approche retenue consiste à introduire un nouveau système infini de particules $\mathcal{X} = (\mathcal{X}_i^N)$ qui aura les mêmes propriétés asymptotiques que X en utilisant la dynamique asymptotique.

Définition 3.21. Le SYSTÈME INFINI DE PARTICULES AUXILIAIRE associé au système couplé (X_i^N) du problème 3.1 est : $\mathcal{X}_i^N = \pi_{p^N}(\tilde{\mathcal{X}}_i)$ (où les $\tilde{\mathcal{X}}_i$ sont donnés à la définition 3.16 uniformes sur $[0, 1]$ et indépendants).

Ce système correspond à des tirages avec remise de l'urne U^N .

On peut extraire de $\mathcal{X}_i^N$ une limite $\mathcal{X}_i$ trajectoire à trajectoire où $\mathcal{X}_i = \pi_p(\tilde{\mathcal{X}}_i)$. Ainsi $\mathcal{X}_i$ est une suite de tirages iid.

La définition est la même, mais notre approche (dite «nouvelle version») un peu différente présente plusieurs avantages selon nous :

- nous introduisons à priori $\bar{\mathcal{X}}$ qui est un objet mathématiquement intéressant,
- nous n'avons pas besoin du système de particule auxiliaire qui est un objet dont on peut se passer y compris à l'intérieur des démonstrations,
- les $\mathcal{X}_i$ sont des objets définis de manière constructive et pas comme limite d'une suite extraite d'une suite.
- enfin définir $\bar{\mathcal{X}}$ est exactement de la même difficulté que définir $\mathcal{X}_i^N$ et ne relève pas d'une approche moins intuitive du problème.

3.4.6 Comparaison des méthodes

Nous avons vu brièvement trois approches. La méthode classique est bien balisée mais elle donne en général des résultats un peu trop faibles pour traiter des problèmes où il est nécessaire de connaître les trajectoires comme par exemple lorsqu'il s'agit de prendre des temps d'arrêts sur les trajectoires. Ainsi elle présente l'avantage de donner un certain nombre d'étapes de raisonnement sans pour autant résoudre le point difficile qui est en général l'unicité de la limite.

L'approche par le caractère chaotique donne des résultats très forts tout en étant aussi bien balisée, mais il faut pouvoir la mettre en oeuvre : en particulier l'hypothèse de chaotité implique un caractère quasiment échangeable, ce qui n'est pas le cas dès que les conditions initiales par exemple ne sont pas identiques pour toutes les particules. La dernière méthode, que nous allons utiliser dans la suite est plus difficile parce qu'un grand nombre d'étapes vont rentrer en jeu, et elle demande de deviner la dynamique limite. Cependant elle peut fournir des résultats extrêmement forts parce qu'elle n'oublie pas les particules ce qui est un point très important quand on veut avoir des dépendances retardées et des problèmes de frontières.

3.5 Apports mathématiques de cette thèse

Nous utilisons des méthodes mathématiques relativement récentes et totalement inédites dans le contexte de l'étude des réseaux de télécommunications, à savoir la méthode du «lookdown process». Nous adaptons les théorèmes classiques à notre contexte pour l'étude du champ moyen, dans le prochain chapitre nous verrons la *résolution d'un type d'EDS* (Équation Différentielle Stochastique) qui ne rentre pas dans les cas classiques dont nous étudierons les solutions fortes. Enfin nous étudions des *effets de bord et des retards* dans la convergence en champ moyen.

Le chapitre 4 généralise la méthode pour des processus. L'étude du système auxiliaire est faite dans le chapitre 5. Vient ensuite le chapitre 6 qui commence à traiter le problème de frontières. Nous refermons cette partie mathématique par un exemple d'utilisation de la méthode pour résoudre une question restée jusqu'alors ouverte de la convergence du modèle Baccelli-Hong en champ moyen ce qui fait l'objet du chapitre 7. Cette partie est écrite sans se préoccuper de manière générale du problème des retards, ceci surtout par souci de clarté plus que pour des difficultés techniques.

Dans la seconde partie nous reprenons les démonstrations de cette première partie pour le modèle central de la thèse. Ceci pour plusieurs raisons. La méthode utilisée est très proche mais diffère quelque peu dans l'introduction du système auxiliaire plus proche de la méthode originale de «lookdown process». Par ailleurs cette étude aurait été difficile à inclure dans le cas général car les délais ont une forme très particulière d'une part, et d'autre part ces mêmes délais sont à l'origine de tous les théorèmes d'existence de limites et de traversées de frontières. La suite de la thèse, c'est à dire la fin de la deuxième partie et la troisième partie sont consacrées à l'exploitation pratique des résultats mathématiques.

Chapitre 4

Problème du champ moyen dans $\mathbb{D}$

Dans ce chapitre nous prolongeons la méthode du champ moyen par système auxiliaire dans le cas où l'espace d'états est $E = \mathbb{D}([0, T], \mathbb{R}^k)$. La dynamique des particules est seulement couplée par une ressource partagée.

Sommaire

4.1	Rappels de mathématiques générales	58
4.1.1	Prévisibilité	58
4.1.2	Topologie de Skorokhod	61
4.1.3	Module de continuité et compacts de $\mathbb{D}$	62
4.1.4	Tension d'une suite de processus càdlàg au sens de Skorokhod	63
4.2	Première présentation de l'équation différentielle stochastique que nous allons étudier	63
4.3	Éléments mathématiques spécifiques	64
4.3.1	Un mode de définition d'un processus de Poisson doublement stochastique	64
4.3.2	Marginales	65
4.3.3	Crochet marginal double	66
4.3.4	Fonctions prévisibles de E	67
4.4	Formulation du problème	67
4.4.1	Notations et définitions	67
4.4.2	Une première formulation du problème	68
4.4.3	Hypothèses techniques	69
4.5	Couplage et norme	69
4.5.1	Norme de convergence d'un système infini de particules	69
4.5.2	Couplage du système initial	70
4.5.3	Reformulation du problème	71
4.6	Discussion autour du problème étudié	71

4.1 Rappels de mathématiques générales

Références Dans cette partie, nous supposons connus les définitions et théorèmes classiques sur les mesures et la convergence faible. Nous effectuons quelques rappels sans démonstrations sur la prévisibilité puis sur la topologie de Skorokhod de $[0, 1]$ dans $\mathbb{R}$. Pour plus de détails, on pourra se référer au livre «Convergence of Probability Measures» de Patrick Billingsley [23] et à «Point Processes and Queues : Martingale Dynamics» de Pierre Brémaud [27]. Des informations complémentaires se trouvent dans «Markov Processes ; Characterization and Convergence» de Stewart Ethier et Thomas Kurtz [44], «Limit theorems for Stochastic Processes» de Jean Jacod et Albert Shiryaev [77] et «Brownian motion and stochastic calculus» de Ioannis Karatzas et Steven Shreve [79].

4.1.1 Prévisibilité

Dans cette section nous rappelons rapidement quelques points techniques autour de la notion de processus stochastique.

Donnons-nous les objets suivants :

- un espace muni d’une probabilité $(\Omega, \mathbb{P})$, on notera $\mathcal{T}$ le domaine de $\mathbb{P}$, c’est une tribu,
- un ensemble d’états muni d’une tribu $(E, \mathcal{E})$,
- un ensemble de temps T (généralement un segment de $\mathbb{R}$ ou $\mathbb{Z}$), muni (le plus souvent) d’une structure d’ordre total.

Discussion Nous formalisons d’abord la notion d’observation. Lorsque l’observation est unique nous parlons de VARIABLE ALÉATOIRE, lorsqu’elle est répétée tout au long du temps, nous parlons de PROCESSUS STOCHASTIQUE. Ensuite nous étudions les rapport entre continuité et information disponible.

Processus stochastique La notion d’observation ponctuelle est alors formalisée ainsi :

Définition 4.1. Une VARIABLE ALÉATOIRE X sur l’espace $(\Omega, \mathbb{P})$ et à valeurs dans l’espace d’états $(E, \mathcal{E})$ est une fonction mesurable de $(\Omega, \mathcal{T}) \rightarrow (E, \mathcal{E})$.

La fonction d’ensemble X^{-1} transporte la probabilité $\mathbb{P}$ de Ω vers E . Cette probabilité sur E se nomme la LOI de la variable aléatoire X , notée $\mathcal{L}(X) := \mathbb{P}X^{-1}$.

La première tâche est de donner un sens à l’évolution du temps T . Pour ce faire on introduit la notion de FILTRATION qui représente l’information disponible au temps t . La «flèche du temps» est alors formalisée par une propriété de croissance de l’information disponible par rapport au temps (c’est pourquoi T a été muni d’un ordre total).

Définition 4.2. Une FILTRATION DE L’ESPACE PROBABILISÉ $(\Omega, \mathbb{P})$ POUR L’ENSEMBLE DE TEMPS T est une famille de sous-tribus de $\mathcal{T}$, $\mathcal{F}_t$ croissante par rapport au temps $t \in T$.

L’ensemble des variables aléatoires mesurables de $\Omega \rightarrow E$ par rapport à $\mathcal{F}_t$ et $\mathcal{E}$ représente alors l’ensemble de observations possibles au temps $t \in T$.

Un ensemble de temps T et une filtration $\mathcal{F}_t$ de Ω étant donnés, intéressons-nous à une famille d'observations sur le même espace $(X_t)_{t \in T}$. Cette famille forme une fonction de $T \times \Omega$ dans E . Plusieurs notions de mesurabilité sont alors envisageables.

Définition 4.3. Une famille $(X_t)_{t \in T}$ ou plus simplement X_t de fonctions de Ω dans E étant donnée, on peut avoir :

1. X_t mesurable par rapport à $\mathcal{T}$ et à $\mathcal{E}$ pour tout t , c'est-à-dire que X_t est une collection de variables aléatoires. X_t est alors un PROCESSUS STOCHASTIQUE ou plus simplement processus.
2. pour tout $t \in T$, X_t est mesurable par rapport à $\mathcal{F}_t$ et $\mathcal{E}$. X_t est en particulier un processus stochastique ; il est dit PROCESSUS ADAPTÉ À LA FILTRATION $(\mathcal{F}_t)$.
3. pour tout temps t ¹³ $X : [0, t] \times \Omega \rightarrow E$ est mesurable par rapport aux tribus $\mathcal{B}([0, t]) \otimes \mathcal{F}_t$ et $\mathcal{E}$. X_t est en particulier un processus stochastique adapté ; il est dit PROCESSUS PROGRESSIF par rapport à la filtration $\mathcal{F}_t$.
4. la fonction $X : T \times \Omega \rightarrow E$ est mesurable par rapport à la tribu $\mathcal{G} := \sigma([s, \infty[\times F, s \in T, F \in \mathcal{F}_s)$ et X_0 est $\mathcal{F}_0$ -mesurable. X_t est en particulier un processus progressif ; il est dit PROCESSUS PRÉVISIBLE.

Si les trois premières définitions sont à peu près évidentes à comprendre, la dernière est un peu plus subtile. En fait on retiendra qu'il s'agit d'une généralisation du concept de filtration continue à gauche. Regardons ce qu'elle signifie si $T = \mathbb{N}$. Dans ce cas $\mathcal{G} := \sigma([k+1, \infty[\times F_k, \forall n \in \mathbb{N}, \forall k, 0 \leq k \leq n \text{ et } \forall F_k \in \mathcal{F}_k)$. En particulier $\mathcal{G}$ contient les ensembles $\{n\} \times F_{n-1}$ où $F_{n-1} \in \mathcal{F}_{n-1}$, ce qui signifie en particulier que X_n est $\mathcal{F}_{n-1}$ -mesurable. En fait dans le cas de $\mathbb{N}$, être $\mathcal{F}_{n-1}$ -adapté est une définition équivalente de la prévisibilité. Cette réciproque est aussi vraie pour $T = \mathbb{R}^+$ bien qu'il existe des processus prévisibles qui ne soient pas continus à gauche.

Proposition 4.1. ([27] p. 9) Parmi les processus de $T = \mathbb{R}^+$ dans $\mathbb{R}^k$ les classes suivantes sont prévisibles dès qu'elles sont adaptées :

1. les processus déterministes,
2. les processus continus à gauche (c'est-à-dire dont toutes les trajectoires sont continues à gauche).

Beaucoup de théorèmes s'appliquent pour les processus prévisibles, ceci mène naturellement à chercher à étudier des trajectoires continues à gauche.

Filtration engendrée La méthode canonique pour construire une filtration est de donner une famille de processus et de construire la FILTRATION ENGENDRÉE :

Définition 4.4. La FILTRATION $\mathcal{F}_t$ ENGENDRÉE PAR UNE FAMILLE DE PROCESSUS $(Y^k)_{k \in K}$ définis sur une même probabilité $\mathbb{P}$ et un même espace de temps T , segment de $\mathbb{Z}^+$ ou $\mathbb{R}^+$ contenant l'origine est $\mathcal{F}_t = \sigma((Y^k_{[0,t]})_{k \in K})$, la tribu engendrée par les processus restreints aux temps $[0, T]$.

¹³ $[0, t]$ représente les éléments de T classés entre 0 et t au sens large ; $\mathcal{B}(\cdot)$ est l'ensemble des Boréliens.

Par définition la famille de processus $(Y^k)_{k \in K}$ est donc adaptée à la filtration $\mathcal{F}_t$ qu'ils engendrent. De plus il ne fait aucun doute que la famille est bien croissante et incluse dans la tribu $\mathcal{T}$, domaine de $\mathbb{P}$.

Définition 4.5. On nomme ensembles négligeables d'une tribu $\mathcal{F}$ l'ensemble de parties : $\mathcal{N}(\mathcal{F}, \mathbb{P}) := \{G \subset \Omega, \exists F \in \mathcal{F} G \subseteq F \text{ et } \mathbb{P}(F) = 0\}$.

On nomme filtration COMPLÉTÉE de $(\mathcal{F}_t)$ la filtration définie par : $\overline{\mathcal{F}}_t^{\mathbb{P}} := \sigma(\mathcal{F}_t \cup \mathcal{N}(\mathcal{F}_t, \mathbb{P}))$.
On nomme filtration AUGMENTÉE de $(\mathcal{F}_t)$ la filtration définie par : $\mathcal{F}_t^{\mathbb{P}} := \sigma(\mathcal{F}_t \cup \mathcal{N}(\mathcal{F}_{\infty}, \mathbb{P}))$.

Continuité de la filtration Dans un monde idéal, les trajectoires sont toutes continues à gauche, ce que assure la prévisibilité de tout processus adapté. Par ailleurs, une autre condition technique importante est que la filtration soit continue à droite. À partir de maintenant, T est un segment de $\mathbb{R}$.

Définition 4.6. Une filtration $\mathcal{F}_t$ étant donnée pour des temps T , définissons la tribu LIMITE À DROITE (resp. limite à gauche) $\mathcal{F}_{t+} := \bigcap_{s>t} \mathcal{F}_s$ (resp. $\mathcal{F}_{t-} := \bigcup_{s<t} \mathcal{F}_s$).
Une filtration est CONTINUE À DROITE (resp. à gauche) si pour tout temps t , $\mathcal{F}_{t+} = \mathcal{F}_t$ (resp. $\mathcal{F}_{t-} = \mathcal{F}_t$).

Intuitivement, une filtration est continue à droite si les événements arrivent exactement en t est pas juste après.

Définition 4.7. On dit qu'une filtration $\mathcal{F}_t$ satisfait les CONDITIONS USUELLES si

- $\mathcal{F}_0$ contient les négligeables de $\mathcal{T}$,
- $\mathcal{F}_t$ est continue à droite.

On a les propriétés suivantes qui expliquent l'intérêt de la continuité à droite :

Proposition 4.2. 1. La limite à droite d'une filtration est une filtration continue à droite.
2. Le supréum de temps d'arrêt est toujours un temps d'arrêt, si la filtration est continue à droite, un infimum de temps d'arrêt est aussi un temps d'arrêt ([27] p.10).
3. Pour une filtration satisfaisant aux conditions usuelles, le temps de première entrée d'un processus progressif dans un borélien est un temps d'arrêt ([44] p. 55).

Continuité des processus et de la filtration engendrée Intéressons-nous à une TRAJECTOIRE du processus X , c'est-à-dire pour un $\omega \in \Omega$, à la fonctions $t \rightarrow X_t(\omega)$.

Définition 4.8. Un processus sera dit continu/continu à droite/càdlàg... si toutes ses trajectoires le sont.

Voici les liens qui existent entre la continuité des trajectoires est la continuité de la filtration engendrée :

Proposition 4.3. 1. Si X_t est continu à gauche alors la filtration engendrée et la filtration engendrée augmentée sont continues à gauche.

2. Si X_t est un processus de Markov fort à k dimensions de distribution initiale μ , alors la filtration augmentée $\mathcal{F}_t^\mu$ est continue à droite ([79] p.90).

Dans ce qui suit nous étudierons toujours des filtrations issues de processus ayant la propriété de Markov fort et avec des trajectoires càdlàg nous sous-entendrons alors l'augmentation des tribus par rapport aux conditions initiales. Ceci explique pourquoi nous serons toujours dans les conditions usuelles avec des processus prévisibles.

4.1.2 Topologie de Skorokhod

Définition 4.9. Soit F , un espace topologique et (a, b) un intervalle de $\overline{\mathbb{R}}$. On note $\mathbb{D}((a, b), F)$ l'espace des fonctions càdlàg c'est-à-dire continues à droites avec des limites à gauche en tout point de (a, b) . On note $\mathbb{D} = \mathbb{D}([0, 1], \mathbb{R})$.

$\mathbb{D}((a, b), F)$ doit être muni d'une topologie qui en fasse un espace Polonais. Pour cela la convergence uniforme est trop restrictive.

Dans ce qui suit nous traitons le cas de $\mathbb{D}$. Dans la suite nous ne nous intéresserons qu'à des segments où la définition de la topologie de Skorokhod est directement déduite de celle que nous allons définir. Cependant rappelons que le passage à $\mathbb{R}$ pose certaines difficultés techniques.

Avant la définition rappelons qu'un homéomorphisme est un isomorphisme d'espaces topologiques c'est-à-dire une application bi-continue.

Définition 4.10. On nomme CHANGEMENT DE TEMPS un homéomorphisme croissant de $[0, 1]$. Notons Λ l'ensemble des changements de temps.

Définition 4.11. La TOPOLOGIE DE SKOROKHOD est définie par les distances équivalentes :

$$d(x, y) := \inf_{\lambda \in \Lambda} \max(\|\lambda - Id_{[0,1]}\|_\infty, \|x - y \circ \lambda\|_\infty)$$

$$d_0(x, y) := \inf_{\lambda \in \Lambda} \max(\|\lambda\|, \|x - y \circ \lambda\|_\infty)$$

où $\|\lambda\| := \sup_{s \neq t} \left| \log \frac{\lambda(s) - \lambda(t)}{s - t} \right|$

d_0 se nomme la DISTANCE DE SKOROKHOD, on la notera dans la suite aussi $\|\cdot\|_s$.

Proposition 4.4. ([44] page 112 et suivantes) $(\mathbb{D}, d_0)$ est un espace Polonais.

Notons que $(\mathbb{D}, d)$ n'est pas complet, cependant, la distance d est plus simple à manipuler pour les propriétés topologiques.

Notons au passage puisqu'on en est aux propriétés d'ordre topologique que

Proposition 4.5. ([23] chapitre 3)

- $\mathcal{Bor}(\mathcal{O}_{(\mathbb{D}, d)}) = \sigma(\{\pi_t, t \in [0, 1]\})$ où π_t est la fonction qui à $x \in \mathbb{D}$ associe $\pi_t(x) = x(t) \in \mathbb{R}^k$.
- la fonction $\sup$ qui à $x \in \mathbb{D}$ associe $\sup_{t \in [0, 1]} |x(t)|$ est continue pour la topologie de Skorokhod sur l'ensemble des fonctions continues. En particulier une limite au sens de Skorokhod vers une fonction continue est une limite uniforme.

4.1.3 Module de continuité et compacts de $\mathbb{D}$

Le but ici est d'indiquer rapidement comment on peut caractériser les compacts de $\mathbb{D}$, il s'agit donc d'une généralisation du théorème d'Ascoli-Arzelà pour les fonctions continues (les compacts des fonctions continues pour la norme uniforme sont les ensembles de fonctions uniformément équi continues) dont nous rappelons une formulation adaptée.

Une première étape est de trouver la bonne caractérisation des éléments de $\mathbb{D}$, cela se fait en s'inspirant de la notion de module de continuité.

Définition 4.12. Soit $x : [0, 1] \rightarrow \mathbb{R}$. Le MODULE DE CONTINUITÉ w de x est la fonction

$$w(x, \delta) = \sup_{|s-t| \leq \delta} |x(s) - x(t)|.$$

Proposition 4.6. x est continue si et seulement si $w(x, \delta) \xrightarrow{\delta \rightarrow 0} 0$.

Théorème 4.1. ASCOLI-ARZELÀ

HYPOTHÈSES : A est une partie de $\mathcal{C}([0, 1], \mathbb{R})$

CONCLUSION : A est relativement compact si et seulement si les deux conditions suivantes sont satisfaites :

- les fonctions de A sont uniformément bornées,
- $\lim_{\delta \rightarrow 0} \sup_{x \in A} w(x, \delta) = 0$.

De manière similaire

Définition 4.13. Soit $x : [0, 1] \rightarrow \mathbb{R}$. Le MODULE DE CONTINUITÉ w' de x est la fonction

$$w'(x, \delta) = \inf \left\{ \max_{s, t \in [t_i, t_{i+1}[} |x(s) - x(t)|; p \in \mathbb{N}, 0 = t_0 < t_1 < \dots < t_{p+1} = 1, t_{i+1} - t_i > \delta \right\}.$$

Proposition 4.7. $x \in \mathbb{D} \Leftrightarrow w'(x, \delta) \xrightarrow{\delta \rightarrow 0} 0$.

Théorème 4.2. HYPOTHÈSES : A est une partie de $\mathbb{D}$

CONCLUSION : A est relativement compact si et seulement si les deux conditions suivantes sont satisfaites :

- les fonctions de A sont uniformément bornées,
- $\lim_{\delta \rightarrow 0} \sup_{x \in A} w'(x, \delta) = 0$.

Une mise en garde Rappelons rapidement les faits suivants :

- $\mathbb{D}$ a une structure naturelle d'espace vectoriel sur $\mathbb{R}$,
- Cependant la topologie de Skorokhod n'est pas en général une topologie d'espace vectoriel.

Voici un exemple typique de cette difficulté sur $\mathbb{D}([0, 1], \mathbb{R})$: $f_n = \chi_{[0, \frac{1}{2} - \frac{1}{n+1}[}$ et $f_\infty = \chi_{[1, \frac{1}{2}[}$. $f_n \rightarrow^{J_1} f_\infty$ pourtant $f_n + f_\infty \not\rightarrow^{J_1} 2f_\infty$. Remarquons au passage que le choix de d_0 permet à la suite $f_n + f_\infty$ de ne pas être une suite de Cauchy car le $\log \frac{\lambda(s) - \lambda(t)}{s-t}$ ne permet pas au changement de temps de stationner au point $\frac{1}{2}$.

4.1.4 Tension d'une suite de processus càdlàg au sens de Skorokhod

Théorème 4.3. HYPOTHÈSES : La suite X_i de variables aléatoires de $(\Omega_i, \mathcal{F}_i, \mathbb{P}_i)$ à valeurs dans $\mathbb{D}$

CONCLUSION 1 : $(X_i)_i \in \mathbb{N}$ est tendue si et seulement si les deux conditions suivantes sont réalisées :

- (i) $\forall \epsilon > 0, \exists i \in \mathbb{N}, \mathbb{P}_i(\sup_{t \in [0,1]} |X_i(t)| > M) \leq \epsilon,$
- (ii) $\forall \epsilon > 0, \forall \eta > 0, \exists \delta, \exists i_0 \in \mathbb{N}, \forall i \geq i_0, \mathbb{P}_i(w'(X_i, \delta) > \eta) < \epsilon.$

CONCLUSION 2 : si on remplace w' par w dans (ii), alors les limites extraites ne chargent que les fonctions continues.

Un grand nombre de critères de tension existent déduits de cette réécriture de la définition (voir dans [23] par exemple). De plus rappelons encore une fois que sur un espace Polonais toute variable aléatoire X est tendue ([42] p.311), ce qui est la moindre des choses avant de se lancer dans l'étude de la tension d'une famille de variables aléatoires.

Voici maintenant un théorème qui explique l'intérêt particulier pour l'étude de la tension sur $\mathbb{D}$. Il dit que la convergence faible sur $\mathbb{D}$ pour les suites tendues est caractérisée par la convergence fini-dimensionnelle des mesures¹⁴.

Proposition 4.8. NOTATIONS : P est une mesure de probabilité sur $\mathbb{D}$. $T_P \subset [0, 1]$ est l'ensemble des points t où π_t est continue P -presque sûrement.

CONCLUSION : T_P a un complémentaire au plus dénombrable.

Théorème 4.4. ([23] Thm 15.1 p. 124)

NOTATIONS : P^N est une suite de mesures de probabilité sur $\mathbb{D}$ qui tend pour la convergence faible sur la topologie de Skorokhod vers une probabilité P .

HYPOTHÈSES :

- pour toute famille finie $(t_i)_{i=1..k}$ d'éléments de T_P ,
 $(P^N \pi_{t_1}^{-1}, \dots, P^N \pi_{t_k}^{-1}) \xrightarrow[N \rightarrow \infty]{*} (P \pi_{t_1}^{-1}, \dots, P \pi_{t_k}^{-1})$
pour la convergence faible sur la topologie produit.
- (P^N) est tendue

CONCLUSION : Alors $P^N \xrightarrow[N \rightarrow \infty]{*} P$ pour la convergence faible sur la topologie de Skorokhod.

4.2 Première présentation de l'équation différentielle stochastique que nous allons étudier

Nous souhaitons étudier l'ÉDS suivante :

$$X_i^N(t) = x_i + \int_{s \in [0, t]} f(X_i^N, M^N)(s^-) ds + \int_{s \in [0, t]} l(X_i^N, M^N)(s^-) dJ_i^N(\Lambda_i^N(s^-)) \quad (4.1)$$

¹⁴si une mesure est toujours caractérisée par l'ensemble de ses projections fini-dimensionnelles (c'est-à-dire ses marginales), il n'est pas toujours vrai que la convergence projection par projection implique la convergence faible.

où le processus d'intensité des sauts Λ_i^N vaut :

$$\Lambda_i^N(s) = \int_{[0,s]} I(X_i^N, M^N)(u^-) du, \quad (4.2)$$

et où

- $X = (X^N) = (X_i^N)$ est un système infini de particules à valeur dans $\mathbb{R}^k$, c'est-à-dire que pour tout N fixé, les particules $(X_1^N, \dots, X_N^N)$ interagissent,
- M^N est la mesure empirique de X^N ,
- f modélise un «drift» qui prend en compte la particule et la mesure empirique,
- J_i^N est un processus de sauts pur et l une taille de sauts.

L'objectif des paragraphes qui viennent est d'expliquer comment formaliser cette étude et quelles hypothèses mathématiques faire sur les objets en présence.

4.3 Éléments mathématiques spécifiques

Dans tout ce qui suit, $E = \mathbb{D}([0, T], \mathbb{R}^k)$. Les processus que nous considérerons seront des variables aléatoires à valeurs dans E avec la conversion : si X est une variable aléatoire à valeurs dans E , $X_t(\omega) := X(\omega)(t)$ est le processus associé.

4.3.1 Un mode de définition d'un processus de Poisson doublement stochastique

Nous supposons connues les définitions et propriétés classiques sur les processus de Poisson homogènes et inhomogènes.

Définition 4.14. Une probabilité sur les ensembles dénombrables de points de $\mathbb{R}^2$, Υ étant donnée. Notons λ , la mesure de Lebesgue sur $\mathbb{R}^2$ et $N_B(\omega)$ le nombre de points de $\Upsilon(\omega) \cap B$.

Υ est un PROCESSUS DE POISSON DANS LE PLAN¹⁵ uniforme d'intensité $I \in \mathbb{R}_+^*$ s'il vérifie les deux conditions :

- pour tout borélien B de $\mathbb{R}^2$ $N_B = \mathcal{L}$ Poisson($I\lambda(B)$),
- si B_1 et B_2 sont deux Boréliens d'intersection vide, N_{B_1} et N_{B_2} sont indépendants.

Proposition 4.9. (Voir [140] p.42) Il existe des processus de Poisson dans le plan d'intensité 1.

Proposition 4.10. Soit Υ un processus de Poisson dans $\mathbb{R}^2$ d'intensité 1 et $I(t)$ un processus càdlàg à valeurs dans $\mathbb{R}$. Alors

$$J(t) = \int \int_{(u,v) \in [0,t] \times \mathbb{R}} \chi_{[0,I(u)]}(v) d\Upsilon(u, v)$$

est un processus de Poisson doublement stochastique d'intensité $I(t)$. De plus $J(t)$ est un processus càdlàg.

Ceci est illustré par la figure (4.1).

¹⁵Notons que le processus de Poisson dans le plan est une variable aléatoire et pas un processus au sens habituel, nous conservons le nom car c'est l'usage.

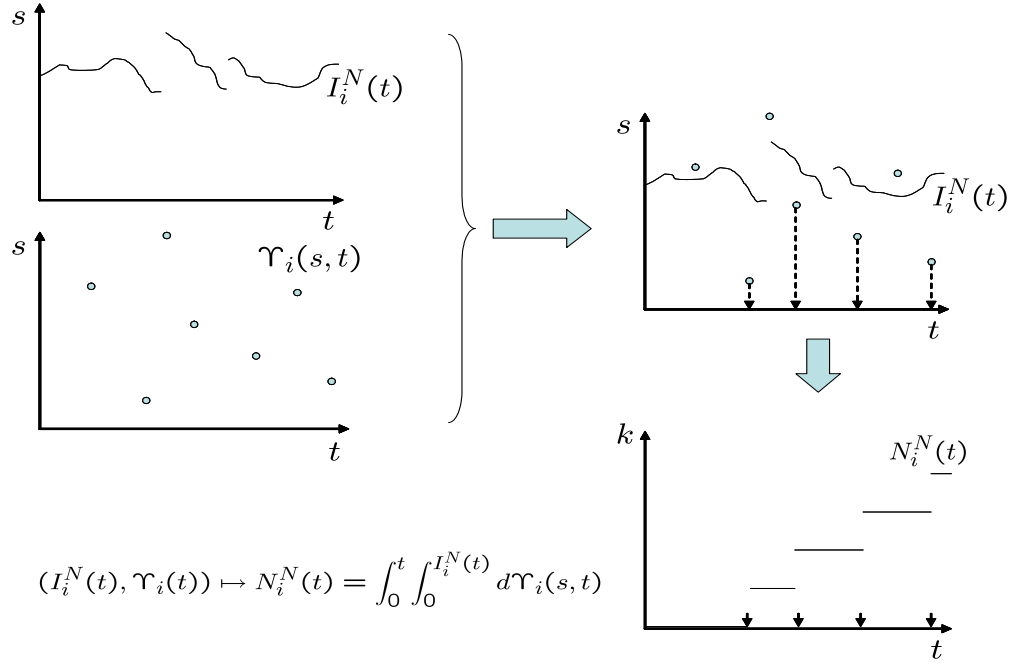


FIG. 4.1 – Génération d'un processus de Poisson d'intensité donnée à partir d'un processus de Poisson dans le plan d'intensité 1.

4.3.2 Marginales

Rappels Notons $\mathcal{F}(F, G) := G^F$, les fonctions de F dans G . Rappelons que $X = (X^N)_{N \in \mathbb{N}} = (X_i^N)_{i \leq N}$ est un système infini de particules où les X_i^N sont des variables aléatoires cadlag. De plus on note M^N la mesure empirique de X^N , c'est une mesure aléatoire. Enfin notons Ω , l'espace de probabilité sous-jacent.

Rappelons la notation des crochets de dualité introduite dans le chapitre 3 : pour une séquence de variables aléatoire $X^N = (X_1^N, \dots, X_N^N)$ (dont la mesure empirique est $M^N := \frac{1}{N} \sum_{i=1}^N \delta_{X_i^N}$) à valeurs dans E et une fonction continue bornée g de E dans $\mathbb{R}$:

$$\langle g, X^N \rangle = \langle g, M^N \rangle = \frac{1}{N} \sum_{i=1}^N g(X_i^N).$$

Nous allons prolonger cette notation en passant au processus marginal.

Définition

Définition 4.15. MESURE MARGINALE Une mesure aléatoire M de Ω sur $E = \mathbb{D}([0, T], \mathbb{R}^k)$ étant donnée, on note M_t le processus marginal c'est-à-dire la distribution dans $\mathcal{P}(\mathbb{R}^k)$ de $M(\cdot)(t)$ pour t fixé.

Commençons par une propriété à peu près évidente dans la mesure où la convergence presque sûre implique la convergence en loi.

Proposition 4.11. *La mesure marginale d'un processus de $\mathbb{D}([0, T], \mathbb{R}^k)$ appartient à $\mathbb{D}([0, T], \mathcal{P}(\mathbb{R}^k))$.*

En particulier, si la mesure est une mesure empirique, $M^N = \frac{1}{N} \sum_{i=1}^N \delta_{X_i^N}$, N étant fixé, alors la mesure empirique marginale M_t^N vaut $\frac{1}{N} \sum_{i=1}^N \delta_{X_i^N(t)}$; c'est un processus dont les trajectoires sont dans $\mathbb{D}([0, T], \mathcal{P}(\mathbb{R}^k))$.

Nous disposons maintenant de deux notions possibles de crochet de dualité :

- $\langle g, M \rangle$ où $g \in \mathcal{C}_b(E, \mathbb{R})$ et $M \in \mathcal{P}(E)$,
- pour chaque t fixé $\langle g, M_t \rangle$ où $g \in \mathcal{C}_b(\mathbb{R}^k, \mathbb{R})$ et $M_t \in \mathcal{P}(\mathbb{R}^k)$.

Définition 4.16. *Nous nommons cette deuxième version le CROCHET MARGINAL SIMPLE, que nous noterons par extension pour X un système fini de particules fixé :*

$$\langle g, X^N \rangle(t) := \langle g, M_t^N \rangle.$$

Discussion Nous voyons alors que

$$\langle g, X^N \rangle(t) = \frac{1}{N} \sum_{i=1}^N g(X_i^N(t)).$$

Ainsi si les $X_i^N(t)$ sont iid de distribution commune représentée par une particule que nous noterons par exemple $\bar{\mathcal{X}}(t)$, alors dans un cas borné par exemple la loi forte des grands nombres s'applique et il vient :

$$\langle g, X^N \rangle(t) \xrightarrow{N \rightarrow \infty} \mathbb{E} (g(\bar{\mathcal{X}}(t))) .$$

Dans l'optique de «lookdown process» qui consiste à regarder la dynamique limite, il est donc utile d'introduire la définition du crochet marginal double qui suit qui va remplacer les crochets simples dans les dynamiques à étudier.

4.3.3 Crochet marginal double

La définition du crochet double ne servira pas avant le chapitre suivant et l'introduction du système auxiliaire.

Discussion (suite) Nous souhaitons maintenant introduire un remplaçant de $\langle g, M^N \rangle(t)$ lorsque N va tendre vers l'infini que nous allons noter $\langle \langle \cdot, \cdot \rangle \rangle(\cdot)$. Ce remplaçant doit prendre les mêmes arguments, c'est-à-dire une fonction continue g , une mesure $\mathcal{M}$ sur E et un temps t ; de plus si $\mathcal{M}$ est la loi de la variable aléatoire $\bar{\mathcal{X}}$ sur E , le crochet double doit retourner le réel $\langle \langle g, \mathcal{M} \rangle \rangle(t) := \mathbb{E} (g(\bar{\mathcal{X}}(t)))$. On voit bien qu'il s'agit du même objet, le but du changement de notation est de bien mettre en évidence quand le crochet s'applique à une mesure empirique et quand le crochet s'applique à une loi.

Notons finalement que le crochet double que nous allons définir est un objet plus simple puisqu'il agit sur une mesure alors que le crochet simple agit sur une mesure empirique, c'est-à-dire sur une variable aléatoire à valeur mesure.

Définition du crochet double Comme nous venons de le voir, le crochet simple agit sur les mesures empiriques ou bien sur les systèmes de particules dont elles sont issues. Nous définissons un nouveau crochet que nous ne ferons agir que sur les mesures lois.

Définition 4.17. Soit $\overline{\mathcal{X}}$ une particule (c'est-à-dire une variable aléatoire) fixée dont la loi vaut $\mathcal{M}$.

Nous nommons CROCHET MARGINAL DOUBLE avec pour argument une fonction et une mesure :

$$\langle\langle g, \mathcal{M} \rangle\rangle(t) := \langle g, \mathcal{M} \rangle(t).$$

De plus le CROCHET MARGINAL DOUBLE avec pour argument une fonction et une variable aléatoire vaut :

$$\langle\langle \cdot, \cdot \rangle\rangle := \left(\begin{array}{ccc} \mathcal{C}_b(\mathbb{R}^k, \mathbb{R}) \times \mathbb{L}_1(\Omega, E) & \rightarrow & \mathbb{D}([0, T], \mathbb{R}) \\ (g, \overline{\mathcal{X}}) & \mapsto & t \mapsto \mathbb{E}(g(\overline{\mathcal{X}}(t))) \end{array} \right) \quad (4.3)$$

c'est-à-dire

$$\langle\langle g, \overline{\mathcal{X}} \rangle\rangle(t) := \langle g, \mathcal{L}(\overline{\mathcal{X}}) \rangle(t) = \mathbb{E}(g(\overline{\mathcal{X}}(t))).$$

Discussion Nous avons laissé le temps hors du crochet alors qu'on pourrait penser plus simple d'écrire $\langle\langle g, \overline{\mathcal{X}}(t) \rangle\rangle$. La raison est que le fait que le crochet dépende de $\overline{\mathcal{X}}(t)$, $\overline{\mathcal{X}}(t-r)$ ou tout autre dépendance du passé ne changerait pas grand chose dans les démonstrations. Aussi nous avons souhaiter garder une forme plus générale pour la notation.

4.3.4 Fonctions prévisibles de E

Définition Nous allons avoir besoin de fonctions d'éléments de E . Il sera nécessaire que ces fonctions vérifient une certaine notion de prévisibilité, c'est à dire que « $f(x, t)$ où $x \in E$ et $t \in \mathbb{R}$ est en fait uniquement fonction de la restriction $x|_{[0, t]}$ ». Une manière formelle d'écrire ceci tient dans la définition suivante :

Définition 4.18. Une fonction $f : E \times [0, T] \rightarrow \mathbb{R}^k$ est dite PRÉVISIBLE si pour tout processus prévisible X , $Y_t(\omega) = f(X(\omega), t)$ est un processus prévisible.

Par exemple $f(x, t) = x(t)$ et $f(x, t) = x(t - 1/2 \vee 0)$ sont prévisibles au sens de la définition précédente, mais pas $f(x, t) = x(t + 1/2 \wedge 1)$.

4.4 Formulation du problème

4.4.1 Notations et définitions

Condition initiale et génération du hasard Considérons le système infini de particules X_i^N . Les particules vont évoluer en fonction de leur trajectoire passée et d'une ressource partagée selon l'équation différentielle stochastique que nous avons présentée au deuxième paragraphe que nous allons récrire dans le problème 4.1.

Définition 4.19. Nous nous donnons les variables indépendantes suivantes :

- une CONDITION INITIALE : $(x_i)_{i \in \mathbb{N}}$, des variables aléatoires à valeurs dans $\mathbb{R}^k$,
- des SAUTS À INTENSITÉ : (J_i^N) , des processus de Poisson réels d'intensité 1.

En réalité pour un i fixé il n'est pas nécessaire que les $(J_i^N)_{N \geq i}$ soient indépendants car ils n'interviendront pas dans les mêmes systèmes d'équations : par exemple on pourrait dire que $J_i^i = \dots = J_i^N = \dots$ pour tout entier i dans la mesure où les particule numéro (i, N_1) et (i, N_2) n'interagissent pas pour $N_1 \neq N_2$.

Définition 4.20. Soit $(\Omega, \mathbb{P})$, un espace supportant l'ensemble des $(x_i)_{i \in \mathbb{N}}$ et des (J_i^N) . Nommons la filtration engendrée $\mathcal{F}_t = \sigma(x_i, J_i^N)_{|[0, t]}$.

Forme particulière des fonctionnelles étudiées Nous voulons que la particule que nous étudions connaisse uniquement sa trajectoire et celle du groupe (c'est-à-dire de la ressource partagée) jusqu'au temps t . Ceci nous mène à poser :

Définition 4.21. Une fonction mesurable $f : E \times \mathbb{D}([0, T], \mathcal{P}(\mathbb{R}^k)) \times [0, T] \rightarrow \mathbb{R}^k$ est dite FONCTION DE LA MARGINALE si elle est de la forme :

$$f(x, M, t) = f_e(x, \langle f_i, M \rangle(\cdot))(t).$$

où $f_i \in \mathcal{C}_b(\mathbb{R}^k, \mathbb{R})$ et f_e est une fonction déterministe prévisible au sens de la définition 4.18 qui prend en argument deux trajectoires de $[0, T] \rightarrow \mathbb{R}$ et a des valeurs dans E .

Nous avons posé dans cette définition que la fonction f devait prendre ses valeurs dans E , c'est-à-dire être càdlàg par rapport au temps si les deux premières variables sont càdlàg.

4.4.2 Une première formulation du problème

Problème 4.1. Nous étudions le système infini de particules (X_i^N) à valeurs dans $E = \mathbb{D}([0, T], \mathbb{R}^k)$. Notons $(M^N)_{N \in \mathbb{N}}$ la mesure empirique de X^N .

Les processus et variables aléatoires (x_i) et (J_i^N) sont donnés conformément à la définition 4.19.

Définissons les fonctions $f, l, I : E \times \mathcal{P}(E) \rightarrow E$. Ces fonctions sont supposées fonctions de la marginale au sens défini plus haut et toutes L -Lipschitz pour un $L \geq 0$ donné (c'est-à-dire que les fonctions $f_i, f_e, l_i, l_e, I_i, I_e$ sont L -Lipschitz pour la norme uniforme sur les trajectoires et la norme infinie sur $\mathbb{R}^k$)

(X_i^N) vérifie les équations différentielles stochastiques couplées suivantes : pour $t \geq 0$

$$X_i^N(t) = x_i + \int_0^t f(X_i^N, M^N)(s^-) ds + \int_{s=0}^t l(X_i^N, M^N)(s^-) dJ_i^N(\Lambda_i^N(s)) \quad (4.4)$$

où le processus d'intensité des sauts Λ_i^N vaut :

$$\Lambda_i^N(s) = \int_0^s I(X_i^N, M^N)(u^-) du \quad (4.5)$$

Remarquons que le Théorème 9.1 dans [73] assure l'existence et l'unicité des solutions au sens trajectoriel (fortes) pour ce système fini d'équations différentielles stochastiques.

Nous souhaitons démontrer un théorème du type :

Conjecture 4.1. CONVERGENCE FAIBLE DES MESURES EMPIRIQUES VERS UNE LIMITE DÉTERMINISTE *Il existe $\mu \in \mathcal{P}(E)$ telle que $M^N \rightarrow^{\mathcal{L}} \delta_\mu$. En particulier les lois des marginales M_t^N convergent vers $\delta_{\mu_t} \in \mathcal{P}(\mathcal{P}(\mathbb{R}^k))$.*

De plus nous souhaitons donner un sens trajectoriel à cette convergence.

4.4.3 Hypothèses techniques

Nous ajoutons deux hypothèses particulières, la première évite un grand nombre de difficultés techniques même si elle est extrêmement restrictive. La deuxième assure le minimum que l'on peut vouloir, c'est-à-dire une convergence pour la condition initiale.

Hypothèse 4.1. CARACTÈRE BORNÉ DES ÉVOLUTIONS *Il existe une valeur b_X telle que les évolutions ne peuvent sortir avant T de la boule de $\mathbb{R}^k$ de centre 0 et de diamètre b_X (pour la norme uniforme).*

Hypothèse 4.2. CONVERGENCE DE LA CONDITION INITIALE DU SYSTÈME INFINI DE PARTICULES *Les $(x_i)_{i \in \mathbb{N}}$ vérifient les hypothèses suivantes :*

- conformément à l'hypothèse précédente $\forall i \in \mathbb{N}, \|x_i\|_\infty \leq b_X$
- il existe une mesure limite μ_0 telle qu'on ait la convergence presque sûre des mesures empiriques $\mathcal{P}(\mathbb{R}^k)$:

$$\frac{1}{N} \sum_{i=1..N} \delta_{x_i(\omega)} \rightarrow^* \mu_0 \text{ ps.}$$

4.5 Couplage et norme

4.5.1 Norme de convergence d'un système infini de particules

Définition 4.22. X un système infini de particules et $\mathcal{X} = (\mathcal{X}_1, \dots, \mathcal{X}_i, \dots)$ une suite de variables aléatoires de E . On dit que X converge vers $\mathcal{X}$ pour la CONVERGENCE CESARÒ-AUXILIAIRE sur $[0, T]$ si et seulement si :

$$\|X^N - \mathcal{X}\|_{[0, T]} := \mathbb{E} \left(\frac{1}{N} \sum_{i=1}^N \sup_{0 \leq \tau \leq T} |X_i^N(\tau) - \mathcal{X}_i(\tau)| \right) \xrightarrow{N \rightarrow \infty} 0.$$

On écrira : $X \xrightarrow{\|\cdot\|_{[0, T]}} \mathcal{X}$.

Par l'intermédiaire de la norme $\mathbb{L}^1$ -Cesarò sur $[0, T]$, $\mathbb{D}$ est étudié pour la norme infinie et non pour la norme de Skorokhod.

Proposition 4.12. CONVERGENCE CESARÒ-AUXILIAIRE ET CONVERGENCE DES MESURES MARGINALES

HYPOTHÈSES : Soit X uniformément borné et $\mathcal{X} = (\mathcal{X}_1, \dots, \mathcal{X}_i, \dots)$ tels que $\|X^N - \mathcal{X}\|_{[0,T]} \rightarrow 0$.

On suppose que pour tout t , $\mathcal{M}_t^N = \frac{1}{N} \sum_{i=1}^N \delta_{\mathcal{X}_i(t)}$ converge faiblement en probabilité vers une limite $\mathcal{M}_t$ déterministe.

CONCLUSION : de manière uniforme pour $t \in [0, T]$, $M_t^N \rightarrow^* \mathcal{M}_t$ en probabilité (pour la distance d^* qui engendre la topologie de la convergence faible).

Preuve : Soit une fonction g L_g -Lipschitz bornée. Notons pour commencer que par convergence dominée :

$$\lim_{N \rightarrow \infty} \mathbb{E} \frac{1}{N} \sum_{i=1}^N g(\mathcal{X}_i(t)) = \mathbb{E} \langle g, \mathcal{M}_t \rangle.$$

En notant $\overline{\lim}$ la limite supérieure, il vient :

$$\begin{aligned} & \overline{\lim}_{N \rightarrow \infty} |\mathbb{E} \langle g, M_t^N \rangle - \mathbb{E} \langle g, \mathcal{M}_t \rangle| \\ & \leq \overline{\lim}_{N \rightarrow \infty} \left| \frac{1}{N} \sum_{i=1}^N \mathbb{E} [g(X_i^N(t)) - g(\mathcal{X}_i(t))] \right| + \overline{\lim}_{N \rightarrow \infty} \left| \frac{1}{N} \sum_{i=1}^N \mathbb{E} g(\mathcal{X}_i(t)) - \mathbb{E} \langle g, \mathcal{M}_t \rangle \right| \\ & \leq L_g \overline{\lim}_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N \mathbb{E} |X_i^N(t) - \mathcal{X}_i(t)| + 0 \\ & \leq L_g \overline{\lim}_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N \mathbb{E} \sup_{\tau \leq t} |X_i^N(\tau) - \mathcal{X}_i(\tau)| \\ & = L_g \overline{\lim}_{N \rightarrow \infty} \|X^N - \mathcal{X}\|_{[0,T]} = 0. \end{aligned}$$

■

4.5.2 Couplage du système initial

Nous avons défini deux types de variables aléatoires :

- la CONDITION INITIALE : $(x_i)_{i \in \mathbb{N}}$,
- les SAUTS À INTENSITÉ : (J_i^N) ,

Nous aurons donc deux volets de couplage. Nous allons reprendre le problème tenant compte de ces nouvelles notations.

Condition initiale

La condition initiale est déjà couplée puisque $X_i^N = x_i$.

Sauts à intensité

Soit Υ_i une suite de processus de Poisson indépendants dans le plan d'intensité 1. Nous allons modifier la définition des J_i^N du problème 4.1. À la place de $J_i^N(\Lambda(t))$, nous allons faire figurer :

$$J_i^N(t) = \int \int_{(u,v) \in [0,t] \times \mathbb{R}} \chi_{[0, I_i^N(u)]}(v) d\Upsilon(u, v)$$

que nous noterons dans ce qui suit ainsi :

$$J_i^N(t) = \int_0^t \int_0^{I(u)} d\Upsilon(u, v).$$

4.5.3 Reformulation du problème

Problème 4.2. *Nous étudions le système infini de particules (X_i^N) à valeurs dans $E = \mathbb{D}([0, T], \mathbb{R}^k)$. Notons $(M^N)_{N \in \mathbb{N}}$ la mesure empirique de X^N .*

Nous nous donnons les processus et variables aléatoires indépendants suivants, (x_i) à valeurs dans $\mathbb{R}^k$, (Υ_i) Poisson du plan. Les (x_i) vérifient l'hypothèse 4.2.

Soient les fonctions $f, l, I : E \times \mathcal{P}(E) \rightarrow E$. Ces fonctions sont supposées fonctions de la marginale au sens défini plus haut et toutes L -Lipschitz pour un $L \in \mathbb{R}$ donné.

(X_i^N) vérifie les équations différentielles stochastiques couplées suivantes : pour $t \geq 0$

$$X_i^N(t) = x_i + \int_0^t f(X_i^N, M^N)(s^-) ds + \int_{s=0}^t \int_0^{I_i^N(s^-)} l(X_i^N, M^N)(s^-) d\Upsilon_i(u, s). \quad (4.6)$$

où l'intensité instantanée des sauts vaut :

$$I_i^N(s) = I(X_i^N, M^N)(s).$$

4.6 Discussion autour du problème étudié

Le problème posé est celui de l'étude de l'interaction d'un grand nombre de particules qui sont uniquement interdépendantes par l'intermédiaire de leur distribution empirique. Typiquement, on peut imaginer le problème de l'interaction gravitationnelle où des particule interagissent avec le centre de gravité du système qu'elles forment.

Dans la suite, nous utiliserons les résultats de cette partie pour modéliser le comportement d'utilisateurs d'Internet qui interagissent par le partage de la bande passante disponible sur le réseau. De plus nous serons amenés à traiter des problèmes de délais et de frontières. Il ne serait pas difficile de permettre des délais fixes. Par contre dans le cas de délais variables, les problèmes deviennent très spécifiques surtout dans le cas de traversée de frontières. C'est pourquoi nous ne parlerons que brièvement de cette question pour nous concentrer sur les problèmes de convergences et d'effets de bords. Nous traiterons de manière exhaustive cette difficulté dans le cas spécifique du notre modèle de TCP.

Chapitre 5

Système auxiliaire dans $\mathbb{D}$

Nous démontrons ici les théorèmes attendus pour résoudre le problème 4.2. Dans un premier temps nous introduisons la particule du champ moyen comme solution d'une équation différentielle stochastique associée au problème. Dans un deuxième temps nous introduisons un système auxiliaire où les particules évoluent de manière indépendante et interagissent seulement avec la particule du champ moyen.

Nous étudions ensuite le système auxiliaire et démontrons en particulier la loi des grands nombres qu'il vérifie. Enfin nous prouvons que le système original tend vers le système auxiliaire qui est asymptotiquement distribué comme la particule du champ moyen.

Dans ce chapitre nous avons montré certains résultats qui n'étaient pas strictement nécessaires mais qui méritaient cependant de figurer. Nous avons fait suivre les sections concernées par des (*).

Sommaire

5.1	Existence et étude de la particule du champ moyen	74
5.1.1	Définition de la particule du champ moyen	74
5.1.2	Un lemme technique	74
5.1.3	Existence et unicité forte des solutions de l'EDS du champ moyen	76
5.1.4	Compatibilité par rapport aux représentations*	78
5.2	Système auxiliaire : définition et propriétés	79
5.2.1	Équations distributionnelles et continuité du champ moyen*	79
5.2.2	Flot de solutions et définition du système auxiliaire	81
5.2.3	Quelques caractéristiques du flot	81
5.2.4	Loi des grands nombres pour le système auxiliaire	82
5.3	Convergence du système original vers le système auxiliaire	85
5.3.1	Un lemme technique	85
5.3.2	Le théorème de convergence	86
5.3.3	Conséquences de la convergence	86
5.4	Discussion et prolongements*	86
5.4.1	Particules couplées k à $k^{(*)}$	87
5.4.2	Dépendance dans le passé*	87

5.1 Existence et étude de la particule du champ moyen

5.1.1 Définition de la particule du champ moyen

Dans la définition suivante M et $\mathcal{M}$ sont des mesures sur E , mais l'emploi de la notation M sous-entend qu'il s'agit d'une mesure empirique alors que $\mathcal{M}$ correspond à la loi d'une variable aléatoire.

Définition 5.1. Une fonction f de la forme $f(X, M) = f_e(X, \langle f_i, M \rangle)$ étant donnée, on définit $\bar{f}$, la FONCTION AUXILIAIRE ASSOCIÉE À f :

$$\bar{f} := (X, \mathcal{M}) \mapsto f_e(X, \langle \langle f_i, \mathcal{M} \rangle \rangle).$$

Rappelons que nous avons défini au chapitre précédent le crochet marginale double (cf. 4.3.3), $\langle \langle f_i, \mathcal{M} \rangle \rangle := \int f d\mathcal{M}$. Ce crochet peut aussi être appliqué à une variable aléatoire $\bar{\mathcal{X}}$ dont la loi est $\mathcal{M} : \langle \langle f_i, \bar{\mathcal{X}} \rangle \rangle := \mathbb{E}(f_i(\bar{\mathcal{X}})) = \langle \langle f_i, \mathcal{M} \rangle \rangle$.

Deux notations sont alors possibles pour $\bar{f} : \bar{f}(X, \mathcal{M})$ ou bien $\bar{f}(X, Y)$ pour une particule Y . Nous utiliseront plutôt la dernière notation quand elle n'est pas ambiguë. Par contre quand Y est de même loi que X , cela mènerait à écrire $\bar{f}(X, X)$ où le premier X désigne $X(\omega)$ et le second désigne X (c'est à dire toutes les réalisations), ce qui n'est pas particulièrement clair. Aussi dans ce cas nous noterons plutôt $\bar{f}(X, \mathcal{L}(X))$ ou bien $\bar{f}(X, \mathcal{M})$ pour bien différencier les objets quand cela est nécessaire.

Définition 5.2. Soient les variables aléatoires et processus indépendants :

- Υ un processus de Poisson dans le plan,
- $\bar{\mathcal{X}}(0)$ une variable aléatoire à valeurs dans $\mathbb{R}^k$ de loi μ_0 ,

L'ÉQUATION DU CHAMP MOYEN («MASTER EQUATION») associée au problème (4.2) est l'équation différentielle stochastique :

$$\bar{\mathcal{X}}(t) = \bar{\mathcal{X}}(0) + \int_0^t \bar{f}(\bar{\mathcal{X}}, \mathcal{M})(s^-) ds + \int_0^t \int_0^{\bar{I}(\bar{\mathcal{X}}, \mathcal{M})(s^-)} \bar{l}(\bar{\mathcal{X}}, \mathcal{M})(s^-) d\Upsilon(u, s). \quad (5.1)$$

où $\mathcal{M}$ est la loi de $\bar{\mathcal{X}}$.

Cette équation est assez différente du système original, en effet elle est plus simple parce qu'on n'étudie plus qu'une unique particule, mais elle pose une difficulté technique particulière car les solutions de cette équation n'ont pas de sens trajectorien.

Définition 5.3. Par définition la PARTICULE DU CHAMP MOYEN $\bar{\mathcal{X}}$ est la solution de l'équation du champ moyen (5.1) lorsque celle-ci existe et est unique.

Nous allons rechercher maintenant les solutions de l'équation.

5.1.2 Un lemme technique

Ce lemme permet de comparer deux itérations de la méthode de Picard-Banach du point fixe qui va nous permettre plus tard de prouver l'existence et l'unicité de solutions pour l'équation (5.1). Mais avant tout précisons la norme employée, $\bar{\mathcal{X}}$ étant une seule particule, il n'est pas choquant de conserver la notation $\|\cdot\|_U$ pour la norme suivante :

Définition 5.4. NORME $\mathbb{L}_1$ -AUXILIAIRE Une particules X (cad une variable aléatoire) à valeurs dans E étant fixée, notons :

$$||X||_U = \mathbb{E}(\sup_{\tau \in U} |X(\tau)|),$$

où $|\cdot|$ désigne la norme uniforme sur $\mathbb{R}^k$.

Notons que la norme $\mathbb{L}_1$ -auxiliaire est une instantiation de la norme Cesarò-auxiliaire avec une seule particule.

Prouvons maintenant un analogue de (4.12). Nous utilisons ici les suites de Cauchy pour caractériser les suites convergentes sans faire référence à leur limite. Il faut bien garder à l'esprit que les espaces que nous manipulons ne sont pas complets.

Proposition 5.1. CONVERGENCE $\mathbb{L}^1$ -UNIFORME ET CONVERGENCE DES MESURES MARGINALES

HYPOTHÈSES : U étant fixé, soit (X_k) une suite de Cauchy pour $||\cdot||_U$.

CONCLUSION : Pour tout $t \in U$, pour toute fonction g , L_g -Lipschitz bornée, $\langle\langle g, X_k \rangle\rangle(t)$ est une suite de Cauchy pour la norme uniforme sur $\mathbb{R}^k$.

Preuve : Pour tout couple $(k, p) \in \mathbb{N}^2$, il vient :

$$|\mathbb{E}(g(X_k(t))) - \mathbb{E}(g(X_{k+p}(t)))| \leq \mathbb{E}|g(X_k(t)) - g(X_{k+p}(t))| \leq L_g \mathbb{E}|X_k - X_{k+p}|$$

Ce qui démontre notre propriété. ■

Lemme 5.1. HYPOTHÈSES : Soit $t_0 < t$, deux réels positifs. Notons :

$$\mathcal{F}(X)(t) = X(t_0) + \int_{t_0}^t \bar{f}(X, \mathcal{L}(X))(s^-) ds + \int_{t_0}^t \int_0^{\bar{I}(X, \mathcal{L}(X))(s^-)} \bar{l}(X, \mathcal{L}(X))(s^-) d\Upsilon(u, s).$$

CONCLUSION :

$$||\mathcal{F}(X) - \mathcal{F}(Y)||_{[t_0, t]} \leq B(X, Y) + C \int_{t_0}^t ||X - Y||_{[t_0, s]} ds$$

avec $B(X, Y) = \mathbb{E}|X(t_0) - Y(t_0)|$ et $C = 2L(1 + \sup(I) + b_X)$.

Preuve : Pour simplifier l'écriture prenons $T_0 = 0$, et omettons les s^- en convenant par exemple que les fonctions de la marginale renvoient toujours la limite à gauche. Nous avons la majoration :

$$\begin{aligned} |\mathcal{F}(X)(\tau) - \mathcal{F}(Y)(\tau)| &\leq |X(0) - Y(0)| + \int_0^\tau |dX(z) - dY(z)| \\ &\leq |X(0) - Y(0)| + \int_0^\tau |f_e(X, \langle\langle f_i, X \rangle\rangle)(z) - f_e(Y, \langle\langle f_i, Y \rangle\rangle)(z)| dz \\ &\quad + \int_0^\tau \int_{\mathbb{R}^+} |\chi_{y \leq \bar{I}(X, \mathcal{L}(X))(z)} l_e(X, \langle\langle l_i, X \rangle\rangle)(z) - \chi_{y \leq \bar{I}(Y, \mathcal{L}(Y))(z)} l_e(Y, \langle\langle l_i, Y \rangle\rangle)(z)| d\Upsilon(y, z). \end{aligned}$$

En appliquant $\mathbb{E}(\sup_{\tau \in [0,t]}(\cdot))$ aux deux membres de l'inégalité, puis à l'aide d'un théorème de martingale (cf [27] page 25) pour la fonction prévisible $z \rightarrow \bar{l}(X, \mathcal{L}(X))(z)$ et l'intensité prévisible $t \rightarrow \bar{I}(X, \mathcal{L}(X))(z)$:

$$\begin{aligned} \|\mathcal{F}(X) - \mathcal{F}(Y)\|_{[0,t]} &\leq \mathbb{E}|X(0) - Y(0)| + \int_0^t L\|X - Y\|_{[0,z]}dz \\ &+ \mathbb{E} \int_0^t \int_0^\infty |\chi_{y \leq I_e(X, \langle \langle I_i, X \rangle \rangle)(z)} l_e(X, \langle \langle l_i, X \rangle \rangle)(z) - \chi_{y \leq I_e(Y, \langle \langle I_i, Y \rangle \rangle)(z)} l_e(Y(z), \langle \langle l_i, Y \rangle \rangle)| dy dz \end{aligned}$$

La deuxième ligne est majorée en introduisant le terme $\chi_{y \leq \bar{I}(Y, Y)(z)} \bar{l}(X, \mathcal{L}(X))(z)$

$$\begin{aligned} &\mathbb{E} \int_0^t \int_0^\infty |\chi_{y \leq \bar{I}(X, \mathcal{L}(X))(z)} \bar{l}(X, \mathcal{L}(X))(z) - \chi_{y \leq \bar{I}(Y, \mathcal{L}(Y))(z)} \bar{l}(Y, \mathcal{L}(Y))(z)| dy dz \\ &\leq \mathbb{E} \int_0^t |\bar{I}(X, \mathcal{L}(X))(z) - \bar{I}(Y, \mathcal{L}(Y))(z)| |\bar{l}(X, \mathcal{L}(X))(z)| dz \\ &\quad + \mathbb{E} \int_0^t |\bar{I}(Y, \mathcal{L}(Y))(z)| |\bar{l}(X, \mathcal{L}(X))(z) - \bar{l}(Y, \mathcal{L}(Y))(z)| dz \\ &\leq 2(\sup(I)L + b_X L) \int_0^t \|X - Y\|_{[0,z]} dz. \end{aligned}$$

D'où le résultat. ■

5.1.3 Existence et unicité forte des solutions de l'EDS du champ moyen

La vue du lemme indique qu'il va falloir introduire des hypothèses supplémentaires pour traiter le problème des temps d'arrêt. Nous repoussons ce problème au chapitre 6.

Théorème 5.1. L'EDS «MASTER EQUATION» DU CHAMP MOYEN (5.1) associée au problème (4.2) a une solution forte unique $\bar{\mathcal{X}}$.

L'existence et l'unicité des solutions de notre EDS repose sur la démonstration classique du théorème de Cauchy-Lipschitz avec suite de Cauchy et point fixe de Picard-Banach (dite aussi méthode des approximations successives) à ceci près que nous devons faire avec des espaces qui ne sont pas complets pour les normes considérées. Nous ramenons donc la convergence dans un espace complet sur l'espérance de la marginale et étendons cette convergence jusqu'à obtenir une solution forte.

Preuve :

La méthode classique Étudions la suite de processus à valeurs dans E ainsi définie :

$$\begin{cases} X_0(t) &= \overline{\mathcal{X}}(0) \text{ pour } t \in \mathbb{R} \\ X_{k+1}(t) &= \mathcal{F}(X_k)(t). \end{cases}$$

où $\overline{\mathcal{X}}(0)$ est une constante du temps qui a pour loi μ_0 donnée dans l'hypothèse 4.2.

Pour $t \leq 0$, la condition initiale est identique pour les deux systèmes. En considérant les processus pour $t > 0$ fixé, d'après le lemme 5.1, où $B(X_k, X_{k+1}) = 0$:

$$\|X_{k+1} - X_k\|_{[0,t]} \leq (Ct)^k \|X_1 - X_0\|_{[0,t]} \quad (5.2)$$

Se ramener à un espace complet pour extraire une limite En particulier X_k est une suite de Cauchy pour $t < \frac{1}{C}$. Malheureusement $\mathbb{D}$ n'est pas complet pour notre norme (il ne l'est même pas pour la norme uniforme). Cependant par la propriété 5.1 intitulée convergence $\mathbb{L}^1$ -Uniforme et convergence des mesures marginales, on voit que toute $\langle\langle g, X_k \rangle\rangle(t)$ est de Cauchy dans $\mathbb{R}^k$. Ceci assure l'existence et l'unicité ponctuelle de la limite de $t \rightarrow \langle\langle g, X_k \rangle\rangle(t)$ sur $[0, T]$ (l'argument n'est valable en toute rigueur que pour T suffisamment petit, ce qui donnera une existence et unicité locale de la limite. Ensuite, le prolongement à une solution sur l'intervalle entier est un problème classique).

Notons $\tilde{f}_i$ (resp. $\tilde{l}_i$ et $\tilde{I}_i$) la limite de $\langle\langle f_i, X_k \rangle\rangle(t)$ (resp. $\langle\langle l_i, X_k \rangle\rangle(t)$ et $\langle\langle I_i, X_k \rangle\rangle(t)$).

De la limite sur $\langle\langle \cdot, \cdot \rangle\rangle$ vers une limite faible Prouvons que les fonctions $g_k := t \rightarrow \langle\langle g, X_k \rangle\rangle(t)$ sont équi-Lipschitz.

$$\begin{aligned} g_{k+1}(t) &= \mathbb{E}(\mathcal{F}(X_k)) = g_k(0) + \int_0^t \mathbb{E}(f_e(X_k, \mathbb{E}(f_i(X_k)))(s))ds + \\ &\quad \int_0^t \mathbb{E}(\bar{I}(X_k, \mathcal{L}(X_k))(s))\bar{l}(X_k, \mathcal{L}(X_k))(s))ds \end{aligned}$$

qui est une intégrale de fonctions bornées (à cause de l'hypothèse (4.1)) donc $(Lb_X + L^2b_X^2)$ -Lipschitz.

Or une limite simple de fonctions L_0 -Lipschitz sur un compact est L_0 -Lipschitz et converge uniformément, ce qui assure que $\tilde{f}_i$, $\tilde{l}_i$ et $\tilde{I}_i$ sont des limites uniformes sur $[0, T]$.

De la limite marginale faible au candidat limite Définissons $\overline{\mathcal{X}}(t)$ comme solution de l'EDS :

$$\overline{\mathcal{X}}(t) = \overline{\mathcal{X}}(0) + \int_0^t f_e(\overline{\mathcal{X}}, \tilde{f}_i)(s)ds + \int_0^t \int_0^{\bar{I}(\overline{\mathcal{X}}, \tilde{I}_i)(s)} \bar{l}(\overline{\mathcal{X}}, \tilde{l}_i)(s)d\Upsilon(u, s). \quad (5.3)$$

Nous nous sommes donc ramenés à une EDS classique qui a une solution forte c'est à dire trajectorielle presque sûre ([73] p. 200).

Validation du candidat limite Vérifions que $X_k \rightarrow \overline{\mathcal{X}}$ au sens de la norme $\mathbb{L}_1$ -auxiliaire. Comme dans la démonstration du Lemme nous avons (mutatis mutandis) :

$$\begin{aligned} \|X_{k+1} - \overline{\mathcal{X}}\|_{[0,t]} &= 0 + \mathbb{E} \int_0^t L \left(|X_k(s) - \overline{\mathcal{X}}(s)| + \left| \langle \langle f_i, X_k \rangle \rangle(s) - \tilde{f}_i(s) \right| \right) ds + \dots \\ &\leq B_k + C \int_0^t \|X_k - \overline{\mathcal{X}}\|_{[0,s]} ds \end{aligned}$$

où

$$B_k = C \mathbb{E} \int_0^T \left(\left| \langle \langle f_i, X_k \rangle \rangle(s) - \tilde{f}_i(s) \right| + \left| \langle \langle l_i, X_k \rangle \rangle(s) - \tilde{l}_i(s) \right| + \left| \langle \langle I_i, X_k \rangle \rangle(s) - \tilde{I}_i(s) \right| \right) ds$$

et

$$C = L(1 + b_X + \sup(I)).$$

Or par le principe de convergence dominée, $B_k \xrightarrow[k \rightarrow \infty]{} 0$, ce qui permet de conclure la démonstration puisqu'on prouve ainsi que $\overline{\mathcal{X}}$ est une limite extraite de X_k pour notre convergence $\mathbb{L}_1$ -auxiliaire.

Reste un mot à dire pour compléter l'argument car la convergence $\mathbb{L}_1$ -auxiliaire n'assure pas à elle seule la convergence simple trajectorielle, mais seulement une convergence uniforme en probabilité. Cependant l'unicité de la limite est valable presque sûrement, et d'autre part nous avons bien défini un système qui a un sens point à point.

Nous venons donc de démontrer que (5.1) a une unique solution forte $\overline{\mathcal{X}}$. ■

Remarque 5.1. Du point de vue pratique pour traiter le problème original, il n'est pas utile de montrer l'existence d'une solution forte, on pouvait se contenter d'extraire les limites $\tilde{f}_i$, $\tilde{I}_i$ et $\tilde{l}_i$.

5.1.4 Compatibilité par rapport aux représentations*

Le premier fait à remarquer est que les solutions données par le théorème 5.1 ont la propriété suivante :

Proposition 5.2. Soit $\overline{\mathcal{X}}(0)$ et $\overline{\mathcal{Y}}(0)$ deux conditions initiales de même loi de l'équation (5.1) et Υ_X et Υ_Y deux Poisson dans le plan d'intensité 1, toutes choses étant égales par ailleurs¹⁶. Alors les solutions $\overline{\mathcal{X}}$ et $\overline{\mathcal{Y}}$ données par le théorème 5.1 ont même loi.

Preuve : Il suffit de remarquer que les fonctions $\tilde{\square}$ ¹⁷ sont identiques pour $\overline{\mathcal{X}}$ et $\overline{\mathcal{Y}}$. C'est le cas parce que si X et Y ont même loi alors $\mathcal{F}(X)$ et $\mathcal{F}(Y)$ ont même loi (ils sont définis par des dynamiques aux mêmes lois). Ainsi les suites qui permettent d'extraire les limites sont en faite exactement identiques. D'où l'identité de leurs limites.

¹⁶c'est-à-dire que f , l et I sont identiques.

¹⁷ $\tilde{\square}$ désigne l'argument d'une fonction, par exemple ici la fonction tilde qui à une fonction f associe $\tilde{f}$ est désignée par $\tilde{\square}$

Mais alors les équations (5.3) sont identiques à la condition initiale près. Il est alors connu que les solutions fortes sont à fortiori des solutions faibles indépendantes du représentant de la condition initiale ([73] p198).

■

Corollaire 5.1. *Une loi initiale μ étant fixée, l'équation du champ moyen (5.1) associée au problème 4.2 a une unique solution faible.*

5.2 Système auxiliaire : définition et propriétés

Rappelons que les sections suivies d'une * sont facultatives et ne font pas partie du parcours critique de la preuve.

5.2.1 Équations distributionnelles et continuité du champ moyen*

Nous donnons dans cette section une méthode qui peut permettre de déduire les équations différentielles vérifiées par la distribution du champ moyen. Cependant le problème dépasse l'objectif de ces chapitres, et nous n'approfondirons pas plus ici la méthode et nous nous contentons d'un exemple simple. Le travail est fait dans le cas de TCP au chapitre 10. D'autre part une fois le champ moyen établi d'autres méthodes peuvent permettre d'obtenir les équations parfois plus facilement (voir [9] par exemple).

Calcul rapides Soit $g \in \mathcal{C}_b^1(\mathbb{R}^k, \mathbb{R})$ telle que $g(\partial\mathbb{R}^k) = 0$ (pour les conditions aux bords de l'IPP), les fonctions considérées forment une partie dense du domaine du générateur du semi-groupe associé à notre évolution pour la norme $\mathbb{L}_1$ (on consultera [44] chapitre 8 pour cela). On suppose de plus que la fonction $\bar{f}(\mathcal{L}(\bar{\mathcal{X}}))(t)$ ne dépend pas de la première variable, que la fonction $\bar{l} = l$ est constante. Enfin supposons que la mesure $\mathcal{M}$ admet une densité différentiable pour tout temps que l'on note $\mu(t, x)$:

Voici l'EDP vérifiée par $\mathbb{E}(g(\bar{\mathcal{X}}))$:

$$\begin{aligned} \mathbb{E}(g(\bar{\mathcal{X}}(t))) - \mathbb{E}(g(\bar{\mathcal{X}}(0))) &= \int_0^t \mathbb{E} \left((Dg)_{\bar{\mathcal{X}}(s)} (\bar{f}(\bar{\mathcal{X}}, \mathcal{L}(\bar{\mathcal{X}}))(s^-)) \right) ds \\ &+ \int_0^t \mathbb{E} \int_0^{\bar{l}(\bar{\mathcal{X}}, \mathcal{L}(\bar{\mathcal{X}}))(s^-)} (g[\bar{\mathcal{X}}(s-) + l] - g[\bar{\mathcal{X}}(s-)]) d\Upsilon(u, s). \end{aligned} \quad (5.4)$$

En appliquant la Fubini et une IPP à la partie continue :

$$\begin{aligned}
& \int_0^t \mathbb{E} \left((Dg)_{\overline{\mathcal{X}}(s)} (\overline{f}(\overline{\mathcal{X}}, \mathcal{L}(\overline{\mathcal{X}}))(s^-)) \right) ds \\
&= \int_0^t \int_{\mathbb{R}^k} ((Dg)_x (\overline{f}(\mathcal{M}_s)(s^-)) d\mathcal{M}_s(x)) ds \\
&= \int_0^t \int_{\mathbb{R}^k} \left(\sum_{i=1}^k \frac{\partial g}{\partial x_i}(x) (\overline{f}_i(\mathcal{M}_s)(s^-)) \mu(s, x_1, \dots, x_k) dx_1 \dots dx_k \right) ds \\
&= - \int_{\mathbb{R}^k} g(x) \int_0^t \sum_{i=1}^k \left(\overline{f}_i(\mathcal{M}_s)(s^-) \frac{\partial \mu}{\partial x_i}(t, x_1, \dots, x_k) \right) ds dx_1 \dots dx_k.
\end{aligned}$$

En appliquant un théorème sur l'intensité du Poisson et en réordonnant la partie à sauts :

$$\begin{aligned}
& \int_0^t \mathbb{E} \int_0^{\overline{I}(\overline{\mathcal{X}}, \mathcal{L}(\overline{\mathcal{X}}))(s^-)} (g[\overline{\mathcal{X}}(s-) + l] - g[\overline{\mathcal{X}}(s-)]) d\Upsilon(u, s) \\
&= \int_0^t \int_{\mathbb{R}^k} \overline{I}(x, \mathcal{M}_s) (g[x + l] - g[x]) d\mathcal{M}_s(x) ds \\
&= \int_{\mathbb{R}^k} g(x) \int_0^t \overline{I}(x - l, \mathcal{M}_s) \mu(s, x - l) - \overline{I}(x, \mathcal{M}_s) \mu(s, x) ds dx_1 \dots dx_k.
\end{aligned}$$

Comme on le voit avec l'espérance on obtient sous l'intégrale des quantités réelles bornées, $\mathbb{E}(g(\overline{\mathcal{X}}))(\cdot)$ est donc continue. Mais alors les quantités sous l'intégrale sont du coup toutes continues, on a le caractère continûment dérivable de $\mathbb{E}(g(\overline{\mathcal{X}}))(\cdot)$ (comme dans la théorie classique des équations différentielles, c'est la fonctionnelle uniquement qui impose la limite de dérivabilité). De plus ajoutant la remarque finale que :

$$\begin{aligned}
\mathbb{E}(g(\overline{\mathcal{X}}(t))) - \mathbb{E}(g(\overline{\mathcal{X}}(0))) &= \int_{\mathbb{R}^k} g(x) (\mu(t, x) - \mu(0, x)) dx \\
&= \int_{\mathbb{R}^k} g(x) \int_0^t \frac{\partial \mu}{\partial t}(s, x) ds dx
\end{aligned}$$

(sous réserve de différentiabilité de μ par rapport à t).

Ce qui donne finalement l'EDP suivante :

$$\frac{\partial \mu}{\partial t}(t, x) = \overline{I}(x - l, \mathcal{M}_t) \mu(t, x - l) - \overline{I}(x, \mathcal{M}_t) \mu(t, x) - \sum_{i=1}^k \overline{f}_i(\mathcal{M}_t)(t) \frac{\partial \mu}{\partial x_i}(t, x).$$

Conjecture 5.1. *Soit g une fonction continûment dérivable, si les fonctionnelles sont «dérivables», alors $t \rightarrow \mathbb{E}(g(\overline{\mathcal{X}}(t)))$ et tous les $(t \rightarrow \mathbb{E}(g(\mathcal{X}_i(t))))_{i \in \mathbb{N}}$ sont continûment dérivables et en particulier Lipschitz du temps sur tout borné.*

La dérivabilité des $\mathbb{E}\mathcal{X}_i$ s'obtient à partir de leur équation qui est analogue à (5.4) et de la dérivabilité de $\mathbb{E}(g(\overline{\mathcal{X}}))(\cdot)$. De plus il faut donner un sens correct au mot dérivable pour transformer cette conjecture en théorème.

On peut sans doute conjecturer des forme d'équations aux dérivées partielles en fonction des besoins. Cependant le problème général dépend beaucoup de la forme exacte de la fonctionnelle et paraît difficilement abordable.

5.2.2 Flot de solutions et définition du système auxiliaire

Définition 5.5. *La particule du champ moyen $t \rightarrow \bar{\mathcal{X}}(t)$ étant donnée, définissons le FLOT DU CHAMP MOYEN $\bar{\mathcal{X}}_t(y)$ comme la fonctionnelle qui à une condition initiale quelconque y (une variable aléatoire sur $\mathbb{R}^k$) fait correspondre la solution au temps t de l'équation :*

$$\bar{\mathcal{X}}_t(y) = y + \int_0^t f_e(\bar{\mathcal{X}}_{s-}(y), \langle \langle f_i, \bar{\mathcal{X}} \rangle \rangle) ds + \int_0^t \int_0^{I_e(\bar{\mathcal{X}}_{s-}(y), \langle \langle I_i, \bar{\mathcal{X}} \rangle \rangle)} l_e(\bar{\mathcal{X}}_{s-}(y), \langle \langle l_i, \bar{\mathcal{X}} \rangle \rangle) d\Upsilon(u, s) \quad (5.5)$$

Le corollaire 5.1 permet de définir une notion faible de flot :

Définition 5.6. *Le FLOT FAIBLE DU CHAMP MOYEN $\mathcal{X}_t(\mu)$ est la loi des $\bar{\mathcal{X}}_t(y)$ tels que $\mathcal{L}(y) = \mu$.*

Proposition 5.3. *Le flot de l'équation différentielle du champ moyen vérifie : $\bar{\mathcal{X}}(t) = \bar{\mathcal{X}}_t(\bar{\mathcal{X}}(0))$ et $\mathcal{M} := \mathcal{L}(\bar{\mathcal{X}}) = \bar{\mathcal{X}}_t(\mu_0)$.*

Définition 5.7. *Le SYSTÈME LIMITE AUXILIAIRE associé au problème (4.2) pour lequel la particule du champ moyen $\bar{\mathcal{X}}$ est définie par le théorème 5.1 est le couple composé de :*

- la suite de variables aléatoires $(\mathcal{X}_i)$ à valeurs dans $\mathbb{D}$:

$$\mathcal{X}_i(t) := \bar{\mathcal{X}}_t(x_i)$$

avec $\Upsilon = \Upsilon_i$,
 - et la loi $\mathcal{M} := \mathcal{L}(\bar{\mathcal{X}}) = \bar{\mathcal{X}}_t(\mu_0)$.

Commentaires Le système auxiliaire limite du champ moyen que nous définissons ici est composé d'une suite de particules $\mathcal{X}_i$ et de la loi de la solution de l'équation maîtresse (5.1) $\mathcal{M}$.

Dans la définition de $\mathcal{X}_i$ nous voyons que la particule est couplée avec les X_i^N pour i fixé (même générateur de saut Υ_i et même condition initiale x_i), mais que la ressource partagée qui était représentée par une mesure empirique est remplacée dans l'équation d'évolution par la mesure limite $\mathcal{M}$ que nous venons de calculer.

5.2.3 Quelques caractéristiques du flot

Théorème 5.2. Caractère Lipschitz du flot par rapport à la condition initiale sur tout borné *Le flot $\bar{\mathcal{X}}_t(y)$ est une fonction Lipschitz de y sur tout borné du temps pour la convergence $\mathbb{L}_1$ (espérance de la norme uniforme sur $\mathbb{R}^k$) au départ et $\mathbb{L}_1$ -auxiliaire à l'arrivée.*

Preuve : Par le raisonnement du lemme 5.1 mutatis mutandis, notant $|\cdot|$, la norme uniforme de $\mathbb{R}^k$, il vient :

$$\|\bar{\mathcal{X}}_\tau(y_2) - \bar{\mathcal{X}}_\tau(y_1)\|_{[0,t]} \leq \mathbb{E}|y_2 - y_1| + C \int_0^t \|\bar{\mathcal{X}}_s(y_2) - \bar{\mathcal{X}}_s(y_1)\| ds. \quad (5.6)$$

Ainsi par le lemme de Gronwall (rappelé en annexe B.1) :

$$\|\bar{\mathcal{X}}_\tau(y_2) - \bar{\mathcal{X}}_\tau(y_1)\|_{[0,t]} \leq e^{Ct} \mathbb{E}|y_2 - y_1|. \quad (5.7)$$

■

Corollaire 5.2. LE FLOT FAIBLE MARGINAL DU CHAMP MOYEN EST UNIFORMÉMENT CONTINU PAR RAPPORT À LA DISTRIBUTION INITIALE *dans le sens : Les mesure μ_k ayant toutes leurs valeurs dans un même support borné,*

$$\mu_k \rightarrow^* \mu_\infty \Rightarrow \sup_{t \in [0,T]} d^*(\bar{\mathcal{X}}_t(\mu_k), \bar{\mathcal{X}}_t(\mu_\infty)) \xrightarrow{k \rightarrow \infty} 0.$$

Preuve : Il suffit de le prouver séquentiellement pour des représentants choisis des lois considérées. Utilisons le théorème de représentation de Skorokhod que nous appliquons aux probabilités μ_k pour les représenter par des variables aléatoires y_k de loi μ_k qui convergent presque sûrement vers une limite y de loi μ . La convergence presque sûre sur les y_k devient intégrale par l'intermédiaire du principe de convergence dominée (l'hypothèse de domination fait partie de l'hypothèse et est standard dans notre problème par l'hypothèse (4.2)). Ceci qui achève la démonstration en utilisant le théorème précédent et le fait que la convergence considérée tout comme la norme Césarè-auxiliaire entraîne la convergence des mesures pour d^* (démontré à la proposition 4.12 dans le cas de mesures aléatoires).

■

5.2.4 Loi des grands nombres pour le système auxiliaire

Notons $\mathcal{M}^N := \frac{1}{N} \sum_{i=1}^N \delta_{\mathcal{X}_i}$ et $\mathcal{M}_t^N$, la mesure marginale.

Discussion Alors que jusqu'à maintenant nous n'avions que des lois dans ce chapitres, nous faisons maintenant le lien avec les mesures empiriques. L'avantage du système auxiliaire est que les trajectoires des particules sont indépendantes les une des autres.

Les particules $\mathcal{X}_i$ ont été définies avec des conditions initiales aléatoires x_i sur le flot $\bar{\mathcal{X}}_t$. Lorsque le nombre de ces particules est grand, nous avons supposé (hypothèse 4.2) une convergence presque sûre sur les x_i , c'est à dire que $\mathcal{M}_0^N(\omega) \rightarrow \mu_0$ ps. Le flot étant continu en loi comme nous venons de le voir, il est naturel d'attendre que la loi forte des grands nombres sur la condition initiale va se propager le long du flot et que la limite de la moyenne des flots $\frac{1}{N} \sum \mathcal{X}_i$ sera la moyenne de la limite $\mathbb{E}\bar{\mathcal{X}}(t)$. Cependant dans la pratique n'écrivons pas une version fonctionnelle de la loi forte des grands nombres, mais seulement une version point à point.

Théorème 5.3. LOI FORTE DES GRANDS NOMBRES *Les hypothèses du problème 4.2 étant vérifiées pour toute fonction continue bornée g , $\forall t \in [0, T]$:*

$$\frac{1}{N} \sum_{i=1..N} g(\mathcal{X}_i(t)) \rightarrow_{N \rightarrow \infty} \mathbb{E}(g(\overline{\mathcal{X}}))(t) \text{ ps.} \quad (5.8)$$

ou bien : $\forall t, \mathcal{M}_t^N \rightarrow^* \mathcal{L}(\overline{\mathcal{X}}(t))$ presque sûrement pour la distance d^* .

Preuve :

Structure de la preuve Soit F_t telle que $g(\mathcal{X}_i(t)) = F_t(\Upsilon_i, x_i)$. F_t est une fonction déterministe ; nous montrons au paragraphe suivant qu'elle est continue presque sûrement (par rapport à la mesure produit de $\mathcal{L}(\Upsilon_i)$ et de μ_0).

Notons $R(A, B)$ l'espace d'états des Υ restreints au pavé $A \times B$ de $\mathbb{R}^2$. Remarquons pour commencer que les Υ_i et les x_i sont globalement indépendants. Ainsi la loi forte des grands nombres peut s'écrire comme l'implications 1 vers 2 :

1. la mesure empirique $\frac{1}{N} \sum_{i=1}^N \delta_{x_i(\omega)}$ tend presque sûrement au sens de la convergence faible vers δ_{μ_0} dans $\mathcal{P}(\mathbb{R}^k)$,
2. la mesure empirique $\frac{1}{N} \sum_{i=1}^N \delta_{(\Upsilon_i(\omega), x_i(\omega))}$ tend presque sûrement au sens de la convergence faible vers $\delta_{\mathcal{L}(\Upsilon) \otimes \mu_0}$ dans $\mathcal{P}(R([0, T] \times [0, b_X]) \times \mathbb{R}^k)$.

Utilisons le théorème dit de l'application continue cf. [44] p. 103¹⁸).

Il vient en particulier que $\frac{1}{N} \sum g(\mathcal{X}_i(t)) = \frac{1}{N} \sum F_t(\Upsilon_i, x_i)$ tend vers une limite déterministe presque sûrement. Reste à montrer que la limite est bien $\mathbb{E}(\overline{\mathcal{X}}(t))$ ce qui se fait en inversant une limite et une intégrale et en utilisant le fait que la limite est déterministe. Ceci est formalisé dans le dernier paragraphe de la démonstration.

Preuve de continuité Définissons donc à ω fixé la fonction F_t qui à la mesure de Radon sur $\mathbb{R}^2$, $\Upsilon_i(\omega)$ et au réel $x_i(\omega)$ associe $\mathcal{X}_i(t) = \overline{\mathcal{X}}_t(x_i)(\omega)$.

L'intensité des sauts est donnée par le graphe de $(t, \lambda(t))$ où $\lambda(t) = \overline{I}(\mathcal{X}_i(\omega)(t), \tilde{I}(t))$. Les fonctions étant Lipschitz, on est assuré pour commencer que ce graphe a une mesure nulle dans $\mathbb{R}^2$. Si (a_p, b_p) les points du processus de Poisson tels que $b_p \leq Lb_X$ rangés par ordre lexicographique croissant (L est la constante de Lipschitz et b_X est la borne des trajectoires). Les points (a_p, b_p) sont les points potentiels de sauts, la surface considérée étant $[0, T] \times [0, b_X]$ qui est de mesure finie dans $\mathbb{R}^2$ ces sauts potentiels sont en nombre fini. Quelle est la probabilité que la trajectoire rencontre l'un de ces points ?

¹⁸Théorème de l'application continue

Théorème 5.4. NOTATIONS : Soient (S, d) et (S', d') deux espaces Polonais, et $f : S \rightarrow S'$ mesurable. μ_n étant une suite de probabilités qui tendent faiblement vers μ dans $\mathcal{P}(S)$. Soient $\nu_n = \mu_n f^{-1}$ et $\nu = \mu f^{-1}$ (rappelons que si μ est une mesure sur S , μf^{-1} est la mesure sur S' telle que $\mu f^{-1}(B) = \mu(f^{-1}(B))$)

HYPOTHÈSES : C_f , l'ensemble des points où f est continue vérifie $\mu(C_f) = 1$.

CONCLUSION : Alors $\nu_n \xrightarrow[N \rightarrow \infty]{*} \nu$.

On remarque habituellement que la continuité n'est à vérifier que par rapport à la mesure limite et rien par rapport aux éléments de la suite.

Précisons que rencontrer signifie que l'adhérence de la trajectoire contient le point considéré. Par récurrence Si la probabilité est 0 jusqu'au temps t sur l'ensemble des ω' identiques à ω jusqu'à t , alors comme les incréments sont indépendants du prochain point (a_p, b_p) rencontré (c'est l'hypothèse de prévisibilité *et* le fait que les sauts ont une taille qui ne dépend pas de b_p), la trajectoire passe donc par ce point avec une probabilité nulle.

Regardons alors la distance entre l'adhérence du graphe $(t, \lambda(t))$ et les points. C'est la distance de deux fermés bornés de $\mathbb{R}^2$. Ces ensembles sont donc compacts et la distance strictement positive est atteinte. Notons ϵ cette distance. Les instants de sauts étant fixés, le flot est continu par rapport à la condition initiale donc la trajectoire reste dans le couloir uniforme de taille $\epsilon/2$ autour de $(t, \lambda(t))$ pour un voisinage V_x de la condition initiale. D'autre part pour la convergence faible sur les mesures de Radon dans $\mathbb{R}^2$, notons V_Υ le voisinage où les variations font bouger les points d'au plus $\epsilon/2$. Changer $\Upsilon_i(\omega)$ par n'importe quel point de V_Υ laisse la dynamique inchangée (les instants de sauts sont exactement les mêmes).

Nous avons donc montré que F_t est continue presque sûrement.

Inversion des limites Soit

$$h(t) := \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N g(\mathcal{X}_i(t)).$$

Nous venons de voir que cette limite existe et qu'elle est déterministe. Ceci signifie en prenant l'espérance des deux cotés de l'égalité :

$$h(t) = \mathbb{E}(h(t)) = \mathbb{E} \left(\lim_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N g(\mathcal{X}_i(t)) \right).$$

Comme g et $\mathcal{X}_i$ sont toutes deux bornées, la quantité intégrée est dominée ; ce qui permet d'inverser la limite et l'intégrale et :

$$\begin{aligned} h(t) &= \lim_{N \rightarrow \infty} \mathbb{E} \left(\frac{1}{N} \sum_{i=1}^N g(\mathcal{X}_i(t)) \right) \\ &= \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N \mathbb{E}(g(\mathcal{X}_i(t))). \end{aligned}$$

Une deuxième expression de la limite Remarquons que

$$\mathcal{X}_t(\mathcal{M}_0^N) \stackrel{\mathcal{L}}{=} \frac{1}{N} \sum_{i=1}^N \mathcal{X}_t(x_i),$$

en effet si les x_i sont distincts par exemple, $\mathbb{E}(\mathcal{X}_t(\mathcal{M}_0^N) | \text{condition initiale} = x_i)$ vaut $\mathcal{X}_t(x_i)$, et ceci intervient exactement avec probabilité $1/N$.

Nous avons alors par hypothèse $\mathcal{M}_0^N$ qui tend presque sûrement vers μ_0 (hypothèse 4.2). Alors le corollaire 5.2 dit que $\mathcal{X}_t(\mathcal{M}_0^N) \rightarrow^* \mathcal{L}(\overline{\mathcal{X}}(t))$ uniformément par rapport à t

(en particulier pour tout temps). Cette convergence en loi implique donc par définition que $\langle g, \mathcal{X}_t(\mathcal{M}_0^N) \rangle \xrightarrow[N \rightarrow \infty]{} \mathbb{E}(\overline{\mathcal{X}}(t))$.

Finalement en égalant les deux limites, et en notant que il vient :

$$h(t) = \mathbb{E} (g(\overline{\mathcal{X}}(t))) .$$

■

5.3 Convergence du système original vers le système auxiliaire

5.3.1 Un lemme technique

Définition 5.8. Soit $g = (g_e, g_i)$, de la forme que nous étudions. Notons

$$\mathcal{S}_g^N = \mathbb{E} \int_0^T |\langle g_i, \mathcal{X}_{1\dots N} \rangle(s) - \langle \langle g_i, \overline{\mathcal{X}} \rangle \rangle(s)| ds.$$

Lemme 5.2. Soit $0 < T$, alors $\|X^N - \mathcal{X}\|_{[0,T]} \leq B^N + C \int_{T_{min}}^t \|X^N - \mathcal{X}\|_{[0,s]} ds$ avec

$$\begin{aligned} B^N &= L\mathcal{S}_f^N + \sup(I)L\mathcal{S}_I^N + b_X L\mathcal{S}_I^N, \\ C &= 2(L + L \sup(I) + b_X L). \end{aligned}$$

Preuve : La preuve est calquée sur celle du lemme 5.1 en utilisant de temps en temps $(\mathcal{X}_i)_i$ comme intermédiaire.

$$\begin{aligned} |X_i^N(\tau) - \mathcal{X}_i(\tau)| &\leq \int_0^\tau |f_e(X_i^N, \langle f_i, X^N \rangle)(z) - f_e(X_i^N, \langle \langle f_i, \overline{\mathcal{X}} \rangle \rangle)(z)| dz \\ &\quad + \int_0^\tau \int_0^\infty \left| \chi_{y \leq I_e(X_i^N, \langle I_i, X^N \rangle)}(z) l_e(X_i^N(z), \langle l_i, X^N \rangle) - \chi_{y \leq \bar{I}(\mathcal{X}_i, \overline{\mathcal{X}})} \bar{l}(\mathcal{X}_i(z), \overline{\mathcal{X}}) \right| \Upsilon_i(dy, dz) \end{aligned}$$

En appliquant $\frac{1}{N} \sum_{i=1}^N \mathbb{E} \sup_{[0,t]}$ aux deux membres de l'inégalité et en intercalant $\langle f_i, (\mathcal{X})_{1\dots N} \rangle$ entre $\langle f_i, X^N \rangle$ et $\langle \langle f_i, \overline{\mathcal{X}} \rangle \rangle$:

$$\|X^N - \mathcal{X}\|_{[0,t]} \leq \int_0^t L \|X^N - \mathcal{X}\|_{[0,z]} dz + L \int_0^t \mathbb{E} |\langle f_i, (\mathcal{X})_{1\dots N} \rangle(z) - \langle \langle f_i, \overline{\mathcal{X}} \rangle \rangle(z)| dz + \dots$$

Le reste est une copie mutatis mutandis du lemme 5.1.

■

5.3.2 Le théorème de convergence

Théorème 5.5. CONVERGENCE VERS LE CHAMP MOYEN

HYPOTHÈSES : Soit X un système infini de particules associé au problème 4.2.

CONCLUSION : Alors $X \rightarrow (\mathcal{X}_i)$ au sens Cesarò-auxiliaire.

Preuve : En utilisant le lemme 5.2

$$\|X^N - \mathcal{X}\|_t \leq (L\mathcal{S}_f^N + \sup(I)L\mathcal{S}_I^N + b_X L\mathcal{S}_I^N)e^{Ct}. \quad (5.9)$$

De plus les $\mathcal{S}^N$ tendent vers 0 par convergence dominée (appliqué deux fois après avoir interverti l'espérance et l'intégrale) car toutes nos fonctions sont bornées définies sur des bornés et qu'on a la limite presque sûre de la loi des grands nombres démontrée au théorème 5.3. ■

5.3.3 Conséquences de la convergence

Le théorème précédent et les propriétés précédentes peuvent alors s'énoncer ainsi en faisant référence au problème original :

Corollaire 5.3. HYPOTHÈSES : dans le cadre du problème 4.2, les conclusions sur le système initial sont les suivantes :

CONCLUSION 1 : soit f continue bornée, fixée, alors $\langle f, X^N \rangle(t) \xrightarrow[N \rightarrow \infty]{} \langle \langle f, \overline{\mathcal{X}} \rangle \rangle(t)$ dans

$\mathbb{L}_1$ pour tout t . De plus, la limite $\langle \langle f, \overline{\mathcal{X}} \rangle \rangle(t)$ est continue et déterministe.

CONCLUSION 2 : la mesure empirique marginale $M^N(t)$ de X converge en probabilité (pour d^*) vers la limite déterministe $\mathcal{M}(t) := \mathcal{L}(\overline{\mathcal{X}}(t))$: c'est-à-dire $\|M^N - \mathcal{M}\|_w \xrightarrow[N \rightarrow \infty]{} 0$ en probabilité.

CONCLUSION 3 : la limites $\langle \langle f, \overline{\mathcal{X}} \rangle \rangle(t)$ et $\mathcal{M}$ satisfont des équations (5.1) et (5.4)

5.4 Discussion et prolongements*

Nous ne parlerons pas pour l'instant de passages de frontières ni de temps d'arrêts sur les trajectoires que permet une solution forte à notre problème. Ces questions sont en effet envisagées au chapitre suivant. La discussion portera ici plus sur la forme très particulière que nous avons choisie pour les fonctions $f = (f_e, f_i)$. En effet les autres restrictions de type Lipschitz ou borné, même si elles peuvent être un peu relaxées, sont les «bonnes» hypothèses à faire. Ce n'est pas le cas pour les fonctions d'interaction entre le groupe et une particule.

5.4.1 Particules couplées k à $k^{(*)}$

Écriture formelle Soit une fonction de k variables $G(x_1, \dots, x_k)$. Cherchons ce qui se passe si on remplace $\langle \square, X^N \rangle$ par

$$\langle G, X^N \rangle_k := \frac{1}{N!} \sum_{\sigma \in \mathfrak{S}_N} G(X_{\sigma(1)}^N, \dots, X_{\sigma(k)}^N)$$

où $\mathfrak{S}_N$ est l'ensemble des permutations de $[1, N]$.

Pour faire simple si $G(x_1, x_2) = (x_1 - x_2)^2$, alors on peut définir l'analogue $\langle \langle \cdot, \cdot \rangle \rangle_2$ de $\langle \langle \cdot, \cdot \rangle \rangle$ ainsi :

$$\langle \langle G, X \rangle \rangle_2 := 4\mathbb{E}(X^2) - 2\mathbb{E}(X)^2$$

en remarquant que¹⁹ :

$$\begin{aligned} \langle G, X^N \rangle_2 &= \frac{1}{N(N-1)} \sum (x_1 - x_2)^2 \\ &= \frac{2}{N} \sum x_1^2 - \frac{2}{N(N-1)} \sum x_1 x_2 \\ &= 2 \left(\frac{1}{N} \sum x_1^2 \right) - 2 \left\{ \left(\frac{1}{N} \sum x_1 \right)^2 - \left(\frac{1}{N} \sum x_1^2 \right) \right\} \\ &= 4 \left(\frac{1}{N} \sum x_1^2 \right) - 2 \left(\frac{1}{N} \sum x_1 \right)^2. \end{aligned}$$

Explications Ce type de dépendance est beaucoup plus compliqué que la simple intégrale par rapport à la mesure marginale. Cependant nous voyons que dans les cas polynomiaux la simplification est simple à réaliser. Dans le cas général la transformation est beaucoup plus délicate à écrire. Nous laissons la recherche d'applications de ceci ainsi que la généralisation aux soins du lecteur.

5.4.2 Dépendance dans le passé*

$f = (f_e, f_i)$ pouvait être gardé de la même forme, mais avec f_i telle que $\langle f_i, M^N \rangle(\cdot)$ est prévisible Lipschitz et que la fonction $\langle \langle f_i, X \rangle \rangle(\cdot)$ est Lipschitz du temps. Ce sont en effet les deux seules propriétés que nous utilisons dans la démonstration. Une illustration simplifiée de la dépendance dans le passé quand elle est exclusive (on ne dépend pas du passé proche) est donnée dans la partie sur RED.

¹⁹conformément à la notation traditionnelle des sommations symétriques, nous écrivons $\sum x_1 x_2$ par exemple pour désigner $\sum_{\sigma \in \mathfrak{S}_N / \mathfrak{S}_{N-2}} x_{\sigma(1)} x_{\sigma(2)}$. De sorte que $\sum x_1 x_2$ désigne finalement $\sum_{i \neq j} x_i x_j$

Chapitre 6

Problème aux frontières dans $\mathbb{D}$

Dans ce chapitre nous prolongeons le problème 4.2 en introduisant des frontières. Lorsqu'une frontière est touchée, nous disons qu'un saut synchrone intervient : c'est-à-dire qu'un événement exceptionnel se produit, et les particules ont une action coordonnée.

Nous commençons par prolonger la définition du problème, puis nous adaptons les démonstrations dans certains cas où il est possible de conclure.

Sommaire

6.1	Spécification du problème	90
6.1.1	Discussion	90
6.1.2	Premières définitions	90
6.1.3	Couplage des sauts	92
6.1.4	Problème avec frontière	93
6.2	Particule du champ moyen et convergence	93
6.2.1	Définition de la particule du champ moyen	93
6.2.2	Loi des grands nombres pour le champ moyen	94
6.2.3	Passage de la convergence après un saut	95
6.2.4	Convergence par l'étude du candidat limite	97
6.3	Un exemple d'utilisation : convergence du modèle Baccelli-Hong	98
6.3.1	Adaptation du modèle à N particules	98
6.3.2	Preuve de convergence en utilisant les résultats des chapitres précédents	100

6.1 Spécification du problème

6.1.1 Discussion

Rappel Pour l’instant nous avons introduit des particules X_i^N dont l’équation d’évolution est :

$$X_i^N(t) = x_i + \int_0^t f(X_i^N, M^N)(s^-)ds + \int_{s=0}^t \int_0^{I(s)} l(X_i^N, M^N)(s^-)d\Upsilon_i(u, s).$$

Cette équation représente un drift et des sauts réalisés de manière quasiment indépendante par les particules. Les seules sources de corrélation sont des ressources partagées. L’utilisation d’une ressource partagée peut être représentée sous la forme $\sigma(M^N) = \sigma_e(\langle \sigma_i, M^N \rangle)$. C’est-à-dire qu’une ressource partagée dépend uniquement de la mesure empirique de X^N (mesure que nous avons nommé M^N dans les chapitres précédents). Dans la notation $f(X_i^N, M^N)$ par exemple contient une dépendance dans une ressource partagée (ainsi que dans la trajectoire).

Saturation d’une ressource Nous souhaitons maintenant nous intéresser à ce qui peut se passer si la ressource que nous considérons vient à saturer. L’instant de saturation est caractérisé par le fait que l’utilisation de la ressource dépasse un certain seuil C que nous nommons CAPACITÉ de la ressource. Ainsi nous pouvons définir un premier instant de congestion T_1^N pour la ressource partagée qui vérifie $T_1^N = \inf\{t \in \mathbb{R}^+, \sigma(M^N)(t) \geq C\}$. À cet instant un changement a lieu, ce peut être une notification de congestion ou un changement de dynamique ou tout autre phénomène ponctuel. Ce phénomène est caractérisé par le fait que chacun des X_i^N va faire un saut que nous nommerons SYNCHRONES au même instant. Ce saut peut être un changement d’état, toutes choses étant égales par ailleurs (pour modéliser un changement de dynamique dû au fait que du fluide est perdu ou retenu à cause de la saturation), ou une diminution brutale de l’utilisation de la ressource.

Itération Nous souhaitons dans ce qui suit nous intéresser à l’évolution du système sur des temps longs. Lorsqu’une ressource est saturée, un mécanisme de contrôle rentre en jeu, ce qui peut amener à «dé-saturer» la ressource, puis une autre ressource partagée peut venir à saturer et ainsi de suite. Nous souhaitons nommer T_j^N ces temps où des sauts synchrones interviennent (j pour «jump»). Dans ce qui suit nous allons considérer qu’une seule ressource peut saturer que que juste après l’instant de saturation, le mécanisme de rétroaction est instantané et la ressource n’est plus saturée.

6.1.2 Premières définitions

Nouvelle dynamique, plan du chapitre Dans ce chapitre nous traitons le cas où un saut synchrone est ajouté à la dynamique initiale. C’est-à-dire que la dynamique originale est enrichie et devient :

$$\begin{aligned}
 X_i^N(t) = & x_i + \int_0^t f(X_i^N, M^N)(s^-) ds + \int_{s=0}^t \int_0^{I(s)} l(X_i^N, M^N)(s^-) d\Upsilon_i(u, s) \\
 & + \sum_{j \in \mathbb{N}} h_{i,j}^N \left(X_i^N \left((T_j^N)^- \right) \right) \chi_{s \geq T_j^N} ds.
 \end{aligned}$$

où les fonctions $h_{i,j}^N$ donnent la taille du saut synchrone réalisé par la particule X_i^N à des instants T_j^N définis à partir de la trajectoire.

Nous allons donner un sens à cette équation définie de manière incrémentale de saut à saut, puis nous nous intéresserons à la convergence du système X^N vers une limite que nous définirons à partir de la particule du champ moyen du chapitre précédent. Enfin nous montrerons, sous des hypothèses assez fortes sur les sauts synchrones, la convergence du système original vers le système auxiliaire.

Hasard dans le nouveau problème Définissons maintenant le hasard, et la filtration associée avant un saut.

Définition 6.1. Nous nous donnons les variables **indépendantes** suivantes :

- une CONDITION INITIALE : $(x_i)_{i \in \mathbb{N}}$, des variables aléatoires à valeurs dans $\mathbb{R}^k$,
- des SAUTS À INTENSITÉ : (Υ_i) , des processus de Poisson dans le plan d'intensité 1,
- des TAILLES DE SAUTS SYNCHRONES $h_{i,j}^N$, variables aléatoires $\mathbb{R}^k \rightarrow \mathbb{R}^k$ (à priori quelconques qui vérifient l'hypothèse (6.2)) et $h_{i,j}$, variables aléatoires $\mathbb{R}^k \rightarrow \mathbb{R}^k$ iid.

Définition 6.2. Nous notons $(\mathcal{F}_{t,j})$, la FILTRATION AVANT LE j -ÈME SAUT, la filtration engendrée par $(x_i)_{i \in \mathbb{N}}$, $(\Upsilon_i([0, t] \times \mathbb{R}^+))_{i \in \mathbb{N}}$, $(h_{i,j'}^N)_{i \leq N, j' < j}$ et $(h_{i,j'})_{i \leq N, j' < j}$.

Cette filtration peut être rendue continue à droite par un procédé d'augmentation (se reporter à [93] pour une mise en oeuvre de l'augmentation par rapport aux conditions initiales). Remarquons que le nom de «filtration avant le j -ème saut» est anticipatif (on ne sait pas encore que des sauts auront lieu), mais l'objet est bien défini en lui-même.

Dynamique avant le premier saut Nous allons laisser évoluer la dynamique comme dans le problème 4.2 jusqu'au moment où la ressource partagée que nous étudions va saturer. Ceci nous conduit à étudier les trajectoires de 4.2 jusqu'à un temps aléatoire que nous allons définir maintenant.

Définition 6.3. RESSOURCE PARTAGÉE

Soit $\sigma : \mathbb{D}(\mathbb{R}, \mathcal{P}(E)) \times \mathbb{R} \rightarrow \mathbb{R}^k$ (déterministe) continue de la forme suivante :

$$\sigma(M^N)(t) := \sigma_e(\langle \sigma_i, M_t^N \rangle),$$

où σ_i et σ_e sont toutes deux continues.

Par analogie aux chapitre précédents, notons

$$\bar{\sigma}(t) := \sigma_e(\langle \langle \sigma_i, M_t^N \rangle \rangle).$$

$\sigma(M^N)(t)$ est la CONSOMMATION COURANTE de la ressource partagée dont la CAPACITÉ vaut C .

Remarquons que si M_t^N est un processus prévisible alors $Y_t := \sigma(M^N)(t)$ est aussi prévisible. Si nous supposons qu'au temps initial, la ressource n'est pas saturée, nous allons pouvoir définir un temps T_1^N par :

Définition 6.4. $T_1^N = \inf\{t : \sigma(M^N)(t) \in [C, \infty[\}$.

Proposition 6.1. T_1^N est un temps d'arrêt pour $\mathcal{F}_{t,1}$.

Preuve : Nous utilisons le théorème 1.6 p. 55 dans [44] pour le temps de première entrée T_1^N . Les trois hypothèses sont vérifiées ici :

- la filtration $\mathcal{F}_{t,1}$ est complète et continue à droite (car on a pris une version augmentée de la filtration),
- $\sigma(M^N)(t)$ est un processus prévisible pour la filtration (il est adapté et continu à gauche),
- $[C, \infty[$ est un borélien de $\mathbb{R}$,

entraînant que T_1^N est un temps d'arrêt. ■

Itération aux sauts suivants Définissons ensuite de manière itérative les temps d'arrêt et les trajectoires associées. Nous pouvons étudier les trajectoires en partant de la condition initiale $X_i^N((T_1^N)^-) + h_{i,1}(X_i^N((T_1^N)^-))$ dans la filtration $\mathcal{F}_{t,2}$. Les processus solution sont alors adaptés, et sont bien uniques solutions du problème, ce qui permet de définir de la même façon que précédemment un temps d'arrêt T_2^N pour peu que l'hypothèse suivante soit vérifiée :

Hypothèse 6.1. Soit T un instant où un saut synchrone intervient, alors $\sigma(M^N, T) < C$ ps.

Définition 6.5. $T_1^N < T_2^N < \dots < T_j^N < \dots$ est la suite strictement croissante de temps d'arrêt (à valeurs dans $\overline{\mathbb{R}}^+$) définis de manière itérative sur j où T_j^N est le premier temps de saturation après T_{j-1}^N pour la filtration $(\mathcal{F}_{t,j})$.

6.1.3 Couplage des sauts

En plus du couplage des chapitres précédents nous introduisons un nouveau couplage pour les sauts synchrones.

Hypothèse 6.2. COUPLAGE DES SAUTS SYNCHRONES. Définissons les tailles de sauts synchrones du système à N particules $(h_{i,j}^N)$ ainsi : ce sont des variables aléatoires à valeurs dans $\mathbb{R}^k \rightarrow \mathbb{R}^k$:

$$\forall j, \sum_{i=1}^N \mathbb{P}(h_{i,j}^N \neq h_{i,j}) \xrightarrow{N \rightarrow \infty} 0$$

Nous n'imposons pas que les $h_{i,j}^N$ soient indépendants.

6.1.4 Problème avec frontière

Problème 6.1. Nous étudions le système infini de particules (X_i^N) à valeurs dans $E = \mathbb{D}([0, T], \mathbb{R}^k)$. Notons $(M^N)_{N \in \mathbb{N}}$ la mesure empirique de X^N .

Nous nous donnons les processus et variables aléatoires indépendants suivants, (x_i) à valeurs dans $\mathbb{R}^k$, (Υ_i) Poisson du plan et $(h_{i,j})$ à valeurs dans les fonctions L -Lipschitz de $\mathbb{R}^k \rightarrow \mathbb{R}^k$. Les x_i vérifient l'hypothèse 4.2 et on se donne des fonctions $h_{i,j}^N$ qui vérifient 6.2 qui sont à valeurs dans les fonctions L -Lipschitz.

Soient les fonctions $f, l, I : E \times \mathcal{P}(E) \rightarrow (\mathbb{R} \rightarrow \mathbb{R}^k)$ et $\sigma : \mathcal{P}(E) \times \mathbb{R} \rightarrow \mathbb{R}$ de capacité C . Ces fonctions sont supposées dépendantes de la marginale et toutes L -Lipschitz pour un $L \in \mathbb{R}$ donné.

Donnons-nous de plus un seuil $C \in \mathbb{R}$, que nous nommerons CAPACITÉ. La fonction σ sera nommée FONCTION D'UTILISATION DES RESSOURCES.

(X_i^N) vérifie les équations différentielles stochastiques couplées suivantes : pour $t \geq 0$

$$\begin{aligned} X_i^N(t) = & x_i + \int_0^t f(X_i^N, M^N)(s^-) ds \\ & + \int_{s=0}^t \int_0^{I(s)} l(X_i^N, M^N)(s^-) d\Upsilon_i(u, s) + \sum_{j \in \mathbb{N}} h_{i,j}^N(X_i^N(s^-)) \chi_{t \geq T_j^N} \end{aligned} \quad (6.1)$$

et les temps de sauts synchrones T_j^N ont lieu à la première entrée de $t \rightarrow \sigma(M^N)(t)$ dans $[C, \infty]$.

Remarquons que la définition forte trajectoire par trajectoire du système précédent ne fonctionne que jusqu'au temps $T_\infty^N = \sup_{j \in \mathbb{N}} T_j^N$. Nous voyons donc bien la nécessité qu'il y aura de contrôler la progression de ces temps d'arrêt.

Nous souhaitons dériver le même type de théorèmes que précédemment, c'est à dire une loi des grands nombres fonctionnelle.

6.2 Particule du champ moyen et convergence

6.2.1 Définition de la particule du champ moyen

Nous avons montré au chapitre précédent l'existence et l'unicité du champ moyen sans frontière, c'est à dire pour $C = +\infty$. Nous allons raisonner de saut en saut pour définir la particule avec les sauts. Nous introduisons en même temps une notion qui va s'avérer utile, la particule du champ moyen à laquelle on ne laisse faire que les j premiers sauts et qu'on laisse libre ensuite.

Définition 6.6. CHAMP MOYEN AUX j PREMIERS SAUTS ET TEMPS D'ARRÊT DU CHAMP MOYEN Définissons les temps d'arrêt $\mathcal{T}_j$ et la particule $\bar{\mathcal{X}}_j$ du champ moyen limité aux j premiers sauts par la récurrence suivante :

- $\bar{\mathcal{X}}_0$ est l'unique solution du théorème 5.1 sur $\mathbb{R}$ sans temps d'arrêt,
- pour $j \geq 1$, $\bar{\mathcal{X}}_j := \bar{\mathcal{X}}_{j-1}$ jusqu'au premier temps après $\mathcal{T}_{j-1}$ que l'on nomme $\mathcal{T}_j$ où $\bar{\sigma}(\bar{\mathcal{X}}_{j-1})(t) = 0$ et pour $t \geq \mathcal{T}_j$, $\bar{\mathcal{X}}_j$ est l'unique solution sans temps d'arrêt donnée par le théorème 5.1 avec condition initiale $h_j(\bar{\mathcal{X}}_{j-1}(\mathcal{T}_j^-))$.

Définition 6.7. CHAMP MOYEN AVEC SAUTS Définissons la particule $\overline{\mathcal{X}}_j$ du champ moyen avec sauts ainsi : pour $t \leq \sup_{j \in \mathbb{N}} \mathcal{T}_j$:

$$\overline{\mathcal{X}}(t) = \text{valeur commune} \{ \overline{\mathcal{X}}_j(t), \mathcal{T}_j \geq t \}. \quad (6.2)$$

Définition 6.8. La particule du champ moyen $t \rightarrow \overline{\mathcal{X}}(t)$ étant donnée, définissons le FLOT DU CHAMP MOYEN $\overline{\mathcal{X}}_t(y)$, fonctionnelle qui à une condition initiale quelconque y (une variable aléatoire sur $\mathbb{R}^k$) fait correspondre la solution au temps t de l'équation :

$$\begin{aligned} \overline{\mathcal{X}}_t(y) = & y + \int_0^t f_e(\overline{\mathcal{X}}_{s-}(\mu), \langle \langle f_i, \overline{\mathcal{X}} \rangle \rangle) ds + \int_0^t \int_0^{I_e(\overline{\mathcal{X}}_{s-}(\mu), \langle \langle I_i, \overline{\mathcal{X}} \rangle \rangle)} l_e(\overline{\mathcal{X}}_{s-}(\mu), \langle \langle l_i, \overline{\mathcal{X}} \rangle \rangle) d\Upsilon(u, s) \\ & + \sum_{j \in \mathbb{N}} h_j(\overline{\mathcal{X}}_t(y)) \chi_{s \leq \mathcal{T}_j} ds \end{aligned} \quad (6.3)$$

Cette définition ne pose pas de problème car les $\mathcal{T}_j$ sont des réels ce qui en fait une EDS tout à fait classique.

Définition 6.9. La particule du champ moyen $\overline{\mathcal{X}}$ étant définie, le SYSTÈME AUXILIAIRE avec sauts synchrones vaut :

$$\mathcal{X}_i(t) = \overline{\mathcal{X}}_t(y)$$

avec $\Upsilon = \Upsilon_i \ h_j = h_{i,j}$.

6.2.2 Loi des grands nombres pour le champ moyen

Théorème 6.1. Soit $T < \mathcal{T}_\infty$, alors $\forall t < T, \forall g \in \mathcal{C}_b(\mathbb{R}^k, \mathbb{R})$,

$$\sum_{i=1}^N g(\mathcal{X}_i(t)) \xrightarrow[N \rightarrow \infty]{} \mathbb{E}(g(\overline{\mathcal{X}}(t))).$$

Preuve : Ce théorème est identique à celui du chapitre précédent, en effet la fonction $F_t(\Upsilon_i, x_i, h_j)$ qui à $\mathcal{X}_i(\omega)(0)$ associe $\mathcal{X}_i(\omega)(t)$ est continue presque sûrement. Pour compléter l'argument, il suffit de voir qu'une fois la fonction $h_{i,j}$ fixée L-Lipschitz, le temps $\mathcal{T}_j$ est un réel donné. Ainsi si $\{F_t, t < \mathcal{T}_j\}$ sont continues alors $\{F_t, t \leq \mathcal{T}_j\}$ aussi en notant que

$$\begin{aligned} \mathcal{X}_i(\omega)(\mathcal{T}_j) &= \mathcal{X}_i(\omega)(\mathcal{T}_j^-) + h_{i,j}(\mathcal{X}_i(\omega)(\mathcal{T}_j^-)) \\ &= F_{\mathcal{T}_j^-}(\mathcal{X}_i(\omega)(0)) + h_{i,j}(F_{\mathcal{T}_j^-}(\mathcal{X}_i(\omega)(0))) \end{aligned}$$

et que $\mathcal{X}_i(\omega)(t)$ est une fonction presque sûrement localement Lipschitzienne en $\mathcal{T}_j^-$ (cad en particulier que x étant fixé, $t \rightarrow F_t(\dots)(x)$ est continue).

■

6.2.3 Passage de la convergence après un saut

Discussion Pour assurer la convergence après un saut nous devons trouver un temps auquel la nouvelle condition initiale après le saut converge au sens Cesarò-auxiliaire. Dans ce qui suit nous allons nous intéresser uniquement au passage du premier saut. Par récurrence nous obtiendrions alors la convergence jusqu'à $\mathcal{T}_\infty := \lim \mathcal{T}_j$ comme l'illustre le raisonnement ci-dessous.

Une hypothèse forte Il est nécessaire de rajouter une hypothèse technique assez forte pour éviter que les particules ne fassent prématurément de saut synchrone.

Hypothèse 6.3. RÉGULARITÉ DE L'UTILISATION Notons $Y_t^N := \sigma(M^N(\omega))(t)$ et $Y_t := \bar{\sigma}(\mathcal{M}(\omega))(t)$ des versions du problème sans temps temps d'arrêt; pour tout réel T positif, on peut montrer que

$$\sup_{t \in [0, T]} |Y_t^N - Y_t| \xrightarrow[N \rightarrow \infty]{} 0 \text{ en probabilité.}$$

Lemme 6.1. L'hypothèse (6.3) est automatiquement vérifiée si $P^N := \mathcal{L}(Y_t^N)$ est tendue dans $\mathcal{P}(\mathbb{D})$.

Preuve : On sait que Y_t est une fonction déterministe et continue (déterministe parce qu'il y a une espérance dans sa définition et continue par le corollaire 5.3). Le caractère déterministe implique que prouver la convergence en probabilité ou en loi revient au même. Par ailleurs la fonction sup est Skorokhod-continue (cf : proposition 4.5) et la limite continue, c'est-à-dire que le problème revient à passer d'une convergence fini-dimensionnelle connue (on sait en particulier que pour tout t , $|Y_t^N - Y_t| \xrightarrow[N \rightarrow \infty]{} 0$ en probabilité par le corollaire 5.3) à une convergence dans $\mathbb{D}$ pour la norme de Skorokhod. Par le théorème 4.4 ceci est en particulier réalisé dès que la suite des mesures est tendue. ■

Un point technique : Nous savons que hors des sauts $t \rightarrow \mathbb{E}(g(\bar{\mathcal{X}}(t)))$ est une fonction continue, ceci assure que $\bar{\sigma}(\bar{\mathcal{X}})$ est aussi une fonction continue; ce qui veut dire qu'elle atteint son maximum sur tout compact. Ainsi sur tout intervalle $A = [0, \mathcal{T}_1 - \epsilon]$, il existe un $\eta > 0$ tel que $\bar{\sigma}(\bar{\mathcal{X}}) < C - \eta$.

Par ailleurs sur l'intervalle A , nous savons par les arguments du chapitre 5 que le système original converge en norme cesarò-auxiliaire vers le champ moyen, ce qui assure en particulier que pour tout t ,

$$\mathbb{E}(\sigma(X^N)) \rightarrow \bar{\sigma}(\bar{\mathcal{X}}).$$

L'hypothèse (6.3) renforce cela en disant que la convergence a lieu en probabilité, mais pour la norme uniforme sur A . Ceci signifie qu'avec une grande probabilité, qu'aucun saut n'aura lieu dans A .

En particulier cela nous donne la convergence (car dans le cas peu probable où passe C en avance, le temps est borné) :

Proposition 6.2. Les temps d'arrêt du système original T_1^N vérifient :

$$\mathbb{E}(\mathcal{T}_1 \wedge T_1^N - \mathcal{T}_1) \rightarrow 0.$$

Condition initiale après un saut Revenons au passage d'un saut.

Nous notons $X_i^N(t \wedge (T_1^N)^-)$ le processus arrêté en $(T_1^N)^-$ juste avant le saut.

Lemme 6.2. HYPOTHÈSES : Supposons la convergence jusqu'au premier saut : X^N et son système auxiliaire $(\mathcal{X}_i)$ vérifient les équations sans sauts synchrones. D'autre part supposons qu'il existe $t_1 \in \mathbb{R}$ tel que $\mathcal{T}_1 \leq t_1 < \mathcal{T}_2$ et que avec probabilité 1, $T_1^N \leq t_1 < T_2^N$ pour tout N suffisamment grand.

CONCLUSION : Une nouvelle condition initiale en t_1 qui correspond à la valeur atteinte en t_1 vérifie :

$$\|X^N - \mathcal{X}\|_{[t_1, t_1]} \leq \|X^N(s \wedge (T_1^N)^-) - \mathcal{X}(s \wedge (T_1^N))\|_{[0, \mathcal{T}_1]} + C(|t_1 - \mathcal{T}_1| + \mathbb{E}|T_1^N - \mathcal{T}_1|) + b_X V^N,$$

avec $C = 2(Lb_X + L^2 b_X^2)$ et $V^N = \sum_{i=1}^N \mathbb{P}(h_{i,j}^N \neq h_{i,j})$.

Attention ! V^N figure sans le $\frac{1}{N}$, ce n'est pas une erreur.

Preuve : Notons $m = T_1^N \wedge \mathcal{T}_1$ et $M = T_1^N \vee \mathcal{T}_1$. Dans le cas $m = \mathcal{T}_1$ et $M = T_1^N$ (le cas inverse se traiterait exactement de la même façon). Après M , les deux particules X^N et $(\mathcal{X}_1, \dots, \mathcal{X}_N)$ ont fait le même nombre de sauts synchrones (à cause de la deuxième hypothèse). Nous disons donc que si $h_{i,j}^N(\omega) = h_{i,j}(\omega)$:

$$\begin{aligned} |X_i^N(M) - \mathcal{X}_i(M)| &\leq |X_i^N(M) - \mathcal{X}_i(m)| + |\mathcal{X}_i(m) - \mathcal{X}_i(M)| \\ &\leq L|X_i^N(M^-) - \mathcal{X}_i(m^-)| + LX_m|M - m| \\ &\quad + \int_m^M \left| \int_{[0, \bar{I}(\mathcal{X}_i, \bar{\mathcal{X}})(z)]} \bar{l}(\mathcal{X}_i, \bar{\mathcal{X}})(z) \right| dz, \end{aligned}$$

en remontant aux instants où les particules auxiliaires ont sauté. Puis en remontant au temps avant le premier saut :

$$\begin{aligned} |X_i^N(M) - \mathcal{X}_i(M)| &\leq L|X_i^N(m^-) - \mathcal{X}_i(m^-)| + 2LX_m|M - m| \\ &\quad + \int_m^M \left| \int_{[0, \bar{I}(\mathcal{X}_i, \bar{\mathcal{X}})(z)]} \bar{l}(\mathcal{X}_i, \bar{\mathcal{X}})(z) \right| + \left| \int_{[0, \bar{I}(Y, Y)(z)]} \bar{l}(Y, Y)(z) dz \right|. \end{aligned}$$

Si nous appliquons partout $\mathbb{E} \frac{1}{N} \sum_{i=1}^N$ en tenant compte des cas où au moins l'un des $h_{i,j}^N(\omega) \neq h_{i,j}(\omega)$ par une majoration brutale par b_X au lieu de $L|\dots|_{[0, t]}$ et en faisant évoluer les trajectoires entre M et t_1 , ce qui introduit la constante $C|t_1 - \mathcal{T}_1|$, nous obtenons le résultat. ■

Discussion Nous voyons dans le lemme précédent la difficulté essentielle introduite par des sauts à temps d'arrêts : le temps $\tau(C)$ défini par l'équation implicite $f(\tau(C)) = C$ n'a aucune raison d'être Lipschitz par rapport à C . Par exemple il peut se faire que le système auxiliaire frôle la frontière et qu'une bonne partie des particules originales ne touchent jamais cette frontière ou la touchent beaucoup plus tard. Il n'est alors par possible de

prolonger la convergence. Aussi on voit pourquoi nous allons être amenés à introduire une hypothèse qui implique que

$$X^N \rightarrow (\mathcal{X}_i) \Rightarrow T_1^N \rightarrow \mathcal{T}_1.$$

Dans ce cas le lemme 6.2 permet de prolonger la convergence après le temps de saut. Par récurrence ensuite, la convergence se poursuivra jusqu'à la limite de la suite croissante $\mathcal{T}_j$, qui peut être ∞ . La sous-section suivante indique une manière de faire.

6.2.4 Convergence par l'étude du candidat limite

Nous avons défini au début de ce chapitre le champ moyen en faisant référence à des particule du champ moyen $\overline{\mathcal{X}}_j$ limitées aux j premiers sauts. Nous allons réutiliser cette définition, en effet l'étude de cet objet au delà de la traversée de la frontière où la particule $\overline{\mathcal{X}}$ fait un saut nous renseigne sur la convergence des trajectoires des particules du processus original qui n'a pas encore fait son saut.

Hypothèses

Hypothèse 6.4. Traversée franche au j -ième saut *On suppose que pour le saut numéro j , $\overline{\sigma}\overline{\mathcal{X}}_j(t)$ est localement strictement monotone dans le voisinage du temps d'arrêt $\mathcal{T}_j$. De plus on suppose que la frontière suivante pour $\overline{\mathcal{X}}$ se trouve à une distance strictement positive.*

Hypothèse 6.5. Non accumulation uniforme des sauts dans le système original *Il existe une suite η_j de réels strictement positifs telle que :*

$$\mathbb{P}(T_{j+1}^N \geq T_j^N + \eta_j) \xrightarrow{N \rightarrow \infty} 1$$

Théorème de convergence

Théorème 6.2.

HYPOTHÈSES : *Dans le cadre du problème 6.1, $\overline{\mathcal{X}}_j$, les particules du champ moyen limitées aux j premiers sauts sont définies. Rajoutons les hypothèses (6.4) et (6.5).*

CONCLUSION : *Sur tout borné jusqu'à la valeur $\mathcal{T}_\infty$ on a $X^N \xrightarrow{N \rightarrow \infty} (\mathcal{X}_i)$ au sens Cesarò-auxiliaire.*

Preuve : Le raisonnement général serait une récurrence, nous nous contenterons de montrer la convergence à travers le premier temps d'arrêt, la récurrence exacte se calquerait sur le même raisonnement.

Nous connaissons la convergence jusqu'au premier temps d'arrêt en utilisant le théorème 5.5. Ceci donne en corollaire que

$$\liminf T_1^N \geq \mathcal{T}_1,$$

il reste à s'assurer que les trajectoires passent suffisamment vite la frontière pour pouvoir utiliser le lemme 5.2. Mais ce point est désormais facile sous l'hypothèse (6.4), en effet on

sait que les trajectoires qui n'ont pas fait leur saut se comportent comme dans $\overline{\mathcal{X}}_1$ et la convergence sans le saut assure pour tout $\epsilon > 0$ que toutes les trajectoires dépassent le seuil avant $\mathcal{T}_1 + \epsilon$ avec une probabilité qui tend vers 1 à cause de la convergence uniforme en probabilité des crochets simples vers les crochets doubles qui déterminent le temps d'arrêt.

De plus une fois le premier saut passé, l'hypothèse de l'espace borné permet de ne pas se soucier de la faible probabilité des trajectoires qui n'ont pas fait le saut. Ainsi on a en fait :

$$\lim T_1^N = \mathcal{T}_1.$$

Mais ce n'est pas encore suffisant pour appliquer le lemme car il ne faudrait pas que $\mathcal{T}_1$ soit un point d'accumulation de sauts avec une probabilité non nulle.

Or l'hypothèse (6.5) assure qu'il existe η_1 tel que le deuxième saut a lieu à moins à η_1 du premier avec une grande probabilité, les trajectoires qui ne vérifient pas cela restent toujours bornées. Ceci permet d'appliquer le lemme 6.2. Voyant que $V^N \rightarrow 0$ à cause de l'hypothèse (6.2), en probabilité, on récupère une condition initiale analogue à l'hypothèse (4.2) qui permet d'itérer jusqu'au prochain saut.

■

6.3 Un exemple d'utilisation : convergence du modèle Baccelli-Hong

Dans cette section nous adaptons et prouvons le modèle HTTP-AIMD de Baccelli-Hong en utilisant les résultats précédents. Ce modèle a été décrit à la section 2.2.1 du chapitre 2. Le but est de prouver le théorème 2.1 en donnant les hypothèses techniques appropriées.

6.3.1 Adaptation du modèle à N particules

Nous conservons le modèle à N particules décrit dans la section 2.2.1, en précisant les hypothèses. En particulier nous expliquons ce qu'on doit entendre par "silencieux pendant une durée $1/\beta$ en moyenne" ou bien "l'utilisateur transmet un fichier de taille $1/\mu$ en moyenne".

Les débuts et fins de transmission La variable d'état de l'utilisateur X contient non seulement le débit de l'utilisateur, mais aussi la taille de fichier déjà transmise et/ou le temps depuis lequel il est en sommeil. On rajoute enfin une quatrième variable d'état qui indique si l'utilisateur est en train de transmettre ou bien s'il est silencieux. Un utilisateur sera dit ON ou ACTIF pendant les périodes d'envois de fichier, et OFF ou EN SOMMEIL sinon. Nous noterons $X.1$ l'état de transmission, $X.2$, le débit, $X.3$, la taille de fichier téléchargé et $X.4$, le temps passé silencieux.

Les chances pour un utilisateur qui au temps t est dans l'état $X = (1 : \text{utilisateur on}, x : \text{débit}, i : \text{non significatif}, f : \text{taille de fichier transmis})$ de terminer sa transmission

sont supposées intervenir selon un processus de Poisson d'intensité $X.2(t) \frac{d_f(f)}{\int_{s=f}^{\infty} d_f(s)ds}$ où d_f la est densité (on suppose qu'elle existe) des tailles de fichiers. Lorsqu'une transmission s'achève on saute au vecteur d'état $(0 : \text{utilisateur off}, 0, 0 : \text{temps d'attente}, 0)$.

De même un utilisateur dont la variable d'état vaut $(0 : \text{utilisateur off}, x : \text{non significatif}, i : \text{temps d'attente}, f : \text{non significatif})$ recommence à émettre selon un processus de Poisson d'intensité $\frac{d_i(i)}{\int_{s=i}^{\infty} d_i(s)ds}$ où d_i est la densité (on suppose qu'elle existe) des temps d'attente. Lorsque la transmission reprend on effectue un saut vers le vecteur d'état $(1 : \text{utilisateur on}, 0 : \text{débit}, 0, 0 : \text{fichier transmis})$.

Remarque 6.1. *On voit que rien n'interdirait de reprendre la transmission à une fenêtre plus grande que 0 ou bien d'introduire d'autres états. Cependant il est bon si possible de rester modéré dans la complexité des états si on veut être capable de simuler les équations par la suite.*

Remarque 6.2. *Nous allons laisser la précision de la condition initiale dans le vague, ceci ne présente aucun problème. Par simplicité nous supposons dans ce qui suit que tout le monde commence à $(0, 0, 0, 0)$.*

Modèle de pertes et couplage des sauts synchrones Soient $b_{i,j}$, des tirages de Bernoulli iid qui valent 1 avec la probabilité p . Définissons ensuite :

$a_{i,j}^N = b_{i,j}$ si $\sum_{i=0}^N b_{i,j} \chi_{X_i^N(1)=1} > 0$.

Sinon tirons uniformément un des $\{a_{i,j}^N\}_{i,X_i^N(1)=1}$ qui vaut 1 (on a vu au chapitre 3 comment écrire cette idée dans le langage mathématique).

On pourrait tout aussi bien dire que la première particule qui transmet fait un saut dans le cas où personne ne saute parmi les utilisateurs en activité, le cas étant rare, le seul problème est d'assurer que la définition dans ce cas ait un sens. Remarquons au passage que par rapport à ce que nous avons prévu, $\mathbb{E}(a_{i,j}^N) > p$, le modèle fini n'a pas exactement le taux de synchronisation qui avait été annoncé.

Précisions sur les processus de Poisson Comme précédemment il faut que nous ayons des versions des processus de Poisson qui supportent aisément le couplage. Pour ce faire nous prenons un processus de Poisson à deux dimensions fixé pour chaque utilisateur Υ_i d'intensité 1. Et nous utilisons le processus de sauts suivant pour une intensité $I_i^N(t)$:

$$J_i^N : t \rightarrow \int_{u=0}^t \int_{s=0}^{I_i^N(u)} d\Upsilon_i(s, u).$$

Nous conservons dans les deux états le même Υ_i , mais l'intensité des sauts fait elle-même des sauts (ce qui ne nous pose pas vraiment de problèmes). Grâce à cette méthode que nous avons expliquée au chapitre 4, les sauts auront lieu soit exactement au même instant soit se décaleront. On peut voir une illustration de la création de $J_i^N(t)$ dans la figure (4.1).

6.3.2 Preuve de convergence en utilisant les résultats des chapitres précédents

Définition de la limite Un processus de Poisson 2D Υ étant fixé, la limite $\overline{\mathcal{X}}(t) = (\overline{e}(t), \overline{x}(t), \overline{f}(t), \overline{i}(t))$ se présente comme notre système original à une particule ; à ceci près que les sauts synchrones n'interviennent pas lorsque $\overline{x}(t)$ atteint la frontière C , mais lorsque $\mathbb{E}_{\Upsilon, h_1, h_2, \dots}(\overline{x})$ atteint cette valeur. La deuxième frontière est atteinte à un temps déterministe, alors que la première est atteinte pour un temps aléatoire. Quand cette frontière est atteinte, des sauts interviennent selon le tirage de variables de Bernoulli h_j avec la probabilité p de valoir 1. De plus $(\Upsilon, h_1, h_2, \dots, h_j, \dots)$ sont indépendants.

Notons maintenant $\mathcal{X}_i$ une copie de $\overline{\mathcal{X}}$ pour $(\Upsilon, h_1, h_2, \dots, h_j, \dots) = (\Upsilon_i, b_{i,1}, b_{i,2}, \dots, b_{i,j}, \dots)$. Les $b_{i,j}$ sont les mêmes que ceux qui ont permis de définir les $a_{i,j}^N$, c'est donc ceci qui réalise le couplage des sauts.

Convergence hors des temps d'arrêt Les théorèmes 5.1 et 5.5 assurent la définition de $\overline{\mathcal{X}}$ et des $\mathcal{X}_i$ et la convergence du système à N particules hors des passages de frontières²⁰. En effet, toutes les fonctions qui définissent la dynamique ont le caractère de régularité suffisant d'une part et d'autre part la croissance bornée des variables d'état du système assure bien que toutes les trajectoires restent dans une partie bornée pour tout compact du temps (si la condition initiale est contenue dans une partie bornée). Notons enfin que dans le cas de ce modèle simple, on pourrait appliquer directement une loi des grands nombres, les évolutions étant indépendantes hors des instants de sauts, on aurait même sans doute un théorème central limite. Cependant, rappelons que le but de cette section n'est pas de donner la meilleure démonstration pour ce cas-ci mais d'illustrer la méthode de couplage-système auxiliaire que nous avons développé dans les chapitres précédents. De plus la présence d'effets de bords rend l'approche trajectorielle plus visuelle.

Il nous faut régler le problème des frontières maintenant.

Convergence uniforme en probabilité du débit Pour pouvoir vérifier l'hypothèse (6.3), il faut que nous transformions la convergence pour tout t en probabilité que nous avons sur le débit en une convergence uniforme en probabilité. Pour cela l'argument est que le débit ne fait pas de saut vers le haut (on voit l'importance de dire que les utilisateurs partent au débit nul après une période de silence) et qu'il est à croissance bornée (à cause du caractère AIMD).

Comme la limite du débit est déterministe et continue, on peut appliquer un argument direct de compacité inspiré des lemmes de Dini (cf : corollaire B.1 de la section B.2). Une autre méthode est de prouver la tension des processus comme nous l'avons rappelé dans le lemme 6.1. Ici la tension est évidente puisque les trajectoires ne font des grands sauts qu'avec une faible probabilité (l'argument est contenu dans la démonstration du lemme 8.1 du chapitre suivant).

²⁰Hors des sauts synchrones, l'EDS a une forme classique puisqu'on a pas de couplage par la moyenne. Aussi il n'est pas utile de faire appel à des raisonnements compliqués comme c'était le cas précédemment pour pouvoir définir le champ moyen.

Non-accumulation des sauts synchrones Pour commencer traitons le problème de l'accumulation des points de sauts pour le système original.

Lemme 6.3. Notons S_j^N , la taille du j -ième saut du système original à N particules ($S_j^N > 0$ pour un saut vers le bas), et D_j^N le temps pour atteindre la prochaine frontière. Alors :

$$\begin{aligned} - \mathbb{P}(S_j^N \geq \frac{pC^2}{16X_m}) &\xrightarrow{N \rightarrow \infty} 1, \\ - D_j^N &\geq \frac{N}{r} S_j^N, \end{aligned}$$

Le système original vérifie l'hypothèse (6.5).

Preuve : Pour le premier point, remarquons que $\mathbb{E}(S_j^N) \geq p\frac{C}{2}$, ce serait une égalité si les tirages étaient indépendants. La valeur X_m (le diamètre de l'espace étudié) majore en particulier le débit maximal. Notons N_{on} , le nombre d'utilisateurs actifs, et p_{on} leur proportion. Le débit moyen par utilisateur actif au moment de la congestion est donc $\frac{C}{p_{on}}$. Notons p_0 la proportion d'utilisateurs actifs qui ont un débit supérieur à la moitié de la moyenne, comme ce débit est par ailleurs majoré par la constante X_m , il vient $p_0 p_{on} X_m + (1 - p_0) \frac{C}{2} \geq C$, ainsi :

$$p_0 \geq \frac{C}{2p_{on}X_m - C}.$$

Or $p_{on}X_m \leq C$, ainsi :

$$p_0 \geq \frac{1}{2}.$$

On voit donc que plus de la moitié des utilisateurs actifs ont un débit supérieur à $C/2$ au moment de la congestion. Mais ce nombre d'utilisateurs actifs est par ailleurs minoré par $\frac{C}{X_m}$, par conséquent $\frac{C}{2X_m}$ utilisateurs ont un débit supérieur à $C/2$ au temp de congestion. À l'instant de la congestion, une proportion p de ces utilisateurs sera touchée (les pertes sont indépendantes du passé), ainsi lorsque N devient grand, avec une probabilité qui tend vers 1, plus de $p/2$ utilisateurs dont le débit dépasse $C/2$ sont touchés. Finalement, on est assuré que la taille du saut a la taille :

$$S_j^N \geq \frac{pC^2}{16X_m}$$

avec une probabilité qui tend vers 1.

La deuxième formule résulte du fait que le rythme de croissance est limité. La conclusion est alors directe. ■

Remarque 6.3. La démonstration précédente n'est valable sous sa forme que si la probabilité de coupure des fenêtres est indépendante de la taille de la fenêtre ce qui est le cas dans le modèle mais constitue un point qui pourrait être amélioré.

En reprenant le raisonnement que nous venons de faire pour le système auxiliaire, on voit que les sauts n'admettent pas de point d'accumulation lorsque $j \rightarrow \infty$, ainsi le système auxiliaire sera défini sur tout borné du temps donc sur $\mathbb{R}^+$ en entier.

Couplage aux instants de sauts Il nous faut vérifier que les trajectoires sont bien couplées aux instants de sauts synchrones. Le problème que nous pourrions avoir serait du type $\frac{N_{on}}{N} \rightarrow 0$ auquel cas on pourrait avoir une trajectoire donnée qui ne tend pas vers sa limite. Ceci n'arrive pas comme nous l'avons vu dans la démonstration du lemme précédent.

Ainsi $\frac{N_{on}}{N} > \eta > 0$ lors d'une époque de congestion. Par conséquent $a_{i,j}^N \neq b_{i,j}$ avec probabilité $(1-p)^{N_{on}}$, ce qui est une condition de couplage suffisante puisque $N(1-p)^{N\eta}$ est une série sommable. Ceci permet de vérifier l'hypothèse (6.2).

Condition de convergence vers le champ moyen aux passage des frontières D'après le théorème 6.2, on aura convergence vers la limite $\bar{\mathcal{X}}$ (qui coïncide avec le procédé de construction de la particule du champ moyen à travers les sauts) si l'hypothèse suivante sur les limites $\bar{\mathcal{X}}_j$ est vérifiée :

Hypothèse 6.6. *Pour chaque j , la pente du débit $\bar{\mathcal{X}}_{j,2}(t)$ au voisinage du temps d'arrêt $\mathcal{T}_j$ est strictement positive.*

Moyennant cette hypothèse qui peut ne pas être vérifiée pour certains seuils, le champ moyen de l'article [9] existe donc bien comme il est annoncé. Il est l'unique limite du système à N particules que nous avons décrit dans la section 2.2.1. Ainsi le théorème 6.2 compte tenu des hypothèses vérifiées précédemment devient dans notre cas :

Théorème 6.3. *Le théorème 2.1 est vrai avec la modélisation du paragraphe 7 et l'hypothèse (7.1). On a le résultat supplémentaire :*

CONCLUSION 3 : *les temps d'arrêt T_j^N convergent vers les temps déterministes $\mathcal{T}_j$ dans $\mathbb{L}_1$ (c'est-à-dire en espérance). Ces temps $\mathcal{T}_j$ n'admettent aucun point d'accumulation, la convergence est valable jusqu'à $t = \infty$ (ou jusqu'au premier point où (7.1) cesse d'être vérifiée).*

Chapitre 7

Un premier exemple : champ moyen pour le modèle HTTP-AIMD de Baccelli-Hong

Dans ce chapitre, nous adaptons et prouvons le modèle HTTP-AIMD de Baccelli-Hong en utilisant le chapitre 6. Ce modèle a été décrit à la section 2.2.1. Pour que ce chapitre soit auto-contenu, nous nous autorisons certaines redites dans la première section où nous rappelons le modèle.

Sommaire

7.1	Rappel des modèles et problèmes à résoudre pour prouver la convergence du champ moyen	103
7.1.1	Le modèle AIMD original	104
7.1.2	Les équations du champ moyen intuitives pour HTTP sur AIMD	105
7.2	Adaptation du modèle à N particules	106
7.2.1	Les débuts et fins de transmission	106
7.2.2	Le modèle de pertes et couplage des sauts synchrones . . .	107
7.2.3	Précisions sur les processus de Poisson	107
7.3	Preuve de convergence en utilisant les résultats des chapitres précédents	108
7.3.1	Définition de la limite	108
7.3.2	Convergence hors des temps d'arrêt	108
7.3.3	Traversée des frontières	108

7.1 Rappel des modèles et problèmes à résoudre pour prouver la convergence du champ moyen

Quand nous parlerons dans la suite du modèle HTTP-AIMD nous ferons référence au modèle décrit dans la section 7.2. Dans cette première section nous rappelons les modèles

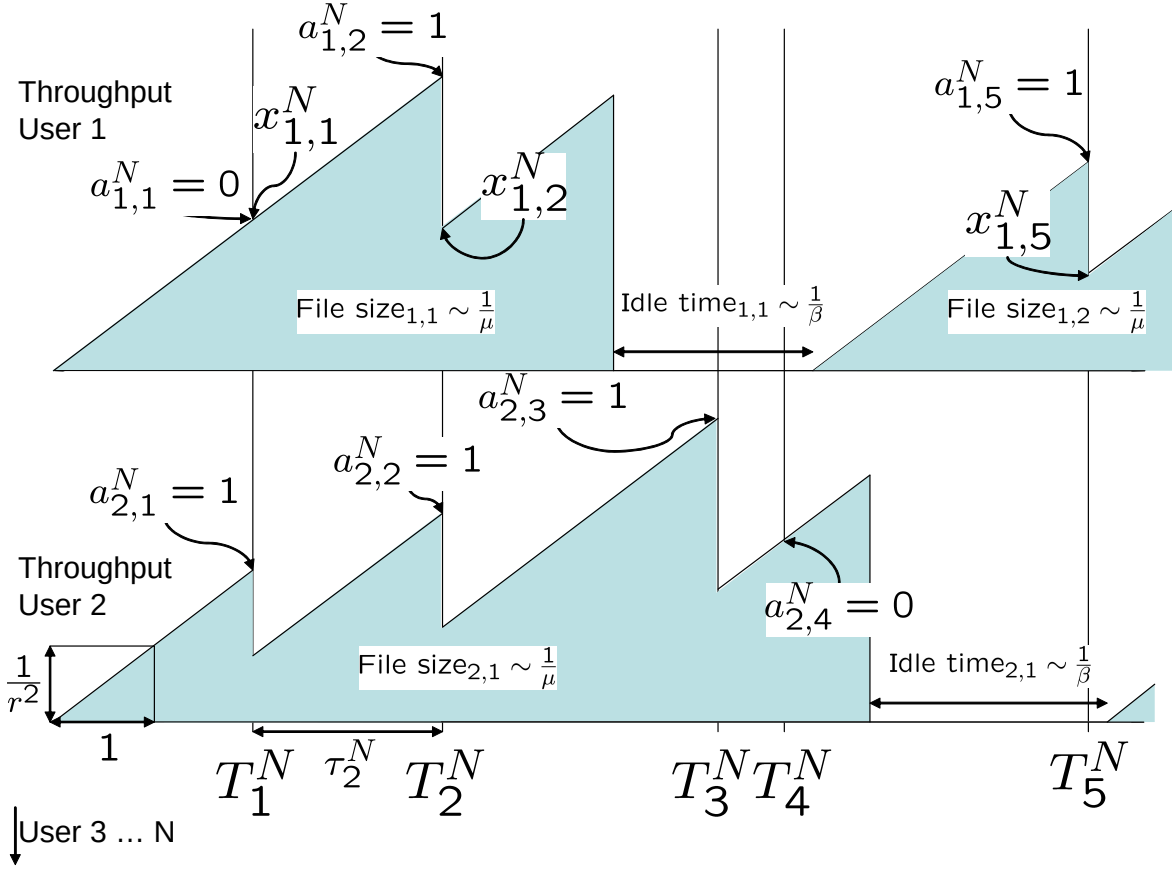


FIG. 7.1 – Illustration du modèle d'utilisateurs de HTTP, les utilisateurs adaptent leur débit en fonction des indications de congestion enchaînant période d'activité et périodes de sommeil.

qui ont été proposés et nous expliquons pourquoi nous sommes amenés à réaliser certaines petites modifications qui permettent la preuve de la convergence du champ moyen.

7.1.1 Le modèle AIMD original

Nous reprenons ici [11] en l'adaptant à HTTP de la manière préconisée dans [9]. La plupart des notations sont illustrées sur les figures (7.1) et (7.2).

Pour faire simple, nous supposons comme d'habitude que N utilisateurs partagent l'utilisation d'un routeur avec un petit buffer (de taille petite, mais suffisante pour absorber les fluctuations probabilistes du débit entrant) et ont tous le même RTT, r . De plus, nous allons adapter et simplifier les parties qui ne changent pas la démonstration et celles qui perturberaient l'homogénéité de notations de ce mémoire.

Pour modéliser HTTP nous supposons que chaque flot est silencieux pendant une durée $1/\beta$ en moyenne selon une certaine distribution de temps de silence. Après l'expiration de la période de silence, l'utilisateur transmet un fichier de taille $1/\mu$ en moyenne selon une certaine distribution de tailles de fichiers. Par défaut on utilise TCP Reno ainsi le

taux de transmission augmente à la vitesse $1/r^2$ (en effet la fenêtre augmente de 1 tous les r et le débit vaut W/r en moyenne). À la fin de la transmission, le débit revient à 0 et l'utilisateur rentre de nouveau dans une période de silence.

Pendant la durée de la transmission les débits des flots évoluent au gré des pertes imposées par le routeur de la manière suivante. Par définition la j -ième temps de congestion, T_j^N correspond à la j -ième époque de congestion au cours de laquelle des pertes de paquets quasiment simultanées vont se produire. On notera $\tau_j = T_j - T_{j-1}$. Les autres notations utilisées sont les suivantes :

- $C^N = NL$ est la capacité du routeur et L la capacité ramenée au nombre d'utilisateurs,
- $x_{i,j}^N$ et $W_{i,j}^N$ sont le débit et la fenêtre de l'utilisateur i après la j -ième époque de congestion,
- $a_{i,j}^N$ sont des variables aléatoires qui ne dépendent pas de N , valent 0 ou 1 selon que l'utilisateur i perd un paquet au temps de congestion T_j ou pas,
- les $a_{i,j}^N$ étant supposés identiquement distribués et indépendants par rapport au numéro j du saut ; p , le taux de synchronisation vaut $p := \mathbb{E}(a_{i,j}^N)^{21}$,
- le taux de croissance des fenêtres vaut $1/r$ et le taux de réduction des fenêtres par saut vaut $\frac{1}{2}$,

Les $a_{i,j}^N$ ne sont pas indépendants en i , mais issus d'un tirage de variables indépendantes conditionné par le fait que la somme des $a_{i,j}^N$ qui correspondent à des utilisateurs en cours de transmission vaut au moins 1. En effet une période de congestion est caractérisée par le fait qu'au moins une perte est constatée.

Ainsi les équations pendant la période de transmission sont :

$$x_{i,j+1}^N = (1 - \frac{1}{2}a_{i,j+1})(x_{i,j}^N + \frac{1}{r}\tau_{j+1}^N), \quad (7.1)$$

$$C^N = \sum_{i=1}^N x_{i,j}^N + \frac{N}{r}\tau_{j+1}^N. \quad (7.2)$$

7.1.2 Les équations du champ moyen intuitives pour HTTP sur AIMD

Nous reprenons ici le modèle proposé dans [9] à quelques modifications de notation près.

Soit $\bar{x}(t)$ le taux de transmission moyen d'un flot dans le régime stationnaire. Le champ moyen supposé est que $\bar{x}(t)$ croît à la vitesse $1/r^2$ multiplié par l'espérance d'être actif. Des sauts vers le bas viennent contrebalancer cette croissance. Ceci quand le fichier se finit ou bien par les réductions de fenêtres aux époques de congestion. L'hypothèse principale est que le taux de synchronisation des réductions de fenêtres reste p .

²¹En fait comme on le redira plus loin, le tirage est biaisé en réalité et c'est avant conditionnement sur des variables iid que l'on souhaite imposer l'espérance p . Nous donnerons plus loin une définition plus précise

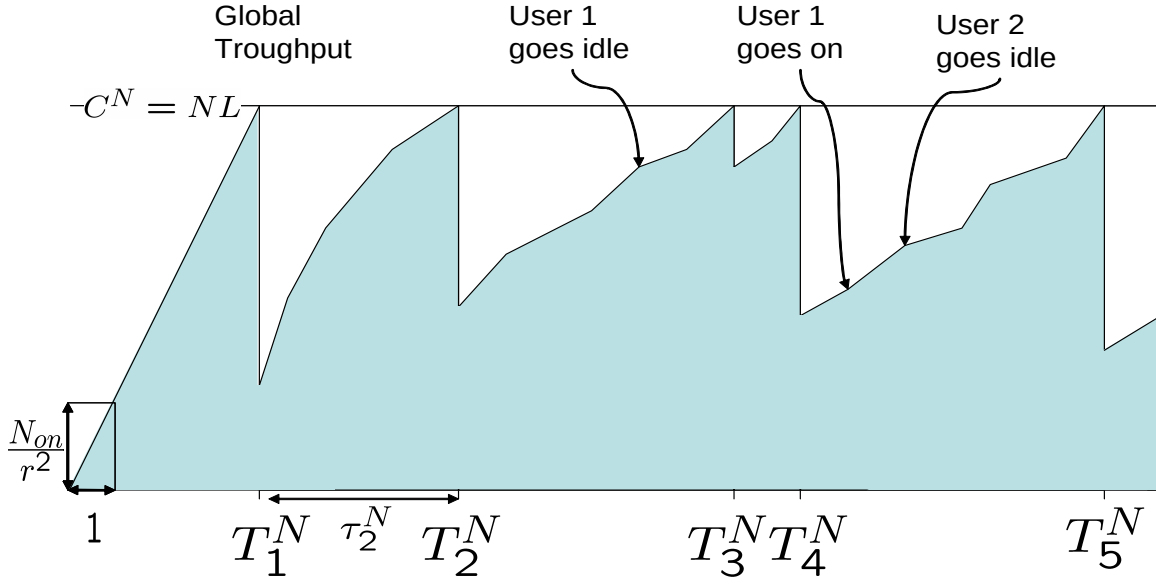


FIG. 7.2 – Illustration du modèle d'utilisateurs de HTTP, les temps d'arrêts T_j^N interviennent lorsque la condition de saturation $\sum_{i=1}^N x_i^N(t) = C^N = NL$ où $x_i^N(t)$ est le débit de la source i au temps t est vérifiée.

7.2 Adaptation du modèle à N particules

Nous conservons le modèle à N particules décrit dans la section 7.1.1, en précisant les hypothèses. En particulier nous expliquons ce qu'on doit entendre par "silencieux pendant une durée $1/\beta$ en moyenne" ou bien "l'utilisateur transmet un fichier de taille $1/\mu$ en moyenne".

7.2.1 Les débuts et fins de transmission

La variable d'état de l'utilisateur X contient non seulement le débit de l'utilisateur, mais aussi la taille de fichier déjà transmise et/ou le temps depuis lequel il est en sommeil. On rajoute enfin une quatrième variable d'état qui indique si l'utilisateur est en train de transmettre ou bien s'il est silencieux. Un utilisateur sera dit ON ou ACTIF pendant les périodes d'envois de fichier, et OFF ou EN SOMMEIL sinon. Nous noterons $X.1$ l'état de transmission, $X.2$, le débit, $X.3$, la taille de fichier téléchargé et $X.4$, le temps passé silencieux.

Les chances pour un utilisateur qui au temps t est dans l'état $X = (1 : \text{utilisateur on}, x : \text{débit}, i : \text{non significatif}, f : \text{taille de fichier transmis})$ de terminer sa transmission sont supposées intervenir selon un processus de Poisson d'intensité $X.2(t) \frac{d_f(f)}{\int_{s=f}^{\infty} d_f(s) ds}$ où d_f la est densité (on suppose qu'elle existe) des tailles de fichiers. Lorsqu'une transmission s'achève on saute au vecteur d'état $(0 : \text{utilisateur off}, 0, 0 : \text{temps d'attente}, 0)$.

De même un utilisateur dont la variable d'état vaut (0 : utilisateur off, x : non significatif, i : temps d'attente, f : non significatif) recommence à émettre selon un processus de Poisson d'intensité $\frac{d_i(i)}{\int_{s=i}^{\infty} d_i(s)ds}$ où d_i est la densité (on suppose qu'elle existe) des temps d'attente. Lorsque la transmission reprend on effectue un saut vers le vecteur d'état (1 : utilisateur on, 0 : débit, 0, 0 : fichier transmis).

Remarque 7.1. *On voit que rien n'interdirait de reprendre la transmission à une fenêtre plus grande que 0 ou bien d'introduire d'autres états. Cependant il est bon si possible de rester modéré dans la complexité des états si on veut être capable de simuler les équations par la suite.*

Remarque 7.2. *Nous allons laisser la précision de la condition initiale dans le vague, ceci ne présente aucun problème. Par simplicité nous supposons dans ce qui suit que tout le monde commence à (0, 0, 0, 0).*

7.2.2 Le modèle de pertes et couplage des sauts synchrones

Soient $b_{i,j}$, des tirages de Bernoulli iid qui valent 1 avec la probabilité p . Définissons ensuite $a_{i,j}^N = b_{i,j}$ si $\sum_{i=0}^N b_{i,j} \chi_{X_i^N(1)=1} > 0$. Sinon tirons uniformément un des $\{a_{i,j}^N\}_{i,X_i^N(1)=1}$ qui vaut 1 (on a vu au chapitre 3 comment écrire cette idée dans le langage mathématique). On pourrait tout aussi bien dire que la première particule qui transmet fait un saut dans le cas où personne ne saute parmi les utilisateurs en activité, le cas étant rare, le seul problème est d'assurer que la définition dans ce cas ait un sens. Remarquons au passage que par rapport à ce que nous avons prévu, $\mathbb{E}(a_{i,j}^N) > p$, le modèle fini n'a pas exactement le taux de synchronisation qui avait été annoncé.

7.2.3 Précisions sur les processus de Poisson

Comme précédemment il faut que nous ayons des versions des processus de Poisson qui supportent aisément le couplage. Pour ce faire nous prenons un processus de Poisson à deux dimensions fixé pour chaque utilisateur Υ_i d'intensité 1. Et nous utilisons le processus de sauts suivant pour une intensité $I_i^N(t)$:

$$J_i^N : t \rightarrow \int_{u=0}^t \int_{s=0}^{I_i^N(u)} d\Upsilon_i(s, u).$$

Nous conservons dans les deux états le même Υ_i , mais l'intensité des sauts fait elle-même des sauts (ce qui ne nous pose pas vraiment de problèmes). Grâce à ce subterfuge que nous avons exploité déjà au chapitre 9, les sauts auront lieu soit exactement au même instant soit se décaleront. On peut voir une illustration de la création de $J_i^N(t)$ dans la figure (4.1).

7.3 Preuve de convergence en utilisant les résultats des chapitres précédents

7.3.1 Définition de la limite

Un processus de Poisson 2D Υ étant fixé, la limite $\overline{\mathcal{X}}(t) = (\overline{e}(t), \overline{x}(t), \overline{f}(t), \overline{i}(t))$ se présente comme notre système original à 1 particule à ceci près que les sauts synchrones n'interviennent pas lorsque $\overline{x}(t)$ atteint la frontière C , ce qui est un temps d'arrêt, mais lorsque $\mathbb{E}_{\Upsilon, h_1, h_2, \dots}(\overline{x})$ qui est une valeur déterministe (en toute rigueur c'est aussi un temps d'arrêt). Ces sauts interviennent selon le tirage de variables de Bernoulli h_j avec la probabilité p de valoir 1. De plus $(\Upsilon, h_1, h_2, \dots, h_j, \dots)$ sont indépendants.

Notons maintenant $\mathcal{X}_i$ une copie de $\overline{\mathcal{X}}$ pour $(\Upsilon, h_1, h_2, \dots, h_j, \dots) = (\Upsilon_i, b_{i,1}, b_{i,2}, \dots, b_{i,j}, \dots)$. Les $b_{i,j}$ sont les mêmes que ceux qui ont permis de définir les $a_{i,j}^N$, c'est donc ceci qui réalise le couplage des sauts.

7.3.2 Convergence hors des temps d'arrêt

Les théorèmes 5.1 et 5.5 assurent la définition de $\overline{\mathcal{X}}$ et des $\mathcal{X}_i$ et la convergence du système à N particules hors des passages de frontières²². En effet, toutes les fonctions qui définissent la dynamique ont le caractère de régularité suffisant d'une part et d'autre part la croissance bornée des variables d'état du système assure bien que toutes les trajectoires restent dans une partie bornée pour tout compact du temps (si la condition initiale est contenue dans une partie bornée). Notons enfin que dans le cas de ce modèle simple, on pourrait appliquer directement une loi des grands nombres, les évolutions étant indépendantes hors des instants de sauts, on aurait même sans doute un théorème central limite. Cependant, rappelons que le but de cette section n'est pas de donner la meilleure démonstration pour ce cas-ci mais d'illustrer la méthode de couplage-système auxiliaire que nous avons développé dans les chapitres précédents. De plus la présence d'effets de bords rend l'approche trajectorielle plus visuelle.

Il nous faut régler le problème des frontières maintenant.

7.3.3 Traversée des frontières

Non-accumulation des points de sauts synchrones Pour commencer traitons le problème de l'accumulation des points de sauts pour le système original.

Lemme 7.1. *Notons S_j^N , la taille du j -ième saut du système original à N particules ($S_j^N > 0$ pour un saut vers le bas), et D_j^N le temps pour atteindre la prochaine frontière. Alors :*

- $\mathbb{P}(S_j^N \geq \frac{pC^2}{16X_m}) \xrightarrow{N \rightarrow \infty} 1,$
- $D_j^N \geq \frac{N}{r} S_j^N,$

²²Hors des sauts synchrones, l'EDS a une forme classique puisqu'on a pas de couplage par la moyenne. Aussi il n'est pas utile de faire appel à des raisonnements compliqués comme c'était le cas précédemment pour pouvoir définir le champ moyen.

Le système original vérifie l'hypothèse (6.5).

Preuve : Pour le premier point, remarquons que $\mathbb{E}(S_j^N) \geq p \frac{C}{2}$, ce serait une égalité si les tirages étaient indépendants. La valeur X_m (le diamètre de l'espace étudié) majore en particulier le débit maximal. Ainsi si nous notons N_{on} , le nombre d'utilisateurs actifs, et p_{on} la proportion des utilisateurs actifs. Le débit moyen moyen par utilisateur actif au moment de la congestion est donc $\frac{C}{p_{on}}$. Notons p_0 la proportion d'utilisateurs actifs qui ont un débit supérieur à la moitié de la moyenne, comme ce débit est par ailleurs majoré par la constante X_m , il vient $p_0 p_{on} X_m + (1 - p_0) \frac{C}{2} \geq C$, ainsi :

$$p_0 \geq \frac{C}{2p_{on}X_m - C}.$$

Or $p_{on}X_m \leq C$, ainsi :

$$p_0 \geq \frac{1}{2}.$$

On voit donc que plus de la moitié des utilisateurs actifs ont un débit supérieur à $C/2$ au moment de la congestion. Mais ce nombre d'utilisateurs actifs est par ailleurs minoré par $\frac{C}{X_m}$, par conséquent $\frac{C}{2X_m}$ utilisateurs ont un débit supérieur à $C/2$ au temp de congestion. À l'instant de la congestion, une proportion p de ces utilisateurs sera touchée (les pertes sont indépendantes du passé), ainsi lorsque N devient grand, avec une probabilité qui tend vers 1, plus de $p/2$ utilisateurs dont le débit dépasse $C/2$ sont touchés. finalement, on est assuré que la taille du saut a la taille :

$$S_j^N \geq \frac{pC^2}{16X_m}$$

avec une probabilité qui tend vers 1.

La deuxième formule résulte du fait que le rythme de croissance est limité. La conclusion est alors directe.

■

Remarque 7.3. La démonstration précédente n'est valable sous sa forme que si la probabilité de coupure des fenêtres est indépendante de la taille de la fenêtre ce qui est le cas dans le modèle mais constitue un point qui pourrait être amélioré.

En recopiant le raisonnement que nous venons de faire pour le système auxiliaire, on voit que les sauts n'admettent pas de point d'accumulation lorsque $j \rightarrow \infty$, ainsi le système auxiliaire sera défini sur tout borné du temps donc sur $\mathbb{R}^+$ en entier.

Couplage aux instants de sauts Il nous faut vérifier que les trajectoires sont bien couplées aux instants de sauts synchrones. Le problème que nous pourrions avoir serait du type $\frac{N_{on}}{N} \rightarrow 0$ auquel cas on pourrait avoir une trajectoire donnée qui ne tend pas vers sa limite. Ceci n'arrive pas comme nous l'avons vu dans la démonstration du lemme précédent.

Ainsi $\frac{N_{on}}{N} > \eta > 0$ lors d'une époque de congestion. Par conséquent $a_{i,j}^N \neq b_{i,j}$ avec probabilité $(1 - p)^{N_{on}}$, ce qui est une condition de couplage suffisante puisque $N(1 - p)^{N_{on}}$ est une série sommable. Ceci permet de vérifier l'hypothèse (6.2).

Condition de convergence vers le champ moyen aux passage des frontières

D'après le théorème 6.2 du chapitre 6, on aura la convergence vers la limite $\bar{\mathcal{X}}$ (qui ici est d'ailleurs définie de manière unique) qui coïncide avec le procédé de construction de la particule du champ moyen à travers les sauts si l'hypothèse suivante sur les limites $\bar{\mathcal{X}}_j$ du champ moyen où on retire tous les sauts à partir du numéro j (se reporter au chapitre 6 pour les précisions de construction) :

Hypothèse 7.1. *Pour chaque j , la pente du débit $\bar{\mathcal{X}}_{j,2}(t)$ au voisinage du temps d'arrêt $\mathcal{T}_j$ est strictement positive.*

Moyennant cette hypothèse qui peut ne pas être vérifiée pour certains seuils, le champ moyen de l'article [9] existe donc bien comme il est annoncé. Il est l'unique limite du système à N particules que nous avons décrit dans la section 7.1.1 :

Théorème 7.1. *Sous l'hypothèse supplémentaire 7.1, le modèle de 7.1.1 admet comme unique limite $(\mathcal{X}_1, \mathcal{X}_2, \dots)$.*

Le mot limite doit être pris au sens Cesarò-auxiliaire hors des décalage de sauts. De plus on a convergence en espérance des temps d'arrêt T_j^N vers les temps déterministes $\mathcal{T}_j$. Ces temps $\mathcal{T}_j$ n'admettent aucun point d'accumulation, aussi la convergence est valable jusqu'à $t = \infty$ ou bien au premier temps où l'hypothèse (7.1) n'est pas vérifiée.

Ce théorème est une conséquence du théorème 6.2 compte tenu des hypothèses vérifiées précédemment.

Deuxième partie

Application du champ moyen à la modélisation de flots TCP

Description Cette partie commence avec la traduction compilée en Français des articles [109] et [18]. Ensuite nous développons ces résultats et adaptons le modèle à des sources intermittentes pour modéliser le comportement d'utilisateurs de HTTP.

Le chapitre 8 présente la modélisation de TCP dans le cas de flots persistants avec l'introduction de la structure de retards. Le chapitre 9 reprend en partie ce qui a été dit dans la partie précédente mais sans se référer directement à ce qui a été fait. Nous expliquons au paragraphe suivant ce choix.

Le chapitre 10 présente la simplification et l'étude des EDP qui résulte de l'application du champ moyen à des flots TCP persistants. C'est là que se trouvent les résultats exploitables du point de vue pratique, comme certains compléments sur la stabilité de différents contrôles au niveau de la file d'attente.

Enfin le chapitre 11 referme cette partie en présentant une adaptation de notre modèle pour des sources intermittentes «HTTP». Ce modèle est une alternative qui s'inspire de celui introduit dans le chapitre 7.

Pour accompagner la lecture, on pourra se reporter aux transparents [108, 107].

Différences entre la démonstration du chapitre 9 et celle de la première partie

Le modèle que nous allons introduire pour commencer est en un sens est moins puissant que celui qui était présenté dans la partie précédente puisqu'il n'inclue pas de sources on-off. Cependant, il contient une structure de délais plus riche. Ces délais sont une difficulté de modélisation mais ils simplifient fortement les preuves ; ils sont le moteur principal de moins deux points cruciaux de la preuve :

- pour l'existence du système auxiliaire,
- pour traiter les passages de frontières au moment où la file atteint sa valeur maximale.

Dans le deuxième cas ils forment même un élément indispensable pour réaliser la preuve.

La preuve est donc basée sur la même architecture et contient beaucoup d'ingrédients communs, mais elle n'est pas directement équivalente à ce qui a été fait auparavant : on notera en plus des délais que les files d'attentes sont étudiées directement et que le modèle est multiclasse.

Chapitre 8

Une modélisation de TCP

Ce chapitre explique la modélisation que nous avons adopté pour comprendre le fonctionnement du partage de bande passante réalisé par TCP. Nous y détaillons les simplifications par rapport à la réalité et nous définissons les grandeurs que nous allons étudier.

Sommaire

8.1	Le modèle et ses équations d'évolution : processus de Markov à N particules	116
8.1.1	Modélisation de TCP : hypothèses et équations d'évolution	116
8.1.2	Modélisation des débits à partir de la valeur des fenêtres	117
8.1.3	Génération et prise en compte des pertes de paquets	117
8.1.4	Équation différentielle sur la taille de la file d'attente :	118
8.2	Discussion et explication du modèle	120
8.2.1	Modélisation de RED et réductions de fenêtres	120
8.2.2	Équation différentielle pour les fenêtres :	121
8.2.3	Limites de la modélisation	122
8.3	Hypothèses techniques	122
8.3.1	Couplage des processus de Poisson	122
8.3.2	Hypothèses portant sur l'état initial	122
8.3.3	Borne pour la taille de la fenêtre au temps t	123
8.3.4	Relation entre le RTT et la taille de la file d'attente	123
8.4	Mathématisation du modèle	124
8.4.1	Énoncé du problème	124
8.4.2	Reformulation en termes de processus prenant pour valeur des mesures	124
8.4.3	Énoncé du théorème	126
8.4.4	Discussion	127
8.5	Vers une preuve du champ moyen	128
8.5.1	Plan de la preuve	128
8.5.2	Théorèmes sur l'équation de la limite et convergence	130
8.5.3	Difficulté pour mener à bien la démarche classique par tension*	131

8.1 Le modèle et ses équations d'évolution : processus de Markov à N particules

Dans tout ce qui suit l'unité sera le paquet dont on suppose la taille fixée dans le réseau étudié. Nous reprenons les notations et le modèle de base défini dans la section 2.3.1 où nous avons comparé différents modèles de champ moyen pour TCP.

Nous étudions N connexions caractérisées par des temps de transmission hors file d'attente $(T_i)_{i=1..N}$ et des fenêtres $(W_i^N(t))_{i=1..N}$. Ces connexions sont réparties (indépendamment de N) en d CLASSES $\mathfrak{C}_c$ avec $c \in [1, 2, \dots, d]$. Dans une classe $T_i^N = T_c$.

Ces connexions partagent une file d'attente de taille $Q^N(t) = Nq^N(t)$, de débit NL et de capacité B^N . Cette file implémente des pertes anticipées avec probabilité $F(q^N)$.

8.1.1 Modélisation de TCP : hypothèses et équations d'évolution

Nous étudierons TCP dans son mode d'évitement de congestion. Il existe trois autres modes, l'attente de l'expiration d'un timeout, le «slow-start» et le plafonnement de la fenêtre de réception W_{max} .

Dans la phase d'évitement de congestion, quand une perte est constatée, la fenêtre est divisée par deux (on renvoie un paquet pour deux ACK reçus). Nous supposons que la perte peut être constatée par la méthode du Duplicate ACK (recevoir 3 ACK dupliqués), dans le cas contraire on sortirait du mode d'évitement de congestion (Congestion avoidance) après l'expiration du Timeout. Dans la pratique lors de la détection d'une perte on recommence à envoyer des paquets à partir de celui qui n'a pas été reçu par la destination. Ce sous mode de l'évitement de congestion se nomme Fast retransmit/Fast recovery (en effet on retransmet rapidement le paquet suspecté d'avoir été perdu, et on reprend comme si de rien n'était à partir de là).

Hypothèse 8.1. MODÉLISATION FLUIDE *Nous supposons que les valeurs de fenêtre, de débit et de RTT sont des réels définis à tout instant. Tout se passe donc approximativement comme si la fenêtre de la connexion i augmentait à la vitesse $1/R_i^N(t)$ paquets par seconde.*

Hypothèse 8.2. *Conformément à l'hypothèse (8.1), nous supposons que lorsque l'information selon laquelle un paquet est détruit arrive à la source, le débit sortant et la fenêtre sont instantanément divisés exactement par deux.*

Nous donnons tout de suite les équations qui correspondent à TCP. Nous justifierons les choix et la définition implicite du RTT dans la discussion de la section suivante.

Définition 8.1. *Un routeur avec N connexions étant donné, le repère de temps est pris au niveau du routeur. Nous définissons le RTT (Round Trip Time) noté $R_i^N(t)$ par connexion comme suit. C'est le temps qu'un paquet qui arrive au routeur au temps t met entre son départ de la source (à un temps $s \leq t$) et le temps où l'ACK qui lui correspond²³ revient à la source (à un temps $s' = s + R_i^N(t) \geq t$). Ce temps résulte uniquement du temps de*

²³Si le paquet est perdu on considère un ACK négatif virtuel jusqu'à la source.

transmission T_i et du temps passé dans la file d'attente supposée FIFO. Le RTT vérifie l'équation implicite :

$$R_i^N(t) = T_i + \frac{Q^N(t - R_i^N(t))}{NL} \quad (8.1)$$

où $Q^N(t) = Nq^N(t)$ sont les tailles de la file d'attente absolue et normalisée.

De même nous décidons que $W_i^N(t)$ est la valeur de la fenêtre de la connexion i au moment où le paquet qui arrive au temps t au niveau du routeur est parti de la source. Avant $t = 0$, $W_i^N(t) = w_i \in \mathbb{R}$ un réel donné, ensuite on a l'équation que nous expliquerons plus loin :

$$dW_i^N(t) = \frac{1}{R_i^N(t)} dt - \frac{W_i^N(t^-)}{2} dJ_i^N(t) \quad (8.2)$$

où J_i^N est un processus de Poisson doublement stochastique d'intensité $I_i^N(t)$ qui modélise la manière dont les sources ressentent les pertes de paquets imposées par le routeur à l'entrée de la file d'attente. Appelons $\mathbf{W}^N(t) = (W_i^N(t))_{i=1..N}$, le vecteur des tailles de fenêtres.

Par définition du concept de fenêtre de congestion, la source i a $W_i^N(t)$ en cours d'envoi au temps t . Par définition de $R_i^N(t)$, ces paquets "en vol" sont les paquets envoyés entre $t - R_i^N(t)$ et t par la source.

8.1.2 Modélisation des débits à partir de la valeur des fenêtres

Hypothèse 8.3. DÉBIT TYPE LITTLE Nous supposons que la VITESSE DE TRANSMISSION OU DÉBIT DE LA SOURCE i AU TEMPS t vaut

$$X_i^N(t) = W_i^N(t) / R_i^N(t),$$

c'est-à-dire la fenêtre au temps t divisée par le RTT du dernier ACK arrivé.

Cette définition est une approximation de la réalité, dans la mesure où tous les RTT, W paquets sont envoyés.

8.1.3 Génération et prise en compte des pertes de paquets

Modélisation des pertes générées par RED Dans ce qui suit nous supposons que les liens sont parfaits, et que les pertes qui ont lieu sont uniquement dues à la congestion au routeur. Ce routeur peut avoir sa file d'attente remplie et perdre mécaniquement des paquets (tail-drop), ou bien appliquer une politique de contrôle AQM (Active Queue Management) consistant à détruire volontairement certains paquets. Dans ce qui suit nous supposons que la file d'attente applique RED (Random Early Detection [51]). Nous allons nous placer dans le cas de tampons dans la file d'attente suffisamment grands pour pouvoir négliger les pertes qui auraient lieu alors que l'utilisation est faible.

Nous notons $B^N = Nb$ (où b est une constante) la taille de ce buffer. Une fois cette capacité d'accueil dépassée, la file d'attente détruit les nouveaux paquets. Ce système "tail-drop" vient en plus du mécanisme de destruction de paquets implémenté par RED.

Un paquet qui arrive a une probabilité d'être détruit sans même arriver dans la file d'attente. Cette probabilité vaut 0 si la taille de la file d'attente est inférieure à $Q_{min}^N = Nq_{min}$ (où q_{min} est une constante); elle croît ensuite linéairement de la valeur 0 à p_{max} lorsque la file passe de Q_{min}^N à $Q_{max}^N = Nq_{max}$ (où q_{max} est une constante), puis vaut 1 si on dépasse Q_{max}^N .

Notons que ceci n'est pas exactement l'algorithme RED tel que spécifié dans [51] : la fonction F qui sert à générer les pertes est identique mais pas appliquée de la taille courante de la file d'attente, mais à moyenne des valeurs jusqu'au temps présent (avec des poids exponentiels). Cette moyenne étant une adaptation d'un filtre de Kalman qui donne la moyenne de la fonction dans les cas périodiques. Cependant, cette moyenne n'est pas nécessaire, et elle fait empirer le retard d'action dans le système qui est la principale raison de son instabilité. En tout cas elle ne simplifie pas l'analyse mathématique d'une part, et ne marche pas mieux d'autre part, ce qui nous conduit à ne pas l'étudier.

Si chacune des N connexions est en phase d'évitement de congestion, nous pouvons reformuler cette probabilité de pertes en fonction de q^N , la taille de la file d'attente divisée par N . En effet parler de la perte d'un paquet dans un modèle fluide n'a plus aucun sens, nous devons donc bien nous intéresser à la probabilité qu'une indication de perte soit reçue par la source.

Hypothèse 8.4. *Nous supposons que la source reçoit des indications de pertes sous la forme $K^N(t) := F(q^N(t))$, où F est une fonction qui vaut 0 en dessous de $q_{min} = Q_{min}/N$ et croît linéairement à p_{max} pour $q_{max} = Q_{max}/N$ puis saute à 1 au point q_{max} .*

Bien sûr, le cas tail drop est un cas particulier de ceci avec $q_{min} = 0$, $q_{max} = b$ ($B^N = Nb$) et $p_{max} = 0$.

Prise en compte des pertes par les connexions

Hypothèse 8.5. *L'intensité stochastique des pertes ressenties par la connexion i du système de taille N vaut :*

$$I_i^N(t) := \frac{W_i^N(t - R_i^N(t))}{R_i^N(t - R_i^N(t))} K^N(t - R_i^N(t))$$

(rappelons que $W_i^N(t) = w_i$ et $K^N(t) = 0$ pour $t < 0$, de plus $I_i^N(t)$ est l'intensité stochastique du processus de Poisson $J_i^N(t)$ des divisions de fenêtres par deux).

Nous allons voir que la file d'attente quant à elle ressent directement l'intensité des pertes et pas un processus de destruction de paquets à l'entrée.

8.1.4 Équation différentielle sur la taille de la file d'attente :

Si f est une fonctions à valeurs dans $\mathbb{R}$, nous noterons $f^+ = \max(f, 0)$ et $f^- = \min(f, 0)$.

Pour $q(t) < q_{max}$, $K^N(t) = F(q^N(t))$ et la vitesse de changement du buffer fluide est donnée par

$$\begin{aligned} \frac{dQ^N(t)}{dt} &= \sum_{i=1}^N \frac{W_i^N(t)}{R_i^N(t)} (1 - K^N(t)) - NL \\ &+ \left(\sum_{i=1}^N \frac{W_i^N(t)}{R_i^N(t)} (1 - K^N(t)) - NL \right)^- \chi\{Q^N(t) = 0\}. \end{aligned}$$

La proportion $K^N(t) := F(q^N(t))$ du total du fluide est perdue, ce qui correspond à la quantité :

$$\sum_{i=1}^N \frac{W_i^N(t)}{R_i^N(t)} K^N(t).$$

Le second terme empêche que la file d'attente prenne des valeurs négatives. En effet la file d'attente peut rester à 0 en attendant qu'un nombre suffisant de connexions augmentent leur fenêtre.

En divisant par N , on obtient l'équation sur la variable normalisée q^N :

$$\begin{aligned} \frac{dq^N(t)}{dt} &= \frac{1}{N} \sum_{n=1}^N \frac{W_i^N(t)}{R_i^N(t)} (1 - K^N(t)) - L \\ &+ \left(\frac{1}{N} \sum_{i=1}^N \frac{W_i^N(t)}{R_i^N(t)} (1 - K^N(t)) - L \right)^- \chi\{q^N(t) = 0\}, \end{aligned} \quad (8.3)$$

avec $q^N(0) = q(0)$.

Si $q^N(t)$ **atteint** q_{max} **et**

$$(1 - p_{max}) \frac{1}{N} \sum_{n=1}^N \frac{W_i^N(t)}{R_i^N(t)} > L,$$

alors la file d'attente q^N rentre dans un état que nous appelons «la gigue»²⁴ à la valeur q_{max} , nous définissons la probabilité de pertes $K^N(t)$ effective par

$$(1 - K^N(t)) \frac{1}{N} \sum_{n=1}^N \frac{W_i^N(t)}{R_i^N(t)} = L.$$

En d'autres mots, si $q^N(t) = q_{max}$ alors

$$K^N(t) = \max\{p_{max}, 1 - L \left(\frac{1}{N} \sum_{i=1}^N \frac{W_i^N(t)}{R_i^N(t)} \right)^{-1}\}.$$

²⁴Comme nous en avons discuté dans le chapitre 2, la dynamique originale ne permet pas de définir un taux de pertes qui soit une fonction. C'est pourquoi nous introduisons une nouvelle dynamique qui correspond à la moyenne intuitive de la distribution solution du problème. Le nom gigue est utilisé en référence aux fluctuations qui ont lieu dans le système réel au niveau paquets.

Explications sur l'état de gigue Pour justifier la définition de $K^N(t)$ nous souhaiterions bien montrer que la probabilité de pertes d'un modèle paquet qui fait la gigue à q_{max} converge faiblement vers $K^N(t)$. Au lieu de cela nous montrons que la probabilité de pertes de *Gentle RED* (défini à la section 1.3.4) converge vers $K^N(t)$ quand Gentle RED converge vers RED (voir la section [137]). Nous devrions aussi prendre en compte les pertes provoquées par de petits buffers (c'est-à-dire si B^N est constant quand $N \rightarrow \infty$), étudié dans [129, 130]. Pour les petits buffers, les fluctuations vont provoquer des pertes de paquets bien longtemps avant que le taux de transmission n'ait atteint la capacité du lien NL . Essentiellement on peut modéliser K^N comme $L_B(\frac{1}{LN} \sum_{i=1}^N \frac{W_i^N(t)}{R_i^N(t)})$ où L_B peut être calculé en trouvant la distribution d'équilibre d'une chaîne de Markov comme dans [95]. Cependant, comme notre buffer est à l'échelle de N , les fluctuations de ce type peuvent être négligées. Il n'est pas vraiment nécessaire que nous puissions prouver une convergence en champ moyen pour le cas de petites files d'attente. La file d'attente et le RTT seraient des constantes. Pour une discussion sur ce sujet on pourra consulter [150].

8.2 Discussion et explication du modèle

8.2.1 Modélisation de RED et réductions de fenêtres

Nous nous occupons ici de la manière dont les pertes générées par RED à la file d'attente sous forme d'une fonction $K^N(t)$ impactent les fenêtres des connexions.

Notre modèle prend en compte le délai d'un RTT entre le temps où un paquet est tué et le moment où le buffer reçoit un débit moindre. Nous supposons que les réductions de fenêtres de la connexion n arrivent à cause d'une perte un RTT dans le passé. Au premier ordre, la probabilité qu'une réduction de fenêtre ait lieu entre le temps t et $t + h$ est

$$\begin{aligned} & \int_{t-R_i^N(t)}^{t+h-R_i^N(t+h)} \frac{W_i^N(s)}{RTT_i^N(s)} K^N(s) ds \\ & \sim [1 - \frac{d}{dt} R_i^N(t)] \frac{W_i^N(t - R_i^N(t))}{R_i^N(t - R_i^N(t))} K^N(t - R_i^N(t)) h. \end{aligned} \quad (8.4)$$

Comme la probabilité qu'un paquet soit détruit est proportionnelle à $I_i^N(t) := W_i^N(t - R_i^N(t)) / R_i^N(t - R_i^N(t))$, la vitesse de transmission un RTT dans le passé est multipliée par $K^N(t - R_i^N(t))$, la probabilité de pertes un RTT dans le passé. Le terme d'effet Doppler $[1 - \frac{d}{dt} R_i^N(t)]$ est une petite correction que nous négligeons mais qui avait été étudiée dans l'article [18]. Le problème est qu'elle n'apporte pas une correction spectaculaire et qu'elle reste une approximation d'un ordre difficile à estimer.

Il y a beaucoup de manières d'implémenter une destruction de paquet une fois une probabilité de pertes $p = K^N(t)$ fixée au temps t . On pourrait jeter de manière déterministe un paquet tous les $1/p$. Cependant cela pourrait introduire une sorte de synchronisation qui n'est pas souhaitée. En fait dans [51] deux méthodes sont proposées. Une première est simplement de générer une variable aléatoire Bernoulli avec probabilité p de jeter un paquet. Une deuxième méthode consiste à choisir un paquet uniformément parmi les $1/p$ paquets qui suivent.

En fait nous ne nous soucions pas trop de la méthode exacte employée. En effet quand $N \rightarrow \infty$, la contribution de chaque flot devient négligeable dans le débit total. Par conséquent, l'arrivée de paquets de la i ième connexion est très étalée parmi les autres paquets. Tant qu'on s'intéresse à une connexion en particulier, les paquets sont jetés au hasard avec la probabilité $p = K^N(t)$ au temps t (c'est en quelque sorte un principe des événements rares qui s'applique). C'est pour cette raison que nous modéliserons le processus de réduction des fenêtres par un processus ponctuel de Poisson avec l'intensité stochastique $I_i^N(t)$.

Bien sûr, la seconde méthode proposée par [51] introduirait une petite dépendance entre les différents processus de Poisson des connexions. Cependant l'interaction entre les flots via la moyenne des tailles de fenêtres et la très faible dépendance n'empêcheront pas la moyenne de converger vers une limite déterministe (ce point est du domaine de la conjecture mais il n'est pas très mystérieux). Nous supposons que la première méthode est utilisée, mais compte tenu des covariances qui tendent vers 0 très vite avec N , il serait possible d'étendre l'argument des sections 9.1 et 10.1 pour tenir compte de cette petite dépendance.

8.2.2 Équation différentielle pour les fenêtres :

Pendant la phase d'évitement de congestion, lorsque la file d'attente n'est ni vide ni pleine, la taille de la fenêtre augmente d'un paquet chaque fois qu'une fenêtre complète est reçue avec succès. Comme [18], [109] et [67] entre autres, nous supposons que ce terme prend simplement la valeur $1/R_i^N(t)$ (c'est-à-dire un incrément d'un paquet par RTT). Notons que ce fait ignore le fait que les ACK reviennent aux sources au débit NL quand la file d'attente n'est pas vide (on se reportera à [69] à ce sujet).

Si une source détecte une perte au temps $t - R_i^N(t)$ à cause de 3 ACK dupliqués, la source divise sa fenêtre courante $W_i^N(t^-)$ par deux vers la valeur $W_i^N(t^-)/2$. Le seuil de slow start (sssthresh) vaut alors (si la fenêtre n'arrive pas en dessous de la valeur 2) $H_i^N(t) = W_i^N(t^-)/2$. La source fait son fast retransmit et fast recovery. Le paquet perdu est retransmis et on continue à envoyer des paquets à peu près comme si la taille de la fenêtre était constante (pour une description précise se reporter à la section 1.1). Nous ignorons cet effet et supposons que la taille de la fenêtre croît à la vitesse $1/R_i^N(t)$ même pendant cette phase. Dans [18], nous avons inclus un terme $(1 - \chi_{S_i^N(t)})$ pour essayer de tenir compte de cet effet. Nous avons vu depuis qu'il n'y a pas d'effet très significatif à le supprimer. L'évitement de congestion continue normalement dès la réception de l'ACK du paquet perdu. Ainsi l'évolution de la fenêtre dans la phase d'évitement de congestion est décrite par l'équation différentielle stochastique suivante

$$dW_i^N(t) = \frac{1}{R_i^N(t)}dt - \frac{W_i^N(t^-)}{2}dJ_i^N(t), \quad (8.5)$$

avec $W_i^N(0) = w_i$, $i = 1, \dots, N$ donnés.

Nous ne tenons pas compte dans l'évolution de la taille de la fenêtre de la fenêtre maximale de réception. Ceci amènerait à écrire $dW_i^N(t) = \max\left(W_{max,i}; \frac{1}{R_i^N(t)}dt - \frac{W_i^N(t^-)}{2}dJ_i^N(t)\right)$ à la place de 8.5.

8.2.3 Limites de la modélisation

Ce modèle ne prend pas en compte les paquets, aussi les sources devront appliquer des pertes à certains instants avec une intensité stochastique $K^N(t)$. Ceci signifie par exemple qu'une intensité de pertes 1 n'entraîne pas qu'aucun paquet n'est reçu, mais qu'en moyenne, la fenêtre est divisée par deux toutes les unités de temps, et qu'entre temps une partie du fluide passe dans le réseau. Par ailleurs, comme nous allons le voir, le taux de perte n'a pas le même sens pour les connexions et pour la file d'attente. Les connexions voient $K^N(t)$ comme une intensité stochastique de pertes, alors que la file d'attente va détruire de manière déterministe cette proportion du débit qui lui arrive.

D'autre part la supposition selon laquelle le débit évolue en suivant l'hypothèse Little (8.3) qui est une approximation stationnaire est particulièrement mauvaise quand la taille de la file d'attente et donc les temps de transmissions varient rapidement. Cependant cette hypothèse pourrait sans mal être amendée sans avoir à modifier en profondeur les démonstrations qui suivent.

8.3 Hypothèses techniques

8.3.1 Couplage des processus de Poisson

Conformément à la démarche suivie dans les chapitres précédents, introduisons le couplage suivant sur les sauts :

Définition 8.2.

$$J_i^N(t) = \int_0^t \int_0^\infty \chi_{[0, I_i^N(v)]}(u) \Upsilon_i(du, dv)$$

où les $\Upsilon_i(u, v)$ sont des processus de Poisson à deux dimensions sur $[0, T] \times [0, \infty)$ d'intensité 1.

$\{\Upsilon_i\}_{i=1..N}$ est donc une séquence iid de processus de Poisson à deux dimensions d'intensité 1 définis sur un espace de probabilités $(\Omega, \mathcal{F}, P)$. De plus le processus de Poisson en $\{(t, u) : 0 \leq u \leq \bar{I}(t), 0 \leq t \leq T\}$ avec $\bar{I}$ défini d'après l'hypothèse 8.6 donnera une borne inférieure a priori sur le taux de transmission.

Définition 8.3. Définissons $\mathcal{F}_t = \sigma\{\Upsilon_n(u, v); v \leq t, n = 1, 2, \dots\}$.

Dans la mesure où le hasard n'est généré que par les Υ_i et les conditions initiales w_i et $q(0)$, la formule ci-dessus et le fait que $W_i^N(0) = w_i$ et $q^N(0) = q(0)$ définissent entièrement le couplage de notre problème.

8.3.2 Hypothèses portant sur l'état initial

Hypothèse 8.6. Nous supposons que

- i) Avant le temps 0, la taille de la fenêtre de la connexion n parmi les N du système vaut la constante w_i indépendante de N où $0 \leq w_i \leq W_{max}$.

- ii) Le temps de transmission T_i de la connexion i satisfait $T_{min} \leq T_i \leq T_{max}$ pour tout i .
- iii) $q^N(0) = q(0)$ est constant.

L'hypothèse i) est bien vérifiée si on tient compte de l'existence d'une fenêtre maximale de réception.

8.3.3 Borne pour la taille de la fenêtre au temps t

Notons $a(t) := W_{max} + \frac{t}{T_{min}}$. De l'hypothèse i) il vient pour tout i ,

$$a(t) \geq w_i + \frac{t}{T_i} \geq W_i^N(t)$$

à tout temps t . L'intensité stochastique du processus ponctuel de Poisson J_i^N des pertes de la connexion i est $I_i^N(s) \leq a(s)/T_{min} =: \bar{J}(t)$ pour tout $0 \leq t \leq T$.

Par définition de la fonction de pertes de RED, $(1 - K^N(t)) \geq 1 - p_{max}$ tant que $q^N(t) < q_{max}$. Si $q^N(t) = q_{max}$ alors $(1 - K^N(t)) \geq LT_{min}/a(t)$.

Définition 8.4. $k_{max} := \max(p_{max}, 1 - \frac{LT_{min}}{a(t)})$.

Proposition 8.1. Pour tout $t \in [0, T]$, $(1 - K^N(t)) \geq (1 - k_{max}) > 0$.

8.3.4 Relation entre le RTT et la taille de la file d'attente

Définissons $\phi_i^N(s)$ le RTT futur inscrit sur un paquet de type i qui quitte le source au temps s . Dans ce cas $\phi_i^N(s) = T_i + q^N(s)/L$. Notons aussi que $s + \phi_i^N(s)$ est strictement monotone car sa dérivée, si $q^N(t) < q_{max}$, vaut

$$\begin{aligned} 1 + \frac{1}{L} \frac{dq^N(s)}{ds} &= 1 + \frac{1}{L} \left(\frac{1}{N} \sum_{i=1}^N \frac{W_i^N(t)}{R_i^N(t)} (1 - K^N(t)) - L \right) \\ &\geq \frac{1}{L} \frac{1}{N} \sum_{i=1}^N W_i^N(t) \frac{1}{T_{min}} (1 - k_{max}). \end{aligned}$$

Ceci est positif à moins que toutes les tailles de fenêtres ne soient 0 ce qui a une probabilité nulle. Si $q^N(t) = q_{max}$ alors la dérivée vaut 1.

Nous redonnons une définition du RTT et nous expliquons ce choix :

Définition 8.5. Le RTT de la connexion i comme la valeur écrite virtuellement sur les paquets qui arrivent au routeur au temps t par $R_i^N(t) := t - s = \phi_i^N(s)$ avec $s + \phi_i^N(s) = t$.

Ceci est naturel dans la mesure où le paquet qui arrive à l'instant t au routeur est le résultat du contrôle exercé sur le flux à un temps s . On sait que l'information arrive exactement au temps $s + \phi_i^N(s)$. Aussi on a bien $t = s + \phi_i^N(s)$.

Comme $s + \phi_i^N(s)$ est strictement monotone, $R_i^N(t)$ est bien défini et $s = t - \phi_i^N(s) = t - R_i^N(t)$. En substituant ceci dans $\phi_i^N(s) = T_i + q^N(s)/L$ nous prouvons que R_i^N satisfait à l'équation suivante :

$$R_i^N(t) = T_i + q^N(t - R_i^N(t))/L. \quad (8.6)$$

De plus, en prenant la dérivée de (8.6) nous voyons que :

$$(1 - \dot{R}_i^N(t)) = \frac{1}{1 + \dot{q}^N(t - R_i^N(t))/L} \quad (8.7)$$

$$= L \left(\frac{1}{N} \sum_{n=1}^N \frac{W_n^N(t - R_i^N(t))}{R_n^N(t - R_i^N(t))} (1 - K^N(t - R_i^N(t))) \right)^{-1}. \quad (8.8)$$

Ceci signifie que $t - R_i^N(t)$ va croître plus rapidement qu'une fonction de la forme λt si on peut assurer que le débit ne se rapproche pas trop de 0. Ceci nous permettra de montrer que la fonction inverse est Lipschitz.

8.4 Mathématisation du modèle

8.4.1 Énoncé du problème

Le problème que nous étudions peut s'énoncer comme suit.

Problème 8.1. *Des connexions et un routeur vérifient les hypothèses de modélisation (8.1), (8.3), (8.2), (8.4), et (8.5). Leurs constantes d'état sont T_i , le temps de transmission, w_i , la fenêtre, q_0 , la file d'attente initiale (déterministes) et L , le débit par utilisateur. Leurs variables d'état $W_i^N(t)$, $q^N(t)$, $R_i^N(t)$, évoluent selon les équations (8.5), (8.3 et suivantes) et (8.6), couplées par les pertes d'intensité I_i^N générées par les sauts du processus de Poisson J_i^N défini par 8.2.*

Nous souhaitons étudier le comportement de ce système lorsque N tend vers l'infini.

8.4.2 Reformulation en termes de processus prenant pour valeur des mesures

Classe de connexions : Rappelons qu'il y a d classes de connexions $\mathfrak{C}_c$, $c = 1, \dots, d$ dont le temps de transmission est T_c . Immédiatement $R_i^N = R_c^N$ pour tout $i \in \mathfrak{C}_c$. Nous allons aussi supposer que la proportion des connexions dans la classe c est $\mathfrak{c}_c^N := \frac{\text{Card}(\mathfrak{C}_c \cap [1, N])}{N}$. Nous rajoutons à 8.6 l'hypothèse suivante :

Hypothèse 8.7. *Hypothèses sur les classes de connexion :*

- iv) *La proportion des utilisateurs toute la classe c converge : $\mathfrak{c}_c^N \rightarrow \mathfrak{c}_c$ pour $c = 1, \dots, d$ lorsque $N \rightarrow \infty$.*
- v) *Notant μ_c^N l'histogramme des fenêtres des connexions de la classe c à temps 0. Nous supposons que pour toutes les classes c , μ_c^N converge faiblement vers une limite μ_c quand $N \rightarrow \infty$. On rajoute que μ_c doit avoir son support dans $[0, W_{\max}]$.*

Processus à valeur dans l'espace des mesures : Pour pouvoir étudier le comportement limite du système quand le nombre N de connexions tend vers l'infini, nous devons d'abord définir le processus empirique (voir Dawson [35, 36]) de ces connexions dans la classe $\mathfrak{C}_c^N = \mathfrak{C}_c \cap [1, N]$. Pour tout borélien A , nous définissons

$$M_c^N(t, A) := \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N \chi_A(W_i^N(t)) \chi\{i \in \mathfrak{C}_c\} \quad (8.9)$$

le processus à valeur dans l'espace des mesures de probabilités associé. Formellement les valeurs sont prises dans $M_1(\mathbb{R}^+)$, muni de la topologie de la convergence faible.

Pour déterminer le futur : La suite $\{\Upsilon_i(u, v)\}_{i \in \mathbb{N}}$ de processus de Poisson indépendants et d'intensité 1 a été définie sur l'espace de probabilités $\{\Omega, \mathcal{F}, P\}$. $\mathbf{W}^N(t) := (W_1^N(t), \dots, W_N^N(t))$, $q^N(t)$ et $M^N(t) := (M_1^N(t), \dots, M_d^N(t))$ peut être construit chemin par chemin comme des processus de $\{\Omega, \mathcal{F}, P\}$ prenant valeurs dans $(\mathbb{R}^+)^{\infty}$, $\mathbb{R}^+$ et $M_1(\mathbb{R}^+)^d$ où les coordonnées de $\mathbf{W}^N(t)$ après N valent 0. Il suffit de supposer que $W_i^N(t) = w_i$ pour $t \leq 0$ et de construire la solution du système (8.5), (8.3) chemin par chemin de saut en saut de Υ_i .

Récrivons (8.5) : Soit $\langle g, \mu \rangle = \int g(w) \mu(dw)$ et $\langle Id, \mu \rangle = \int w \mu(dw)$, il vient

$$\overline{W}_c^N(s) := \langle Id, M_c^N(s) \rangle = \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N W_i^N(s) \chi\{i \in \mathfrak{C}_c\}.$$

Notons $\mathcal{G} = \{f \in \mathcal{C}_b^1(\mathbb{R}^+, \mathbb{R}), f(0) = 0\}$, l'ensemble de fonctions continues bornées à dérivée continue bornée qui valent 0 au bord. Si $g \in \mathcal{G}$, 8.5 implique :

$$\begin{aligned} & \langle g, M_c^N(t) \rangle - \langle g, M_c^N(0) \rangle \\ &= \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N \chi\{i \in \mathfrak{C}_c\} \int_0^t \left[\frac{dg}{dw}(W_i^N(s)) \frac{ds}{R_c^N(s)} + (g(W_i^N(s^-)/2) - g(W_i^N(s^-))) dJ_i^N(s) \right]. \end{aligned} \quad (8.10)$$

Dans la section 10.1, nous considérons la limite du système précédent lorsque N tend vers l'infini pour obtenir une équation d'évolution des distributions de fenêtres.

Récrivons (8.3) : Pour $q^N(t) < q_{max}$, $K^N(t) = F(q^N(s))$ et

$$\begin{aligned} & q^N(t) - q(0) \\ &= \int_0^t \left[\sum_{c=1}^d \mathfrak{c}_c^N \langle Id, M_c^N(s) \rangle \frac{(1 - K^N(s))}{R_c^N(s)} - L \right. \\ & \quad \left. + \left(\sum_{c=1}^d \mathfrak{c}_c^N \langle Id, M_c^N(s) \rangle \frac{(1 - K^N(s))}{R_c^N(s)} - L \right)^- \chi\{q^N(s) = 0\} \right] ds \end{aligned} \quad (8.11)$$

où

$$R_c^N(t) = T_c + q^N(t - R_c^N(t))/L. \quad (8.12)$$

Si q^N est dans l'état de gigue à q_{max} alors la probabilité $K^N(t)$ est donnée par

$$K^N(t) = \max \left\{ p_{max}, 1 - L \left(\frac{\sum_{c=1}^d \mathbf{c}_c^N \langle Id, M_c^N(t) \rangle}{R_c^N(t)} \right)^{-1} \right\}.$$

8.4.3 Énoncé du théorème

Rappelons brièvement la définition 3.8. Nous pouvons définir une métrique pour la convergence faible pour les mesures de probabilité sur $[0, \infty)$ en définissant la distance entre les probabilités μ et ν comme

$$\|\mu - \nu\|_w := \sum_{k=1}^{\infty} 1 \wedge |\langle f^k, \mu \rangle - \langle f^k, \nu \rangle| \frac{1}{2^k}$$

où $\langle f, \mu \rangle = \int_0^\infty f(w) \mu(dw)$ et $(f^k)_k$ est une famille dénombrable de fonctions qui déterminent la convergence. Enfin comme précédemment nous notons $\|\cdot\|_s$ ou d_0 , la distance de Skorokhod de $D[0, T]$.

Théorème 8.1. HYPOTHÈSES : Nous étudions un système infini de particules défini par le problème 8.1 sous les hypothèses technique $[i]$ - $[v]$ (8.6 et 8.7).

CONCLUSION 1 : CONVERGENCE MARGINALE La mesure aléatoire des tailles de fenêtres des connexions dans chacune des classes c converge en probabilité vers une mesure déterministe $M_c(t, dw)$; c'est-à-dire $\|M_c^N(t, dw) - M_c(t, dw)\|_w \rightarrow 0$ en probabilité quand $N \rightarrow \infty$. $M_c(t, dw)$ est la distribution marginale de $M_c(s - R_c(s), dv; s, dw)$, la limite déterministe de la distribution jointe des tailles de fenêtres aux temps t et $t - R_c(t)$.

CONCLUSION 2 : ÉQUATIONS DE LA MESURE ET DE LA FILE D'ATTENTE Soit $\mathcal{G} = \{g \in C_b^1(\mathbb{R}^+) : g(0) = 0\}$ où $C_b^1(\mathbb{R}^+)$ est l'espace des fonctions bornées à dérivées continues bornées. Pour $g \in \mathcal{G}, c = 1, \dots, d$,

$$\begin{aligned} & \langle g, M_c(t) \rangle - \langle g, M_c(0) \rangle \\ &= \int_0^t \left[\frac{1}{R_c(s)} \left\langle \frac{dg(w)}{dw}, M_c(s, dw) \right\rangle \right. \end{aligned} \quad (8.13)$$

$$\left. + \left\langle (g(w/2) - g(w))v, M_c(s - R_c(s), dv; s, dw) \right\rangle \frac{1}{R_c(s - R_c(s))} K(s - R_c(s)) \right] ds$$

$$= \int_0^t \left[\frac{1}{R_c(s)} \left\langle \frac{dg(w)}{dw}, M_c(s, dw) \right\rangle \right. \quad (8.14)$$

$$\left. + \left\langle (g(w/2) - g(w)), e(s, s - R_c(s), w) M_c(s, dw) \right\rangle \frac{1}{R_c(s - R_c(s))} K(s - R_c(s)) \right] ds,$$

où $\langle g, M_c(t) \rangle = \int_{w=0}^\infty g(w) M_c(t, dw)$, avec $R_c(t) = T_c + q(t - R_c(t))/L$ et $K(t) = F(q(t))$ pour $q(t) < q_{max}$ en notant

$$e(s, s - R_c(s), w) = \left\langle v, \frac{M_c(s - R_c(s), dv; s, dw)}{M_c(t, dw)} \right\rangle$$

l'espérance conditionnelle d'une fenêtre au temps $t - RTT$ sachant qu'elle a atteint la valeur w au temps t .

De plus la taille de file d'attente converge en probabilité vers une limite déterministe $q(t)$ qui satisfait à l'équation :

$$\begin{aligned} \frac{dq(t)}{dt} = & \sum_{c=1}^d \mathbf{c}^c \langle w, M_c(t, dw) \rangle \frac{(1 - K(t))}{R_c(t)} - L \\ & - \left(\sum_{c=1}^d \mathbf{c}^c \langle w, M_c(t, dw) \rangle \frac{(1 - K(t))}{R_c(t)} - L \right)^- \chi\{q(t) = 0\}, \end{aligned} \quad (8.15)$$

où $q(t) = q_{max}$, $K(t)$ est déterminé par l'équation

$$K(t) = \max(p_{max}, \frac{1}{L} \left(\sum_{c=1}^d \mathbf{c}^c \langle w, M_c(t, dw) \rangle \frac{1}{R_c(t)} \right)^{-1}),$$

avec $\langle w, M_c(t, dw) \rangle = \int_w w M_c(t, dw)$ Lipschitz en t pour chacune des classes c .

L'exploitation numérique de ces équations fait l'objet du chapitre ??.

8.4.4 Discussion

Notons qu'il peut y avoir des discontinuités dans $K(t)$ quand $q(t)$ atteint q_{max} . Le problème à q_{max} apparaît du fait que F n'est pas continue et encore moins Lipschitz en ce point. Pour justifier cette définition allons considérer une adaptation inspirée de *Gentle RED* (voir [137]) un AQM adapté de RED pour lequel F croît linéairement de p_{max} à 1 entre q_{max} et $2q_{max}$. F est donc Lipschitz. Dans notre adaptation la borne $2q_{max}$ va être remplacée par $q_{max} + \epsilon$ par exemple et nous allons tenter de vérifier que lorsque ϵ tend vers 0, la limite est la même en distribution que la gigue que nous avons introduit directement pour modéliser le mécanisme de la drop tail.

Il est important de bien remarquer que nous n'avons jamais essayé de justifier que le modèle fluide que nous étudions est une limite de quelque modèle discret au niveau paquets. Nous prouvons beaucoup moins, à savoir que le modèle fluide intuitif introduit dans [113] converge vers une limite en champ moyen. Cette convergence est implicitement supposée dans la bibliographie d'ingénierie, et les limites qui en sont tirées sont ensuite étudiées pour prouver la stabilité de tel ou tel mécanisme de contrôle AQM de la congestion. Nous faisons aussi l'hypothèse simplificatrice en ce qui concerne les taux de pertes qui nous amène à négliger l'effet Doppler ($[1 - \frac{d}{dt} R_i^N(t)]$ dans (8.4)). Mais dans la mesure où nous ne sommes pas capables de donner l'ordre de l'erreur ainsi commise, autant rester sur le modèle le plus simple. Ce problème est résolu numériquement dans le chapitre 12 où on fabrique un vrai transport de fluide qui intègre implicitement l'effet Doppler. On voit alors que cet effet peut être important dans certains cas et ne devrait pas être négligé si la taille de la file d'attente est d'une dimension comparable au temps de transmission. Cependant nous nous tiendrons à la démarche historique dans cet exposé.

En ce qui concerne RED, la conclusion pour le choix des paramètres optimaux est négative. Quel que soit le choix que l'on puisse faire, il existe des cas de nombre d'utilisateurs ou de RTT qui le rendent instable autant globalement que localement autour des points fixes. Cette conclusion a été vérifiée sur notre modèle par simulations, tout comme dans la réalité. C'est la raison pour laquelle alors que RED est une option de tous les routeurs aujourd'hui sur le marché, il n'est quasiment jamais mis en fonction.

Nous sommes plus intéressés par la preuve mathématique de la convergence vers le champ moyen que par les détails de modélisation. Aussi nous ne nous gênerons pas pour négliger les aspects concernant les timeouts, le slow start ou les détails précis dans la phase de fast retransmit. L'exposé dans ce qui suit est focalisé sur les idées principales. Cependant nous allons expliquer comment le modèle pourrait être amélioré. La preuve peut fonctionner pour prouver la convergence en champ moyen dans bien d'autres cas que RED.

8.5 Vers une preuve du champ moyen

Le modèle de TCP que nous étudions peut être vu comme N systèmes dynamiques (les tailles des N fenêtres $\mathbf{W}^N := (W_1^N, \dots, W_N^N)$) qui évoluent indépendamment sauf par l'intermédiaire d'une ressource partagée (la file q^N). La dynamique de la ressource partagée ne dépend que de la distribution des systèmes dynamiques (dans le cas présent de la taille moyenne de la fenêtre). C'est-à-dire que nous nous trouvons dans un cadre proche de celui du problème 6.1 que nous avons posé au chapitre 6.

Nous ne sommes pas en mesure d'appliquer directement les résultats précédents, trois raisons expliquent cela :

- principalement, la structure particulière des délais, $R_i^N(t)$, dont la dépendance n'est pas explicite,
- la ressource partagée, q^N , ne dépend de plus de choses que la seule valeur instantanée du système,
- la dynamique de sauts synchrones ne rentre pas dans les hypothèses du théorème 6.2 et nous ne serons pas en mesure d'utiliser la méthode de convergence par étude du candidat limite.

8.5.1 Plan de la preuve

L'approche standard est de prouver l'existence puis l'unicité de la limite. C'est ce que nous faisons mais l'innovation mathématique (qui a été reprise dans la première partie de cette thèse) est que nous définissons d'abord un système auxiliaire $(\mathcal{W}^N, \mathcal{Q}^N)$ qui converge mieux que le système original et on le verra vers la même limite. La modification est faite uniquement dans la dynamique de la ressource partagée qui au lieu de dépendre de la moyenne de la distribution dépend de l'espérance de cette moyenne.

Comme pour le système auxiliaire nous avons une ressource partagée déterministe, les équations des individus deviennent indépendantes. Ainsi il est facile de prendre une sous-suite de la ressource partagée qui va converger (ici $\mathcal{Q}^N \rightarrow \mathcal{Q}$). Il est ensuite facile de prouver que les individus $\mathcal{W}_i^N$ convergent vers une limite $\mathcal{W}_n$ le long de l'extraction

et pour chaque n . Ceci prouve l'existence d'un système auxiliaire infini (ici $(\mathcal{W}, \mathcal{Q})$ où $\mathcal{W} = (\mathcal{W}_1, \mathcal{W}_2, \dots)$). Ensuite la remarque clef est que par la loi des grands nombres

$$\overline{\mathcal{W}}(t) := \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N \mathcal{W}_i(t) = E\overline{\mathcal{W}}(t).$$

C'est-à-dire que le système auxiliaire infini est exactement une limite du système initial (ce qui est exactement l'idée de *lookdown process*) !

Nous pouvons donc récrire $(\mathcal{W}, \mathcal{Q})$ comme $(\mathbf{W}, q)$ avec $\mathbf{W} := (W_1, W_2, \dots)$ dans la mesure où le terme d'interaction est $\overline{\mathbf{W}}(t)$; i.e. l'interaction se fait à travers le fenêtre moyenne sans besoin sans l'espérance.

Pour poursuivre nous utilisons un argument de couplage pour montrer que le système original (ici chaque W_i^N) converge presque sûrement dans la norme de Skorokhod vers la limite infinie (ici W_i^N converge vers W_i). Ceci prouve la propagation du chaos, chaque système dynamique (W_i) est indépendant des autres et interagit avec les autres uniquement par l'intermédiaire d'une ressource qui est déterministe. On peut déduire de cette dernière propriété que le modèle du champ moyen converge (ici que q^N tend vers q et que l'histogramme des $\mathbf{W}^N$ tend vers le champ moyen). Nous avons utilisé l'approche de couplage suggérée par Thomas Kurtz ([41]) de *bring back the particles* (remettre les particules en selle) en utilisant son *lookdown process* ; c'est-à-dire de ne pas projeter tout sur les histogrammes des tailles de fenêtres comme dans les arguments standard. Il s'agit d'une approche directe qui est une bonne alternative à la méthode de tension-limites-unicité devenue l'approche standard de nos jours (ce point était illustré dans le chapitre 3).

Comparaison par rapport à la première partie Nous aurions pu calquer la démarche des chapitres précédents. C'est à dire introduire un couplage sur le système original (ce que nous avons fait en 8.2). Puis modifier la dynamique en remplaçant dans l'équation d'une particule la moyenne empirique du système par sa propre loi, en partant de la condition initiale limite du système. Ceci nous aurait donné la particule du champ moyen $(\overline{\mathcal{W}}, \overline{\mathcal{Q}})$ (il serait alors nécessaire d'introduire à l'avance la norme Cesarò-auxiliaire). Ensuite nous aurions pu définir les candidats limite $((\mathcal{W}_i)_{i \in \mathbb{N}}, \mathcal{Q})$, indépendants pour chaque indice de particule i avec $\mathcal{Q} = \overline{\mathcal{Q}}$ déterministe. Ceci nous aurait permis de prouver une loi des grands nombres sur les candidats limites puis la preuve de convergence aurait suivi.

Comme nous le voyons la démarche est donc légèrement différente car on introduit un système auxiliaire $\mathcal{W}_i^N, \mathcal{Q}^N$ dont on extrait une limite qui en fait est le même objet que ce que nous définissons directement dans l'approche de la première partie. Plusieurs raisons expliquent ce choix :

- historiquement c'est ainsi que nous avons prouvé la convergence de notre modèle,
- cette méthode est exactement celle préconisée par Thomas Kurtz et elle mérite d'être comparée à ce que nous avons fait dans la première partie,
- nous utilisons les délais pour prouver l'existence du système auxiliaire et la convergence, ce qui peut être une méthode extrêmement puissante,
- cette méthode donne un résultat plus fort concernant la limite puisque l'extraction est réalisée «d'un seul coup» à travers les sauts synchrones, nous n'avons donc pas de raccords de convergences à faire.

L'intérêt de l'extraction à travers les sauts synchrones est le suivant, si le système ne converge pas effectivement vers la limite du champ moyen (ce qui pourrait se produire sous des hypothèses moins fortes concernant les conditions initiales par exemple), il reste vrai que la limite extraite est une limite possible du système auxiliaire, alors que notre approche de la première partie échouera et ne donnera pas de limite.

8.5.2 Théorèmes sur l'équation de la limite et convergence

Nous souhaitons montrer que $(\mathbf{W}^N(t), q^N(t))$, l'unique solution du système à N corps converge quand $N \rightarrow \infty$. Dans la section 9.1 nous commençons par prouver l'existence d'une telle limite.

Théorème 8.2. *Si les hypothèses 8.6 et 8.7 sont vérifiées alors il existe une unique solution forte $(\mathbf{W}, q, (M_1, \dots, M_d))$ au système suivant. Pour $q(t) < q_{max}$, $K(t) = F(q(t))$ et*

$$\begin{aligned} & q(t) - q(0) \\ &= \int_0^t \left[\sum_{c=1}^d \mathfrak{c}^c \langle Id, M_c(s) \rangle \frac{(1 - K(s))}{R_c(s)} - L \right. \\ & \quad \left. + \left(\sum_{c=1}^d \mathfrak{c}^c \langle Id, M_c(s) \rangle \frac{(1 - K(s))}{R_c(s)} - L \right)^- \chi\{q(s) = 0\} \right] ds. \end{aligned} \quad (8.16)$$

Quand $q(t) = q_{max}$, $K(t)$ satisfait

$$K(t) = \max\{p_{max}, 1 - L(\sum_{c=1}^d \mathfrak{c}^c \langle Id, M_c(s) \rangle \frac{1}{R_c(t)})^{-1}\}.$$

Chaque fenêtre évolue selon

$$dW_i(t) = \frac{1}{R_i(t)} dt - \frac{W_i(t^-)}{2} dJ_i(t), \quad (8.17)$$

où $W_i(0) = w_i$, $n = 1, \dots$ sont spécifiés, où

$$J_i(t) = \int_0^t \int_0^\infty \chi_{[0, I_i(t)]}(v) d\Upsilon_i(u, v) \text{ où } I_i^N(s) = \frac{W_i(s - R_i(s))}{R_i(s - R_i(s))} K(s - R_i(s))$$

avec $M_c(t)$ défini par

$$\langle g, M_c(t) \rangle = \lim_{N \rightarrow \infty} \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N g(W_i(t)) \chi\{i \in \mathfrak{C}_c\}$$

ce qui permet de définir

$$\overline{W}_c(s) = \langle Id, M_c(t) \rangle = \lim_{N \rightarrow \infty} \frac{1}{N \mathfrak{c}_c^N} \sum_{i=1}^N W_i(s) \chi\{i \in \mathfrak{C}_c\}.$$

Ainsi de (8.16)

$$\begin{aligned} q(t) - q(0) = & \int_0^t \left[\sum_{c=1}^d \mathfrak{c}_c \frac{\overline{W}_c(s)}{R_c(s)} (1 - K(s)) - L \right. \\ & \left. + \left(\sum_{c=1}^d \mathfrak{c}_c \frac{\overline{W}_c(s)}{R_c(s)} (1 - K(s)) - L \right)^- \chi\{q(s) = 0\} \right] ds. \end{aligned} \quad (8.18)$$

Les solutions $q(t)$ et $\mathbf{R}(t) = (R_1(t), R_2(t), \dots)$ ainsi que les $M_c(t)$ sont déterministes. Finalement les composantes de $\mathbf{W}$ sont des processus indépendants.

Définissons

$$\overline{\mathbf{S}}_N(s) = \frac{1}{N} \sum_{n=1}^N \frac{W_n^N(s)}{R_n^N(s)} = \sum_{c=1}^d \mathfrak{c}_c^N \frac{\overline{W}_c^N(s)}{R_c^N(s)}. \quad (8.19)$$

où $\overline{W}_c^N(s)$ est la taille moyenne des fenêtre de classe $\mathfrak{C}_c$ parmi les premières $W_1^N, \dots, W_N^N$. Définissons $\overline{\mathbf{S}}_N(s)$ de manière analogue à $\mathbf{W}$. Du théorème 8.2 nous pouvons définir

$$\overline{\mathbf{S}}(s) = \lim_{N \rightarrow \infty} \overline{\mathbf{S}}_N(s) = \lim_{N \rightarrow \infty} \frac{1}{N} \sum_{n=1}^N \frac{W_n(s)}{R_n(s)}. \quad (8.20)$$

Dans la section 9.2 nous allons prouver le théorème 8.3 et montrer qu'il y a une unique solution forte de (8.18), (8.17) et que la solution de (8.5) et (8.3) converge vers cette solution forte.

Théorème 8.3. *Si les hypothèses 8.6 et 8.7 sont vérifiées, alors $\|M_c^N(t) - M_c(t)\|_w$, $\|q^N(t) - q(t)\|_s$ et $\|K^N(t) - K(t)\|_s$ convergent vers 0 en probabilité où $M_c(t)$, $q(t)$ et $K(t)$ sont des fonctions déterministes de $t \in \mathbb{R}^+$ dans $M_1(\mathbb{R}^+)$, $\mathbb{R}^+$ et $\mathbb{R}^+$ respectivement donnés dans le théorème 8.2.*

8.5.3 Difficulté pour mener à bien la démarche classique par tension*

Soit P^N la mesure induite sur $D^d[0, T] \times C[0, T]$ par $(\overline{W}_1^N, \dots, \overline{W}_d^N), q^N$.

Lemme 8.1. *Sous les hypothèses 8.6 et 8.7, les mesures P^N sont tendues.*

Preuve : Nous vérifions les conditions du théorème 12.3 dans [23] que nous avons rappelées dans le théorème 4.3. La condition (i) est immédiate comme les suites $((\overline{W}_1^N, \dots, \overline{W}_d^N), q^N(t))$ sont bornées.

Dans la condition (ii) nous nous donnons des constantes positives ϵ et η . Nous pouvons choisir Δ suffisamment petit pour que l'accroissement maximum d'une fenêtre pendant une durée Δ soit moindre que $\epsilon/2$; c'est-à-dire $\Delta < T_{\min} \epsilon/2$. Δ doit aussi être suffisamment petit pour que la probabilité de l'événement B qu'un processus de Poisson d'intensité $\overline{J}(T)$ saute deux fois pendant la durée Δ soit plus petite que $\eta \epsilon a(T)/2$.

Notons que, par la construction des fenêtres W_n^N , l'événement que la fenêtre soit coupée en deux fois pendant Δ jusqu'au temps T est contenue dans B_n , l'événement

où $\Upsilon_n(t, \bar{J}(T))$ a deux sauts dans un intervalle de durée Δ . Notons aussi que la pire oscillation qu'une fenêtre puisse faire durant Δ est $a(T)$; c'est-à-dire le plus grand saut possible. Aussi le module $w'_c(\Delta)$ de toute trajectoire de $\bar{W}_c^N, c = 1, \dots, d$ comme défini dans (12.6) dans [23] satisfait

$$w'_c(\Delta) \leq \frac{1}{N} \sum_{n=1}^N (a(T) \chi_{B_n} + \epsilon/2).$$

Ainsi,

$$\begin{aligned} P^N(w'_c(\Delta) \geq \epsilon) &\leq P(a(T) \frac{1}{N} \sum_{n=1}^N \chi_{B_n} \geq \epsilon/2) \\ &\leq \frac{2}{\epsilon a(T)} P(B) \leq \eta. \end{aligned}$$

Comme nous pouvons faire cette estimation pour chacune des d classes, nous vérifions bien la condition (ii) pour les oscillations de $(\bar{W}_1^N, \dots, \bar{W}_d^N)$.

Les oscillations de $q^N(t)$ et de $(R_1^N(t), \dots, R_d^N(t))$ sont bornées uniformément dans $C[0, T]$ car les trajectoires sont uniformément Lipschitz. Ceci montre que P^N est tendue.

■

Utilisant le lemme 8.1 nous pouvons extraire une sous suite N_k telle que q^N et $(\bar{W}_1^N, \dots, \bar{W}_d^N)$ convergent presque sûrement vers q^∞ et $(\bar{W}_1^\infty, \dots, \bar{W}_d^\infty)$ dans la norme de Skorokhod. La convergence des composantes W_n^N s'en suit. Malheureusement nous n'avons même pas été capables de prouver la première étape : c'est-à-dire que les limites q^∞ et $(\bar{W}_1^\infty, \dots, \bar{W}_d^\infty)$ sont déterministes.

Utilisant la méthode ci-dessus et le critère de Jakubowski²⁵ (cf. [36], Théorème 3.6.4) nous serions peut-être capables de vérifier que les processus à valeur dans les mesures $M_1^N(t), \dots, M_d^N(t)$ dans $\mathbb{D}([0, T], M_1(\mathbb{R}^+))$ sont tendus. Nous pourrions peut-être même vérifier que les processus de mesures à deux dimensions $M_1^N(t - R^N(t), dv; t, dw), \dots, M_d^N(t - R^N(t), dv; t, dw)$ dans $\mathbb{D}([0, T], M_1(\mathbb{R}^+)^2)$ sont tendues. Nous pourrions prendre une sous suite et faire l'analyse de 10.1. Nous obtiendrions une solution limite au théorème 8.1. Malheureusement nous ne savons pas que la solution est unique et déterministe parce que

²⁵Ce critère ramène le problème de la tension dans un espace quelconque E à celui de la tension dans $\mathbb{R}$.

Lemme 8.2. CRITÈRE DE JAKUBOWSKI (1986)

HYPOTHÈSES : (E, d) Polonais, F une famille de fonctions $E \rightarrow \mathbb{R}$ qui sépare les points et close sous addition sont donnés. Pour $f \in F$, on définit $\tilde{f}$ qui à x càdlàg dans E associe $(\tilde{f}x)(t) = f(x(t))$ dans les fonctions càdlàg de $\mathbb{R}$.

CONCLUSION : Une suite de mesures (P_n) sur $\mathbb{D}(\mathbb{R}^+, E)$ est tendue c'est-à-dire relativement séquentiellement compacte pour la topologie faible $M_1(\mathbb{D}(\mathbb{R}^+, E))$ si et seulement si les deux conditions suivantes sont réalisées :

1. $\forall T > 0, \forall \epsilon > 0, \exists K \in E$ compact, $P_n(\mathbb{D}([0, T], K)) > 1 - \epsilon$
2. (P_n) est faiblement tendue par rapport à F , c'est-à-dire : $\forall f \in F, (P_n \circ (\tilde{f}^{-1}))$ est tendue.

les équations (8.13) et (10.8) ne déterminent pas la solution comme $(M_1(t), \dots, M_d(t))$ et $q(t)$ ne forment pas un état du système (à cause du délai).

Il serait éventuellement possible de rectifier cela en définissant l'état au temps t comme la trajectoire entière de toutes les particules $W_1^N(t), \dots, W_d^N(t)$ et $q(t)$ un RTT au moins dans le passé. La procédure numérique de chapitre 12 montre comment maintenir cette information pour résoudre numériquement (8.13) et (10.8). Cependant plutôt que d'essayer de montrer la tension dans un espace aussi gros, nous allons procéder directement dans le chapitre suivant.

Chapitre 9

Démonstration du champ moyen

Ce chapitre apporte les preuves des théorèmes énoncés dans le chapitre précédent.

Dans la section 8.5.2, nous avons expliqué en détail ce qui devait être réalisé dans cette partie, ce plan est légèrement différent de ce que nous avons fait dans la partie mathématique pour tenir compte des délais. Loin d’être une difficulté supplémentaire, nous allons voir comment les délais nous aident à déployer la méthode de couplage-système auxiliaire.

Pour débiter la preuve nous introduisons notre système auxiliaire et démontrons l’existence d’une limite pour le système original (section 9.1). Puis dans la section 9.2 nous prouvons l’unicité de la limite. Alors que l’existence d’une limite est obtenue «d’un seul coup» à travers toutes les singularités, la partie unicité se fait pas à pas et nous abordons le problème délicat des changements de dynamique en utilisant les délais du système pour passer les frontières.

La section 10.1 du chapitre où nous établissons les équations du champ moyen, achève la démonstration de 8.1 et clôt en même temps les parties mathématiques de cette thèse pour laisser progressivement la place aux applications.

Sommaire

9.1	Existence d’une limite	135
9.1.1	Système infini de particules auxiliaire	136
9.1.2	Premières propriétés du système auxiliaire	137
9.1.3	Existence d’une limite pour le système auxiliaire	139
9.1.4	Équations de la limite $(\mathcal{W}, Q)$:	142
9.2	Unicité de la solution forte	143
9.2.1	Convergence loin de la frontière	147
9.2.2	Quand on traverse ou qu’on frôle le bord	153
9.2.3	Convergence en champ moyen sur la frontière	156
9.3	RED est la limite faible de Gentle RED*	159

9.1 Existence d’une limite

Dans cette section nous prouvons l’existence d’une solution à (8.18), (8.17).

9.1.1 Système infini de particules auxiliaire

Nous introduisons maintenant le système auxiliaire que nous avons évoqué plus tôt où on force q^N à être déterministe en modifiant l'équation d'évolution de la file d'attente vers (9.1). Ensuite nous extrayons un candidat limite qui se valide bien dans le système initial.

Soit $\mathcal{K}^N(t) = F(\mathcal{Q}^N(t))$ si $\mathcal{Q}(t) < q_{max}$ où $\mathcal{Q}^N$ est défini par

$$\begin{aligned} & \mathcal{Q}^N(t) - q(0) \\ &= \int_0^t \left[\sum_{c=1}^d \mathfrak{c}_c^N (\mathbb{E} \langle Id, \mathcal{M}_c^N(s) \rangle) \frac{(1 - \mathcal{K}^N(s))}{\mathcal{R}_c^N(s)} - L \right. \\ & \quad \left. + \left(\sum_{c=1}^d \mathfrak{c}_c^N (\mathbb{E} \langle Id, \mathcal{M}_c^N(s) \rangle) \frac{(1 - \mathcal{K}^N(s))}{\mathcal{R}_c^N(s)} - L \right)^- \chi\{\mathcal{Q}(s) = 0\} \right] ds \end{aligned} \quad (9.1)$$

où $\mathcal{W}^N = (\mathcal{W}_1^N, \dots, \mathcal{W}_N^N)$ satisfait l'analogue de (8.5), où $\mathcal{W}_i(0) = w_i$, où

$$\mathcal{R}_c^N(t) = T_c + \mathcal{Q}^N(t - \mathcal{R}_c^N(t))/L, \quad (9.2)$$

et où

$$\langle Id, \mathcal{M}_c^N(s) \rangle = \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N \mathcal{W}_i^N(s) \chi\{i \in \mathfrak{c}_c\}.$$

Si $\mathcal{Q}^N(t) = q_{max}$ alors la probabilité de pertes $\mathcal{K}^N(t)$ satisfait

$$\mathcal{K}^N(t) = \max \left\{ p_{max}, 1 - L \left(\sum_{c=1}^d \mathfrak{c}_c^N \frac{\mathbb{E} \langle Id, \mathcal{M}_c^N(t) \rangle}{\mathcal{R}_c^N(t)} \right)^{-1} \right\}. \quad (9.3)$$

Remarquons que quand $\mathcal{Q}^N$ touche q_{max} , la probabilité $\mathcal{K}^N(t)$ saute de p_{max} à une valeur qui maintient $\dot{\mathcal{Q}}^N(t) = 0$ tant que

$$\sum_{c=1}^d \mathfrak{c}_c^N \frac{\mathbb{E} \langle Id, \mathcal{M}_c^N(t) \rangle}{\mathcal{R}_c^N(t)} (1 - p_{max}) \geq L.$$

Nous montrerons ensuite que $\mathbb{E} \langle Id, \mathcal{M}_c^N(t) \rangle$ est une fonction Lipschitz du temps, donc en fait $\mathcal{K}^N(t) = p_{max}$ juste avant $\mathcal{Q}^N(t)$ de quitter la frontière.

Solution pour N donné : Soit $t^c(0) = 0$, $t^c(k+1) = t^c(k) + T_c + \mathcal{Q}^N(t^c(k))/L$ tel que $t^c(k+1) - \mathcal{R}_c^N(t^c(k+1)) = t^c(k)$. Tant que nous pouvons la définir, la suite $(t^c(k))$ est croissante en k . Soit $\Phi^c(t) =$ le premier k tel que $t^c(k) > t$. Nous allons construire notre solution par récurrence du temps t_i au temps t_{i+1} en définissant $t_{i+1} = \min_c t^c(\Phi^c(t_i))$ partant de $t_0 = 0$.

À t_0 , supposons que $\mathcal{W}^N(t)$ et $\mathcal{Q}^N(t)$ sont donnés (peut-être constants) pour $t \leq 0 = t_0$. Supposons que $(\mathcal{W}^N(t), \mathcal{Q}^N(t))$ est bien défini pour $t \leq t_i$ où t_i est le temps tel que

$t_i = t^c(k)$ pour un c et un k . C'est vrai à temps t_0 . Alors $\Phi^c(t_i)$ et $t^c(\Phi^c(t_i))$ sont définis pour toutes les classes ainsi que t_{i+1} .

Ainsi si $t \leq t_{i+1}$,

$$\mathcal{I}_n^N(t) = \frac{\mathcal{W}_n^N(t - \mathcal{R}_n^N(t))}{\mathcal{R}_n^N(t - \mathcal{R}_n^N(t))} \mathcal{K}^N(t - \mathcal{R}_n^N(t))$$

peut être défini, car pour chaque classe $t - \mathcal{R}_c^N(t) \leq t_i$ pour $s \leq t_{i+1}$ par définition de t_{i+1} (se remettre (8.8) en mémoire). Ainsi le processus ponctuel $\mathcal{I}_n^N(t)$ et les trajectoires $\mathcal{W}^N(t)$ sont définis sur $[0, t_{i+1}]$. Elles sont bornées et mesurables. Ainsi on peut définir les espérances puis $\mathcal{Q}^N(t)$. On a ainsi vérifié l'hypothèse de récurrence jusqu'à t_{i+1} .

Pour conclure nous devons montrer que $t_i \rightarrow \infty$. Remarquons que $t_{i+d+1} \geq \min\{t^c(\Phi^c(t_i) + 1) : c = 1, \dots, d\}$. Sinon les $d + 1$ valeurs $t_{i+j}, j = 1, \dots, d + 1$ doivent être choisies parmi d valeurs $\{t^c(\Phi^c(t_i))\}$ ce qui est impossible. Concluons que $t_{i+d+1} \geq t_i + T_{min}$ et ainsi que $t_i \rightarrow \infty$ lorsque $i \rightarrow \infty$.

9.1.2 Premières propriétés du système auxiliaire

Caractère Lipschitz de $\mathbb{E}\langle Id, \mathcal{M}_c^N(t) \rangle$ et de $\mathcal{Q}^N$: les sommations qui suivent sont de la forme $\sum_{i=1}^N \dots \chi\{i \in \mathfrak{C}_c\}$, ce qui signifie que nous sommes en fait sur $\mathfrak{C}_c \cap [1, N]$, les indices des éléments de classe c dans le système à N particules. Rappelons de plus que $a(T)$ est la borne sur la valeur fenêtres avant le temps T (cf. section 8.3.3).

$$\begin{aligned} & |\mathbb{E}\langle Id, \mathcal{M}_c^N(t+h) \rangle - \mathbb{E}\langle Id, \mathcal{M}_c^N(t) \rangle| \\ & \leq \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N \mathbb{E} |\mathcal{W}_i^N(t+h) - \mathcal{W}_i^N(t)| \chi\{i \in \mathfrak{C}_c\} \\ & \leq \frac{h}{T_{min}} \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N P(\mathcal{W}_i^N(s^-) = \mathcal{W}_i^N(s), t \leq s \leq t+h) \chi\{i \in \mathfrak{C}_c\} \end{aligned} \quad (9.4)$$

$$+ a(T) \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N P(\mathcal{W}_i^N(s^-) \neq \mathcal{W}_i^N(s), \text{ pour un } t \leq s \leq t+h) \chi\{i \in \mathfrak{C}_c\}. \quad (9.5)$$

Le deuxième terme vient car même de nombreux sauts ne vont pas dépasser la fenêtre maximale en longueur cumulée.

(9.4) vaut moins que h/T_{min} et tend vers zéro lorsque $h \rightarrow 0$. (9.5) est bornée par la probabilité qu'une fenêtre fasse un saut dans un intervalle de longueur h . La fonction d'intensité $I^N(t)$ est bornée par $\bar{I}(t) = a(t)/T_{min}$. Ainsi la probabilité d'un saut dans un intervalle de longueur h est bornée par $1 - \exp(-a(T)h/T_{min})$. Donc, $\mathbb{E}\langle Id, \mathcal{M}_c^N(t) \rangle$ est équi-Lipschitz sur $t \in [0, T]$ pour $N \in \mathbb{N}$. De (9.1) il découle que $\mathcal{Q}^N$ est aussi Lipschitz.

Soit $m_c^N(t) = \mathbb{E}\langle Id, \mathcal{M}_c^N(t) \rangle$.

Borne inférieure pour $\dot{m}_c^N(t)$:

Lemme 9.1. *Les dérivées de $m_c^N(t)$ et $\mathcal{Q}^N(t)$ sont uniformément bornées en N .*

Nous savons déjà que ces quantités sont bornées supérieurement (les débits sont bornés et les accroissements de fenêtres aussi), l'information supplémentaire est la borne inférieure.

Preuve : Prenons l'espérance,

$$\begin{aligned} \mathbb{E}\mathcal{W}_i^N(t) - w_i &= \mathbb{E} \int_0^t \left[\frac{1}{\mathcal{R}_i^N(s)} ds - \frac{\mathcal{W}_i^N(s^-)}{2} dI_i^N(s) \right] \\ &= \int_0^t \frac{1}{\mathcal{R}_i^N(s)} ds - \mathbb{E} \int_0^t \left[\frac{\mathcal{W}_i^N(s^-)}{2} \frac{\mathcal{W}_i^N(s - \mathcal{R}_i^N(s))}{\mathcal{R}_i^N(s - \mathcal{R}_i^N(s))} \mathcal{K}^N(s - \mathcal{R}_i^N(s)) \right] ds, \end{aligned}$$

alors,

$$\begin{aligned} \frac{d\mathbb{E}\mathcal{W}_i^N(t)}{dt} &= \frac{1}{\mathcal{R}_i^N(t)} - \mathbb{E} [\mathcal{W}_i^N(t^-) \mathcal{W}_i^N(t - \mathcal{R}_i^N(t))] \frac{1}{2\mathcal{R}_i^N(t - \mathcal{R}_i^N(t))} \mathcal{K}^N(t - \mathcal{R}_i^N(t)) \\ &\geq \frac{1}{T_{max} + q_{max}/L} - \frac{\mathbb{E}\mathcal{W}_i^N(t^-)}{2} \frac{a(t)}{T_{min}}. \end{aligned}$$

Comme $\mathbb{E}\mathcal{W}_i^N(t) \leq a(t)$, on voit que $d\mathbb{E}\mathcal{W}_i^N(t)/dt$ est strictement plus grand que $-C$ où C est une constante positive indépendante de N et de n . On prend la moyenne des $\mathfrak{c}_c N$ fenêtres avec le temps de transmission T_c . Ceci montre que les $m_c^N(t)$ sont bornés inférieurement uniformément en N .

Or $\mathcal{R}_i^N(t) \leq T_{max} + q_{max}/L$, la dérivée de $\frac{1}{N} \sum_{i=1}^N \frac{\mathbb{E}\mathcal{W}_i^N(t)}{\mathcal{R}_i^N(t)}$ est donc bornée inférieurement par la même constante. Le fait que $\mathcal{Q}^N(t)$ soit uniformément borné inférieurement en N provient de (9.1). ■

Nous souhaitons réaliser des extractions valables indépendamment des problèmes de frontières. C'est ce qui motive le lemme suivant :

Lemme 9.2. *La suite de fonctions $\mathcal{K}^N(t)$ sur $[0, T]$ est compacte (preuve séquentielle) pour la métrique de Skorokhod (cf. chapitre 4).*

Preuve : Tant que $\mathcal{Q}^N(t) < q_{max}$, $\mathcal{K}^N(t) = F(\mathcal{Q}^N(t))$ et donc $\mathcal{K}^N(t)$ est uniformément Lipschitz par rapport à N . De la même manière, tant que $\mathcal{Q}^N(t) = q_{max}$, $\mathcal{K}^N(t) = 1 - L \left(\sum_{c=1}^d \mathfrak{c}_c^N m_c^N(t) / \mathcal{R}_c^N(t) \right)^{-1}$. Dans ces deux cas la possibilité d'extraction d'une limite provient du caractère uniformément Lipschitz en N . Le problème provient des sauts quand $\mathcal{Q}^N(t)$ touche la frontière.

Nous vérifions les conditions du théorème 12.3 dans [23] (qui est rappelé en théorème 4.3 avec les définitions du module de continuité dans le chapitre 4). La condition (i) est triviale, en effet, $\mathcal{K}^N(t)$ est uniformément borné. Pour vérifier (ii) il faut montrer que les oscillations sur de petits intervalles après avoir enlevé les grands sauts sont aussi petites que nous le voudrions. Pour tout $\epsilon > 0$ définissons $A^N = \{j_i\}$ l'ensemble des temps associés aux sauts plus grand que $\epsilon/2$. Le nombre de points dans cet ensemble est uniformément borné en N et l'espace entre les points de A^N est borné inférieurement

uniformément en N par un δ_0 . Ceci suit immédiatement du lemme 9.1 car, après un saut de $\mathcal{K}(t)$ lorsque $\mathcal{Q}(t)$ touche q_{max} de taille plus grande qu' ϵ , le temps nécessaire pour décroître jusqu'à p_{max} est strictement borné inférieurement. Choisissons $\delta < \delta_0$ tel que $\mathcal{Q}^N(t)$ et $1 - L \left(\sum_{c=1}^d \mathbf{c}_c^N m_c^N(t) / \mathcal{R}_c^N(t) \right)^{-1}$ oscille moins qu' $\epsilon/4$ sur les intervalles de longueur moindre que δ . Poursuivons en considérant une partition A^N par les points $0 = t_0 < t_1 < \dots < t_v = T$ qui sont δ -rares ; c'est à dire tels que $\min_{1 \leq i \leq v} (t_i - t_{i-1}) > \delta$ qui inclue les points de A^N . Ceci est possible parce que ces points sont séparés par plus que δ_0 .

Considérons maintenant les oscillations maximales sur tout intervalle $[t_{i-1}, t_i]$. Il n'y a aucun saut de taille supérieure à $\epsilon/2$ comme $\mathcal{K}^N$ est continue à droite et les grands sauts font partie des points finaux à gauche en t_{i-1} . Comme les sauts ne vont vers le haut que de p_{max} ils n'ajoutent finalement pas beaucoup. La plus grande oscillation possible est donc $\epsilon/2$ pour un saut plus $\epsilon/4 + \epsilon/4$ pour les oscillations de $F(\mathcal{Q}^N(t))$ et $1 - L \left(\sum_{c=1}^d \mathbf{c}_c^N m_c^N(t) / \mathcal{R}_c^N(t) \right)^{-1}$. Ainsi $w'_{\mathcal{K}^N}(\delta)$ comme dans la définition de [23] est plus petit qu' ϵ , ce qui établit la condition (ii). ■

9.1.3 Existence d'une limite pour le système auxiliaire

Dans cette sous section nous allons extraire des sous suites. Pour plus de clarté nous sous entendrons ce fait dans les notations jusqu'à la fin de la sous section.

Extraction d'une limite de $\mathcal{Q}^N$ et $\mathbb{E}\langle Id, \mathcal{M}_c^N(t) \rangle$: $\mathcal{Q}^N(t)$ est déterministe. De plus, l'intégrande dans (9.1) est bornée par une constante B parce que les tailles des fenêtres jusqu'au temps t sont bornées par $a(T)$ et le RTT est plus grand que $T_{min} > 0$. Ainsi $\mathcal{Q}^N$ est uniformément Lipschitz pour $N \in \mathbb{N}$ et $t \in [0, T]$. Il en découle l'existence d'une sous suite et d'une fonction Lipschitz $\mathcal{Q}(t)$ telles que $\mathcal{Q}^N(t) \rightarrow \mathcal{Q}(t)$ uniformément par le théorème d'Ascoli-Arzelà et le fait qu'une limite (uniforme) de fonctions Lipschitz est Lipschitz.

Nous avons montré au dessus que $\mathbb{E}\langle Id, \mathcal{M}_c^N(t) \rangle$ est Lipschitz. En utilisant encore le théorème d'Ascoli-Arzelà, nous pouvons extraire une sous sous suite qui vérifie pour toutes les classes c , $m_c^N(t) = \mathbb{E}\langle Id, \mathcal{M}_c^N(t) \rangle$ converge uniformément vers une fonction Lipschitz $m_c(t)$.

Notons que prendre la limite $N \rightarrow \infty$ dans le lemme 9.1 donne que la dérivée de $m_c(t)$ et $\mathcal{Q}(t)$ est bornée inférieurement.

Convergence du RTT : C'est une conséquence directe de la convergence de $\mathcal{Q}^N, \mathcal{R}_i^N(t)$ converge uniformément vers $\mathcal{R}_i(t)$ où $\mathcal{R}_i(t) = T_i + \mathcal{Q}(t - \mathcal{R}_i(t))/L$.

Extraction de la limite de $\mathcal{K}^N(t)$: Par le lemme 9.2 nous pouvons extraire une sous suite telle que $\mathcal{K}^N$ converge en topologie de Skorohod vers la limite $\mathcal{K}$ de $D[0, T]$. Si

$\mathcal{Q}(t) < q_{max}$ alors $\mathcal{Q}^N(t) < q_{max}$ pour N assez grand. Comme $\mathcal{K}^N(t) = F(\mathcal{Q}^N(t))$, on a $\mathcal{K}(t) = F(\mathcal{Q}(t))$.

Si $\mathcal{Q}(t) = q_{max}$ nous souhaitons montrer que $\mathcal{K}(t)$ vaut

$$\max \left\{ p_{max}, 1 - L \left(\sum_{c=1}^d \mathbf{c}_c \frac{m_c(t)}{\mathcal{R}_c(t)} \right)^{-1} \right\}. \quad (9.6)$$

Supposons d'abord que t n'est pas un point où $\mathcal{K}$ saute. Alors il existe un petit intervalle $I = (t - \delta_0, t]$ où $\mathcal{Q}(s) = q_{max}$ pour $s \in I$. Il y a deux possibilités :

$$\sum_{c=1}^d \mathbf{c}_c m_c(t) \frac{1 - p_{max}}{\mathcal{R}_c(t)} > L \text{ ou } \sum_{c=1}^d \mathbf{c}_c m_c(t) \frac{1 - p_{max}}{\mathcal{R}_c(t)} = L.$$

Dans le premier cas $\sum_{c=1}^d \mathbf{c}_c m_c(t) \frac{1 - F(\mathcal{Q}(t))}{\mathcal{R}_c(t)} > L$ pour $t \in [t - \delta, t]$ et δ , $0 < \delta < \delta_0$, assez petit. Or $\mathcal{Q}^N(t)$ converge à $\mathcal{Q}(t)$, pour N assez grand, $\mathcal{Q}^N(t)$ est dans un petit tube autour de $\mathcal{Q}(t)$. De plus,

$$\frac{\sum_{c=1}^d \mathbf{c}_c^N m_c^N(t)}{\mathcal{R}_c^N(t)} (1 - F(\mathcal{Q}^N(t))) \rightarrow \frac{\sum_{c=1}^d \mathbf{c}_c m_c(t)}{\mathcal{R}_c(t)} (1 - F(\mathcal{Q}(t))).$$

La partie de droite est strictement plus grande que L pour $t \in [\rho - \delta, \rho]$ si δ est suffisamment petit. Il en est de même pour la partie gauche pour N assez grand. Ceci montre que pour N grand, $\mathcal{Q}^N(t)$ va toucher la frontière dans un temps compris dans $[t - \delta, t]$. Aussi, pour N grand,

$$\mathcal{K}(t) = \lim_{N \rightarrow \infty} \mathcal{K}^N(t) = \lim_{N \rightarrow \infty} \left(1 - L \left(\sum_{c=1}^d \mathbf{c}_c^N \frac{m_c^N(t)}{\mathcal{R}_c^N(t)} \right)^{-1} \right) = \left(1 - L \left(\sum_{c=1}^d \mathbf{c}_c \frac{m_c(t)}{\mathcal{R}_c(t)} \right)^{-1} \right).$$

Dans le second cas, $\mathcal{K}(t) = p_{max} = F(q_{max})$ donc

$$\mathcal{K}(t) = \max \left\{ p_{max}, 1 - L \left(\sum_{c=1}^d \mathbf{c}_c \frac{m_c(t)}{\mathcal{R}_c(t)} \right)^{-1} \right\}.$$

Nous avons donc établi que $\mathcal{K}(t) = F(\mathcal{Q}(t))$ si $\mathcal{Q}(t) < q_{max}$ et

$$\mathcal{K}(t) = \max \left\{ p_{max}, 1 - L \left(\sum_{c=1}^d \mathbf{c}_c \frac{m_c(t)}{\mathcal{R}_c(t)} \right)^{-1} \right\} \quad (9.7)$$

si $\mathcal{Q}(t) = q_{max}$ à tous les temps t où $\mathcal{K}(t)$ ne saute pas. Comme la dérivée de $m_c(t)$ et $\mathcal{Q}(t)$ sont bornées inférieurement, il existe un intervalle de temps sans saut à la droite de chaque saut. Donc $\mathcal{K}(t)$ est continu à droite.

Il s'ensuit que

$$\mathcal{K}(t) = 1 - L \left(\sum_{c=1}^d \mathbf{c}_c \frac{m_c(t)}{\mathcal{R}_c(t)} \right)^{-1}$$

au temps de saut.

Convergence uniforme de $\mathcal{W}_i^N$: Fixons la coordonnée i . Clairement $\mathcal{W}_i^N(t) \leq a(t)$ pour tout $0 \leq t \leq T$ il vient donc que $\mathcal{I}_i^N(t)$ est uniformément bornée par $\bar{I}(t) = a(t)/T_{\min}$ pour tout $0 \leq t \leq T$ et tout $N \geq n$. Par conséquent $\mathcal{J}_i^N(t) \leq \bar{J}_i(t)$ où $\bar{J}_i(t) = \int_0^t \int_0^\infty 1_{[0, \bar{I}(s)]}(u) \Upsilon_n(du, ds)$. Donc, si pour une trajectoire ω de Υ_i , $\bar{J}_i(T) \leq m$ alors $\mathcal{I}_i^N(T) \leq m$ pour tout i et tout N .

$\mathcal{K}$ étant donné, pour chaque trajectoire ω de Υ_i nous pouvons résoudre le système :

$$\mathcal{W}_i(t) - w_i = \int_0^t \left[\frac{1}{\mathcal{R}_i(s)} ds - \frac{\mathcal{W}_i(s^-)}{2} d\mathcal{J}_i(s) \right] \quad (9.8)$$

où $\mathcal{W}_i(t) = w_i$, pour $t \leq 0$, et $i = 1, \dots$ avec

$$\mathcal{J}_i(t) = \int_0^t \int_0^\infty 1_{[0, \mathcal{I}_i(s)]}(u) \Upsilon_i(du, ds)$$

et

$$\mathcal{I}_i(s) = \frac{\mathcal{W}_i(s - \mathcal{R}_i(s))}{\mathcal{R}_i(s - \mathcal{R}_i(s))} \mathcal{K}(s - \mathcal{R}_i(s)).$$

Considérons la solution de (9.8) d'un point de saut de $\mathcal{I}_i$. Soient $T_m; m = 1, 2, \dots$ les sauts de $\mathcal{I}_i$ dans $[0, T]$ correspondant aux points de saut $(X_m, Y_m); m = 1, 2, \dots$ de Υ_i satisfaisant $T_m = X_m$ et $Y_m \leq I_i(X_m)$. La solution de (9.8) pour $t \geq T_m$, est l'incrément positif déterministe de la taille de la fenêtre jusqu'au temps T_{m+1} et ceci n'a aucune chance d'atteindre (X_{m+1}, Y_{m+1}) qui est choisi selon un processus de Poisson. Donc, avec probabilité 1, la trace $(t, \mathcal{I}_i(t))$ pour $0 \leq t \leq T$ évitera les points de $\Upsilon_i(du, ds)$ pour un ω fixé. Ainsi nous pouvons placer une bande de largeur ϵ autour de la trace $(t, \mathcal{I}_i(t))$. Maintenant, aussi longtemps que $\mathcal{I}_i^N(t)$ reste dans cette bande, les processus ponctuels $\mathcal{I}_i$ et $\mathcal{I}_i^N$ sont identiques. Ce sera le cas si $(\mathcal{R}_i^N(t))^{-1}$ est assez proche de $(\mathcal{R}_i(t))^{-1}$ parce que les solutions $\mathcal{W}_i^N(t)$ et $\mathcal{W}_i(t)$ partant du même point w_i vont rester ensemble dans la bande entre les sauts, et seront coupés ensemble aux instants de saut. Ainsi et avec probabilité 1, $\mathcal{W}_i^N(t)$ converge uniformément vers $\mathcal{W}_i(t)$ quand $N \rightarrow \infty$.

Équation de la limite $\mathcal{W}_i$: Comme $\mathcal{R}_i^N(s)$, $\mathcal{K}^N(s)$ et $\mathcal{W}_i^N(s)$ convergent uniformément vers $\mathcal{R}_i(s)$, $\mathcal{K}(s - \mathcal{R}_i(s))$ et $\mathcal{W}_i(s)$ respectivement, il s'ensuit que $\mathcal{I}_i^N(t)$ converge uniformément vers

$$\mathcal{I}_i(t) = \frac{\mathcal{W}_i(s - \mathcal{R}_i(s))}{\mathcal{R}_i(s - \mathcal{R}_i(s))} \mathcal{K}(s - \mathcal{R}_i(s))$$

sur $[0, T]$. Les deux cotés de

$$\mathcal{W}_i^N(t) - w_i = \int_0^t \frac{1}{\mathcal{R}_i^N(s)} ds - \int_0^t \int_{u=0}^\infty \frac{\mathcal{W}_i^N(s^-)}{2} 1_{[0, \mathcal{I}_i^N(s)]}(u) \Upsilon_i(du, ds)$$

convergent donc uniformément sur $D[0, T]$ vers (9.8).

9.1.4 Équations de la limite $(\mathcal{W}, Q)$:

Détermination de $\overline{\mathcal{W}}_c$: Comme $\mathcal{Q}(t)$ est déterministe, les équations de (9.8) sont indépendantes. Ceci signifie que

$$\mathbb{E}\overline{\mathcal{W}}_c(s) = \overline{\mathcal{W}}_c(s) = \lim_{N \rightarrow \infty} \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N \mathcal{W}_i(s) \chi\{i \in \mathfrak{C}_c\}.$$

Notons que cette limite est valable pour tous les N sans qu'il ne soit nécessaire de réaliser une extraction.

Ceci est essentiellement une conséquence de la loi des grands nombres. Considérons les $\mathcal{W}_i(s) = f(\mathcal{W}_i(0), \Upsilon_i)$ définis par (9.8). f a pour domaine $E = [0, \infty) \times \mathbb{D}([0, T], \mathbb{R}^+)$ où $\mathcal{W}_i(0) \in \mathbb{R}^+ = [0, \infty)$ et $\Upsilon_i \in \mathbb{D}([0, T], \mathbb{R}^+)$, l'espace des fonctions cadlag. E est un espace métrique muni de la distance $m = e \oplus d$ où e est la distance Euclidienne sur $[0, \infty)$ et d la distance de Skorohod sur $\mathbb{D}([0, T], \mathbb{R}^+)$.

Nous avons montré ci-dessus que, pour tout point initial w_0 , l'espace des trajectoires $\Upsilon_i^{\omega_0}$ ($\Upsilon_i^{\omega_0}$ est le processus ponctuel Υ_i évalué aux points ω_0) telles que le graphe associé $(t, I_i(t))$ ne touche aucun des points de saut de $\Upsilon_i^{\omega_0}$ avec probabilité 1. f est continue sur cet ensemble par les arguments du dessus. Si le graphe $(t, I_i(t))$ évite les points de $\Upsilon_i^{\omega_0}$ alors suivant [23] ou bien [44], pour un point $q = (w_1, \Upsilon_i^{\omega_1}) \in E$ proche de $p = (w_0, \Upsilon_i^{\omega_0})$ on peut trouver une application strictement croissante θ de $[0, T]$ dans lui-même avec $\sup_{[0, T]} |\theta(t) - t| \leq \gamma(\theta)$ où $\gamma(\theta) \rightarrow 0$ quand $q \rightarrow p$ tel que $\Upsilon_i^{\omega_1}(\theta(t))$ est uniformément proche de $\Upsilon_i^{\omega_0}(t)$. Notons que cela signifie que les sauts se produisent en même temps.

La solution de (9.8) pour $\mathcal{W}_i$ et I_i pour p et pour $(w_1, \Upsilon_i^{\omega_1}(\theta(t)))$ sont donc uniformément proches. Donc pour tout temps fixé s où $\Upsilon_i^{\omega_0}$ ne saute pas, nous aurons $\mathcal{W}_i^{\omega_1}(s)$ et $I_i^{\omega_1}(s)$ uniformément proches de $\mathcal{W}_i^{\omega_0}(\theta(s))$ et $I_i^{\omega_0}(\theta(s))$. Comme $\theta(s)$ est arbitrairement proche de s et comme la probabilité d'un saut proche de s tend vers zéro, il vient que pour q proche de p , $\mathcal{W}_i^{\omega_1}(s)$ est arbitrairement proche de $\mathcal{W}_i^{\omega_0}(s)$. Ceci prouve que f est continue en $(w_0, \Upsilon_i^{\omega_0})$ tant que $\Upsilon_i^{\omega_0}$ ne saute pas en s . Ceci a probabilité 1.

Par hypothèse,

$$\lim_{N \rightarrow \infty} \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N \delta_{\mathcal{W}_i(0)} \chi\{i \in \mathfrak{C}_c\} \rightarrow \mu_c$$

et les Υ_i sont des processus de Poisson i.i.d. Ainsi la mesure empirique des paires $(\mathcal{W}_i(0), \Upsilon_i)$ converge ; c'est-à-dire

$$\frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N \delta_{(\mathcal{W}_i(0), \Upsilon_i)} \chi\{i \in K_c^N\} \rightarrow \mu_c \otimes \nu$$

où ν est la distribution de Υ_i sur $\mathbb{D}([0, T], \mathbb{R}^+)$. Le résultat en découle comme f est bornée et que l'ensemble des discontinuités a une probabilité nulle.

Détermination de $\mathcal{M}_c$: Ce qu'il y a au dessus signifie en plus que $\mathcal{M}_c^N(t)$ converge faiblement vers $\mathcal{M}_c(t)$ P -presque sûrement. Pour toute fonction continue à support compact,

définissons la mesure limite $\mathcal{M}_c(t)$ par

$$\langle g, \mathcal{M}_c(t) \rangle = \lim_{N \rightarrow \infty} \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N g(\mathcal{W}_i^N(t)) \chi\{i \in \mathfrak{C}_c\}.$$

La limite déterministe existe presque sûrement par l'argument précédent. On a donc aussi que $m_c(t) = \langle Id, \mathcal{M}_c(t) \rangle$ donc quand $\mathcal{Q}(t) = q_{max}$, $\mathcal{K}(t)$ satisfait

$$\mathcal{K}(t) = \max \left\{ p_{max}, 1 - L \left(\sum_{c=1}^d \mathfrak{c}_c^N \langle Id, \mathcal{M}_c(s) \rangle \frac{1}{R_c(t)} \right)^{-1} \right\}.$$

En particulier,

$$\begin{aligned} \overline{\mathcal{W}}_c(t) &:= \lim_{N \rightarrow \infty} \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N \mathcal{W}_i^N(t) \chi\{i \in \mathfrak{C}_c\} \\ &= \langle Id, \mathcal{M}_c(t) \rangle. \end{aligned}$$

Existence de solutions fortes : Le long des extractions, nous obtenons donc un point limite presque sûre $(\mathcal{W}, \mathcal{Q}, (\mathcal{M}_1, \dots, \mathcal{M}_d))$ satisfaisant le système auxiliaire : (9.8) et

$$\begin{aligned} &\mathcal{Q}(t) - \mathcal{Q}(0) \\ &= \int_0^t \left[\sum_{c=1}^d \mathfrak{c}_c \frac{\overline{\mathcal{W}}_c(s)}{\mathcal{R}_c(s)} (1 - \mathcal{K}(s)) - L \right. \\ &\quad \left. + \left(\sum_{c=1}^d \mathfrak{c}_c \frac{\overline{\mathcal{W}}_c(s)}{\mathcal{R}_c(s)} (1 - F(\mathcal{Q}(s))) - L \right)^- \chi\{\mathcal{Q}(s) = 0\} \right] ds \end{aligned} \tag{9.9}$$

où

$$\begin{aligned} \overline{\mathcal{W}}_c(t) &:= \lim_{N \rightarrow \infty} \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N \mathcal{W}_i(t) \chi\{i \in \mathfrak{C}_c\} \\ &= \langle Id, \mathcal{M}_c(t) \rangle. \end{aligned} \tag{9.10}$$

Nous nommons cette solution, une solution forte associée à l'extraction. Rappelons que (9.10) est vraie hors extraction, ce qui signifie que la solution du système auxiliaire (9.9) et (9.8) est une solution forte du système initial du théorème 8.2.

9.2 Unicité de la solution forte

Nous avons construit une solution forte $(\mathbf{W}, \mathbf{R}, q, \mathbf{M}) := ((\mathcal{W}_i), (\mathcal{R}_c), \mathcal{Q}, (\mathcal{M}_c))$ de (8.18) et (8.17)²⁶. Cette solution sera notre candidat limite. Le but de cette section est de

²⁶Nous adoptons une nouvelle notation pour bien mettre en évidence que la solution que nous étudions est bien solution du système original infini et pas uniquement du système auxiliaire. D'autre part pour

prouver la convergence $\mathbb{L}^1$ de $(\mathbf{W}^N, q^N)$ vers la solution forte candidate sur $[0, T]$. Pour cela nous allons nous aider d'une «norme» qui prolonge les normes Cesàro-auxiliaires que nous avons utilisé dans la première partie de cette thèse.

Proposition 9.1. *Sous les hypothèses 8.6 et 8.7, $D^N(T) \rightarrow 0$ quand $N \rightarrow \infty$ où*

$$D^N(t) := \mathbb{E} \sup_{\tau \leq t} |q^N(\tau) - q(\tau)| + \mathbb{E} \|\mathbf{W}^N(t) - \mathbf{W}(t)\|$$

où

$$\|\mathbf{W}^N(t) - \mathbf{W}(t)\| = \frac{1}{N} \sum_{i=1}^N \sup_{0 \leq \tau \leq t} |W_i^N(\tau) - W_i(\tau)|.$$

Nous concluons que q^N et donc R_i^N convergent en probabilité vers q et R_i (la convergence en probabilité est un sous-produit de la convergence de D^N). Ainsi, pour chaque fenêtre W_i^N solution de (8.17), l'intensité $I_i^N(s)$ converge vers

$$I_i(s) = \frac{W_i(s - R_i(s))}{R_i(s - R_i(s))} K((s - R_i(s))).$$

Ce qui implique que chaque fenêtre W_i^N converge en probabilité vers W_i dans la norme de Skorokhod.

On a aussi un résultat plus fort que la convergence faible de $\mathbf{W}^N$ à $\mathbf{W}$:

Corollaire 9.1. *Si les hypothèses 8.6 et 8.7 sont vérifiées, alors, $\|M_c^N(t) - M_c(t)\|_w \rightarrow 0$ en probabilité pour tout $t \leq T$.*

Rappelons que la définition de la norme $\|\cdot\|_w$ dont dérive la distance d^* se trouve en définition 3.8. Cette distance engendre en particulier la topologie faible sur les mesures.

Preuve : Pour toute fonction g Lipschitz bornée,

$$\begin{aligned} & \limsup_{N \rightarrow \infty} |\mathbb{E}\langle g, M_c^N(t) \rangle - \mathbb{E}\langle g, M_c(t) \rangle| \\ &= \limsup_{N \rightarrow \infty} \left| \mathbb{E} \left[\frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N g(W_i^N(t)) \chi\{i \in \mathfrak{C}_c\} \right] - \lim_{N \rightarrow \infty} \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N g(W_i(t)) \chi\{i \in \mathfrak{C}_c\} \right| \\ &\leq \limsup_{N \rightarrow \infty} \left| \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N \mathbb{E} [g(W_i^N(t)) - g(W_i(t))] \chi\{i \in \mathfrak{C}_c\} \right| \\ &\quad + \limsup_{N \rightarrow \infty} \left| \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N g(W_i(t)) \chi\{i \in \mathfrak{C}_c\} - \mathbb{E}\langle g, M_c(t) \rangle \right| \\ &\leq \limsup_{N \rightarrow \infty} \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N \mathbb{E} |g(W_i^N(t)) - g(W_i(t))| \\ &\leq C_g \limsup_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N \mathbb{E} |W_i^N(t) - W_i(t)| \end{aligned}$$

respecter la logique suivie jusqu'ici, une variable mise à l'échelle s'écrit en minuscule ; ceci explique que q qui est la limite supposée de $q^N = \frac{Q^N}{N}$ soit en caractère minuscule alors que Q qui est une limite de files d'attentes auxiliaires n'est pas mise à l'échelle dans le même sens et garde une lettre majuscule. Pour être complet sur les notations nous avons conservé dans ce chapitre uniquement la convention d'écriture des vecteurs en lettres grasses de l'article [109].

où C_g est la constante de Lipschitz divisée par $\min\{\mathbf{c}_c^N\}$. Ainsi,

$$\begin{aligned}
 & \limsup_{N \rightarrow \infty} |\mathbb{E}\langle g, M_c^N(t) \rangle - \mathbb{E}\langle g, M_c(t) \rangle| \\
 & \leq C_g \limsup_{N \rightarrow \infty} \frac{1}{N} \sum_{i=1}^N \mathbb{E} \sup_{\tau \leq t} |W_i^N(\tau) - W_i(\tau)| \\
 & \leq C_g \limsup_{N \rightarrow \infty} \mathbb{E} \|\mathbf{W}^N(t) - \mathbf{W}(t)\| \\
 & = 0.
 \end{aligned}$$

D'où le résultat. ■

Un argument similaire montre que $\|M_c^N(t - R_c^N(t), \cdot; t, \cdot) - M_c(t - R_c(t), \cdot; t, \cdot)\|_w \rightarrow 0$ en probabilité pour tout $t \leq T$.

Pour prouver que $D^N(t) \rightarrow 0$ lorsque $N \rightarrow \infty$ nous allons établir une inégalité de Gronwall : $D^N(t) \leq B_N(t) + C \int_0^t D^N(s) ds$ où $B_N(t) \rightarrow 0$ quand $N \rightarrow \infty$. Il est plus facile de prouver l'inégalité de Gronwall pour Gentle RED où les probabilités de pertes gardent une dépendance Lipschitz du temps. Pour Gentle RED les résultats de la sous section 9.2.1 s'appliquent avec $\rho = T$. De plus, nous allons voir plus bas que $B^N(t) = \mathbb{E} \int_0^t |\bar{\mathbf{S}}_N(s) - \bar{\mathbf{S}}(s)| ds$ où $\bar{\mathbf{S}}_N$ et $\bar{\mathbf{S}}$ ont été définis par (8.19) et (8.20). Ceci signifie que nous pouvons même récupérer des informations sur la vitesse de convergence comme $D^N(t) \leq B^N(t) + C \int_0^t B^N(s) \exp(C(t-s)) ds$.

Malheureusement pour RED, quand q^N touche la frontière q_{max} , la dynamique change parce que F n'est pas Lipschitz à q_{max} . Notre solution pour démontrer la proposition 9.1 est de prouver la convergence sur $[0, \rho]$ où ρ est le temps (déterministe) de première atteinte par $q(t)$ du niveau q_{max} . Ceci est fait dans la sous section 9.2.1. Dans la section 9.2.2 nous prouvons que la convergence des taux de transmission de la pré limite s'étend pour un temps T_{min} au delà de ρ (et en fait T_{min} est arbitraire). Ceci tient parce que les taux de transmission sont déterminés un RTT dans le passé. Ceci nous permet d'étendre la preuve à travers la frontière quand q atteint q_{max} . Enfin, dans la sous section 9.2.3, nous prouvons la convergence sur l'intervalle $[0, \sigma]$ où $\rho < \sigma$ et σ est le premier temps où $q(t)$ quitte la frontière ; c'est-à-dire $q(t) < q_{max}$ pour $t \in (\sigma, \sigma + \delta]$ pour un $\delta > 0$.

Voici quelques lemmes dont nous allons faire usage ainsi que leurs preuves.

Lemme 9.3. $\dot{q}(t) > -L + \delta$ où $\delta > 0$ et $\dot{q}(t) \leq a(t)/T_{min} - L$.

Preuve : On prend l'espérance de (8.17) et en suivant la démarche du lemme 9.1 nous trouvons :

$$\mathbb{E}W_i(t) - w_i \geq \frac{1}{T_{max} + q_{max}/L} - \frac{\mathbb{E}W_i(t^-)}{2} \frac{a(t)}{T_{min}}.$$

Comme cette borne est strictement positive quand $\mathbb{E}W_i(t)$ est petit, il suit que $\mathbb{E}W_i(t)$ est uniformément strictement supérieure à 0 pour tous les temps t et les indices n . Ainsi $\bar{W}_c(t)$ est uniformément strictement supérieur à 0 en t , n et c . De (8.18), il vient $\dot{q}(t) > -L$.

La deuxième inégalité est une conséquence directe de (8.18).

■

Lemme 9.4. Pour $0 \leq s \leq t \leq T$,

$$\begin{aligned} & \sup_{0 \leq s \leq t} |R_i^N(s) - R_i(s)|, \sup_{0 \leq s \leq t} |R_i^N(s - R_i^N(s)) - R_i(s - R_i(s))|, \\ & \text{et } \sup_{0 \leq s \leq t} |q^N(s - R_i^N(s)) - q(s - R_i(s))| \end{aligned}$$

sont bornés par $C \sup_{0 \leq s \leq t - T_i} |q^N(s) - q(s)|$.

Preuve : Soit $\phi_i(s) = T_i + q(s)/L$, $g_i(s) = s + \phi_i(s)$, $\phi_i^N(s) = T_i + q^N(s)/L$ et $g_i^N(s) = s + \phi_i^N(s)$. Pour tout t soit $g_i(u) = t$ et soit $g_i^N(u^N) = t$. Notons que $u \leq t - T_i$ et $u^N \leq t - T_i$. Aussi, $R_i^N(t) - R_i(t) = \phi_i^N(u^N) - \phi_i(u) = u^N - u$. Supposons que $u \leq u^N$ (ou $u^N \leq u$), la fonction g_i croît (ou décroît) strictement de $g_i(u) = t = g_i^N(u^N)$ à $g_i(u^N)$ (resp. de $g_i(u^N)$ à $g_i^N(u^N)$) d'au moins $\delta(u^N - u)/L$ dans la mesure où par le lemme 9.3 la dérivée de $\phi_i(u)$ est bornée par δ/L . Il vient que $\frac{\delta}{L}|u^N - u| \leq |g_i^N(u^N) - g_i(u)| \leq \sup_{0 \leq s \leq t - T_i} |g_i^N(s) - g_i(s)| \leq \sup_{0 \leq s \leq t - T_i} |q^N(s) - q(s)|/L$. Ainsi

$$|R_i^N(t) - R_i(t)| = |u^N - u| \leq \sup_{0 \leq s \leq t - T_i} |q^N(s) - q(s)|/\delta$$

c'est-à-dire ce que nous cherchions.

De la même manière que nous avons obtenu (8.7) nous avons

$$(1 - \dot{R}_i(t)) = \frac{1}{1 + \dot{q}(t - R_i(t))/L}$$

et du lemme 9.3

$$\begin{aligned} \dot{R}_i(s) &= \frac{\dot{q}(s - R_i(s))/L}{1 + \dot{q}(s - R_i(s))/L} \\ &\leq |a(s)/T_{min} - L|/(\delta/L). \end{aligned}$$

En particulier

$$\begin{aligned} |R_i(s - R_i^N(s)) - R_i(s - R_i(s))| &\leq C |R_i^N(s) - R_i(s)| \\ &\leq C \sup_{0 \leq u \leq t - T_i} |q^N(u) - q(u)|. \end{aligned}$$

Finalement,

$$\begin{aligned} & \sup_{0 \leq s \leq t} |R_i^N(s - R_i^N(s)) - R_i(s - R_i(s))| \\ & \leq \sup_{0 \leq s \leq t} |R_i^N(s - R_i^N(s)) - R_i(s - R_i^N(s))| + \sup_{0 \leq s \leq t} |R_i(s - R_i^N(s)) - R_i(s - R_i(s))| \\ & \leq \sup_{0 \leq s \leq t} |R_i^N(s) - R_i(s)| + \sup_{0 \leq s \leq t} |R_i(s - R_i^N(s)) - R_i(s - R_i(s))| \\ & \leq C \sup_{0 \leq u \leq t - T_i} |q^N(u) - q(u)| \end{aligned}$$

en utilisant le premier résultat et l'inégalité de dessus.

On a ensuite le troisième résultat,

$$\begin{aligned}
& \sup_{0 \leq s \leq t} |q^N(s - R_i^N(s)) - q(s - R_i(s))| \\
& \leq \sup_{0 \leq s \leq t} |q^N(s - R_i^N(s)) - q(s - R_i^N(s))| + \sup_{0 \leq s \leq t} |q(s - R_i^N(s)) - q(s - R_i(s))| \\
& \leq \sup_{0 \leq s \leq t - T_i} |q^N(s) - q(s)| + C |R_i^N(s) - R_i(s)| \text{ comme la dérivée de } q \text{ est bornée et } R_i \geq T_i \\
& \leq C \sup_{0 \leq u \leq t - T_i} |q^N(u) - q(u)| \text{ en utilisant le premier résultat.}
\end{aligned}$$

■

9.2.1 Convergence loin de la frontière

Commençons avec $q^N(0) = q(0)$ à l'intérieur du domaine sans gigue : $0 < q(0) < q_{max}$. Définissons ρ , respectivement ρ^N , les temps d'arrêt de première atteinte par $q(t)$ et $q^N(t)$ de q_{max} . Nous devons définir la distance entre le processus marginal $\mathbf{W}^N(t) \equiv (W_1^N(t), \dots, W_N^N(t))$, $q^N(t)$ et $M^N(t) \equiv (M_1^N(t), \dots, M_d^N(t))$ et la limite du processus jusqu'au temps d'arrêt $\rho \wedge \rho^N$. Pour tout $t \leq \rho$ définissons

$$||\mathbf{W}^N(t) - \mathbf{W}(t)|| = \frac{1}{N} \sum_{i=1}^N \sup_{0 \leq \tau \leq t \wedge \rho^N} |W_i^N(\tau) - W_i(\tau)|$$

où τ est un temps d'arrêt pour la filtration $\mathcal{F}_t$. Soit

$$D^N(t) := \mathbb{E} \sup_{\tau \leq t \wedge \rho^N} |q^N(\tau) - q(\tau)| + \mathbb{E} ||\mathbf{W}^N(t) - \mathbf{W}(t)||.$$

Nous allons établir une inégalité de Gronwall :

Proposition 9.2. $D^N(t) \leq B^N(t) + C \int_0^t D^N(s) ds$ pour $t \leq \rho$ et $\sup_{t \leq \rho} B^N(t) \rightarrow 0$ quand $N \rightarrow \infty$ où C est la constante canonique tout au long des calculs (elle dépendra malheureusement de F'). De plus ρ^N converge vers ρ en probabilité.

Nous regroupons juste les constantes universelles qui ne dépendent pas de N dans C . Nous notons qu'un des facteurs de C est la constante de Lipschitz, $F'(q)$, $q < q_{max}$.

Estimé de q : rappelons que $\bar{\mathbf{S}}_N^N$ et $\bar{\mathbf{S}}$ ont été définis par (8.19) et (8.20).

$$\begin{aligned}
 & \mathbb{E} \sup_{0 \leq \tau \leq t \wedge \rho^N} |q^N(\tau) - q(\tau)| \leq \\
 & \quad + \mathbb{E} \sup_{0 \leq \tau \leq t \wedge \rho^N} \int_0^\tau |\bar{\mathbf{S}}_N^N(s)(1 - K^N(s)) - \bar{\mathbf{S}}(s)(1 - K(s))| ds \quad (9.11) \\
 & \leq \mathbb{E} \sup_{0 \leq \tau \leq t \wedge \rho^N} \int_0^\tau |(\bar{\mathbf{S}}_N^N(s) - \bar{\mathbf{S}}(s))(1 - K^N(s))| ds \\
 & \quad + \mathbb{E} \sup_{0 \leq \tau \leq t \wedge \rho^N} \int_0^\tau |(\bar{\mathbf{S}}_N(s) - \bar{\mathbf{S}}(s))(1 - K^N(s))| ds \\
 & \quad + \mathbb{E} \sup_{0 \leq \tau \leq t \wedge \rho^N} \int_0^\tau |\bar{\mathbf{S}}(s)(K^N(s) - K(s))| ds
 \end{aligned}$$

$$\begin{aligned}
 \mathbb{E} \sup_{0 \leq \tau \leq t \wedge \rho^N} |q^N(\tau) - q(\tau)| & \leq \mathbb{E} \sup_{0 \leq \tau \leq t \wedge \rho^N} \int_0^\tau \left(\left| \frac{1}{N} \sum_{i=1}^N \left[\frac{W_i^N(s)}{R_i^N(s)} - \frac{W_i(s)}{R_i(s)} \right] \right| \right) ds \\
 & \quad + \mathbb{E} \sup_{0 \leq \tau \leq t \wedge \rho^N} \int_0^\tau |\bar{\mathbf{S}}_N(s) - \bar{\mathbf{S}}(s)| ds \\
 & \quad + \mathbb{E} \sup_{0 \leq \tau \leq t \wedge \rho^N} \int_0^\tau |\bar{\mathbf{S}}(s)| |K^N(s) - K(s)| ds \\
 & \leq \int_0^t \frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} \left| \frac{W_i^N(\tau)}{R_i^N(\tau)} - \frac{W_i(\tau)}{R_i(\tau)} \right| \right) ds \\
 & \quad + B^N + \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \rho^N} \int_0^\tau \frac{1}{T_{min}} a(s) |K^N(s) - K(s)| ds \right)
 \end{aligned}$$

où $B^N = \mathbb{E} \int_0^t |\bar{\mathbf{S}}_N(s) - \bar{\mathbf{S}}(s)| ds$.

Puis,

$$\begin{aligned}
 & \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} \left| \frac{W_i^N(\tau)}{R_i^N(\tau)} - \frac{W_i(\tau)}{R_i(\tau)} \right| \right) \quad (9.12) \\
 & \leq \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |W_i^N(\tau) - W_i(\tau)| \frac{1}{R_i^N(\tau)} \right) + \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |W_i(\tau)| \left| \frac{1}{R_i^N(\tau)} - \frac{1}{R_i(\tau)} \right| \right) \\
 & \leq \frac{1}{T_{min}} \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |W_i^N(\tau) - W_i(\tau)| \right) + \frac{a(s)}{T_{min}^2} \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |R_i^N(\tau) - R_i(\tau)| \right) \\
 & \leq \frac{1}{T_{min}} \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |W_i^N(\tau) - W_i(\tau)| \right) + \frac{1}{T_{min}^2} a(s) C \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |q^N(\tau) - q(\tau)| \right)
 \end{aligned}$$

avec le lemme 9.4.

De plus,

$$\begin{aligned}
 & \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \rho^N} \int_0^\tau \frac{1}{T_{min}} a(s) |K^N(s) - K(s)| ds \right) \\
 & \leq C \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \rho^N} \int_0^\tau |K^N(s) - K(s)| ds \right) \\
 & \leq C \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \rho^N} \int_0^\tau |F(q^N(s)) - F(q(s))| ds \right) \leq C \int_0^t \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |q^N(\tau) - q(\tau)| ds \right).
 \end{aligned} \tag{9.13}$$

Donc,

$$\begin{aligned}
 & \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \rho^N} |q^N(\tau) - q(\tau)| \right) \\
 & \leq \int_0^t \frac{1}{T_{min}} \frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |W_i^N(\tau) - W_i(\tau)| \right) ds \\
 & \quad + \int_0^t \frac{1}{T_{min}^2} a(s) C \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |q^N(\tau) - q(\tau)| \right) ds + B_1^N \\
 & \quad + C \int_0^t \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |q^N(\tau) - q(\tau)| \right) ds \\
 & \leq B_1^N + C \int_0^t D^N(s) ds.
 \end{aligned} \tag{9.14}$$

Estimé de W . De (8.17),

$$\begin{aligned}
 & \frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \rho^N} |W_i^N(\tau) - W_i(\tau)| \right) \\
 & \leq \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \rho^N} \int_0^\tau \frac{1}{N} \sum_{i=1}^N \mathbb{E} \left| \frac{1}{R_i^N(s)} - \frac{1}{R_i(s)} \right| ds \right)
 \end{aligned} \tag{9.15}$$

$$+ \frac{1}{2} \frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \rho^N} \left| \int_0^\tau [W_i^N(s^-) dJ_i^N(s) - W_i(s^-) dN_i(s)] \right| \right). \tag{9.16}$$

De nouveau, (9.15) est borné par

$$\begin{aligned}
 & \int_0^t \frac{a(s)}{T_{min}^2} \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |R_i^N(\tau) - R_i(\tau)| \right) ds \\
 & \leq C \int_0^t \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |q^N(\tau) - q(\tau)| \right) ds.
 \end{aligned} \tag{9.17}$$

Par définition de J_i^N ,

$$\int_0^\tau \frac{W_i^N(s^-)}{2} dJ_i^N(s) = \int_{s=0}^\tau \int_{u=0}^\infty \frac{W_i^N(s^-)}{2} \chi_{[0, I_i^N(s))}(u) \Upsilon_i(du, ds).$$

Ainsi

$$\begin{aligned} & \frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \rho^N} \left| \int_0^\tau [W_i^N(s^-) dJ_i^N(s) - W_i(s^-) dJ_i(s)] \right| \right) \\ & \leq \frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \rho^N} \int_0^\tau \int_{u=0}^\infty |W_i^N(s^-) \chi_{[0, I_i^N(s))}(u) - W_i(s^-) \chi_{[0, I_i(s))}(u)| \Upsilon_i(du, ds) \right) \\ & = \frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\int_0^{t \wedge \rho^N} \int_{u=0}^\infty |W_i^N(s^-) \chi_{[0, I_i^N(s))}(u) - W_i(s^-) \chi_{[0, I_i(s))}(u)| \Upsilon_i(du, ds) \right) \\ & = \frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\int_0^{t \wedge \rho^N} \int_{u=0}^\infty |W_i^N(s^-) \chi_{[0, I_i^N(s))}(u) - W_i(s^-) \chi_{[0, I_i(s))}(u)| dud s \right). \end{aligned}$$

Donc,

$$\begin{aligned} & \frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \rho^N} \left| \int_0^\tau [W_i^N(s^-) dJ_i^N(s) - W_i(s^-) dJ_i(s)] \right| \right) \\ & \leq \frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\int_0^{t \wedge \rho^N} |W_i^N(s^-) - W_i(s^-)| I_i^N(s) \wedge I_i(s) ds \right) \\ & \quad + \frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\int_0^{t \wedge \rho^N} |W_i^N(s^-) \vee W_i(s^-)| \cdot |I_i^N(s) - I_i(s)| ds \right) \\ & \leq \int_0^t \frac{a(s)}{T_{min}} \left[\frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |W_i^N(\tau) - W_i(\tau)| \right) \right] ds \end{aligned} \tag{9.18}$$

$$+ \frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\int_0^{t \wedge \rho^N} a(s) |I_i^N(s) - I_i(s)| ds \right) \tag{9.19}$$

où $I_i^N(s)$ et $I_i(s)$ valent moins que $(w_i + s/T_{min})/T_{min} = a(s)/T_{min}$.

On a aussi

$$\begin{aligned} & |I_i^N(s) - I_i(s)| \\ & \leq |W_i^N(s - R_i^N(s)) - W_i(s - R_i^N(s))| \frac{1}{R_i^N(s - R_i^N(s))} K^N(s - R_i^N(s)) \\ & \quad + |W_i(s - R_i^N(s)) - W_i(s - R_i(s))| \frac{1}{R_i^N(s - R_i^N(s))} K^N(s - R_i^N(s)) \\ & \quad + |W_i(s - R_i(s))| \frac{1}{R_i^N(s - R_i^N(s))} - \frac{1}{R_i(s - R_i(s))} |K^N(s - R_i^N(s)) \\ & \quad + W_i(s - R_i(s)) \frac{1}{R_i(s - R_i(s))} |K^N(s - R_i^N(s)) - K(s - R_i(s))|. \end{aligned}$$

Ainsi,

$$|I_i^N(s) - I_i(s)| \leq |W_i^N(s - R_i^N(s)) - W_i(s - R_i^N(s))|/T_i \quad (9.20)$$

$$+ |W_i(s - R_i^N(s)) - W_i(s - R_i(s))|/T_i \quad (9.21)$$

$$+ a(s)|R_i^N(s - R_i^N(s)) - R_i(s - R_i(s))|/T_i^2 \quad (9.22)$$

$$+ a(s)|K^N(s - R_i^N(s)) - K(s - R_i(s))|/T_i. \quad (9.23)$$

Nous devons borner $\mathbb{E} \left(\int_0^{t \wedge \rho^N} |I_i^N(s) - I_i(s)| ds \right)$ donc nous devons borner l'espérance de l'intégrale de chacun des termes du dessus. Le premier (9.20) satisfait

$$\begin{aligned} & \mathbb{E} \left(\int_0^{t \wedge \rho^N} |W_i^N(s - R_i^N(s)) - W_i(s - R_i^N(s))| ds / T_i \right) \\ & \leq \frac{1}{T_{min}} \int_0^t \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |W_i^N(\tau) - W_i(\tau)| ds \right). \end{aligned}$$

Le second (9.21)

$$\begin{aligned} & \mathbb{E} \left(\int_0^{t \wedge \rho^N} |W_i(s - R_i^N(s)) - W_i(s - R_i(s))| ds \right) / T_i \\ & \leq \mathbb{E} \left(\int_0^{t \wedge \rho^N} \int_{[s - R_i^N(s) \wedge s - R_i(s), s - R_i^N(s) \vee s - R_i(s)]} \frac{1}{R_i(u)} du ds \right) \\ & + \frac{1}{2} \mathbb{E} \left(\int_0^{t \wedge \rho^N} \int_{[(s - R_i^N(s)) \wedge (s - R_i(s)), (s - R_i^N(s)) \vee (s - R_i(s))]} W_i(u^-) dN_i(u) ds \right). \end{aligned}$$

Remarquons que

$$\begin{aligned} & \int_0^{t \wedge \rho^N} \chi \{ (s - R_i^N(s)) \wedge (s - R_i(s)) \leq u \leq (s - R_i^N(s)) \vee (s - R_i(s)) \} ds \\ & = (u + \phi_i^N(u) \vee \phi_i(u)) \wedge (t \wedge \rho^N) - (u + \phi_i^N(u) \wedge \phi_i(u)) \wedge (t \wedge \rho^N). \end{aligned}$$

où $\phi_i(u) = T_i + q(u)/L$ et $\phi_i^N(u) = T_i + q^N(u)/L$.

Ainsi,

$$\begin{aligned}
& \mathbb{E} \left(\int_0^{t \wedge \rho^N} |W_i(s - R_i^N(s)) - W_i(s - R_i(s))| ds \right) / T_i \\
& \leq \frac{1}{T_{min}} \mathbb{E} \left(\int_0^{t \wedge \rho^N} |R_i(s) - R_i^N(s)| ds \right) \\
& \quad + \frac{1}{2} \mathbb{E} \left(\int_0^{t \wedge \rho^N} |\phi_i^N(u) \vee \phi_i(u) - \phi_i^N(u) \wedge \phi_i(u)| W_i(u^-) I_i(u) du \right) \\
& \leq \frac{1}{T_{min}} \mathbb{E} \left(\int_0^{t \wedge \rho^N} |R_i(s) - R_i^N(s)| ds \right) + \frac{1}{2} \mathbb{E} \left(\int_0^{t \wedge \rho^N} |\phi_i^N(u) - \phi_i(u)| \frac{a^2(u)}{T_{min}} du \right) \\
& \leq \frac{C}{T_{min}} \int_0^t \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |q(s) - q^N(s)| \right) ds + \int_0^t \frac{1}{2} \frac{a^2(s)}{T_{min} L} \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |q^N(\tau) - q(\tau)| \right) ds \\
& \leq C \int_0^t \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |q^N(\tau) - q(\tau)| \right) ds
\end{aligned}$$

où C est un exemplaire de notre constante universelle (dont la valeur est variable d'une formule à l'autre mais ne dépend jamais de N).

Le troisième terme (9.22) se borne ainsi, pour $s \leq t \wedge \rho^N$

$$|R_i^N(s - R_i^N(s)) - R_i(s - R_i(s))| \leq C \sup_{\tau \leq s} |q^N(\tau) - q(\tau)|.$$

Par le lemme 9.4, en prenant les espérances :

$$\mathbb{E} \left(\int_0^{t \wedge \rho^N} \frac{a_i(s)}{T_i^2} |R_i^N(s - R_i^N(s)) - R_i(s - R_i(s))| ds \right) \leq C \int_0^t \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |q^N(\tau) - q(\tau)| \right) ds.$$

De même, pour le quatrième terme (9.23), pour $s \leq t \wedge \rho^N$

$$\begin{aligned}
& \frac{a(s)}{T_n} |K^N(s - R_i^N(s)) - K(s - R_i(s))| \\
& \leq C [|K^N(s - R_i^N(s)) - K(s - R_i^N(s))| + |K(s - R_i^N(s)) - K(s - R_i(s))|] \quad (9.24) \\
& \leq C [|q^N(s - R_i^N(s)) - q(s - R_i^N(s))| + |q(s - R_i^N(s)) - q(s - R_i(s))|] \\
& \leq C [\sup_{\tau \leq s} |q^N(\tau) - q(\tau)| + |R_i^N(s) - R_i(s)|] \\
& \leq C \sup_{\tau \leq s} |q^N(\tau) - q(\tau)| \quad (9.25)
\end{aligned}$$

comme q a une dérivée bornée.

En prenant l'espérance, le quatrième terme est borné par

$$C \int_0^t \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |q^N(\tau) - q(\tau)| \right) ds.$$

Ainsi nous pouvons borner (9.19) par

$$\begin{aligned} & C \int_0^t \left[\mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |q^N(\tau) - q(\tau)| \right) + \frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\sup_{0 \leq \tau \leq s \wedge \rho^N} |W_i^N(\tau) - W_i(\tau)| \right) \right] ds \\ & \leq C \int_0^t D^N(s) ds. \end{aligned}$$

En réunissant (9.18), (9.19) et (9.17) nous trouvons

$$\begin{aligned} & \frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \rho^N} |W_i^N(\tau) - W_i(\tau)| \right) \\ & \leq C \int_0^t \left[\mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \rho^N} |q^N(\tau) - q(\tau)| \right) + \frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \rho^N} |W_i^N(\tau) - W_i(\tau)| \right) \right] ds. \end{aligned}$$

Ainsi,

$$\mathbb{E} \|\mathbf{W}^N(t) - \mathbf{W}(t)\| \leq C \int_0^t D^N(s) ds.$$

Finalement en remettant dans (9.14) nous tombons sur notre inégalité de Gronwall : $D^N(t) \leq B^N(t) + C \int_0^t D^N(s) ds$.

9.2.2 Quand on traverse ou qu'on frôle le bord

La construction de l'inégalité de Gronwall de la section précédente est standard et suffit à prouver la convergence vers le champ moyen avec Gentle RED. Cette sous section résout le problème fondamental qui intervient lorsque $q(t)$ frôle la frontière q_{max} quand RED est utilisé. Par exemple, si nous changions la dynamique sur la frontière pour que la file d'attente augmente quand q touche q_{max} alors q^N ne convergerait certainement pas vers q sur beaucoup de réalisations. Ces réalisations où q^N a juste manqué q_{max} décroîtraient alors que ceux qui touchent q_{max} augmenteraient. On aurait une bifurcation. Dans le cas d'espèce le champ moyen est déjà dégénéré dans la mesure où après la bifurcation, la file d'attente n'est plus intrinsèquement déterministe.

Heureusement le délai de notre système nous permet de résoudre ce genre de problèmes. Nous montrons ainsi que $\sup_{0 \leq \tau \leq \rho + T_{min}} |\bar{S}_N^N(\tau) - \bar{S}(\tau)| \rightarrow 0$ quand $N \rightarrow \infty$. En d'autre termes on peut étendre la convergence du taux de transmission un RTT à l'avance, et cela au delà des temps d'arrêt. En effet l'évolution de la fenêtre pour $[\rho, \rho + T_{min}]$ est déjà déterminée au temps ρ .

Ceci force q^N à suivre q pour un RTT après avoir touché la frontière. Ainsi si $\bar{S}(t)(1 - p_{max}) > L$ pour $t \in (\rho, \rho + \delta)$ où $\delta > 0$ alors nous sommes assurés que la pré limite q^N touche le bord à un point proche de celui où q touche la frontière et que ρ^N converge vers ρ . De plus nous sommes assurés que si N est grand, il y aura bien un dernier temps de sortie η^N où $\rho^N \leq \eta^N < \rho + \delta$ quand q^N fait la gigue sur la frontière et que $\mathbb{E}|\eta^N - \rho| \rightarrow 0$. Nous pouvons donc définir σ^N sans ambiguïté comme le plus petits de ces temps après $\rho + \delta$ où q^N quitte le bord.

Quand q frôle la frontière au temps ρ alors $\bar{S}(t)(1 - p_{max}) < L$ pour $t \in (\rho, \rho + \delta]$ pour un $\delta > 0$. Ainsi $\sigma = \rho$. Ce cas pose une difficulté mathématique parce que la pré limite q^N touche le bord à un temps proche de ρ ou sinon évite le bord. Il est donc difficile de définir rigoureusement ρ^N et σ^N (ce que nous allons tout de même faire). Nous résolvons ce problème à la fin de cette sous section en enlevant ρ , et en disant que q n'avait en fait pas touché le bord. Alors les pré limites passent un très petit temps sur le bord.

Cette extension de la convergence de $\bar{S}_N^N(t)$ vers $\bar{S}(t)$ pour les temps $t \in [\rho, \rho + T_{min}]$ est en fait valable pour tout temps. La preuve ne change pas. Dans la sous section 9.2.3 nous utilisons ce fait pour étendre la convergence à σ et ensuite par itération jusqu'à T .

Notons que dans le cas de l'étude d'un système modélisé sans délai qui pose des problèmes pour passer des frontières on aura intérêt à introduire un petit délai si c'est possible pour simplifier la preuve de champ moyen.

Pour montrer la convergence un RTT en avant, pour $t \leq \rho + T_{min}$ définissons

$$H(t) = \mathbb{E} \left[\frac{1}{N} \sum_{i=1}^N \sup_{0 \leq \tau \leq t} |W_i^N(\tau) - W_i(\tau)| \right].$$

Pour $t \leq T_{min}$ nous pouvons utiliser les mêmes étapes que dans l'obtention des estimés (9.15) et (9.16) pour obtenir

$$\begin{aligned} H(t) \leq & \mathbb{E} \left[\frac{1}{N} \sum_{i=1}^N \sup_{0 \leq \tau \leq \rho} |W_i^N(\tau) - W_i(\tau)| \right] \\ & + \frac{1}{N} \sum_{i=1}^N \left(\int_{\rho}^{\rho+T_{min}} \mathbb{E} \left| \frac{1}{R_i^N(s)} - \frac{1}{R_i(s)} \right| ds \right) \end{aligned} \quad (9.26)$$

$$+ \frac{1}{2} \frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\sup_{\rho \leq \tau \leq t} \left| \int_{\rho}^{\tau} [W_i^N(s^-) dJ_i^N(s) - W_i(s^-) dN_i(s)] \right| \right). \quad (9.27)$$

De nouveau, (9.26) est borné par

$$T_{min} \frac{a(T)}{T_{min}^2} \mathbb{E} \left(\sup_{0 \leq \tau \leq \rho+T_{min}} |R_i^N(\tau) - R_i(\tau)| \right) \leq C \mathbb{E} \left(\sup_{0 \leq \tau \leq \rho} |q^N(\tau) - q(\tau)| \right)$$

en utilisant le lemme 9.4. Nous avons déjà prouvé la convergence jusqu'à ρ , ceci tend donc vers 0 lorsque $N \rightarrow \infty$.

Nous approchons (9.27) comme nous l'avons fait en (9.16) pour obtenir les termes suivants correspondant à (9.18) et (9.19) :

$$\begin{aligned} & \frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\sup_{\rho \leq \tau \leq t} \left| \int_{\rho}^{\tau} [W_i^N(s^-) dJ_i^N(s) - W_i(s^-) dN_i(s)] \right| \right) \\ & \leq \int_{\rho}^t \frac{a(s)}{T_{min}} \left[\frac{1}{N} \sum_{i=1}^N \mathbb{E} \left(\sup_{\rho \leq \tau \leq s} |W_i^N(\tau) - W_i(\tau)| \right) \right] ds \end{aligned} \quad (9.28)$$

$$+ \frac{1}{N} \sum_{i=1}^N \int_{\rho}^t \mathbb{E} (a(s) |I_i^N(s) - I_i(s)| ds). \quad (9.29)$$

On peut borner (9.29) en utilisant la décomposition donnée par (9.20), (9.21), (9.22) et (9.23). Comme chacun de ces termes n'implique que des temps dans le passé plus lointains que T_{min} , il n'est pas difficile de constater que l'intégrale du terme (9.20) est bornée par

$$C\mathbb{E} \left[\frac{1}{N} \sum_{i=1}^N \sup_{0 \leq t \leq \rho} |W_i^N(t) - W_i(t)| \right];$$

que le terme (9.21) est borné par $C\mathbb{E}(\sup_{0 \leq \tau \leq \rho} |q^N(\tau) - q(\tau)|)$ en tant qu'intégrale de (9.22) et (9.23). Toutes ces bornes tendent vers 0 lorsque $N \rightarrow \infty$.

Nous concluons donc que $H(t) \leq B^N(t) + C \int_0^t H(s) ds$ pour $0 \leq t \leq \rho + T_{min}$ où $B^N(t)$ va encore pour représenter un terme qui tend vers 0 si $N \rightarrow \infty$. En utilisant l'inégalité de Gronwall, nous concluons que $H(t) \rightarrow 0$ si $N \rightarrow \infty$. Ainsi nous avons la convergence des fenêtres sur $[\rho, \rho + T_{min}]$. En raffinant l'estimation (12.3) nous trouvons

$$\begin{aligned} & \mathbb{E} \left(\sup_{0 \leq \tau \leq \rho + T_{min}} \left| \frac{W_i^N(\tau)}{R_i^N(\tau)} - \frac{W_i(\tau)}{R_i(\tau)} \right| \right) \\ & \leq \frac{1}{T_{min}} \mathbb{E} \left(\sup_{0 \leq \tau \leq \rho + T_{min}} |W_i^N(\tau) - W_i(\tau)| \right) + \frac{1}{T_{min}^2} a(s) C \mathbb{E} \left(\sup_{0 \leq \tau \leq \rho} |q^N(\tau) - q(\tau)| \right) \end{aligned}$$

par le lemme 9.4. Ces deux estimés tendent vers 0 si $N \rightarrow \infty$ donc nous avons montré que

$$\sup_{0 \leq \tau \leq \rho + T_{min}} |\bar{S}_N^N(\tau) - \bar{S}_N(\tau)| \xrightarrow[N \rightarrow \infty]{} 0.$$

Ce qui complète la preuve.

Une fois qu'on a la convergence du taux de transmission jusqu'à $\rho + T_{min}$, le cas où q traverse la frontière devient clair et ρ^N converge à ρ .

Le cas de frôlement de la frontière par q donne $\rho = \sigma$. Il est aussi résolu parce qu'en probabilité,

$$\bar{S}_N^N(t)(1 - p_{max}) \rightarrow \bar{S}(t)(1 - p_{max}) < L \text{ pour } \rho \leq t \leq \rho + T_{min}.$$

Ainsi si q^N traverse la frontière au temps ρ^N , il part sans tarder donc $\sigma^N - \rho^N \rightarrow 0$ en probabilité. Par conséquent nous pouvons continuer les iterations décrites dans la sous section 9.2.3, jusqu'à la prochaine fois où ρ_1, q traversent réellement la frontière, la contribution aux termes

$$\mathbb{E} \left(\int_0^{t \wedge \rho_1^N} |K^N(s) - K(s)| ds \right) \text{ et } \mathbb{E} \left(\int_0^{t \wedge \rho_1^N} |K^N(s - R_i^N(s)) - K(s - R_i(s))| ds \right)$$

par l'intégrale sur l'intervalle $[\rho^N, \eta^N]$ est négligeable.

Le cas où q reste sur la frontière et $\bar{S}(t)(1 - p_{max}) = L$ pour $t \in [\rho, \sigma]$ est théoriquement possible. Il ne pose pas de problèmes cependant. En effet q peut être considéré à la fois comme ayant passé la frontière ou pas. Ainsi les estimés de la section peuvent fonctionner.

9.2.3 Convergence en champ moyen sur la frontière

Nous prouvons la convergence sur l'intervalle $t \in [0, \sigma]$. Supposons que $\sigma > \rho$. Dans la sous section 9.2.2 nous avons montré comment traiter le cas où la file frôle la frontière. Soit σ^N l'infimum sur les temps plus grands que $\rho + \delta$ tels que q^N est plus petit que q_{max} où $\delta < \sigma - \rho$.

Pour tout $0 \leq t \leq \sigma$, on redéfinit

$$\|\mathbf{W}^N(t) - \mathbf{W}(t)\| = \frac{1}{N} \sum_{i=1}^N \sup_{0 \leq \tau \leq t \wedge \sigma^N} |W_i^N(\tau) - W_i(\tau)|$$

où τ est un temps d'arrêt pour la filtration $\mathcal{F}_t$ et σ^N est la fin du premier séjour sur la frontière de q^N . On redéfinit aussi

$$D^N(t) := \mathbb{E} \sup_{\tau \leq t \wedge \sigma^N} |q^N(\tau) - q(\tau)| + \mathbb{E} \|\mathbf{W}^N(t) - \mathbf{W}(t)\|.$$

Nous allons encore établir une inégalité de Gronwall : $D^N(t) \leq B_1^N(t) + C \int_0^t D^N(s) ds$ pour $t \in [0, \sigma]$ où $\sup_{t \leq \sigma} B_1^N(t) \rightarrow 0$ quand $N \rightarrow \infty$.

Le calcul est presque le même qu'à la sous section 9.2.1. Nous avons seulement besoin d'améliorer les bornes sur les termes (9.13) et (9.19) via (9.24). Si $s \leq t \wedge \sigma^N$ où $t \leq \sigma$, trois choses sont possibles ; $q^N(s)$ et $q(s)$ sont tous deux loins de la frontière, soit ils y sont tous les deux, soit l'un et pas l'autre. Si $q^N(s) < q_{max}$ et $q(s) < q_{max}$ alors

$$|K^N(s) - K(s)| = |F(q^N(s)) - F(q(s))| \leq C \sup_{\tau \leq s \wedge \sigma^N} |q^N(\tau) - q(\tau)|.$$

Si $q^N(s) = q(s) = q_{max}$, $K^N(s)$ et $K(s)$ sont donnés par

$$\bar{S}_N^N(s)(1 - K^N(s)) = L \text{ et } \bar{S}(s)(1 - K(s)) = L$$

avec $K^N(s), K(s) > p_{max}$. Ce qui signifie que

$$\bar{S}_N^N(s) \geq L/(1 - p_{max}) \text{ et } \bar{S}(s) \geq L/(1 - p_{max}). \quad (9.30)$$

Ainsi, si $q^N(s) = q(s) = q_{max}$,

$$\begin{aligned} |K^N(s) - K(s)| &= |L/\bar{S}_N^N(s) - L/\bar{S}(s)| \\ &\leq \frac{(1 - p_{max})^2}{L} |\bar{S}_N^N(s) - \bar{S}(s)| \\ &\leq C(|\bar{S}_N^N(s) - \bar{S}_N(s)| + |\bar{S}_N(s) - \bar{S}(s)|) \\ &\leq CD^N(s) + |\bar{S}_N(s) - \bar{S}(s)| \end{aligned}$$

en utilisant le même estimé que (9.11).

Aussi, pour borner l'expression (9.13), pour $t \leq \sigma$,

$$\begin{aligned}
 & \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \sigma^N} \int_0^\tau |K^N(s) - K(s)| ds \right) \\
 & \leq C \int_0^t D^N(s) ds + \int_0^t |\bar{S}_N(s) - \bar{S}(s)| ds \\
 & \quad + \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \sigma^N} \int_0^\tau [\chi\{q^N(s) < q_{max} = q(s)\}] ds \right) \\
 & \quad + \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \sigma^N} \int_0^\tau [\chi\{q(s) < q_{max} = q^N(s)\}] ds \right) \\
 & \leq C \int_0^t D^N(s) ds + B^N(t) + \mathbb{E}|\rho^N - \rho| + \mathbb{E}|\rho^N - \eta^N|. \tag{9.31}
 \end{aligned}$$

Pour borner (9.19) nous avons besoin d'améliorer un petit peu la borne sur (9.24). Comme précédemment

$$\begin{aligned}
 & |K^N(s - R_i^N(s)) - K(s - R_i(s))| \\
 & \leq |K^N(s - R_i^N(s)) - K(s - R_i^N(s))| + |K(s - R_i^N(s)) - K(s - R_i(s))|. \tag{9.32}
 \end{aligned}$$

Si $q(s - R_i^N(s)) < q_{max}$ et $q(s - R_i(s)) < q_{max}$, alors comme $F(q)$ est Lipschitz pour $q < q_{max}$,

$$\begin{aligned}
 |K(s - R_i^N(s)) - K(s - R_i(s))| &= F(q(s - R_i^N(s))) - F(q(s - R_i(s))) \\
 &\leq C|q(s - R_i^N(s)) - q(s - R_i(s))| \\
 &\leq C|R_i^N(s) - R_i(s)| \text{ parce que } q \text{ est différentiable} \\
 &\leq C \sup_{\tau \leq s \wedge \sigma^N} |q^N(\tau) - q(\tau)|.
 \end{aligned}$$

Cependant, si $q(s - R_i^N(s)) = q(s - R_i(s)) = q_{max}$,

$$\begin{aligned}
 & |K(s - R_i^N(s)) - K(s - R_i(s))| \\
 &= |L/\bar{S}(s - R_i^N(s)) - L/\bar{S}(s - R_i(s))| \\
 &\leq \frac{(1 - p_{max})^2}{L} |\bar{S}(s - R_i^N(s)) - \bar{S}(s - R_i(s))| \\
 &\leq C|R_i^N(s) - R_i(s)| \\
 &\leq C \sup_{\tau \leq s \wedge \sigma^N} |q^N(\tau) - q(\tau)|.
 \end{aligned}$$

Nous avons utilisé le fait que $\bar{S}(s) = \sum_{c=1}^d \mathbf{c}_c \frac{m_c(s)}{R_c(s)}$ est Lipschitz comme nous l'avons vérifié à la sous section 9.1.3.

Nous pouvons borner l'intégrale de la seconde expression dans (9.32) :

$$\begin{aligned} & \mathbb{E} \left(\int_0^{t \wedge \sigma^N} |K(s - R_i^N(s)) - K(s - R_i(s))| ds \right) \leq C \int_0^t D^N(s) ds \\ & + \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \sigma^N} \int_0^\tau [\chi\{q^N(s - R_i^N(s)) < q_{max} = q(s - R_i(s))\}] ds \right) \\ & + \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \sigma^N} \int_0^\tau [\chi\{q(s - R_i(s)) < q_{max} = q^N(s - R_i^N(s))\}] ds \right) \end{aligned}$$

Mais $q(s - R_i(s)) < q_{max} = q^N(s - R_i^N(s))$ implique $s - R_i(s) < \rho$ et $s - R_i^N(s) \geq \rho^N$; c'est-à-dire quand $s \in (\rho^N + \phi_i^N(\rho^N), \rho + \phi_i(\rho))$. Donc,

$$\begin{aligned} & \int_0^\tau [\chi\{q(s - R_i(s)) < q_{max} = q^N(s - R_i^N(s))\}] ds \\ & \leq |\rho^N - \rho| + |\phi_i^N(\rho^N) - \phi_i(\rho)| \leq |\rho^N - \rho| + \frac{1}{L} |q^N(\rho^N) - q(\rho)| \\ & \leq |\rho^N - \rho| + \frac{1}{L} |q(\rho^N) - q(\rho)| + C \sup_{\tau \leq s \wedge \sigma^N} |q^N(\tau) - q(\tau)| \\ & \leq C |\rho^N - \rho| + C \sup_{\tau \leq s \wedge \sigma^N} |q^N(\tau) - q(\tau)|. \end{aligned}$$

Or la situation $q^N(s - R_i^N(s)) < q_{max} = q(s - R_i(s))$ peut se produire quand $s - R_i(s) \leq \rho$ et $s - R_i^N(s) > \rho^N$ ou $\rho^N \leq s - R_i^N(s) \leq \eta^N$; c'est-à-dire quand $s \in (\rho^N + \phi_i^N(\rho^N), \rho + \phi_i(\rho))$ ou quand $s \in (\rho^N + \phi_i^N(\rho^N), \eta^N + \phi_i^N(\eta^N))$. Ainsi,

$$\begin{aligned} & \int_0^\tau [\chi\{q^N(s - R_i^N(s)) < q_{max} = q(s - R_i(s))\}] ds \\ & \leq |\rho^N - \rho| + |\rho^N - \eta^N| + |\phi_i^N(\rho^N) - \phi_i(\rho)| + |\phi_i^N(\eta^N) - \phi_i^N(\rho^N)| \\ & \leq 2|\rho^N - \rho| + 2|\rho^N - \eta^N| + \frac{1}{L} |q^N(\rho^N) - q(\rho)| + \frac{1}{L} |q^N(\eta^N) - q^N(\rho^N)| \\ & \leq 2|\rho^N - \rho| + 2|\rho - \eta^N| + \frac{1}{L} |q(\eta^N) - q(\rho)| + \frac{2}{L} |q(\rho^N) - q(\rho)| + C \sup_{\tau \leq s \wedge \sigma^N} |q^N(\tau) - q(\tau)| \\ & \leq C(|\rho^N - \rho| + |\rho - \eta^N|) + C \sup_{\tau \leq s \wedge \sigma^N} |q^N(\tau) - q(\tau)|. \end{aligned}$$

En additionnant tout cela, il vient

$$\mathbb{E} \int_0^{t \wedge \rho^N} |K(s - R_i^N(s)) - K(s - R_i(s))| ds \leq C(|\rho^N - \rho| + |\rho - \eta^N|) + C \int_0^t D^N(s) ds.$$

Nous pouvons facilement prouver que l'intégrale de la première expression de (9.32) utilise la même estimation que nous avons faite pour (9.31) :

$$\begin{aligned} & \mathbb{E} \left(\sup_{0 \leq \tau \leq t \wedge \rho^N} \int_0^\tau |K^N(s - R_i^N(s)) - K(s - R_i^N(s))| ds \right) \\ & \leq C \int_0^t D^N(s) ds + B^N(t) + C(\mathbb{E}|\rho^N - \rho| + \mathbb{E}|\eta^N - \rho|). \end{aligned}$$

En remettant tout ensemble, il vient $D^N(t) \leq B_1^N(t) + C \int_0^t D^N(s)ds$ où $B_1^N(t) = 2B^N(t) + C(\mathbb{E}|\rho^N - \rho| + \mathbb{E}|\eta^N - \rho|)$. La première itération a montré que ρ^N et η^N convergent en probabilité vers ρ donc $B_1^N(t) \rightarrow 0$ comme $N \rightarrow \infty$. En conséquence $\sup_{t \leq \sigma} D^N(t) \rightarrow 0$ et la convergence sur $[0, \sigma]$ est prouvée.

En utilisant maintenant les résultats de la sous section 9.2.2, nous sommes en mesure de montrer que $\sup_{0 \leq \tau \leq \sigma + T_{min}} |\bar{S}_N^N(\tau) - \bar{S}(\tau)| \rightarrow_{N \rightarrow \infty} 0$. Aussi si $\bar{S}(t)(1 - p_{max}) < L$ pour $t \in (\sigma, \sigma + \delta)$ où $\delta > 0$, alors nous sommes sûrs que la pré limite q^N part de la frontière proche de l'endroit où q l'a fait et que σ^N converge vers σ . On peut maintenant itérer pour montrer la convergence après n'importe quelle suite d'entrées et de sorties de la frontière. Il est tout à fait envisageable qu'il y ait un point où la suite des entrées et des sorties convergent vers un temps Θ . Nous pouvons utiliser ce qui a été dit plus haut pour prouver la convergence en champ moyen aussi proche que nous le voulons de Θ . Là encore, en utilisant l'astuce du délai, on peut montrer : $\sup_{0 \leq \tau \leq \Theta + T_{min}} |\bar{S}_N^N(\tau) - \bar{S}(\tau)| \rightarrow_{N \rightarrow \infty} 0$.

Aussi la convergence traverse bien Θ . Il n'y a pas de dernier point de convergence du champ moyen dans ce cas. On peut donc avoir cette convergence sur $\mathbb{R}$ tout entier par connexité.

9.3 RED est la limite faible de Gentle RED*

Nous nous posons une dernière question : aurait-on pu résoudre le problème en utilisant la convergence sans problèmes de frontière qui intervient pour Gentle RED et en faisant la limite pour une fonction de pertes qui a une pente de plus en plus grande ? La réponse de cette section est oui, mais on aurait obtenu uniquement des renseignements sur les distributions, ce qui est un résultat moins fort. Au moins nous avons une inversion des limites qui se passe bien.

On définit la fonction de pertes de Gentle RED ainsi : $F_\delta(q)$ est linéaire entre 0 et q_{max} et entre q_{max} et $q_{max} + \delta$. $F_\delta(0) = 0$, $F_\delta(q_{max}) = p_{max}$ et $F_\delta(q_{max} + \delta) = 1$. La sous section 9.2.1 prouve la convergence en champ moyen de Gentle RED sur tout compact $[0, T]$, comme F_δ est Lipschitz. Nous allons ici faire tendre $\delta \rightarrow 0$, Gentle RED devient RED et F_δ tend (faiblement) vers F , la probabilité de pertes de RED.

Pour tout couple (δ, N) , on redéfinit la solution du système à N corps (ou particules) de la section 8.1, $(\mathbf{W}^N(t), q^N(t))$ par $(\mathbf{W}^{\delta, N}, q^{\delta, N})$ maintenant $(\mathbf{W}^N(t), q^N(t))$ ne représente plus que la solution pour la fonction de pertes F . Ces processus sont construits itérativement sur le nombre presque sûrement fini des segments définis par les sauts de $\Upsilon_i(s, \bar{I}(T))$; $i = 1, \dots, N$ où $\Upsilon_i, i = 1, 2, \dots$ sont définis sur l'espace de probabilités $(\Omega, \mathcal{F}, P)$. Soit $\mathbf{R}^{\delta, N} = (R_1^{\delta, N}, \dots, R_N^{\delta, N})$ le RTT correspondant. Soit

$$\bar{S}_N^{\delta, N}(t) = \sum_{c=1}^d \mathbf{c}_c^N \frac{\bar{\mathbf{W}}_c^{\delta, N}(t)}{R_c^{\delta, N}(s)}.$$

Soit $P^{\delta, N}$ la mesure induite sur $D[0, T] \times C[0, T]$ par $(\bar{S}_N^{\delta, N}, q^{\delta, N})$ (où les coordonnées plus grandes que N sont 0). En reprenant la méthode du lemme 8.1 on montre que les mesures $P^{\delta, N}$ sont tendues.

En utilisant ce lemme, nous pouvons prouver que :

Lemme 9.5. $(\mathbf{W}^{\delta,N}, q^{\delta,N}, \mathbf{R}^{\delta,N})$ convergent faiblement vers $(\mathbf{W}^N(t), q^N, \mathbf{R}^N)$ quand $\delta \rightarrow 0$. Le résultat est aussi vrai pour $N = \infty$.

Preuve du lemme 9.5 :

preuve intuitive Notons $V_\delta = 1 - K_\delta = 1 - F_\delta(Q_\delta)$.

L'équation vérifiée par V_δ entre les valeurs Q_{max} et $Q_{max} + \delta$ est donc $\dot{V}_\delta = -\dot{F}_\delta \dot{Q}_\delta = -C_\delta((1 - K_\delta)S_\delta - L)$ où S_δ est le débit entrant normalisé dans la file et C_δ le coefficient directeur de F entre Q_{max} et $Q_{max} + \delta$. Ceci donne donc :

$$\dot{V}_\delta(t) = C_\delta(V_\delta(t)S_\delta(t) - L)$$

avec $C_\delta \rightarrow \infty$, et finalement :

$$0 = V_\delta(t)S_\delta(t) - L$$

c'est-à-dire le résultat que nous souhaitons obtenir que

$$K(t) = 1 - V(t) = 1 - \frac{L}{S(t)},$$

où $S(t)$ est le débit.

Retour sur la preuve formelle Par hypothèse, $q^{\delta,N}(t) = q^N(t) = q(0)$ et $W_i^{\delta,N}(t) = W_i^N(t) = w_i$ pour $t \leq 0$. Nous montrons que la probabilité de perte $K^{\delta,N}(t) = F_\delta(q^{\delta,N}(t))$ converge pour $\delta \rightarrow 0$.

Tout d'abord, prenons une sous suite δ_k telle que $P^{\delta_k,N}$ converge faiblement vers $P^{0,N}$ et que de plus $(\mathbf{W}^{\delta_k,N}, \bar{S}_N^{\delta_k,N}, q^{\delta_k,N})$ converge presque sûrement vers $(\mathbf{W}^{0,N}, \bar{S}_N^{0,N}, q^{0,N})$.

Si $q^{\delta_k,N}(t) < q_{max}$ alors $K^{\delta_k,N}(t) = F_{\delta_k}(q^{\delta_k,N}(t)) = F(q^{\delta_k,N}(t))$. D'un autre coté, si $q^{\delta_k,N}(t) > q_{max}$, alors il y aura un temps $\rho^{\delta_k,N}(t) \leq t$ où $q^{\delta_k,N}(t)$ touche q_{max} pour la dernière fois. Considérons la solution de (8.11) sur l'intervalle $[\rho^{\delta_k,N}(t), t]$ et notons :

$$V^{\delta_k,N}(s) = 1 - F_{\delta_k}(q^{\delta_k,N}(s)).$$

Remarquons que pour $\rho^{\delta_k,N}(t) \leq s \leq t$,

$$\frac{dV^{\delta_k,N}(s)}{ds} = -c_{\delta_k} \bar{S}_N^{\delta_k,N}(s) V^{\delta_k,N}(s) + L c_{\delta_k}$$

où $c_{\delta_k} = (1 - p_{max})/\delta_k$. Résolvons cette équation du temps $\rho^{\delta_k,N}(t)$ où $V^{\delta_k,N}(\rho^{\delta_k,N}(t)) = (1 - p_{max})$ jusqu'à t :

$$\begin{aligned} V^{\delta_k,N}(t) &= (1 - p_{max}) \exp \left(- \int_{\rho^{\delta_k,N}(t)}^t c_{\delta_k} \bar{S}_N^{\delta_k,N}(u) du \right) \\ &\quad + \int_{\rho^{\delta_k,N}(t)}^t L c_{\delta_k} \exp \left(- \int_u^t c_{\delta_k} \bar{S}_N^{\delta_k,N}(s) ds \right) du \\ &= (1 - p_{max}) \exp(-v(\rho^{\delta_k,N}(t))) + \int_0^{v(\rho^{\delta_k,N}(t))} e^{-v} \frac{L}{\bar{S}_N^{\delta_k,N}(u(v))} dv \end{aligned}$$

où $v = \int_u^t c_{\delta_k} \bar{S}_N^{\delta_k, N}(s) ds$ et $u(v)$ est l'inverse définie implicitement.

Définissons $D^{\delta_k, N}(t) = 0$ si $q^{\delta_k, N}(t) \leq q_{max}$ et pour $q^{\delta_k, N}(t) > q_{max}$ définissons

$$\begin{aligned} D^{\delta_k, N}(t) &= V^{\delta_k, N}(t) - \frac{L}{\bar{S}_N^{\delta_k, N}(t)} \\ &= (1 - p_{max}) \exp(-v(\rho^{\delta_k, N}(t))) + L \int_0^{v(\rho^{\delta_k, N}(t))} e^{-v} \left(\frac{1}{\bar{S}_N^{\delta_k, N}(u(v))} - \frac{1}{\bar{S}_N^{\delta_k, N}(t)} \right) dv \\ &\quad - e^{-v(\rho^{\delta_k, N}(t))} \frac{L}{\bar{S}_N^{\delta_k, N}(t)}. \end{aligned}$$

Comme $(\bar{S}_N^{\delta_k, N}(t), q^{\delta_k, N}(t))$ converge ps vers $(\bar{S}_N^{0, N}(t), q^{0, N}(t))$ quand $\delta_k \rightarrow 0$, il vient $\rho^{\delta_k, N}(t) \rightarrow \rho^{0, N}(t)$. Comme $\bar{S}_N^{0, N}(t) > 0$ ps, $v(\rho^{\delta_k, N}(t)) \rightarrow \infty$ pour $\delta_k \rightarrow 0$ et pour un v fixé, $v \frac{\delta_k}{(1-p_{max})} = \int_{u(v)}^t \bar{S}_N^{\delta_k, N}(s) ds$ donc $u(v) \rightarrow t$ pour $\delta_k \rightarrow 0$. Nous concluons que $D^{\delta_k, N}(t) \rightarrow 0$ ps (par convergence dominée pour chaque valeur de t).

Si nous prenons maintenant la limite quand $\delta_k \rightarrow 0$ dans (8.11) satisfaite par $(\bar{S}_N^{\delta_k, N}(t), q^{\delta_k, N}(t))$ pour la fonction de pertes F_{δ_k} nous voyons que $(\bar{S}_N^{0, N}(t), q^{0, N}(t))$ satisfait (8.11) avec la fonction de pertes F donc $K^{0, N}(t) = L/\bar{S}_N^{0, N}(t)$ si $q^{0, N}(t) = q_{max}$. De plus, la fenêtre $W_i^{\delta_k, N}$ satisfait (8.5) où l'intensité des réductions de fenêtre est donnée par,

$$I_i^{\delta_k, N}(t) := \frac{W_i^{\delta_k, N}(t - R_i^{\delta_k, N}(t))}{R_i^{\delta_k, N}(t - R_i^{\delta_k, N}(t))} F_{\delta_k}(q^{\delta_k, N}(t - R_i^{\delta_k, N}(t))).$$

Prenons maintenant la limite $\delta_k \rightarrow 0$, on voit que $W_n^{\delta_k, N}$ converge vers $W_n^{0, N}$ qui satisfait (8.5) où les réductions de fenêtre se font avec

$$I_i^{0, N}(t) := \frac{W_i^{0, N}(t - R_i^{0, N}(t))}{R_i^{0, N}(t - R_i^{0, N}(t))} F(q^{0, N}(t - R_i^{0, N}(t))).$$

Tout ceci signifie que $\mathbf{W}^{0, N} = (W_1^{0, N}, \dots, W_N^{0, N})$ et $q^{0, N}$ sont des solutions de (8.5) et (8.11) donc en fait aussi $\mathbf{W}^N = (W_1^N, \dots, W_N^N)$ et q^N . Ainsi la limite le long des sous séquences est unique, ce qui finit de prouver que nous avons une convergence faible.

Chapitre 10

EDP du champ moyen : équations et points fixes

Ce chapitre présente l'exploitation "calculatoire" du modèle étudié. Du point de vue pratique c'est le plus important, dans la mesure où les résultats des calculs que nous faisons ici peuvent être exploités pour du dimensionnement d'infrastructure. Concernant la stabilité, on trouvera des études similaires dans [14, 15]. En règle générale, dans la littérature, on s'intéresse toujours à l'exploitation du champ moyen avant sa preuve ; ici nous avons fait l'inverse. Si le problème dérivé restait trop complexe pour être étudié, la preuve de la limite n'aurait pas d'intérêt du point de vue pratique. Pour résumer, ce chapitre justifie le travail que nous avons fait pour prouver la limite du champ moyen dans le chapitre précédent et il fournit toutes les clés pour permettre à des ingénieurs d'utiliser le modèle.

Sommaire

10.1 Équations aux dérivées partielles du champ moyen	164
10.1.1 La martingale qui tend vers 0	164
10.1.2 L'équation aux dérivées partielles	166
10.1.3 Simplifications de l'EDP	167
10.2 Étude de l'équation simplifiée	168
10.2.1 Point fixe des équations simplifiées du champ moyen	168
10.2.2 Une discussion autour des timeouts	171
10.2.3 Preuve du théorème sur les timeouts	173
10.3 Contrôle et stabilité des points fixes de l'équation ap- proximée	175
10.3.1 Étude des moments de la distribution des fenêtres	175
10.3.2 Hypothèse pour l'étude de la stabilité	176
10.3.3 Stabilité sous les contrôles classiques	177
10.3.4 Un contrôle stable sur les débits	183
10.3.5 Un contrôle stable sur les débits avec utilisation maximale du lien	185

Remarque 10.1. Dans cette section $p(t)$ désignera une distribution de fenêtres et pas un taux de pertes comme on peut être habitué à le voir dans la littérature. Ce taux de pertes est noté $K(t)$ dans les versions non simplifiées de l'équation et $\kappa(t)$, ou κ à partir du moment où on enlève les classes et où on fait la simplification au paragraphe 10.1.3.

10.1 Équations aux dérivées partielles du champ moyen

Cette partie présente les résultats qu'on doit montrer même quand on suppose l'existence et l'unicité du champ moyen. Les résultats présentés ont donc été plus largement étudiés par les physiciens en particulier. Dans un premier temps on étudie uniquement l'équation qui met en jeu des distributions. En effet le problème de la taille de la file d'attente est direct. Il est traité déjà dans ce qui précède. On retrouvera le système complet explicité après les simplifications de (8.17) et (8.18).

10.1.1 La martingale qui tend vers 0

Dans cette section nous prouvons le théorème 8.1. Rappelons brièvement quelques une des notations introduites dans les deux chapitres précédents :

- le générateur $\mathcal{G}$ est composé des fonctions $\{g \in \mathcal{C}_b^1(\mathbb{R}^+), g(0) = 0\}$, c'est-à-dire les fonctions bornées continues dérivables à dérivée continue et bornée qui valent 0 à l'origine des temps,
- les sommations $\sum_{i=1}^N \chi\{i \in \mathfrak{C}_c\}$ sont en fait des sommations sur la classe c : $\sum_{i \in \mathfrak{C}_c \cap [1, N]}$,
- $\mathfrak{c}_c^N := \frac{\text{Card}(\mathfrak{C}_c \cap [1, N])}{N}$ est la proportion des utilisateurs de classe c parmi les N premiers,
- $M_c(t)$ et $M_c^N(t)$ sont les mesures empiriques (marginales au temps t) des utilisateurs de classe c ,
- $M_c^N(t, dv, s, dw)$ est la mesure produit indiquant la probabilité qu'une trajectoire soit en v au temps t et en w au temps s .

On peut reformuler (8.10) pour $g \in \mathcal{G}$:

$$\langle g, M_c^N(t) \rangle - \langle g, M_c^N(0) \rangle \tag{10.1}$$

$$\begin{aligned} &= \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N \int_0^t \left[\frac{dg}{dw}(W_i^N(s)) \frac{1}{R_c^N(s)} ds + \left(g\left(\frac{W_i^N(s^-)}{2}\right) - g(W_i^N(s^-)) \right) dJ_i^N(s) \right] \chi\{i \in \mathfrak{C}_c\} \\ &= \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N \chi\{i \in \mathfrak{C}_c\} \int_0^t \left[\frac{dg}{dw}(W_i(s)) \frac{1}{R_c^N(s)} ds + \left(g\left(\frac{W_i^N(s)}{2}\right) - g(W_i^N(s)) \right) \right. \\ &\quad \cdot \left. \frac{W_i^N(s - R_c^N(s))}{R_c^N(s - R_c^N(s))} K^N(s - R_c^N(s)) ds \right] + \mathcal{E}_c^N(t) \end{aligned} \tag{10.2}$$

où $\mathcal{E}_c^N(t)$ est donné par

$$\frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N \chi\{i \in \mathfrak{C}_c\} \int_0^t \left(g \left(\frac{W_i^N(s^-)}{2} \right) - g(W_i(s^-)) \right) dZ_n^N(s)$$

et

$$Z_n^N(t) - Z_n^N(0) := \int_0^t \left(dJ_i^N(s) - \frac{W_i^N(s - R_c^N(s))}{R_c^N(s - R_c^N(s))} K^N(s - R_c^N(s)) ds \right).$$

Il vient alors,

$$\begin{aligned} & \langle g, M_c^N(t) \rangle - \langle g, M_c^N(0) \rangle \\ &= \int_0^t \left[\frac{1}{R_c^N(s)} \left\langle \frac{dg(w)}{dw}, M_c^N(s) \right\rangle ds \right. \\ & \quad \left. + \langle (g(w/2) - g(w))v, M_c^N(s - R_c^N(s), dv; s, dw) \rangle \frac{1}{R_c^N(s - R_c^N(s))} K^N(s - R_c^N(s)) ds \right] \\ &+ \mathcal{E}_c^N(t). \end{aligned} \tag{10.3}$$

Prouvons d'abord que $\mathcal{E}_c^N$ est asymptotiquement petit lorsque $N \rightarrow \infty$. Rappelons que (nous récrivons les formules pour les mettre exactement dans le format du théorème à appliquer) :

$$\mathcal{E}_c^N(t) = \frac{1}{\mathfrak{c}_c^N N} \sum_{i=1}^N \int_0^t C_{c,i}^N(s) Z_{c,i}^N(ds)$$

où

$$C_{c,i}^N(s) = \chi\{i \in \mathfrak{C}_c\} \left(g \left(\frac{W_i^N(s)}{2} \right) - g(W_i^N(s)) \right)$$

et

$$Z_{c,n}^N(t) - Z_{c,n}^N(0) := \int_0^t \left(dJ_n^N(s) - \frac{W_n^N(s - R_c^N(s))}{R_c^N(s - R_c^N(s))} K^N(s - R_c^N(s)) ds \right) \chi\{n \in \mathfrak{C}_c\}.$$

Si $i \in \mathfrak{C}_c$, $J_i^N(s)$ est un processus ponctuel adapté à $\mathcal{F}_i(t)$ avec l'intensité stochastique :

$$W_i^N(s - R_c^N(s)) \frac{K^N(s - R_c^N(s))}{R_c^N(s - R_c^N(s))}.$$

Par conséquent $Z_i^N(t)$ est une martingale. Or $W_i^N(s)$ est aussi adapté à $\mathcal{F}_i(t)$ donc la version càdlàg est $\mathcal{F}_i(t)$ -prévisible. Le théorème T13 dans [27] s'applique et :

$$\begin{aligned} & E(\mathcal{E}_c^N(t))^2 \\ &= \frac{1}{(\mathfrak{c}_c^N N)^2} E \left[\sum_{i=1}^N \chi\{i \in \mathfrak{C}_c\} \int_0^t (C_{c,i}^N(s))^2 \frac{W_i^N(s - R_c^N(s))}{R_c^N(s - R_c^N(s))} K^N(s - R_c^N(s)) ds \right] \\ &\leq \frac{1}{(\mathfrak{c}_c^N N)^2} E \left[\sum_{i=1}^N \chi\{i \in \mathfrak{C}_c\} \int_0^t \left(g \left(\frac{W_i^N(s)}{2} \right) - g(W_i^N(s)) \right)^2 \frac{W_i^N(s - R_c^N(s))}{R_c^N(s - R_c^N(s))} ds \right] \\ &\leq \frac{C_1}{(\mathfrak{c}_c^N N)^2} \left(\sum_{i=1}^N \chi\{i \in \mathfrak{C}_c\} \int_0^t E[W_i^N(s - R_c^N(s))] ds \right) \end{aligned}$$

où C_1 est une constante qui ne dépend que de $\sup g$ et $\sup(g')$.

On a la borne à priori $W_n^N(t) \leq a(t)$. Donc,

$$E(\mathcal{E}_c^N(t))^2 \leq \frac{tC_1}{(\mathfrak{c}_c^N)^2 N} a(t).$$

Ainsi $\mathcal{E}_c^N(t)$ tend vers 0 dans l'espace $\mathbb{L}^2$. Comme on a de plus $\mathcal{E}_c^N(t)$ qui est une martingale on voit immédiatement que

$$P(\sup_{t \in [0, T]} |\mathcal{E}_c^N(t)| > \lambda) \leq E(\mathcal{E}_c^N(T))^2 / \lambda^2$$

et le processus $\mathcal{E}_c^N(t), t \in [0, T]$ converge en probabilité vers 0.

Les processus $q^N(t)$ et $K^N(t)$ convergent en probabilité vers le processus limite $q(t)$ et $K(t)$ lorsque $(M_1^N(t), \dots, M_d^N(t))$ tend vers $(M_1(t), \dots, M_d(t))$ en probabilité où les processus limites satisfont (8.18) et (8.17). Prenons la limite dans (10.3) et nous avons notre preuve.

10.1.2 L'équation aux dérivées partielles

Première équation Rappelons que nous avons noté dans le théorème 8.1 :

$$e(s, s - R_c(s), w) = \langle v, \frac{M_c(s - R_c(s), dv; s, dw)}{M_c(s, dw)} \rangle.$$

Ceci signifie que $e(s, t, w)$ est la valeur moyenne (c'est-à-dire l'espérance) au temps s des fenêtres dont la valeur était w au temps t .

Ainsi en reprenant (8.14) avec des fonctions test g telles que $g(0) = 0$, et $g(w) \rightarrow 0$ quand $w \rightarrow \infty$, on peut écrire :

$$\begin{aligned} & \langle g, M_c(t) \rangle - \langle g, M_c(0) \rangle \\ &= \int_0^t \left[-\frac{1}{R_c(s)} \langle g(w), D_w M_c(s, dw) \rangle + \frac{1}{R_c(s - R_c(s))} K(s - R_c(s)) \right. \\ & \quad \left. \cdot \langle g(w), e(s, s - R_c(s), 2w) \cdot M_c(s, 2dw) - e(s, s - R_c(s), w) \cdot M_c(s, dw) \rangle \right] ds \end{aligned}$$

où $D_w M_c(s, dw)$ (respectivement $D_t M_c(s, dw)$), est la différentielle de Fréchet²⁷ de la mesure $M_c(s, dw)$ par rapport à w (resp. t). Ainsi,

²⁷Nous rappelons la définition de la différentielle de Fréchet :

Définition 10.1. Une fonction $f : X \rightarrow Y$ où (X, d^X) et (Y, d^Y) sont deux espaces vectoriels métriques est DIFFÉRENTIABLE AU SENS DE FRÉCHET en x_0 s'il existe une application linéaire continue $Df(x_0)$ de X dans Y , telle que :

$$\forall x \in X, \lim_{h \rightarrow 0} d^Y \left(\frac{f(x_0 + hx) - f(x_0)}{h}, Df(x_0)(x) \right) = 0 \text{ } h \text{ étant un réel.}$$

Dans le cas présent nous étudions la différentielle de Fréchet selon une direction donnée, c'est à dire que nous considérons que M_c est fonction des réels w ou de t uniquement. Par exemple, $D_w M_c(t, dw)(w_0)$

Théorème 10.1. *L'équation des fenêtres dans le champ moyen est donnée par*

$$D_t M_c(t, dw) = -\frac{1}{R_c(t)} D_w M_c(t, dw) + \frac{1}{R_c(t - R_c(t))} K(t - R_c(t)) \cdot (e(t, t - R_c(t), 2w) M_c(t, d(2w)) - e(t, t - R_c(t), w) M_c(t, dw)). \quad (10.4)$$

Ceci est une équation qui peut être exploitée numériquement moyennant quelques hypothèses sur la fonction e (voir 12.1.1).

10.1.3 Simplifications de l'EDP

Une hypothèse de densité et d'indépendance de e Pour aller plus loin on va faire une hypothèse supplémentaire sur la distribution initiale, on va supposer qu'elle a une densité continûment dérivable. En fait l'hypothèse qu'il y a une densité est nécessaire dans la mesure où l'équation 10.4 transporte les atomes : il n'y a pas d'espoir ici comme dans l'équation de la chaleur que les atomes soient diffusés instantanément.

De même nous devons ajouter une hypothèse de régularité sur la fonction K et e .

Corollaire 10.1. *$M(s, dw)$ est la distribution de fenêtres qui évolue selon 10.4. Supposons que $\mu(0)$ a une densité $\mathcal{C}^1$ que l'on note $p(0, w)$. On suppose de plus que R_c , K et $e(t, \cdot - R(\cdot), w)$ sont des fonctions continues et dérivables, en particulier K et surtout e sont considérées comme des données externes connues à l'avance.*

Alors $\mu(t, dw)$ a une densité continûment dérivable et vérifie l'EDP :

$$\frac{\partial p}{\partial t} + \frac{a}{R_c} \frac{\partial p}{\partial w} = \frac{K}{R_c} (t - R_c) \left(2e(2w)p(2w) - e(w)p(w) \right). \quad (10.5)$$

Il n'y a rien à démontrer dans la mesure où on a enlevé toutes les singularités en prenant toutes les fonctions dérivables et en rompant la dépendance entre e et μ . C'est simplement une réécriture de 10.4. De plus le fait de supposer que R, K et e sont des données pourrait être à l'origine d'une méthode de point fixe, même si les modalités de mise en oeuvre sont inutilisables dans la pratique, cela justifie une approche itérative si on a montré l'unicité de la limite du champ moyen.

C'est cette équation que nous simulerons dans la pratique. Nous rajouterons même une hypothèse simplificatrice sur e comme nous allons le voir ce qui permet de traiter entièrement le problème et nous supprimons les classes de RTT en notant $r(t)$ le RTT commun.

Une simplification encore plus grande sur e L'astuce consiste à faire l'approximation suivante :

$$e(t, t - R_c, w) \approx w. \quad (10.6)$$

est une mesure (qui dépend linéairement du réel w_0 pour chaque t donné), telle que :

$$\frac{M_c(t, dw + hw_0) - M_c(t, dw)}{h} \xrightarrow[h \rightarrow 0]{*} D_w M_c(t, dw)(w_0),$$

où $M_c(t, dw + hw_0)$ représente la mesure translatée.

C'est clairement peu précis mais si le taux de pertes est très faible, alors la fenêtre à un RTT dans le passé est moindre que maintenant. Surtout, quand le taux de pertes est modéré, les grandes fenêtres ont plus de chances d'avoir été deux fois plus grandes il y a un RTT. On se reportera à 12.1.1 pour voir comment il serait possible de résoudre directement les équations ou de réaliser une simplification moins brusque.

Cependant si le taux de pertes est faible et qu'on est proche d'un point fixe du système (8.13 et 8.14) qui est stable, alors l'espérance de la fenêtre il y a un RTT devrait être proche de sa valeur actuelle, en tout cas on ne peut certainement pas faire la simplification : $e(t, t - R_c, w) \approx \mathbb{E}(w)(t - r)$, car cela revient à dire que les fenêtres sont très peu dispersées autour de la moyenne, ce qui n'est pas vérifié.

Cette approximation donne l'équation de transport simplifiée (avec une seule classe de RTT) :

Corollaire 10.2. *Si la distribution initiale $\mu(t, dw)$ a une densité continue, alors $\mu(t, dw)$ a une densité continue $p(t, w)$ différentiable en t qui satisfait approximativement l'équation de transport suivante :*

$$\frac{\partial p(t, w)}{\partial t} = \left(p(t, 2w) \frac{2w}{r(t - r(t))} \times 2 - p(t, w) \frac{w}{r(t - r(t))} \right) \kappa(t - r(t)) - \frac{1}{r(t)} \frac{\partial p(t, w)}{\partial w} \quad (10.7)$$

et

$$\begin{aligned} \frac{dq(t)}{dt} = & \int_w \frac{w}{r(t)} p(t, w) dw (1 - \kappa(t)) - L \\ & - \left(\int_w \frac{w}{r(t)} p(t, w) dw (1 - \kappa(0)) - L \right)^- \chi\{q = 0\} \end{aligned} \quad (10.8)$$

où $r(t) = T + q(t - r(t))/L$. $\kappa(t) = F(q(t))$ avec F continue et quand $F(q(t)) = 1$ (i.e. quand $q(t) = q_{max}$), $\kappa(t)$ est déterminé par

$$\int_w \frac{w}{r(t)} p(t, w) dw \cdot (1 - \kappa(t)) = L.$$

10.2 Étude de l'équation simplifiée

10.2.1 Point fixe des équations simplifiées du champ moyen

Un point fixe des équations (10.7) et (10.8) vérifiera :

$$\frac{df_\kappa(w)}{dw} = \kappa(2(2w)f_\kappa(2w) - wf_\kappa(w)) \quad (10.9)$$

$$L = (1 - \kappa) \frac{1}{r} \int_w wf_\kappa(w) dw. \quad (10.10)$$

(10.10) est simplement une formule de conservation du débit : le débit entrant égale le débit sortant. Attention il faut bien noter que nous supposons qu'il n'y a pas de pertes quand la file d'attente est vide. On pourrait bien sûr obtenir d'autres points fixes ainsi

qui n'utilisent pas la totalité du débit disponible. C'est ce qui intervient en particulier dans [129] dans le cas de files d'attentes très petites où on a des pertes en saturant la file d'attente même avec une utilisation inférieure à 1.

Théorème 10.2. Soit $\Psi = \sum_{i=0}^{\infty} \frac{2^i}{\prod_{j=1}^i (1-4^j)}$ ($\Psi \approx 0.4194$). On suppose que κ est fixé. Alors l'unique densité $f_{\kappa}(w)$ solution de (10.9) est donnée par

$$f_{\kappa}(w) = \sum_{i=0}^{\infty} a_i \exp\left(-\kappa 4^i \frac{w^2}{2}\right) \quad (10.11)$$

où

$$a_0 = \sqrt{\frac{2}{\pi}} \frac{1}{\Psi} \sqrt{\kappa} ; \quad a_i = a_{i-1} \frac{4}{1-4^i} = a_0 \frac{4^i}{\prod_{j=1}^i (1-4^j)}. \quad (10.12)$$

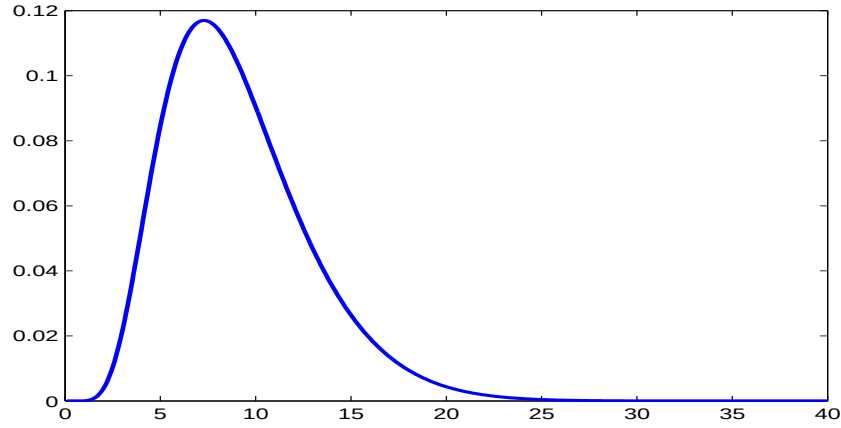


FIG. 10.1 – Histogramme des fenêtres en équilibre dynamique pour l'équation (10.9) avec $\kappa = 2\%$

Cette solution a été obtenue indépendamment dans [1] par Adjih, Jacquet et Vvedenskaya en suivant une autre démarche pour obtenir directement l'équation des points fixes sans étudier la dynamique du système.

Preuve : Le fait que f_{κ} est une solution se constate simplement par substitution, il n'y a aucun problème de convergence (convergence absolue en norme de la série 10.11). Le seul point à élucider est de savoir si f_{κ} est bien une fonction positive. Ce dernier point est plus difficile qu'il n'y paraît, il y a d'ailleurs une erreur sur la question de positivité dans la version initiale de [18], qui est corrigée dans [16] avec une généralisation de l'équation. Ici on pourra se contenter de la démonstration suivante :

Uniquement lorsque w est grand, le module du terme général de la série est décroissant. Nous sommes donc en présence d'une série alternée. Par conséquent $f_{\kappa}(w) \geq f_{\kappa}^2(w) \geq 0, \forall w \geq W_{\kappa}$, où $f_{\kappa}^i(w)$ est la série tronquée au $i^{\text{ème}}$ terme (en tout cas dès que a_0 est positif). Ensuite pour récupérer la positivité des valeurs en dessous de W_{κ} pour un κ donné, il suffit d'utiliser l'équation du point fixe, on sait déjà que f_{κ} la satisfait, et on intègre entre w et l'infini (on voit aisément que $f(+\infty) = 0$) :

$$f_\kappa(W) = - \int_W^\infty \kappa (2(2w)f_\kappa(2w) - wf_\kappa(w)) dw = \kappa \int_W^{2W} f_\kappa(w) dw. \quad (10.13)$$

On sait que f_κ est continue (limite uniforme de fonctions continues), elle prend des valeurs strictement positives pour $w > W_\kappa$. On voit grâce à (10.13) qu'il n'y a pas de plus grande valeur pour laquelle f_κ n'est pas strictement positive. Ainsi $f_\kappa > 0$ pour sur $[0, \infty)$. De plus on voit dans la formule que les valeurs de f_κ aux bords sont 0 et que c'est les valeurs de prolongement $\mathcal{C}^\infty$ (la manière la plus simple de voir que f_κ est $\mathcal{C}^\infty$ est l'équation différentielle où le terme de droite est dérivable au moins une fois de plus que celui de gauche).

La valeur de a_0 est déterminée par la normalisation qui fait de f_κ une densité.

$$\begin{aligned} 1 &= \int_0^\infty f_\kappa(w) ds = \sum_{i=0}^\infty a_i \int_0^\infty e^{-\kappa \frac{w^2}{2} 4^i} dw = \sum_{i=0}^\infty a_i \frac{1}{2} \sqrt{2\pi \frac{1}{\kappa 4^i}} \\ &= \frac{\sqrt{2\pi}}{2} \sum_{i=0}^\infty a_0 \frac{4^i}{\prod_{j=1}^i (1 - 4^j)} \sqrt{\frac{1}{\kappa 4^i}} = \frac{1}{\sqrt{\kappa}} a_0 \sqrt{\frac{\pi}{2}} \sum_{i=0}^\infty \frac{2^i}{\prod_{j=1}^i (1 - 4^j)}. \end{aligned} \quad (10.14)$$

■

Théorème 10.3. *f et toutes ses dérivées valent 0 pour $t = 0$ et tendent de manière hyper-exponentielle vers 0 quand $t \rightarrow \infty$*

Preuve : On voit par récurrence $\sum_{i=0}^N a_i = \frac{a_0}{\prod_{j=1}^N (1 - 4^j)}$. Pour $N = 0$, la formule est vraie. Si elle est vérifiée pour N , alors

$$\sum_{i=0}^{N+1} a_i = \frac{a_0}{\prod_{j=1}^N (1 - 4^j)} + a_{N+1} = \frac{a_0}{\prod_{j=1}^N (1 - 4^j)} \left(1 + \frac{4^{N+1}}{1 - 4^{N+1}}\right) = \frac{a_0}{\prod_{j=1}^{N+1} (1 - 4^j)},$$

elle est donc vérifiée pour $N + 1$. Comme $f(0) = \sum_{i=0}^\infty a_i$ il suit que $f(0) = 0$. Maintenant nous dérivons l'équation du point fixe n fois, et faisons $w \rightarrow 0$. La partie de gauche donne la dérivée $(n + 1)$ -ième de f tandis que la partie de droite donne une somme de produits avec w et des dérivées d'ordre inférieur de f . Comme $f(0) = 0$ nous avons $f^{(n)}(0) = 0$ par récurrence.

■

Il y a une unique solution à (10.9) and (10.10) sous réserve que F est bien choisie. Nous supposerons dans la suite que F est continue et croissante en q et admet un saut à

Q_{max} . Notons tout d'abord que

$$\begin{aligned}
 \int_w w f_\kappa(w) dw &= \sum_{i=0}^{\infty} a_i \int_0^{\infty} w \exp(-\kappa 4^i \frac{w^2}{2}) dw \\
 &= \sum_{i=0}^{\infty} a_i \frac{1}{\kappa 4^i} \\
 &= \frac{1}{\kappa} a_0 \sum_{i=0}^{\infty} \frac{1}{\prod_{j=1}^i (1 - 4^j)} = \frac{1}{\kappa} a_0 \xi \\
 &= \alpha \sqrt{\frac{1}{\kappa}}
 \end{aligned} \tag{10.15}$$

où $\xi = \sum_{i=0}^{\infty} \frac{1}{\prod_{j=1}^i (1 - 4^j)}$ ($\xi \approx 0.6885$) et $\alpha = \sqrt{\frac{2}{\pi} \frac{\xi}{\Psi}}$ ($\alpha \approx 1.310$).

On tire de (10.10) l'égalité :

$$L = \frac{(1 - \kappa)}{r} \alpha \sqrt{\frac{1}{\kappa}} = \frac{\alpha (1 - \kappa)}{r \sqrt{\kappa}}.$$

Ainsi nous avons

$$\frac{(1 - \kappa)^2}{\kappa} = \left(\frac{rL}{\alpha} \right)^2. \tag{10.16}$$

La fonction $(1 - \kappa)^2/\kappa$ est strictement décroissante. Il y a donc un unique point fixe κ satisfaisant (10.16) pour une valeur fixée de r .

Il y a deux possibilités. De (10.16), comme $(1 - F(q))^2/F(q)$ est décroissante en q , la valeur du point fixe pour la taille de la file d'attente est déterminée par l'unique solution de :

$$\frac{(1 - F(q))^2}{F(q)} = \left(\frac{(T + q/L)L}{\alpha} \right)^2. \tag{10.17}$$

Si la solution de (10.17) est plus petite que q_{max} alors $\kappa = F(q)$ donne le point fixe et on a $r = T + q/L$. Sinon $q = q_{max}$ et κ est donné par (10.16) (c'est le cas où $Q_N(t)$ fait la gigue à q_{max}).

10.2.2 Une discussion autour des timeouts

Dans le modèle présenté jusqu'ici, nous avons ignoré les timeouts. Dans un état d'équilibre, les populations qui sortent et qui entrent en timeout s'équilibrent (pour simplifier on supposera que le slow-start est une partie du timeout et qu'on passe directement en congestion avoidance). À tout temps N connexions sont actives sur un total de N' ; $N' - N$ sont en timeout, c'est-à-dire qu'elles ont un débit nul et attendent l'expiration du timer. Nous définissons la proportion de timeout en fonction du taux de pertes $TO(\kappa)$:

$$N' - N = TO(\kappa)N.$$

En état stationnaire avec propagation de chaos, ceci est aussi la part du temps que chaque connexion va passer en timeout. Dans la plupart des cas nous connaissons N' , on résout donc l'équation du dessus pour connaître N , le nombre des connexion actives. La distribution des fenêtres à l'équilibre peut être utilisée pour connaître $TO(\kappa)$ comme nous allons le voir tout de suite.

Une connexion sera en timeout quand aucun mécanisme de récupération ne lui permet d'interpréter une perte avant l'expiration du temps RTO . Pour plus de précisions, on voudra bien se reporter à la section (1.1).

Parmi les grandes causes de timeout existant qui font échouer le procédé de fast retransmit pour diverses raisons que nous ne détaillerons pas, citons :

- un paquet est perdu alors que la taille de la fenêtre est plus petite que 4,
- le paquet retransmis est perdu,
- deux paquets sont perdus dans la même fenêtre.

Ici nous nous occuperons uniquement de la première cause. En effet l'utilisation d'un mécanisme du type SACK qui consiste à renvoyer des ACK contenant l'information exacte des paquets qui ont été reçus permet d'éliminer les deux autres cas de timeout cités.

Si le niveau des pertes est faible, il est peu probable d'avoir deux pertes sur un RTT pour de petites fenêtres. Ainsi les connexions qui iront en timeout après une perte seront celles qui ne peuvent pas recevoir 4 ACK identiques après la perte, c'est-à-dire les connexions dont la fenêtre est plus petite ou égale à 4. Dans la mesure où pour des fenêtres de cette taille faire slow start ou congestion avoidance revient au même, notre hypothèse de négliger le slow start n'aura pas d'impact majeur sur le résultat.

La valeur de RTO est $RTT + 4Var(RTT)$ quand on n'a pas de timeouts répétés, ce qui est le cas pour de faibles niveaux de pertes. Habituellement cela donne $RTO \approx 2RTT$. Aussi la probabilité d'attendre l'expiration du timeout RTO est donnée par :

$$T(\kappa) = \int_0^4 f_\kappa(w)[1 - (1 - \kappa)^w]dw \quad (10.18)$$

$$= \sum_{i=0}^{\infty} a_i T_i \text{ où} \quad (10.19)$$

$$T_i = \frac{\sqrt{\pi}}{2} A_i \left[erf\left(\frac{4}{A_i}\right) - \exp\left(\frac{A_i^2}{4} \ln(1 - \kappa)^2\right) \left(erf\left(\frac{4}{A_i} - \ln(1 - \kappa)\frac{A_i}{2}\right) - erf\left(-\ln(1 - \kappa)\frac{A_i}{2}\right) \right) \right]. \quad (10.20)$$

En notant $A_i = \sqrt{2/(4^i \kappa)}$ et la fonction d'erreur, $erf(z) = (2/\sqrt{\pi}) \int_0^z \exp(-t^2)dt$. La preuve de ceci est donnée dans la section suivante.

Finalement

$$TO(\kappa) \approx \frac{RTO}{RTT} T(\kappa).$$

Théorème 10.4. *Dans la situation de point fixe, la proportion $On = N/N'$ de connexions actives dans le cas TCP persistant sous nos hypothèses d'étude est donnée par $\frac{1}{On} - 1 =$*

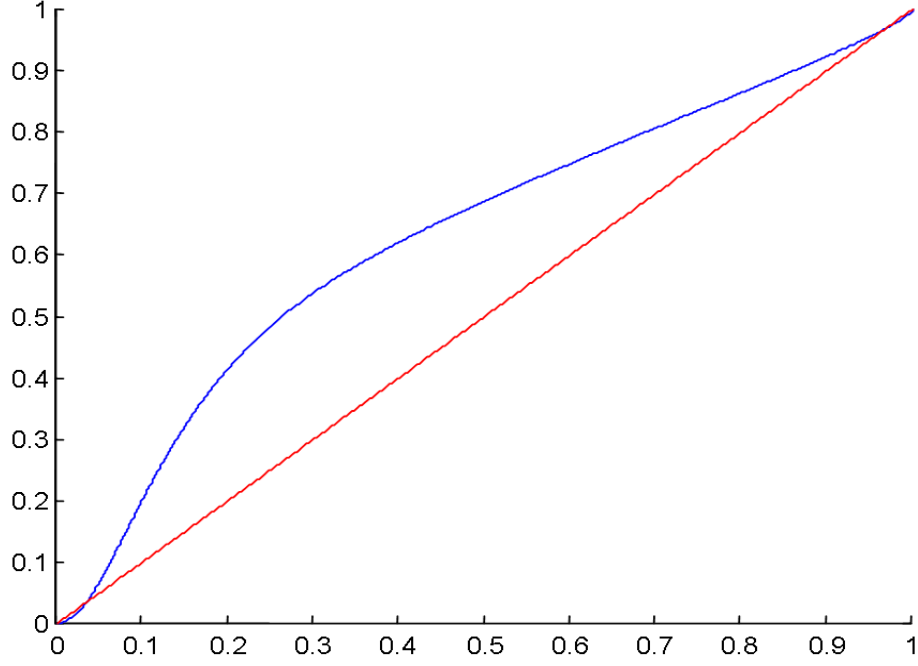


FIG. 10.2 – % timeouts vs taux de pertes par simulation de l'EDP (pour comparaison $x \mapsto x$ en rouge)

$TO(\kappa)$, ie :

$$On = \frac{1}{1 + TO(\kappa)}$$

et

$$TO(\kappa) = \frac{RTO}{RTT} T(\kappa)$$

où $T(\kappa)$ est donnée par (10.19). Un équivalent de $T(\kappa)$ quand κ est petit est donné par :

$$T(\kappa) = 2\alpha\kappa\sqrt{\kappa} + O(\kappa^2\sqrt{\kappa}).$$

La constante devant $\kappa\sqrt{\kappa}$ vaut environ : 2.62

Finalement nous trouvons des résultats très proches de ceux annoncés par [121] avec un peu plus de précision.

10.2.3 Preuve du théorème sur les timeouts

De (10.11),

$$\sum_{i=0}^{\infty} a_i \int_0^4 \left[1 - \exp\left(w \ln(1 - \kappa) - \frac{w^2}{A_i^2}\right) \right] dw = \sum_{i=0}^{\infty} a_i T_i$$

où

$$\begin{aligned} T_i &= \int_0^4 [\exp(-\frac{w^2}{A_i}) - \exp(w \ln(1 - \kappa) - \frac{w^2}{A_i})] dw \\ &= \int_0^4 \exp(-\frac{w^2}{A_i}) dw - \exp(\frac{A_i^2}{4} \ln^2(1 - \kappa)) \int_0^4 \exp(-(\frac{w}{A_i} - A_i \ln(1 - \kappa)/2)^2) dw, \end{aligned}$$

un changement de variable donne (10.19).

$$\text{Soit } g_\delta(x) = \exp(-\delta \frac{x^2}{2}) \text{ et soit } \hat{g}(t) = \int_0^\infty g(x) e^{tx} dx.$$

Nous allons évaluer la transformée de Laplace de f_κ ,

$$T_f(\theta) = \int_0^\infty f_\kappa(w) e^{-\theta w} dw = \sum_{i=0}^\infty a_i \hat{g}_{\kappa 4^i}(-\theta) \text{ par (10.11).} \quad (10.21)$$

On calcule ensuite $\hat{g}_\delta$.

$$\begin{aligned} \hat{g}_\delta'(t) &= \int_0^\infty x g_\delta(x) e^{tx} dx = \int_0^\infty \frac{-1}{\delta} \frac{d(e^{\delta \frac{-x^2}{2}})}{dx} e^{tx} dx \\ &= -\left[\frac{1}{\delta} e^{\delta \frac{-x^2}{2}} e^{tx} \right]_{x=0}^{x=\infty} + \frac{t}{\delta} \int_0^\infty g_\delta(x) e^{tx} dx \\ &= \frac{1}{\delta} + \frac{t}{\delta} \hat{g}_\delta(t), \text{ pour tout } t \in \mathbb{R}. \end{aligned}$$

On résout cette équation en multipliant les deux cotés par $e^{-t^2/(2\delta)}$ ce qui donne

$$\hat{g}_\delta'(t) e^{-t^2/(2\delta)} - \left(\frac{t}{\delta} e^{-t^2/(2\delta)} \right) \hat{g}_\delta(t) = \frac{1}{\delta} e^{-t^2/(2\delta)}.$$

Comme $\hat{g}_\delta(0) = 1/\sqrt{\pi/(2\delta)}$,

$$\hat{g}_\delta(t) = \sqrt{\frac{\pi}{2\delta}} + \frac{1}{\delta} e^{t^2/(2\delta)} \int_0^t e^{-x^2/(2\delta)} dx.$$

Si on substitue dans (10.21) nous trouvons

$$T_f(\theta) = 1 - \sum_{i=0}^\infty a_i \frac{1}{4^i \kappa} e^{\frac{\theta^2}{2} \frac{1}{\kappa 4^i}} \int_0^\theta e^{-\frac{x^2}{2} \frac{1}{\kappa 4^i}} dx$$

en utilisant (10.14).

Maintenant, calculons l'équivalent quand κ tend vers 0. Il serait un peu plus long de montrer rigoureusement que les termes que nous mettons dans le $O()$ sont bien petits, cependant c'est juste calculatoire puisqu'on retrouve des dérivées de f . Sans trop se préoccuper des détails de justification des convergences, on a :

$$T(\kappa) = \sum_{i=0}^{\infty} a_i T_i \text{ où}$$

$$T_i = \frac{\sqrt{\pi}}{2} A_i \left[\frac{2}{\sqrt{\pi}} \frac{4}{A_i} - \left(1 - \frac{A_i^2}{4} \kappa^2\right) \left(\frac{2}{\sqrt{\pi}} \frac{4}{A_i} \right) \right] + O(\kappa^2).$$

Il reste finalement :

$$\begin{aligned} T(\kappa) &= a_0 \sum_{i=0}^{\infty} \frac{a_i}{a_0} [A_i^2 \kappa^2] + O(\kappa^2 \sqrt{\kappa}) \\ &= 2\kappa \sqrt{\frac{2}{\pi}} \frac{1}{\Psi} \sqrt{\kappa} \sum_{i=0}^{\infty} \frac{a_i}{4^i a_0} + O(\kappa^2 \sqrt{\kappa}) \\ &= 2\alpha \kappa \sqrt{\kappa} + O(\kappa^2 \sqrt{\kappa}). \end{aligned}$$

10.3 Contrôle et stabilité des points fixes de l'équation approximée

Discussion sur la notion de stabilité Dans ce qui suit nous allons nous référer à diverses notions de stabilité d'un système dynamique. Dans cette section nous dirons qu'un point fixe d'un système est stable pour dire que le point fixe est localement attractif. Lorsqu'une valeur d'équilibre est v^{eq} , nous noterons $\Delta v(t) := v(t) - v^{eq}$. Par conséquent nous avons $\dot{v}(t) = \Delta \dot{v}(t)$, formule que nous utiliserons plus tard. L'intérêt d'étudier plutôt Δv que v est que la stabilité consiste à vérifier que Δv tend vers 0 quand t tend vers l'infini.

Plan de la section Nous n'allons pas être capables d'étudier la stabilité de l'équation même sous sa forme simplifiée actuelle. Cependant nous commençons l'exploitation de l'équation sous forme générale avant de la simplifier. En particulier nous voyons que les simplifications qui doivent être faites pour retomber sur l'équation fluide de [67], en particulier pour l'étude de l'équilibre sont tout à fait irréalistes. Dans la suite $\dot{X}$ est la dérivée de X par rapport à t . Ensuite nous tentons de faire une étude plus juste de l'équilibre. Mais avant tout rappelons que nous avons déjà réalisé des simplifications qui sont décrites dans 10.1.3, en particulier le moment M_2 que nous verrons apparaître est une version qui suppose de petits taux de pertes.

10.3.1 Étude des moments de la distribution des fenêtres

Équation différentielle des moments Notons $M_m(t)$ le m ème moment de $p(t)$ pour une condition initiale donnée.

Tout d'abord, de (10.7), en multipliant par w^m et en intégrant par rapport à w , il vient

$$\begin{aligned} \frac{\partial}{\partial t} \int_{w>0} w^m p(t, w, q) dw &= 4 \frac{\kappa(t - r(t))}{r(t - r(t))} \int_{w>0} w^{m+1} p(t, 2w, q) dw \\ &\quad - \frac{\kappa(t - r(t))}{r(t - r(t))} M_{m+1} - \int_{w>0} w^m \frac{\partial p}{\partial w}(t, w, q) dw. \end{aligned}$$

En faisant un changement de variable $u = 2x$ dans la première intégrale, et une intégration par partie dans la deuxième, il vient (on se rappelle que p est exponentiellement décroissant) :

$$\dot{M}_m = \frac{\kappa(t - r(t))}{r(t - r(t))} \left(\frac{1}{2^m} - 1 \right) M_{m+1} + \frac{1}{r(t)} m M_{m-1}. \quad (10.22)$$

De (10.8) on a aussi

$$\dot{q} = \frac{1 - \kappa(t)}{r(t)} M_1 - L. \quad (10.23)$$

Une borne sur les moments Dans (10.22)_m, si nous appliquons l'inégalité de Jensen avec la fonction convexe $u \mapsto u^{(m+1)/m}$ à M_{m+1} , on a :

$$\dot{M}_m \leq \frac{\kappa(t - r(t))}{r(t - r(t))} \left(\frac{1}{2^m} - 1 \right) M_m^{(m+1)/m} + \frac{1}{r(t)} m M_{m-1}.$$

Pour $m = 1$, si $M_1 \geq \sqrt{\frac{2r(t-r(t))}{\kappa(t-r(t))r(t)}}$ alors $\dot{M}_1 \leq 0$ ($r \geq d > 0$), et cela décroît.

Supposons maintenant que r et κ sont constants, il vient $\overline{\lim}_{t \rightarrow \infty} (M_1(t)) \leq \sqrt{2/\kappa} = M_1^{Max}$. Par induction :

$$\overline{\lim}(M_{m-1}) \leq M_m^{Max} \text{ avec } M_m^{Max} = \left(\frac{1}{\kappa} \frac{m 2^m}{2^m - 1} M_{m-1}^{Max} \right)^{\frac{m}{m+1}}. \quad (10.24)$$

Dans l'équation (10.23), par le même raisonnement, si $r \geq (1 - \kappa) M_1^{Max} / L$ alors il décroît.

$$\overline{\lim}_{t \rightarrow \infty} r(t) \leq r^{Max} = \frac{1 - \kappa}{L} M_1^{Max}. \quad (10.25)$$

10.3.2 Hypothèse pour l'étude de la stabilité

Hypothèses simplificatrices, mais invérifiables pour étudier la stabilité On va supposer l'égalité dans Jensen, c'est-à-dire que les fenêtres ont une variance très faible (le strict cas d'égalité est la concentration en un point). On connaît les valeurs des moments à l'équilibre (10.15, 10.22), la variance est $\sigma^2 = M_2 - M_1^2 = \frac{2-\alpha^2}{\kappa} \approx \frac{.28}{\kappa}$. Cette approximation tient donc à peu près lorsque $\frac{M_2 - M_1^2}{M_1^2} = \frac{\sigma^2}{M_1^2} = \frac{2-\alpha^2}{\alpha\sqrt{\kappa}} \approx 0.22/\sqrt{\kappa}$ est aussi petit que possible. Globalement ceci amène à supposer que κ est très grand. Cette hypothèse n'est donc absolument pas réalisable dans la mesure où le modèle est fait pour des taux de pertes modérés où TCP fonctionne bien.

En faisant cette hypothèse qui n'est donc essentiellement jamais bonne, on se ramène au système étudié par [67], qui fait en plus l'hypothèse dans certains moments de l'étude que W est grand (W grand = κ petit à l'équilibre). Cette étude n'a donc, à notre sens, à peu près aucune utilité telle qu'elle est réalisée.

Des hypothèses plus réalistes Le but est de ne pas devoir étudier toutes les équations des moments, elles ne tombent pas à notre connaissance dans une classe connue d'équations ; aussi il faut autant que possible se restreindre à l'équation avec la dérivée du premier moment. On va donc dans la suite étudier la stabilité de l'équilibre en imposant une perturbation particulière, toutes les fenêtres sont décalées d'un Δw . C'est une hypothèse réaliste qui correspondrait au cas où on laisse pendant un peu de temps les fenêtres augmenter toutes seules (c'est-à-dire en coupant les pertes).

10.3.3 Stabilité sous les contrôles classiques

Étude simplifiée sous les perturbations Δw de fenêtres avec κ et r fixes sans limite de débit Ce cas est très simple. On part des valeurs d'équilibre avec les exposants $\square^{eq}$. Les conséquences de la perturbation sont au premier ordre :

$$\Delta M_1 = \Delta w, \quad (10.26)$$

$$\Delta M_2 = 2M_1^{eq} \Delta w. \quad (10.27)$$

Une fois injecté dans (10.22)₁ (l'indice 1 réfère à la valeur $m = 1$ dans l'équation sur les moments), il vient :

$$\dot{M}_1 = \Delta \dot{M}_1 = -\frac{\kappa}{2r} (M_2^{eq} + 2M_1^{eq} \Delta w) + \frac{1}{r} = -\frac{\kappa}{r} M_1^{eq} \Delta M_1 = -\frac{\alpha \sqrt{\kappa}}{r} \Delta M_1.$$

L'équilibre est toujours stable. Les informations que nous récupérons au passage sont que l'équilibre est d'autant plus stable que :

- κ est grand (ie : W est petit),
- r est petit.

On peut voir une illustration de ceci dans la figure 10.3

Étude simplifiée sous les perturbations Δw de fenêtres avec κ et r variables selon RED en débit limité Récrivons le système que nous allons étudier ce qui est une reprise de (10.22) et de (10.23) en rajoutant l'équation linéaire de RED et celle du RTT :

$$\begin{aligned} \dot{M}_1 &= -\frac{1}{2} \frac{\kappa(t-r(t))}{r(t-r(t))} M_2 + \frac{1}{r(t)}, \\ \dot{q} &= \frac{1-\kappa(t)}{r(t)} M_1 - L, \\ r(t) &= T + \frac{q(t)}{L}, \\ \kappa(t) &= \kappa^{eq} + \epsilon(q(t) - q^{eq}). \end{aligned}$$

Linéarisé autour de l'équilibre, ceci nous donne :

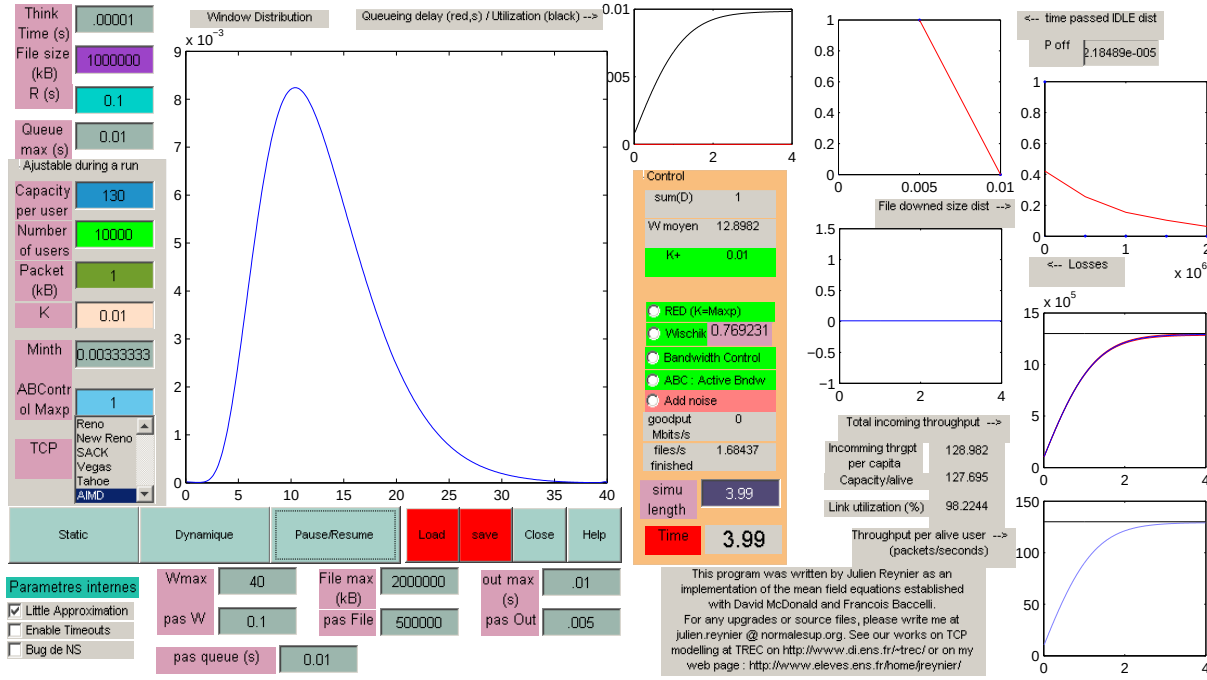


FIG. 10.3 – Cette figure est obtenue à l’aide du simulateur HTTP (décrit au chapitre 12) en mettant des paramètres qui le font simuler les équations que nous considérons. On voit la stabilisation avec le taux de pertes constant fixé à $\kappa = 1\%$ et le RTT $r = 0.1s$. On voit en bas à droite deux figures qui représentent le débit total et le débit par utilisateur. Le grand cadre donne la distribution des fenêtres.

$$\begin{aligned}\Delta \dot{M}_1 &= -\frac{\kappa^{eq}}{r^{eq}} M_1^{eq} \Delta M_1 - \frac{1}{2} \frac{1}{r^{eq}} M_2^{eq} \Delta \kappa(t - r^{eq}) + \frac{1}{2} \frac{\kappa^{eq}}{(r^{eq})^2} M_2^{eq} \Delta r(t - r^{eq}) - \frac{1}{(r^{eq})^2} \Delta r(t - r^{eq}), \\ \Delta \dot{q} &= \frac{1 - \kappa^{eq}}{r^{eq}} \Delta M_1 - \frac{1}{r^{eq}} M_1^{eq} \Delta \kappa(t) - \frac{1 - \kappa^{eq}}{(r^{eq})^2} M_1^{eq} \Delta r(t), \\ \Delta r(t) &= \frac{1}{L} \Delta q(t), \\ \Delta \kappa(t) &= \epsilon \Delta q(t).\end{aligned}$$

Notons que les deux derniers termes dans l’équation sur $\Delta \dot{M}_1$ s’annulent comme $M_2^{eq} = \frac{2}{\kappa^{eq}}$.

En simplifiant :

$$\Delta \dot{M}_1 = -\alpha \frac{\sqrt{\kappa^{eq}}}{r^{eq}} \Delta M_1 - \frac{\epsilon}{r^{eq} \kappa^{eq}} \Delta q(t - r^{eq}), \quad (10.28)$$

$$\Delta \dot{q} = \frac{1 - \kappa^{eq}}{r^{eq}} \Delta M_1 - \frac{\alpha}{r^{eq} \sqrt{\kappa^{eq}}} \left(\epsilon + \frac{1 - \kappa^{eq}}{L r^{eq}} \right) \Delta q(t). \quad (10.29)$$

L’étude de ce genre d’équations avec retard se fait traditionnellement en essayant d’injecter des solutions exponentielles complexes et en regardant pour quel jeu de paramètres

toutes les solutions sont décroissantes (on recherche les modes propres du système). On pourra par exemple consulter [28] où ce problème est traité. Mais avant de résoudre, regardons ce qui se passe intuitivement. L'équation (10.28) paraît stable, mais en fait il faut faire attention au délai, même avec le bon signe il peut produire des oscillations. L'équation (10.29) quant à elle a une tournure assez inquiétante. On voit que si le terme avec le signe + est prépondérant, on risque l'instabilité.

On cherche les solutions sous la forme $X = \Delta M = Ge^{Ht}$ et $Y = \Delta q = e^{Ht}$ (pas la peine de prendre plus général dans la mesure où on doit vérifier une égalité pour tout t et que les solutions sont à constante près), avec G et H complexes. De plus on met le système avec délai sous la forme :

$$\dot{X} = -aX - bY(t - r), \quad (10.30)$$

$$\dot{Y} = cX - dY. \quad (10.31)$$

Ce qui nous donne

$$GH = -aG - be^{-Hr}, \quad (10.32)$$

$$H = cG - d. \quad (10.33)$$

H est donc solution de :

$$H^2 + (a + d)H + ad = -bce^{-Hr},$$

$$(H + a)(H + d) = -bce^{-Hr}. \quad (10.34)$$

On a la condition suffisante de stabilité qui est que H ait une partie réelle strictement négative.

Remarque 10.2. *Voici un raisonnement faux, mais intuitivement intéressant que l'on pourrait faire. Une condition suffisante pour avoir des racines réelles négatives est que la parabole de gauche en son minimum soit plus basse que l'exponentielle de droite (il ne serait pas raisonnable de dériver pour savoir où se trouve exactement le minimum). Dans la mesure où $ad > 0$, cela fera bien deux racines négatives. Une condition suffisante de stabilité serait donc :*

$$(a - d)^2 \geq 4bce^{\frac{r}{2}(a+d)}. \quad (10.35)$$

On linéarise (10.34) pour les petites valeurs de r . Il vient successivement :

$$\begin{aligned} (H + a)(H + d) &= -bc(1 - Hr), \\ H^2 + (a + d - bcr)H + ad + bc &= 0. \end{aligned}$$

On voit alors que la partie réelle des deux racines est strictement négative si et seulement si $bcr < a + d$. Ce qui est la condition de stabilité que nous retiendrons pour les petites valeurs de r .

Théorème 10.5. *Le point fixe du système approximé avec utilisation de RED, sous les hypothèses simplificatrices pour étudier la stabilité et pour les petites valeurs de r , est stable si et seulement si : $rbc < a + d$. En d'autres termes la CNS de stabilité est :*

$$\epsilon(1 - \kappa^{eq} - \alpha\sqrt{\kappa^{eq}}) < (\alpha\sqrt{\kappa^{eq}} + 1)\kappa^{eq}. \quad (10.36)$$

En nous rappelant que κ^{eq} est de l'ordre de $\left(\frac{\alpha}{r^{eq}L}\right)^2$, une condition suffisante de stabilité très simple est que :

$$\epsilon \leq \kappa^{eq}. \quad (10.37)$$

Remarque 10.3. *En regardant la période des oscillations pour r petit et $Lr \gg 1$, on voit que*

$$T \approx 2\pi r. \quad (10.38)$$

Ceci est une pseudo-période. Cela signifie que dans la pratique on converge vers l'équilibre ou qu'on est instable. Cependant tant qu'on reste en dessous de certaines valeurs critiques on observe un comportement périodique amorti ou périodique amplifié. On peut trouver une confrontation de ce résultat à l'expérience dans la section 12.2.1, où on voit que c'est une bonne évaluation du phénomène périodique. De plus cette formule reste assez bien vérifiée dans les cas instables. On peut aussi voir une bonne illustration de ceci dans la figure (10.4).

Attention avant d'exploiter les résultats de ce théorème, ϵ et L sont normalisés par le nombre d'utilisateurs : la capacité vaut NL , la pente de RED en fonction de la file réelle vaut ϵ/N . De plus on rappelle que α est une constante dont la valeur est proche de 1,310.

Preuve : $rbc < a + d$ se récrit :

$$\epsilon(1 - \kappa^{eq}) < \alpha\kappa^{eq}\sqrt{\kappa^{eq}} + \alpha\sqrt{\kappa^{eq}}\left(\epsilon + \frac{1 - \kappa^{eq}}{Lr^{eq}}\right),$$

c'est-à-dire :

$$\epsilon(1 - \kappa^{eq} - \alpha\sqrt{\kappa^{eq}}) < \alpha\kappa^{eq}\sqrt{\kappa^{eq}} + \frac{\alpha}{L}\sqrt{\kappa^{eq}}(1 - \kappa^{eq})\frac{1}{r^{eq}}.$$

En tenant compte du fait que $r^{eq}L = (1 - \kappa^{eq})M_1^{eq} = \frac{\alpha(1 - \kappa^{eq})}{\sqrt{\kappa^{eq}}}$,

$$\epsilon(1 - \kappa^{eq} - \alpha\sqrt{\kappa^{eq}}) < (\alpha\sqrt{\kappa^{eq}} + 1)\kappa^{eq}.$$

■

Comparaison de notre estimation avec celle fournie dans [67]. On prend le même exemple d'un réseau avec un goulot d'étranglement RED : $C = NL = 3750$ paquets/s, $N^- = 60$, $R^+ = 0.2s$. L'évaluation exacte du point de bifurcation avec (10.36) est $L_{RED} = 2.44 \cdot 10^{-4}$ (On reprend ici le paramètre $L_{RED} := \epsilon/N$ défini dans [67]), et avec (10.37), on a la formule :

$$L_{RED} := \frac{\epsilon}{N} \leq N\left(\frac{\alpha}{r^{eq}C}\right)^2 \approx 1.83 \cdot 10^{-4}. \quad (10.39)$$

Ce deuxième résultat est quasiment identique à celui de [67] mais avec une formule largement plus simple. Cependant on remarquera bien évidemment que le fait d'être proche d'un résultat existant n'est pas une fin en soi, même s'il est rassurant de ne pas tomber trop loin sans pouvoir expliquer le phénomène.

Remarque 10.4. Cette formule est extrêmement importante du point de vue de l'utilisation pratique, elle ne fait figurer que des variables connues d'un système réel, son débit et une estimation du RTT maximum moyen des connexions qui le traversent. La linéarité par rapport au nombre d'utilisateurs, qui est à priori une inconnue donne un caractère très intéressant du point de vue de contrôle puisque la variabilité est relativement faible. Ce qui est plus gênant du point de vue du contrôle est que le taux de pertes d'équilibre soit de l'ordre de $\left(\frac{rC}{\alpha N}\right)^{-2}$. Par contre ceci est un point très intéressant d'AIMD puisque de faibles pertes permettent de le contrôler.

Ceci donne un résultat extrêmement proche de [67], notons que cette étude fait intervenir le coefficient $K = 0.005$ pour moyenner l'estimé de q dans le calcul du coefficient de pertes κ . Nous n'avons pas mis cela dans le calcul, mais ceci pourrait être repris. Notons que cet exemple est pris pour une fenêtre moyenne inférieure à 12, 5, ce qui nous met dans le cadre de notre approximation linéaire, il est normal que le résultat soit proche de [67] dans la mesure où l'approximation est à peu près valable pour les très petites fenêtres.

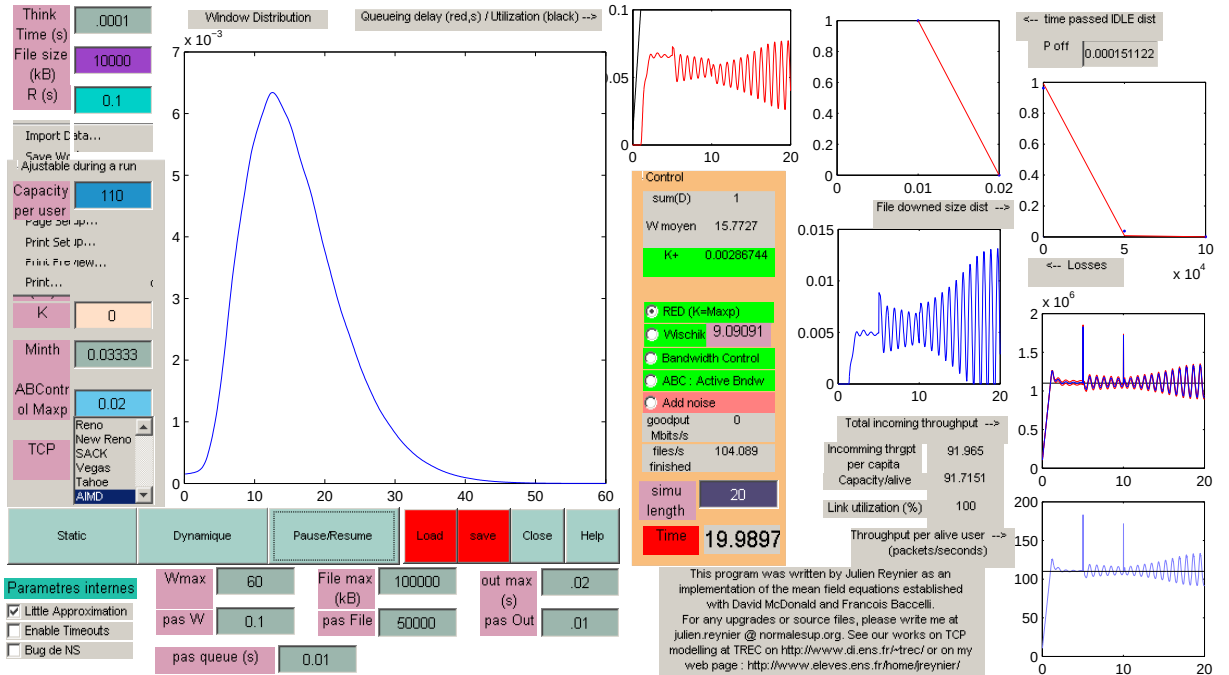


FIG. 10.4 – $Max_p = 1\%$, $Max_{th} = r$, $r = 0.1s$ et $Minth = Max_{th}/3$ jusqu'à $t = 5s$ ($\frac{\epsilon}{\kappa_{eq}} \approx 3$). Ensuite on passe à $Max_{th} = 1.5\%$ jusqu'à $t = 10s$ ($\frac{\epsilon}{\kappa_{eq}} \approx 4.5$), puis on prend $Max_{th} = 1.5\%$ jusqu'à la fin ($\frac{\epsilon}{\kappa_{eq}} \approx 6$). Le premier cas est stable pour RED, alors que le troisième diverge (sous l'approximation Little). Le deuxième cas est à la limite. De plus les périodes sont de l'ordre de $.9s$ ce qui est proche des $.95s$ avec $T = 2\pi r$ et $r_{moyen} = .15s$.

Étude simplifiée sous les perturbations Δw de fenêtres avec κ variable, r fixe avec Tail Drop en débit fixe Récrivons le système que nous allons étudier ce qui est une reprise de (10.22) et de (10.23) en rajoutant l'équation de tail drop sur la frontière (on note que le RTT est constant) :

$$\begin{aligned}\dot{M}_1 &= -\frac{1}{2} \frac{\kappa(t-r)}{r} M_2 + \frac{1}{r(t)}, \\ 0 &= \frac{1-\kappa(t)}{r} M_1 - L, \\ r &= T + \frac{q^{max}}{L}.\end{aligned}$$

Linéarisé autour de l'équilibre, ceci nous donne :

$$\begin{aligned}\Delta \dot{M}_1 &= -\frac{\kappa^{eq}}{r} M_1^{eq} \Delta M_1 - \frac{1}{2r} M_2^{eq} \Delta \kappa(t-r), \\ 0 &= (1-\kappa^{eq}) \Delta M_1 - M_1^{eq} \Delta \kappa(t).\end{aligned}$$

Ce qui nous donne en tenant compte du fait que $\frac{M_2^{eq}}{M_1^{eq}} = \frac{2}{\alpha\sqrt{\kappa^{eq}}}$ et que $M_1^{eq} = \frac{\alpha}{\sqrt{\kappa^{eq}}}$:

$$\begin{aligned}r \Delta \dot{M}_1 &= -\alpha\sqrt{\kappa^{eq}} \Delta M_1 - \frac{1-\kappa^{eq}}{\alpha\sqrt{\kappa^{eq}}} \Delta M_1(t-r), \\ (1-\kappa^{eq}) \Delta M_1 &= M_1^{eq} \Delta \kappa(t).\end{aligned}\tag{10.40}$$

Nous savons maintenant donner l'équation caractéristique de la première équation différentielle avec délai avec la variable complexe H :

$$rH = -\alpha\sqrt{\kappa^{eq}} - \frac{1-\kappa^{eq}}{\alpha\sqrt{\kappa^{eq}}} e^{-Hr}\tag{10.41}$$

qui se linéarise pour les petites valeurs de r :

$$\alpha\sqrt{\kappa^{eq}} rH = -\alpha^2 \kappa^{eq} - (1-\kappa^{eq}) + (1-\kappa^{eq}) Hr,\tag{10.42}$$

ce qui donne :

$$rH = \frac{(\alpha^2 - 1)\kappa^{eq} + 1}{1 - \kappa^{eq} - \alpha\sqrt{\kappa^{eq}}}.\tag{10.43}$$

On remarque qu'au point où le dénominateur s'annule, on change de signe. Il est difficile de dire ce qui se passe exactement en ce point, mais rH est grand et notre linéarisation n'est plus valable. Le point critique est $\kappa^c = \frac{-\alpha + \sqrt{\alpha^2 + 4}}{2} \approx 0.540$.

Théorème 10.6. *Dans les zones où la linéarisation a un sens, c'est-à-dire $\kappa^{eq} \gg \kappa^c$ ou $\kappa^{eq} \ll \kappa^c$, tail drop est instable pour les petites valeurs de κ^{eq} , et stable pour les très grandes. Cependant compte tenu des approximations de notre modèle et du niveau extrêmement élevé de κ^c , tail drop n'a pas d'équilibre stable dans la zone de fonctionnement de TCP.*

Note : si on voulait croire que TCP fonctionne selon notre équation même pour les petites fenêtres, un taux de pertes de κ^c correspondrait à une taille de fenêtre de $\alpha/\sqrt{\kappa^c} \approx 2.4$. C'est exactement ce que l'on observe dans le simulateur de l'équation aux dérivées partielles.

10.3.4 Un contrôle stable sur les débits

Leçon à tirer de l'étude de tail drop Reprenons à partir de l'équation caractéristique en fonction de la variable complexe H :

$$rH = -A - Be^{-Hr} \quad (10.44)$$

qui se linéarise pour les petites valeurs de r , ce qui donne :

$$rH = \frac{A+B}{B-1}. \quad (10.45)$$

Ainsi si le coefficient B devant de terme qui a du retard est plus grand que 1, on n'aura aucune chance d'avoir un équilibre stable. Il est donc inutile d'essayer de modifier TCP pour qu'il réagisse différemment, le problème fondamental est que la fonction de pertes réagit trop violemment, compte tenu du retard, ce n'est pas une bonne chose. Si on garde TCP tel qu'on le modélise, et que le RTT reste fixe, le coefficient A sera toujours positif; en effet une augmentation de la taille des fenêtres ne modifie pas leur manière d'augmenter, mais augmente instantanément la force de la division des fenêtres pour un taux de pertes identique.

Les petits buffers Le problème de TCP avec de petits buffers a aussi fait l'objet d'études. En effet quand le buffer est petit, il y a des pertes anticipées qui ne sont pas dues à la congestion mais aux fluctuations. Ce contrôle automatique produit un peu trop de pertes pour les grandes fenêtres, cependant il est stable dans une large mesure. On pourra se reporter à [130] pour une étude de la stabilité dans le cas des petits buffers. Ceci confirme bien que TCP n'est pas en cause dans les instabilités mais que le contrôle retardé est mal fait. Il nous montre en particulier qu'un contrôle sur les débits peut bien fonctionner. La seule difficulté à savoir bien gérer est de garder une bonne utilisation.

Contrôle stable Reprenons notre équation des moments

$$\dot{M}_1 = -\frac{1}{2} \frac{\kappa(t-r)}{r} M_2 + \frac{1}{r(t)}.$$

Linéarisé autour de l'équilibre, ceci nous donne :

$$\Delta \dot{M}_1 = -\frac{\kappa^{eq}}{r} M_1^{eq} \Delta M_1 - \frac{1}{2r} M_2^{eq} \Delta \kappa(t-r)$$

pour que notre système soit stable, il faut que $0 < \frac{1}{\kappa r} \frac{\Delta \kappa(t-r)}{\Delta M_1(t-r)} < 1$. Soit donc $0 < \epsilon < 1$. Nous souhaitons que le contrôle sur κ vérifie : $\Delta \kappa(t-r) = \epsilon \kappa r \Delta M_1(t-r)$.

On a donc le contrôle stable sur le débit : $\kappa'(M_1) = \epsilon \kappa r M_1'$. Ce qui nous donne le contrôle stable sur les débits :

$$\kappa(M_1) = \kappa(0) e^{\epsilon r M_1}.$$

En variables extensives tout cela nous donne :

$$\kappa(B_i) = \kappa(0) e^{\frac{\epsilon r}{N} B_i}$$

en notant $B_i = N M_1$ le débit entrant.

Ce qui nous permet d'affirmer le théorème suivant :

Théorème 10.7. *Le point fixe de TCP sous nos hypothèses simplificatrices est stable pour le contrôle sur le débit total entrant B_i :*

$$\kappa(B_i) = \kappa(0) e^{\beta \frac{B_i}{C}} \quad (10.46)$$

si

$$\beta < \frac{rC}{N}. \quad (10.47)$$

Pour N, r, B_i ainsi que les coefficients $\kappa(0)$ et β fixés, on a stabilisation pour le débit avec B_i est la solution de

$$X e^{\frac{\beta}{2C} X} = \frac{N\alpha}{\sqrt{\kappa(0)}}. \quad (10.48)$$

On note que $\frac{\kappa(C)}{\kappa(0)} = e^\beta$, il faut que β soit aussi grand que possible si on souhaite avoir de la latitude dans le contrôle.

Bien choisir les paramètres Notons donc que le débit est très peu élastique sous ce contrôle par rapport au nombre d'utilisateurs (qui varie) et à la valeur de contrôle $\kappa(0)$. Ceci est un aspect très intéressant, on peut donc avoir un contrôle robuste sur le débit total. Par la suite on peut lui rajouter un contrôle à échelle de temps longue dans le but de maximiser la bande passante. La formule à utiliser pour mettre à jour $\kappa(0)$ est :

$$\kappa^{new}(0) = \left(\frac{U}{U_{target}} \right)^2 \kappa(0) \quad (10.49)$$

où U représente l'utilisation du lien (cette formule est simplement le fait que $\sqrt{\kappa} M_1 = \alpha = cste$).

En ce qui concerne le choix des paramètres, comme on l'a dit $\kappa(0)$ est ajustable. De plus on peut mettre le contrôle de débit en complément de RED pour s'occuper des petites fenêtres. On ne cherchera pas à obtenir d'équilibre pour les fenêtres de trop petite taille. Aussi nous prendrons $W_{min} = 25$. On pourra toujours tenter ensuite d'utiliser les critères de (10.39) pour les petites fenêtres. Cependant il faut bien noter que RED est à peu près impossible à faire fonctionner dans la pratique (il suffit d'essayer un peu pour

comprendre : augmenter le délai modifie W_{max} . Il est donc très difficile de réaliser les bons réglages).

Sur un routeur de capacité totale $C = 100000$ paquets par seconde avec le temps de transmission $T_{min} = .01$ (2x1500 km à la vitesse de la lumière).

Il faut fixer β . Comme $N_{max} = C/W_{min}$

$$\beta = T_{min}W_{min} = 0.25 \quad (10.50)$$

et $e^\beta \approx 1.28$

Ensuite il faut partir sur une valeur de $\kappa(0)$ raisonnable, par exemple telle que si le débit entrant est C , on ait de l'ordre de $\left(\frac{\alpha}{W_{min}}\right)^2 \approx 0.27\%$ de pertes ce qui correspond à la valeur pour stabiliser W_{min} . Donc on choisit comme valeur initiale :

$$\kappa(0) \geq \left(\frac{\alpha}{W_{min}}\right)^2 e^{-T_{min}W_{min}} \approx .21\%. \quad (10.51)$$

10.3.5 Un contrôle stable sur les débits avec utilisation maximale du lien

Ajustement de la valeur de ϵ Nous allons garder le même contrôle et regarder ce qui se passe si on commence à remplir la file d'attente avec notre contrôle sur les débits. Dans ce cas écrit T et pas r pour le délai de transmission.

$$\begin{aligned} \dot{M}_1 &= -\frac{1}{2} \frac{\kappa(t-r(t))}{r(t-r(t))} M_2 + \frac{1}{r(t)}, \\ \dot{q} &= \frac{1-\kappa(t)}{r(t)} M_1 - L, \\ r(t) &= T + \frac{q(t)}{L}, \\ \kappa(M_1) &= \kappa(0) e^{\epsilon T M_1}. \end{aligned}$$

Linéarisé autour de l'équilibre, ceci nous donne :

$$\begin{aligned} \Delta \dot{M}_1 &= -\frac{\alpha \sqrt{\kappa^{eq}}}{r^{eq}} \Delta M_1 - \frac{1}{r^{eq} \kappa^{eq}} \Delta \kappa(t-r^{eq}), \\ \Delta \dot{q} &= \frac{1-\kappa^{eq}}{r^{eq}} \Delta M_1 - \frac{\alpha}{r^{eq} \sqrt{\kappa^{eq}}} \Delta \kappa(t) - \frac{1-\kappa^{eq}}{(r^{eq})^2} \frac{\alpha}{\sqrt{\kappa^{eq}}} \Delta r(t), \\ \Delta r(t) &= \frac{1}{L} \Delta q(t), \\ \Delta \kappa(t) &= \kappa^{eq} \epsilon T \Delta M_1 \end{aligned}$$

puis

$$\begin{aligned}\Delta \dot{M}_1 &= -\frac{\alpha\sqrt{\kappa^{eq}}}{r^{eq}}\Delta M_1 - \frac{\epsilon T}{r^{eq}}\Delta M_1(t - r^{eq}), \\ \Delta \dot{q} &= \frac{1 - \kappa^{eq}}{r^{eq}}\Delta M_1 - \frac{\epsilon T\alpha\sqrt{\kappa^{eq}}}{r^{eq}}\Delta M_1 - \frac{1 - \kappa^{eq}}{(r^{eq})^2}\frac{\alpha}{L\sqrt{\kappa^{eq}}}\Delta q.\end{aligned}$$

On voit ici que la première équation ne dépend pas de la seconde. Elle est identique à celle du paragraphe précédent, la condition de stabilité que nous avons retenue est :

$$\epsilon < \frac{r^{eq}}{T}.$$

De plus l'équation sur q est stable par rapport aux variations de q . D'où la stabilité.

La stabilité du contrôle avec une file d'attente présente deux intérêts essentiels :

1. comme nous l'avions indiqué au début, on peut avoir un contrôle stable avec une utilisation maximale en remplissant un peu la file d'attente,
2. on gagne un contrôle sur les petites fenêtres pour lesquelles le délai dans la file d'attente va représenter une augmentation de taille significative, ainsi elles reviendront dans la zone contrôlée.

Utilisation pratique Si le routeur a une file d'attente qui produit une attente cumulée avec le lien qui suit T_{full} , on peut prendre $\beta = T_{full}W_{min}$. Dans la pratique on pourra donc toujours prendre $\beta = 1$ avec $Q_{max} = .1s$ pour stabiliser les fenêtres en moyenne plus grande que $W_{min} = 10$. Ensuite on utilisera un algorithme adaptatif qui modifie $\kappa(0)$ sur des temps longs pour atteindre l'objectif de file remplie à $.01s$ par exemple dans le cas où les niveaux de pertes sont bas. Dans les cas pratiques, on se rend compte que cette valeur de $\beta = 1$ est largement plus stable qu'on le prévoit ici.

10.3. Contrôle et stabilité des points fixes de l'équation approximée

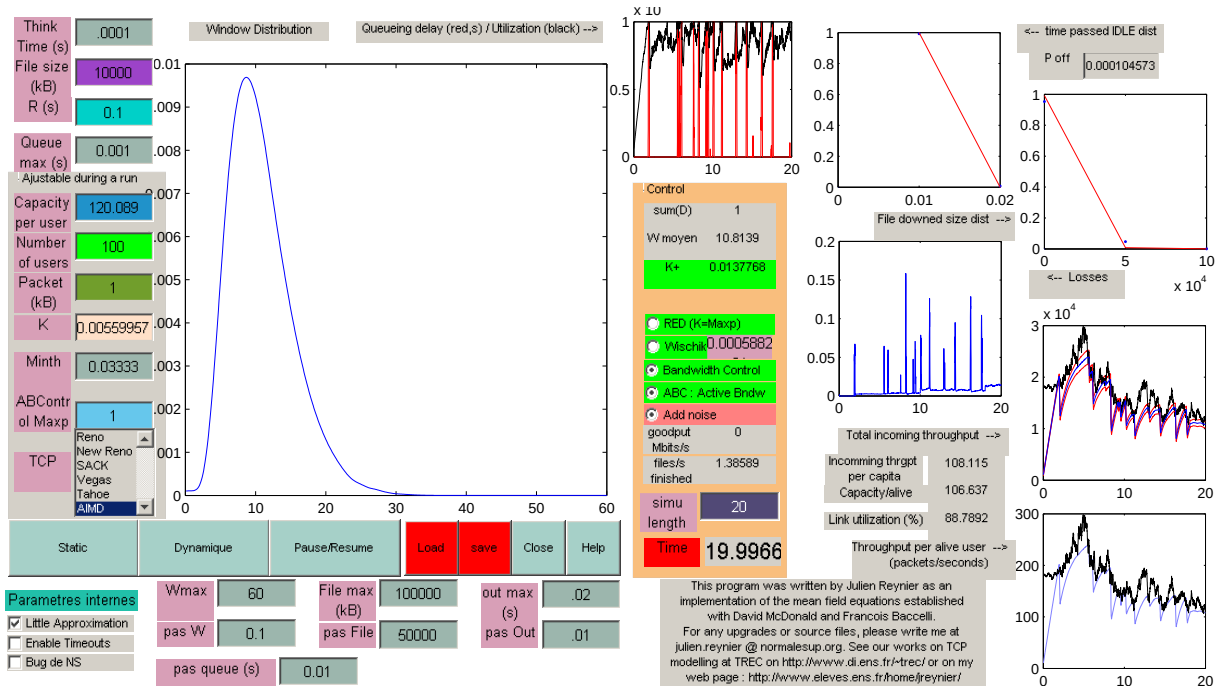


FIG. 10.5 – Illustration du fonctionnement de l'algorithme avec une file d'attente très petite et un débit variable.

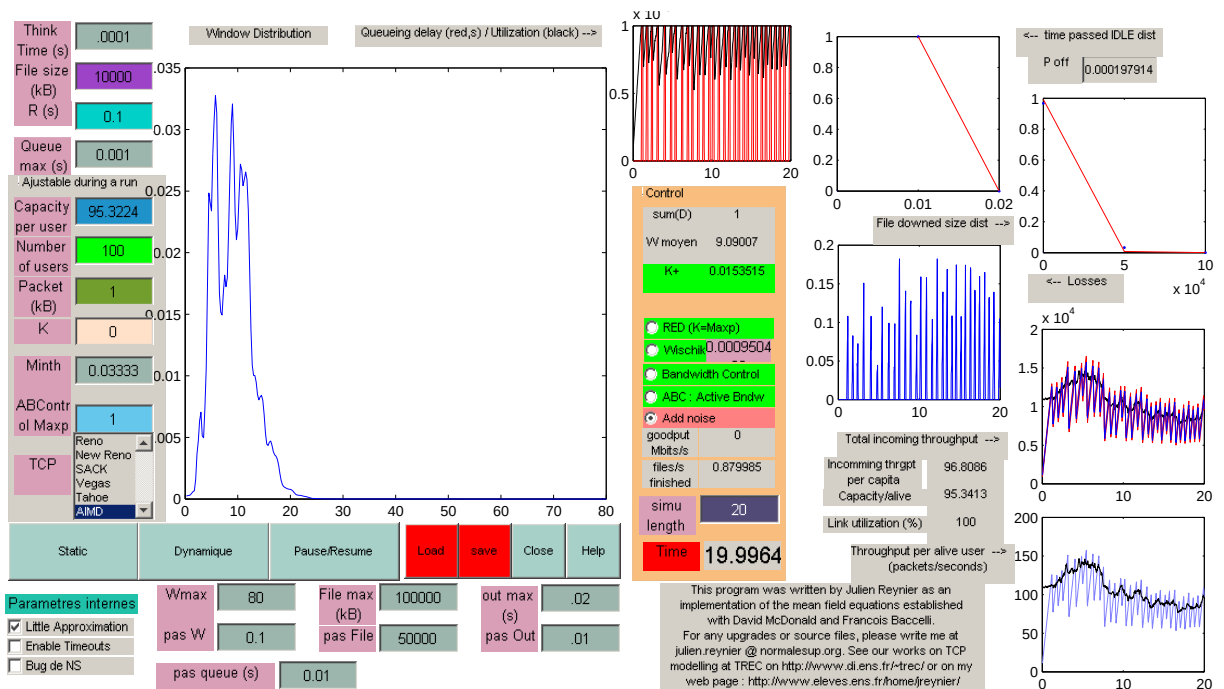


FIG. 10.6 – Ce qui se passerait sans activer le contrôle de débit.

Chapitre 11

Adaptation du champ moyen à HTTP

Sommaire

11.1 Modélisation des particules	189
11.1.1 Environnement	189
11.1.2 Équations d’une particule	190
11.1.3 Les équations du champ moyen	191
11.1.4 Les extensions possibles	191
11.2 Point fixe	191
11.2.1 Le point fixe exact pour les inactifs	192
11.2.2 Point fixe approximé pour les utilisateurs actifs	192
11.2.3 Formules autour du point fixe et normalisation	193
11.2.4 Formule de la racine carrée	195
11.3 Exemple de stabilité dans un cas simple	196
11.3.1 Hypothèse simplificatrice	196
11.3.2 Équation pour M_1	196
11.3.3 Équation de contrôle du débit	197
11.3.4 Petites perturbations	197

11.1 Modélisation des particules

11.1.1 Environnement

On a N utilisateurs de HTTP à travers un routeur congestionné qui permet l’accès à un serveur Web. Nous modélisons HTTP comme une suite de téléchargements de fichiers en utilisant TCP séparés par des temps de réflexion comme dans les articles [9, 17, 30, 88]. La modélisation fine de TCP et l’évolution de la taille de la file d’attente se fait comme dans le chapitre 8 avec les mêmes notations. Ainsi nous faisons l’hypothèse simplificatrice que TCP fonctionne dans le mode d’évitement de congestion (congestion avoidance).

Une des nouveautés est l’arrivée de nouveaux utilisateurs, ils commencent leur téléchargement à la fenêtre 0. D’autre part chaque utilisateur peut être dans l’état actif où

il se comporte comme précédemment en gardant en mémoire la quantité de données qui a été reçue (c'est à dire dont la destination a bien accusé réception) ; ou bien dans l'état d'attente où il ne se passe rien à part qu'on garde en mémoire la durée depuis laquelle on est dans cet état. Le passage d'un état à l'autre se fait ainsi : on se donne $\mathfrak{I}(i)$, la probabilité qu'une source à d'être en attente pendant la durée i . $\mathfrak{F}(f)$ est la probabilité de devoir envoyer fichier de taille f (l'unité de mesure est le paquet). Chaque utilisateur va enchaîner des téléchargements de fichiers de tailles f iid (indépendantes et identiquement distribuées) et des temps d'attente de durée i iid. On se donne un état initial où pour chacun des N utilisateur on connaît :

- son état E (actif ou en attente),
- la taille de fichier qu'il a envoyé et sa fenêtre (W, F) ou bien depuis combien de temps il attend (I) .

De plus on se donne une taille de file d'attente au routeur dont on étudie la congestion, ainsi qu'un historique sur les fenêtres et les RTT (pour pouvoir remonter un peu dans le temps).

En plus de la condition initiale, l'évolution des particules se fait comme au chapitre 8, avec en plus des sauts de changement d'état quand le fichier est fini, ce qui se passe avec la probabilité :

$$\mathbb{P}(F_n^N \text{ est la taille finale sachant que le fichier fait au moins } F_n^N).$$

Les sauts inverses pour revenir à l'état actif se font avec la même probabilité sur le temps d'attente et la distribution des temps d'attente.

11.1.2 Équations d'une particule

L'espace total a 4 états $(W, F, I, \{0, 1\})$, la taille de la fenêtre, le nombre de paquets transmis, le temps d'inactivité et l'état de transmission ou d'inactivité.

Comme d'habitude, on modélise TCP. Ici on rajoute la donnée de la taille du fichier. Donc dans l'espace (W, F) :

$$dW(t) = \frac{1}{r(t)}dt - \frac{W(t)}{2r(t) - r(t)}dN(t) \quad (11.1)$$

$$dF(t) = \frac{W(t)}{r(t)}(1 - \kappa(t))dt \quad (11.2)$$

où N est un processus de Poisson d'intensité $\kappa(t - r)$.

De plus on a un changement d'état puisque l'utilisateur devient inactif après avoir fini son téléchargement. Ceci se produit après la réception du paquet numéro F de la transmission avec la probabilité

$$P_F(f) := \mathbb{P}(F = f | F \geq f) = \frac{\mathfrak{F}(f)}{\int_f^\infty \mathfrak{F}(u)du}.$$

Pendant la période d'inactivité, le temps I progresse (par l'équation $dI(t) = dt$), et on reprend la transmission avec la probabilité :

$$P_I(i) = \mathbb{P}(I = i | I \geq i) = \frac{\mathfrak{I}(i)}{\int_i^\infty \mathfrak{I}(u)du}.$$

11.1.3 Les équations du champ moyen

Comme nos équations ont le degré de régularité suffisant, nous savons alors que si le nombre N de particules tend vers l'infini, alors on a un champ moyen dont les équations sont (après la simplification habituelle $e(t, w) = w$) :

Théorème 11.1. *Nommons $p(f, w, t)$, la projection de la distribution au temps t des utilisateurs actifs à la fenêtre w et qui ont transmis déjà f paquets ; $I(i, t)$, la projection de la distribution des utilisateurs inactifs depuis le temps $t - i$. Alors*

$$\frac{\partial p}{\partial t} = -\frac{1}{r(t)} \frac{\partial p}{\partial w} + \frac{\kappa(t-r)}{r(t-r)} (4wp(2w) - p(w)) \quad (11.3)$$

$$\begin{aligned} & -\frac{(1-\kappa(t))w}{r(t)} \left(\frac{\partial p}{\partial f} + p(w)P_F(f) \right) + B(t)\delta_{(f=0, w=0)} \\ \frac{\partial I}{\partial t} &= -\frac{\partial I}{\partial i} + D(t)\delta_{i=0} - I(i, t)1_{s-1}P_I(i) \end{aligned} \quad (11.4)$$

où

$$D(t) = \int \int_{f,w} p(f, w, t) \frac{(1-\kappa(t))w}{r(t)} P_F(f) df dw \quad (11.5)$$

$$B(t) = \int_i I(i, t) 1_{s-1} P_I(i) di. \quad (11.6)$$

1_{s-1} sert à conserver l'homogénéité et rappelle que tout doit être exprimé dans la même unité de temps.

11.1.4 Les extensions possibles

Comme d'habitude nous n'écrivons pas les équations de slow start et ne considérons pas la fenêtre de réception maximale. Cependant cela ne présente pas de difficulté mathématique particulière, on rajoute deux variables, *sthresh* et *rwnd* à l'espace et les évolutions sont suffisamment régulières pour ne pas poser de problèmes. Par contre du point de vue de la simulation, rajouter des variables n'est pas anodin. La méthode de simulation serait plutôt de supposer que *sthresh* et *rwnd* sont indépendants de la connexion et ont une distribution donnée.

Les équations (11.3,11.4,11.5,11.6) sont simulées dans [135].

11.2 Point fixe

On peut déduire des équations du champ moyen un point fixe pour des valeurs fixées r et κ . Dans le cas d'un point fixe on aura aussi égalité entre les naissances et les morts, ainsi $B(t) = D(t) = B = D$.

11.2.1 Le point fixe exact pour les inactifs

On regarde d'abord le point fixe pour 11.4 avec la constante D de morts en entrée. L'équation d'équilibre est :

$$\frac{dI}{di} = D\delta_{i=0} - I(i, t)P_I(i).$$

Dont la solution pour $i \geq 0$ est :

$$I(i) = D \int_i^\infty \mathfrak{T}(u) du \quad (11.7)$$

et I vaut 0 pour $i < 0$.

11.2.2 Point fixe approximé pour les utilisateurs actifs

L'équation est un peu difficile à résoudre symboliquement. Comme nous sommes surtout intéressés par la simulation, nous allons simplifier la résolution. Pour une résolution dans de nombreux cas particuliers, on pourra se reporter à [17]. En particulier on a une résolution complète dès que la taille des fichiers suit une loi exponentielle. De plus on y trouvera un calcul des moyennes en général. Ici nous donnons un raisonnement par approximation dans le cas général.

L'équation du point fixe est :

$$\begin{aligned} \frac{\partial p}{\partial w} + w \frac{\partial p}{\partial f} &= \kappa(4wp(f, 2w) - p(f, w)) \\ &\quad - wp(f, w)P_F(f) + rD\delta_{(f=0, w=0)}. \end{aligned} \quad (11.8)$$

D'abord notons que si on peut trouver une fonction positive $p_1(f, w)$ solution de l'équation sans le terme $wp(w)P_F(f)$, alors

$$p(f, w) = p_1(f, w) \int_f^\infty \mathfrak{F}(u) du$$

est une solution de (11.8).

Alors l'équation est particulièrement facile à résoudre sur la courbe $\mathcal{C} = (w, \frac{w^2}{2})$, la solution est la densité linéique $p(w, w^2/2) = rDe^{-\kappa w^2/2}\delta_{\mathcal{C}}$, avec $mes(A \cap \mathcal{C}) = Lebesgue(proj_w(A \cap \mathcal{C}))$.

Si nous pouvions nous débarrasser du terme d'arrivées $rD\delta_{(f=0, w=0)}$, il y aurait une solution très simple qui correspond à l'indépendance entre les fenêtres et les tailles de fichier. On aurait donc la solution p_2 indépendante de f . Ainsi il viendrait $p_2(f, w) = \mathcal{F}_\kappa(w)$ où $\mathcal{F}_\kappa$ est la solution du point fixe donnée dans 10.2.1 :

$$\mathcal{F}_\kappa(w) = \sum_{i=0}^{\infty} a_i e^{-\kappa 4^i \frac{w^2}{2}}.$$

Il est nécessaire de rappeler la valeur des deux premiers moments de cette distribution : $M_1 = \frac{\alpha}{\sqrt{\kappa}}$ avec $\alpha \approx 1.310$ et $M_2 = \frac{2}{\kappa}$.

Pour résoudre l'équation nous faisons l'hypothèse heuristique suivante : avant la première perte, une particule se trouve sur la courbe $\mathcal{C}$ (ce qui ne relève pas d'une approximation), et ensuite elle est distribuée selon $\mathcal{F}$ conditionné par le fait qu'au moins une perte s'est produite dans le RTT, c'est à dire selon $\mathcal{F}_{\kappa+\frac{1}{f}}$. Par contre nous ne pouvons rien dire de ces particules à part qu'elles ont téléchargé moins que $f = \frac{w^2}{2}$. Aussi notre approximation ne nous donne pas d'information sur la position en f . Cependant, une fois projeté sur un w fixé, nous avons une bonne approximation de la distribution marginale des fenêtres.

Aussi nous avons le théorème suivant :

Théorème 11.2. *Le point fixe heuristique de 11.8 a la marginale :*

$$\begin{aligned} p_w(w) &= \int_f p(f, w) df \\ &= C \int_f ((1 - e^{-\kappa f}) \mathcal{F}_{\kappa+\frac{1}{f}} + e^{-\kappa f} \delta_{\mathcal{C}}) \frac{\int_f^\infty \mathfrak{F}(u) du}{\sqrt{2f}} df. \end{aligned}$$

Avec la règle d'intégration linéique suivante

$$\int_{w \in \mathbb{R}} \int_{u \in [a, b]} \delta_{u^2/2, w} du dw = b - a.$$

Ici C est une constante de normalisation elle figure car notre approximation change le poids global de la distribution.

Notons que si la queue est plus lourde que $\text{puissance}(\frac{3}{2})$, on ne peut pas résoudre ainsi notre problème.

11.2.3 Formules autour du point fixe et normalisation

Nous allons noter le premier moment d'une variable X :

$$\overline{X} := \int_0^\infty x X(x) dx.$$

Si les queues des distributions ne sont pas trop lourdes, il y a une moyenne finie du temps de téléchargement, et du temps de réflexion.

La normalisation nous donne

$$\begin{aligned} p_{off} &= \int_i I(i) di \\ &= D \int_i \int_i^\infty \mathfrak{F}(u) du di \\ &= D \int_i i \mathfrak{F}(i) di = D \overline{\mathfrak{F}}. \end{aligned} \tag{11.9}$$

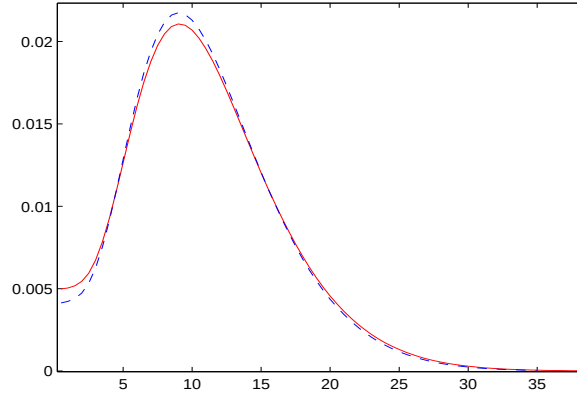


FIG. 11.1 – Point fixe de la densité des fenêtres, simulé (en tirets) et $\int p(u^2/2, w)du$ du théorème 11.2 (en trait plein). $\kappa = 1\%$, $F_{mean} = 1000$ paquets (selon une loi de puissance de coefficient 2).

Comme nous avons fait une approximation qui modifie la masse, nous devons nous garder d'appliquer la formule pour calculer p_{on} en général, cependant dans le cas particulier où $\kappa = 0$ (on a donc $C = rD$) la formule fonctionne et :

$$\begin{aligned}
 p_{on, \kappa=0} &= \int_w \int_u p(u^2/2, w) du dw \\
 &= C \int_u \int_{u^2/2}^{\infty} \mathfrak{F}(s) ds du \\
 &= C \int_u \sqrt{2u} \mathfrak{F}(u) du.
 \end{aligned} \tag{11.10}$$

La fenêtre moyenne peut être calculée pour tous les κ ainsi :

$$\begin{aligned}
 \overline{W} &= \frac{1}{p_{on, \kappa=0}} \int_w \int_u w p(u^2/2, w) du dw \\
 &= \int_u \left[\frac{\alpha(1 - e^{-\kappa u^2/2})}{\sqrt{\kappa + \frac{2}{u^2}}} + u e^{-\kappa u^2/2} \right] \frac{\int_{u^2/2}^{\infty} \mathfrak{F}(v) dv}{\int_0^{\infty} \sqrt{2f} \mathfrak{F}(f) df} du.
 \end{aligned} \tag{11.11}$$

Le débit sortant moyen du routeur par utilisateur est :

$$(1 - \kappa) \frac{\overline{W}}{r}.$$

Ainsi le temps moyen de téléchargement vaut :

$$\overline{Down} = \frac{r \overline{\mathfrak{F}}}{\overline{W}(1 - \kappa)}. \tag{11.12}$$

Théorème 11.3. *De (11.11) et (11.12) il vient*

$$p_{on} = \frac{\overline{Down}}{\mathfrak{T} + \overline{Down}}. \quad (11.13)$$

11.2.4 Formule de la racine carrée

Nous savons maintenant que pour de faibles valeurs de κ , il n'y aura aucune chance d'avoir une formule de la racine carrée (ie : un débit en $\frac{1}{\sqrt{\kappa}}$). En effet on a une borne supérieure exacte sur les débits (elle n'utilise pas d'approximation pour son calcul).

Théorème 11.4. *De manière générale, quand les distributions des tailles de fichiers et de temps de réflexion sont finies, le débit maximal moyen par utilisateur en vie est :*

$$Max_{through} = \frac{\int f \mathfrak{F}(f) df}{r \int \sqrt{2f} \mathfrak{F}(f) df}.$$

Preuve : Le débit maximal se fait pour $\kappa = 0$. Dans ce cas nous connaissons la solution exacte du problème de point fixe :

$$p(f, w) = r D \delta_c \frac{\int_f^\infty \mathfrak{F}(u) du}{\sqrt{2f}}.$$

De plus p_{on} est connu par (11.10). Alors l'équation $p_{on} + p_{off} = 1$ permet de trouver le temps moyen de téléchargement d'un fichier et de conclure.

■

On ne peut pas pour autant rejeter totalement la formule de la racine carrée, en effet pour des taux de pertes grands relativement à la taille des fichiers à télécharger, on est dans un cas proche de téléchargements persistants pour lesquels cette formule s'applique. Par exemple s'il y a une taille minimale de fichier f_{min} telle que $\sqrt{f_{min}} e^{-\kappa f_{min}} \ll 1$, alors (11.11) devient :

$$\overline{W} \approx \frac{\alpha}{\sqrt{\kappa}}. \quad (11.14)$$

Ceci ne saurait être considéré comme une preuve dans la mesure où ce résultat vient directement de l'approximation heuristique que nous avons faite pour résoudre le point fixe.

Par un principe de superposition, le cas général se situera toujours un peu entre les deux cas précédents. Cependant la formule de la racine carrée est aussi une borne supérieure sur les débits, par le raisonnement intuitif qui suit : la différence entre le cas TCP persistant et HTTP est qu'on a, dans le cas de HTTP, des entrées à la fenêtre 0, alors que dans le cas permanent, les entrées équivalentes se font selon la distribution stationnaire. Le cas permanent est donc strictement meilleur que le cas HTTP pour même nombre d'utilisateurs actifs.

11.3 Exemple de stabilité dans un cas simple

Nous allons montrer comment utiliser nos équations dans le cas simple d'un contrôle sur les débits avec un délai fixé. Ceci signifie que $r(t)$ sera une constante que nous noterons r . L'étude menée ici est très semblable à celle de 10.3, mais nous savons maintenant que le cas du contrôle sur le débit est particulièrement intéressant en pratique.

11.3.1 Hypothèse simplificatrice

Commençons par faire quelques simplifications. D'abord nous allons prendre les temps de réflexion exponentiellement distribués. Ceci a le grand mérite d'enlever un terme de délai qui pollue un peu l'étude du système. Ainsi les équations 11.4 et 11.6 sur $I(t)$, la proportion d'utilisateurs en attente deviennent

$$\frac{dI}{dt}(t) = -\frac{1}{T}I(t) + D(t)$$

et

$$B(t) = \frac{I(t)}{T},$$

où T est le temps moyen d'attente.

Tout comme dans 10.3 nous allons aussi restreindre notre étude à des perturbations du type Δw . Ainsi pour $M_i = \int \int w^i p dw df$, ceci nous permet d'écrire $\Delta M_i = 2M_{i-1}\Delta w$ et ainsi

$$\Delta M_i = 2M_{i-1}\Delta M_1. \quad (11.15)$$

Ceci n'est pas vrai pour une perturbation générale, mais l'étude de toutes les équations sur les moments est difficile. Ceci nous donnera des équations de stabilité que nous devons considérer comme heuristiques. Mais nous verrons que cette formule se vérifie bien expérimentalement.

11.3.2 Équation pour M_1

Soit P_i , le i^{me} moment de pP_F :

$$P_i := \overline{w^i p P_f} = \int \int w^i p(f, w) dw P_F(f) df$$

En multipliant (11.3) par w , en prenant les intégrales par rapport à f et w , puis en faisant une intégration par parties (p tend vers 0 quand $|w| \rightarrow \infty$), nous trouvons :

$$r\dot{M}_1 = (1 - \kappa)M_0 - \frac{\kappa(t - r)}{2}M_2 - (1 - \kappa)P_2 \quad (11.16)$$

$$\dot{I} = -\frac{1}{T}I(t) + \frac{1 - \kappa}{r}P_1 \quad (11.17)$$

avec

$$M_0 + I = 1.$$

Notons que le terme $\int \int w \frac{\partial p}{\partial f} df dw = \int w(p(w, 0) - 0) dw$ disparaît parce que les utilisateurs naissent à la fenêtre 0. Pour la même raison, le terme $\int \int w B(t) \delta_{(f=0, w=0)} dw df$ part.

En fait les nouveaux utilisateurs vont modifier la taille moyenne de la fenêtre en changeant le nombre de connexion actives.

Faire les mêmes calculs, mais en multipliant par w^i au lieu de w donnerait l'équation différentielle exacte vérifiée par M_i . Malheureusement ceci ne nous a pas permis d'aboutir à une preuve de stabilité des équilibres. C'est pour cela que nous avons choisi l'approche de perturbations d'un type bien particulier pour trouver des conditions simple.

11.3.3 Équation de contrôle du débit

Le contrôle effectué sur le débit tue certains paquets au niveau du routeur en fonction de la bande passante qui lui arrive. Comme r est constant, le débit entrant vaut exactement $\frac{M_1}{r}$ (rappelons que tout est normalisé par N , le nombre total d'utilisateurs). Dans la mesure où nous nous intéressons à une stabilité locale, nous allons prendre des pertes linéaires du type :

$$\kappa(t) = Cst + \frac{\epsilon M_1}{r}. \quad (11.18)$$

La valeur de la constante n'est pas utile dans la mesure où on connaît la valeur κ^{eq} de κ à l'équilibre.

11.3.4 Petites perturbations

Sous de petites perturbations, $\Delta w = \Delta M_1$:

$$r \Delta \dot{M}_1 = -\Delta I - \frac{M_2^{eq}}{2} \frac{\epsilon}{r} \Delta M_1(t-r) - (\kappa^{eq} M_1^{eq} + 2P_1) \Delta M_1 \quad (11.19)$$

$$\Delta \dot{I} = -(P_1 \frac{\epsilon}{r} + (1 - \kappa^{eq}) P_0) \Delta M_1 - \frac{r}{T} \Delta I \quad (11.20)$$

car :

$$\begin{aligned} - \Delta B(t) &= \frac{1}{T} \Delta I(t) \\ - \Delta \kappa(t) &= \frac{\epsilon}{r} \Delta M_1(t). \end{aligned}$$

Soit $\mathcal{D} = 1 - \frac{\epsilon M_2^{eq}}{2r}$, alors la matrice caractéristique de ce système avec délai linéarisé pour de petites valeurs de r est :

$$\begin{pmatrix} -\frac{1}{r\mathcal{D}} \left(\frac{M_2^{eq}}{2} \frac{\epsilon}{r} + \kappa^{eq} M_1^{eq} + 2P_1 \right) & -\frac{1}{r\mathcal{D}} \\ -\frac{\epsilon P_1}{r^2} + (1 - \kappa^{eq}) \frac{P_0}{r} & -\frac{1}{T} \end{pmatrix}.$$

Nous utilisons ensuite le fait qu'un polynôme réel de degré 2, $X^2 + a_1 X + a_0$ a deux racines de partie réelle strictement négative si et seulement si a_0 et a_1 sont strictement positifs. Ainsi une condition suffisante de stabilité (ie : condition pour avoir des valeurs

propres négatives) sera :

$$\frac{1}{\mathcal{D}} \left(\frac{M_2^{eq}}{2} \frac{\epsilon}{r} + \kappa^{eq} M_1^{eq} + 2P_1 \right) + \frac{r}{T} > 0 \quad (11.21)$$

$$\frac{1}{\mathcal{D}} \left[\left(\frac{M_2^{eq}}{2} \frac{\epsilon}{r} + \kappa^{eq} M_1^{eq} + 2P_1 \right) \frac{r}{T} - \frac{\epsilon}{r} P_1 + (1 - \kappa^{eq}) P_0 \right] > 0. \quad (11.22)$$

Supposons que $\mathcal{D} = 1 - \frac{\epsilon M_2^{eq}}{2r} > 0$. ie :

$$\epsilon < \frac{2r}{M_2^{eq}}. \quad (11.23)$$

Alors 11.21 est vérifiée, et 11.22 devient

$$\left(\kappa^{eq} M_1^{eq} + 2P_1 \right) \frac{r}{T} + (1 - \kappa^{eq}) P_0 > \epsilon \left(\frac{1}{r} P_1 - \frac{1}{T} \frac{M_2^{eq}}{2} \right). \quad (11.24)$$

Le pire des cas est atteint pour de grandes valeurs de T , ce qui est précisément le cas usuel (le temps de lecture d'une page sera de l'ordre de plusieurs secondes alors que le RTT n'en vaut pas plus d'une dans les cas habituels). On peut donc récrire 11.24 avec l'hypothèse que $T \gg r$:

$$(1 - \kappa^{eq}) P_0 > \epsilon \frac{1}{r} P_1. \quad (11.25)$$

La condition de stabilité devient alors :

Théorème 11.5. *Le contrôle sur les débits est stable si*

$$\epsilon < r \min \left[(1 - \kappa^{eq}) \frac{P_0}{P_1}, \frac{2}{M_2^{eq}} \right]. \quad (11.26)$$

Notons que $\frac{P_1}{P_0}$ est la taille moyenne de la fenêtre au moment où le dernier paquet de la transmission du fichier traverse la file d'attente.

Troisième partie

Simulations

Cette partie est organisée en deux chapitres. Le chapitre 12 présente quelques point qui ont mené à la conception du simulateur des équations aux dérivées partielles ainsi que quelques validations. Nous discutons ensuite des améliorations à apporter au modèle et au simulateur dont certaines ont été implémentées.

Le chapitre 13 vient clore cette thèse en présentant quelques exemples d'utilisations du simulateur. Nous y illustrons l'effet de la congestion sur les utilisateurs et un effet de turbulence pour HTTP.

Chapitre 12

Le simulateur

Ce chapitre est dédié aux problèmes liés à la conception d'un simulateur d'équations aux dérivées partielles, et aux validations expérimentales et aux améliorations possibles. Nous commençons en discutant de divers problèmes numériques et d'implémentation, puis nous comparons le comportement de notre simulateur face à OPNET un simulateur à événements discrets du marché et nous vérifions la convergence du système à N utilisateurs vers notre solution numérique. Enfin nous envisageons diverses améliorations, dont une importante sur la modélisation des débits.

Sommaire

12.1 Conception du simulateur	204
12.1.1 Analyse numérique	204
12.1.2 Dans tous les cas : problème des retards et des pertes . . .	206
12.1.3 Pour HTTP : tailles de fichiers et temps de sommeil	208
12.1.4 Difficultés techniques	213
12.2 Validations du simulateur pour des sources persistantes .	213
12.2.1 Comparaison entre le simulateur et OPNET	213
12.2.2 Analyse de la qualité de service	218
12.2.3 Comparaison entre le simulateur et des réalisations du sys- tème à N particules	220
12.3 Amélioration du modèle de débit	222
12.3.1 Diagnostic d'un problème dans la prise en compte des débits	222
12.3.2 Changement du modèle de lien et de file d'attente	224
12.3.3 Implémentation du modèle de débits réel	226
12.4 Intégration des Timeouts	228

Remarque 12.1. *Dans la suite la taille de la file d'attente aura la plupart du temps pour unité la seconde à moins qu'une mention contraire ne le précise.*

12.1 Conception du simulateur

12.1.1 Analyse numérique

La difficulté de l'analyse En partant de (10.4) :

$$D_t M_c(t, dw) = -\frac{1}{R_c(t)} D_w M_c(t, dw) + \frac{1}{R_c(t - R_c(t))} K(t - R_c(t)) \cdot (e(t, t - R_c(t), 2w) M_c(t, d(2w)) - e(t, t - R_c(t), w) M_c(t, dw)). \quad (12.1)$$

Ni $M_c(t, dw)$, ni $M_c(s - R_c(s), dv; s, dw)$ ne sont des états du système. Cependant l'équation fournit suffisamment d'informations pour faire évoluer le système. Soit $\mu_c(t)$ le processus $\{M_c(s, dw); t-1 \leq s \leq t\}$ (si par exemple tous les RTT valent moins que 1). En utilisant (8.13) on peut déduire de $M_c(t, dw)$ au temps t sa valeur au temps $t + \delta t$ dans la mesure où $M_c(t - s + \delta t, dw)$ est obtenu par un décalage de temps. Malheureusement μ_c n'est pas un état utilisable numériquement. Même en discrétisant et en ne gardant que les trajectoires du processus sur une partition que donne $\{M_c(s_i, dw); t-1 = s_0 < s_1 \dots s_n = t\}$ c'est toujours beaucoup par rapport à ce qu'on peut traiter sur un ordinateur.

Une résolution précise On peut éviter ce problème en définissant une séquence de temps t_k^c pour chaque classe de sorte que $t_{k+n+1}^c - R_c(t_{k+n+1}^c) = t_k^c$. En prenant n suffisamment grand, on a une partition qui convient. Si nous définissons $\Phi^c(t) =$ le premier k tel que $t_k^c > t$, nous allons construire notre solution par récurrence du temps t_i au temps t_{i+1} en définissant $t_{i+1} = \min_c(t_{\Phi^c(t_i)}^c)$ partant du temps $t_0 = 0$. Supposons que pour chaque classe nous soyons capables de calculer et de sauver le vecteur $V_c^M(t)$, une version discrète de $M_c(t_k^c)$ pour $k = m - n, \dots, m$ où $m = \Phi^c(t) - 1$ (Ce sont des marginales et pas la distribution jointe en entier). Supposons aussi que nous gardions en mémoire le vecteur des noyaux $V_c^T(t)$ donnés par $T^c(t_k^c)$ pour $k = m - n, \dots, m$ où $M_c(t_k^c) = M_c(t_{k-1}^c) \circ T^c(t_k^c)$ et m défini comme ci-dessus. Supposons finalement que nous sauvions les noyaux $S^c(m) = \prod_{k=m-n}^m T^c(t_k^c)$.

Nous pouvons alors faire évoluer le système au temps t_{i+1} . À chaque pas de temps nous faisons évoluer la file d'attente et la classe c qui vérifie $t_{i+1} = t_{\Phi^c(t_i)}^c$. Le noyau inverse $(S^c(m))^{-1}$ donne la distribution conditionnelle de la fenêtre de la classe c un RTT avant t_i étant donné la fenêtre au temps t_i . Calculons l'espérance conditionnelle $e(t_i, t_i - R_c(t_i), w)$. Avec cela nous pouvons utiliser (8.14) pour calculer $T^c(t_{m+1}^c)$. On se débarrasse de $T^c(t_{m-n}^c)$. On met à jour $S^c(m+1) = (T^c(t_{m-n}^c))^{-1} S^c(m) T^c(t_{m+1}^c)$. Finalement on calcule $Q(t_{i+1})$ en utilisant les $M_c(\Phi^c(t_i) - 1)$ pour $c = 1, \dots, d$.

Résolution numérique de l'équation simplifiée et voies d'amélioration Grâce à la simplification (10.6) on obtient une équation facile à simuler avec la distribution des fenêtres comme état. On peut faire mieux en améliorant l'évaluation de la fonction e . Du point de vue numérique nous avons constaté que l'approximation brutale ne pose en fait pas trop de problèmes et qu'on obtient une solution peu différente de celle plus raffinée qui vient en utilisant ce qui est décrit ici.

On étudie l'évolution d'une fenêtre $W(t)$ d'une connexion canonique dont la distribution est $\mu(t; dw)$. En première approximation, il y a deux manières d'arriver à la fenêtre w au temps t . Cela se produit si $W(t - r(t)) = w - 1$ et qu'il n'y a pas de pertes pendant le RTT précédent. Si les paquets sont répartis uniformément dans le RTT, la probabilité de ne pas avoir de perte est approximativement $H(t - r(t), w - 1)$ où

$$H(t, w) = \prod_{j=1}^w \left(1 - k(t - \frac{r(t)}{w}j) \right).$$

Le deuxième cas est que $W(t - r(t)) = 2w$ et qu'on a une perte et pas plus pendant le dernier RTT. De même si on suppose l'uniformité de répartition des $2w$ paquets de la fenêtre, cette probabilité vaut environ : $1 - H(t - r(t), 2w)$. Dans ce second cas nous allons négliger le cas où les pertes terminent en timeout. Nous allons aussi négliger la possibilité de deux sauts à partir de la fenêtre $4w, 8w, \dots$

Nous avons

$$\langle v, \mu(s - r(s), dv; s, dw) \rangle \quad (12.2)$$

$$\begin{aligned} &= (w - 1)P(W(t) = w | W(t - r(t)) = w - 1)\mu(t - r(t), dw - 1) \\ &\quad + (2w)P(W(t) = w | W(t - r(t)) = 2w)\mu(t - r(t), d(2w)) \\ &= (w - 1)H(t - r(t), w - 1)\mu(t - r(t), dw - 1) \\ &\quad + (2w)(1 - H(t - r(t), 2w))\mu(t - r(t), d(2w)) \quad (12.3) \\ &= (w - 1)H(t - r(t), w - 1)p(t - r(t), w - 1)dw \\ &\quad + (2w)(1 - H(t - r(t), 2w))p(t - r(t), 2w)2dw \\ &= \left[(w - 1)H(t - r(t), w - 1)\frac{p(t - r(t), w - 1)}{p(t, w)} \right. \\ &\quad \left. + (2w)(1 - H(t - r(t), 2w))\frac{2p(t - r(t), 2w)}{p(t, w)} \right] p(t, w)dw \\ &= e(t; w)p(t, w)dw \end{aligned}$$

où $e(t, w)$ vaut

$$(w - 1)H(t - r(t), w - 1)\frac{p(t - r(t), w - 1)}{p(t, w)} + (2w)(1 - H(t - r(t), 2w))\frac{2p(t - r(t), 2w)}{p(t, w)}.$$

Ainsi par le même argument que dans (10.2) nous avons notre amélioration :

$$\frac{\partial p}{\partial t}(t, w) = \left(p(t, 2w)\frac{e(t; 2w)}{r(t - r(t))}2 - p(t, w)\frac{e(t; w)}{r(t - r(t))} \right) k(t - r(t)) - \frac{1}{r(t)}\frac{\partial p}{\partial w}(t, w), \quad (12.4)$$

où $r(t) = T + q(t - r(t))/L$, $q(t)$ satisfait (10.8).

De (12.3) nous avons aussi

$$\begin{aligned}
 & \langle vw, \mu(s - r(s), dv; s, dw) \rangle \\
 &= \langle w, (w - 1)H(t - r(t), w - 1)\mu(t - r(t), dw - 1) \rangle \\
 & \quad + \langle w, (2w)(1 - H(t - r(t), 2w))\mu(t - r(t), d(2w)) \rangle \\
 &= \langle (w + 1), wH(t - r(t), w)\mu(t - r(t), dw) \rangle + \frac{1}{2} \langle w, w(1 - H(t - r(t), w))\mu(t - r(t), d(w)) \rangle \\
 &= \langle (w + 1)wH(t - r(t), w) + \frac{1}{2}w^2(1 - H(t - r(t), w)), \mu(t - r(t), dw) \rangle.
 \end{aligned}$$

Pour le point fixe, (12.4) devient :

$$k \left(f(2w) \frac{e(2w)}{r} 2 - f(w) \frac{e(w)}{r} \right) = \begin{cases} \frac{1}{r} d(f(w - 1)(1 - k)^{w-1})/dw & \text{si } w \geq 4, \\ \frac{1}{r} df(w)/dw & \text{si } w < 4 \end{cases} \quad (12.5)$$

où

$$e(w) = (w - 1)(1 - k)^{w-1} \frac{f(w - 1)}{f(w)} + (2w)(1 - (1 - k)^{2w}) \frac{2f(2w)}{f(w)}.$$

12.1.2 Dans tous les cas : problème des retards et des pertes

Il s'agit de fabriquer une simulation d'une équation aux dérivées partielles. Nous avons une équation de transport ce qui rend les choses relativement simple, car toutes les quantités sont positives et on a le contrôle sur la somme. Les difficultés résident essentiellement dans le calcul des taux de pertes et des retards. Nous donnons informellement quelques indications sur ces difficultés.

Lorsque le RTT est fixe de valeur r Nous avons coupé les valeurs à un certain seuil et appliqué un pas constant ΔW en espace pour les tailles de fenêtres. Ensuite remarquant qu'une évolution du type

$$r \frac{\partial y}{\partial t} + \frac{\partial y}{\partial x} = 0$$

correspond à un déplacement de masse vers la droite à vitesse constante $\frac{1}{r}$, on lie le pas de temps au pas d'espace lorsqu'aucune perte n'est constatée en prenant $\Delta t := r\Delta W$, ceci évite le phénomène de diffusion indésirable qui intervient quand une masse déplacée tombe entre plusieurs cases de discrétisation. Aussi la simulation que nous faisons ne fait aucune erreur dans le déplacement à part le dépassement de la taille maximale. Ceci peut d'ailleurs être contrôlé en regardant la masse qui s'échappe sur le bord droit de l'intervalle.

L'autre partie de la dynamique consiste en des sauts d'une partie de la population. Le saut en lui-même divise l'erreur avant le saut par deux. Ce qui est plus délicat à contrôler est la proportion d'utilisateurs touchés dans l'intervalle $[W, W + \Delta W]$. Cette proportion vaut $K(t - r)W\Delta t$ où $K(t)$ est le taux de pertes au temps t au niveau de la file d'attente. Ceci pose un sérieux problème dans la mesure où cela imposerait que $W\Delta t \ll 1$ pour tous les W que nous serons amenés à considérer. Les pas de calculs seraient alors trop

petits pour la puissance des ordinateurs actuels. Pour nous sortir de cela nous avons pris la proportion partiellement linéarisée :

$$b(t) := 1 - e^{-K(t-r)W\frac{1}{r}\Delta t}.$$

En effet si $k = K(t-r)$ représente la probabilité de perte d'un paquet est fixée pendant la durée Δt (les individus de la population entre W et $W + \Delta W$ ayant le débit W/r) la proportion de connexions qui ne voient aucune perte est $(1 - k)^{\frac{W}{r}\Delta t}$. Cette quantité linéarisée en $k \ll 1$ nous donne bien $1 - b(t)$.

L'erreur sur le calcul de $b(t)$ donne :

$$\begin{aligned} \epsilon_b(t) &= 1 - e^{-(K(t-r)+\epsilon_K)(W+\Delta W)\frac{1}{r}\Delta t} - b(t) \\ &= e^{-K(t-r)W\frac{1}{r}\Delta t} - e^{-(K(t-r)+\epsilon_K)(W+\Delta W)\frac{1}{r}\Delta t} \\ &= e^{-K(t-r)W\frac{1}{r}\Delta t} (1 - e^{-K(t-r)\Delta W\frac{1}{r}\Delta t - \epsilon_K W\frac{1}{r}\Delta t - \epsilon_K \Delta W\frac{1}{r}\Delta t}) \end{aligned}$$

Si W est tel que $K(t-r)W\frac{1}{r}\Delta t \gg 1$, l'erreur est très petite, car elle est la différence de deux quantités très petites, et l'erreur relative est du coup de même ordre puisque $b(t)$ vaut quasiment 1.

Si W est tel que $K(t-r)W\frac{1}{r}\Delta t \ll 1$, l'erreur est petite car on pouvait linéariser au départ $b(t) \sim K(t-r)W\frac{1}{r}\Delta t$ et sommer les erreurs, l'erreur relative vaut alors $\frac{\epsilon_K}{K(t-r)} + \frac{\Delta W}{W}$.

Le cas difficile se produit quand $K(t-r)W\frac{1}{r}\Delta t \sim 1$, on ne peut alors linéariser l'exponentielle qu'aux conditions $\epsilon_K \ll K(t-r)$ et $\Delta W \ll W$. La deuxième condition ne pose pas trop de problème car pour avoir $K(t-r)W\frac{1}{r}\Delta t \sim 1$, il fallait nécessairement que $W\frac{1}{r}\Delta t \gg 1$, comme $r\Delta W$ est de l'ordre de Δt , le problème ne se pose jamais. Par contre la seconde condition $\epsilon_K \ll K(t-r)$ est délicate et intervient aussi dans la valeur relative. Ceci signifie que le simulateur sera assez peu précis aux points où K varie en restant proche de zéro pour les fenêtres extrêmement grandes de taille $W \sim r\frac{1}{K\Delta t}$.

Dans les cas que nous étudierons cette imprécision ne posera pas trop de difficultés. Pour le cas «drop tail», les pertes sont assez fortes ou nulles, et il sera difficile d'atteindre de grandes fenêtres. Pour le cas de RED, on garde un taux de pertes constant. On prendra soin ne pas considérer des cas où les tailles de fenêtres peuvent devenir trop grandes.

Lorsque le retard est variable : résolution directe Le retard $R(t)$ est calculé par la formule implicite $R(t) = T + \frac{Q(t-R(t))}{L}$. Ceci ne pose pas de problème de définition car Q est positive et a une croissance bornée. Ainsi comme nous l'avons déjà vu au milieu d'autres choses dans le chapitre 8, t_0 étant fixé, la fonction qui à x associe $x - (T + \frac{Q(t_0-x)}{L})$ est strictement croissante tend vers ∞ et est négative au moins jusqu'à T . La figure (12.1) illustre le calcul à réaliser dans le cas difficile où $x \mapsto Q(t_0 - x)$ croît au voisinage du point de rencontre. Le calcul du RTT sera donc d'autant plus précis que la file d'attente évite de *décroître* rapidement. Le problème apparaît si le débit peut s'approcher de 0 à un instant donné où la file n'est pas vide.

Or le débit entrant est de la forme :

$$\frac{\bar{W}}{R(t-R)}(1 - K(t-R)),$$

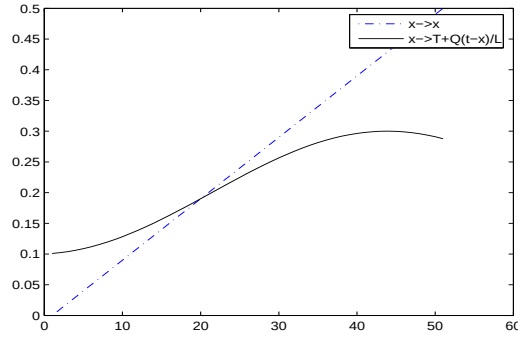


FIG. 12.1 – Calcul du RTT

aussi le problème des petits débits ne se pose que dans le cas de très forts taux de pertes ou bien si les fenêtres sont extrêmement petites. Excluons ce dernier cas et intéressons-nous au problème des forts taux de pertes. Comme nous l'avons déjà dit, le taux de pertes est majoré par le maximum de p_{max} et de $\frac{\bar{D}-L}{\bar{D}}$ où $\bar{D}$ est le débit moyen. Il est très facile de vérifier pendant les calculs que cette valeur ne dépasse pas un seuil fixé au cours de la simulation.

Lorsque le retard est variable : une astuce de calcul En fait il est possible d'écrire le RTT sans faire appel à une équation implicite. La valeur du RTT a en effet été choisie telle que si un paquet passe à la file d'attente au temps $u = t - R(t)$, le signal après renvois successifs qui lui correspond arrivera à la file d'attente au temps $t = u + \left(T + \frac{Q(u)}{L}\right)$. Ceci signifie que l'équation sur le RTT peut aussi être écrite ainsi :

$$R\left(u + T + \frac{Q(u)}{L}\right) = T + \frac{Q(u)}{L}. \quad (12.6)$$

L'avantage de cette approche est que l'on ne commet pas d'erreur d'approximation, de plus on voit que lorsque Q décroît rapidement, on attribue des valeurs assez différentes à R pour des temps très proches les uns des autres ; ceci signifie que pour bien simuler l'équation du champ moyen (qui là n'est pas nécessairement proche de la réalité), il faut baisser les pas de temps dans ces zones.

12.1.3 Pour HTTP : tailles de fichiers et temps de sommeil

Pour implémenter HTTP, une difficulté essentielle est de trouver de bonnes données à mettre en entrée du simulateur. Le rajout d'une variable supplémentaire n'a par contre pas été embêtant, le facteur principal de simplification étant que les évolutions des fenêtres et les tailles de fichiers varient pour les grands fichiers sur des échelles de temps très différentes. Aussi du point de vue des résultats obtenus nous avons constaté que le choix des données est très important.

Nous discutons ici du problème des données qui a mené à une autre difficulté que nous n'avons pas résolu pour l'instant : le simulateur est difficile sous sa forme à faire

fonctionner avec des distributions de tailles de fichiers à queues trop lourdes.

Deux sources d'informations ont été utilisées, des données opérateur, et des logs de sites web.

Données opérateur Ce ne sont pas les données les plus faciles à interpréter dans la mesure où elles sont agrégées. Cependant on peut se faire une première idée de la forme du trafic. Voici donc une première approche des sessions HTTP :

- Flots HTTP définis par (TCP,@source,@dest,portsource,portdest) : moyenne du temps off 20 s, taille du fichier 30 paq soit 34 ko ;
- Flots "usagers" HTTP : 360 paquets, 420 ko (la taille d'une page et des images associées par exemple) ;
- Sessions usagers : temps off de 100 s, échange de 1140 paquets soit 1305 ko en 40 flots simples soit 3 flots "usagers".

(source : statistiques sur des concentrateurs BAS (après les DSLAM) d'accès ADSL en France en juillet 2005 fournies par France Telecom).

On voit donc que la taille des fichiers transportés est très petite d'une part, et d'autre part que la notion de flot doit être précisée.

Études de logs Les logs²⁸ sont très intéressants car ils indiquent quels fichiers ont été téléchargés pendant une période donnée et combien de fois. C'est exactement le type d'information que nous pouvons utiliser. Notons par contre que si on nous donnait la distribution des tailles de fichiers présents sur les disques d'un serveur Web, nous aurions du mal à exploiter les résultats. Enfin, notons que si l'on s'intéresse à améliorer la situation d'un client précis, c'est bien en utilisant ses propres logs que l'on va procéder.

Les tailles de fichiers sont connues pour suivre assez souvent des lois de puissance, aussi on les représente en général en échelle log-log²⁹. On a eu de l'ordre de 3400 requêtes par heure durant la période considérée, et la taille moyenne d'un fichier téléchargé a été de 35 paquets de 576 octets. Si on prend un point (x, y) du graphe dans la figure (12.2), x est le logarithme à base 10 d'un nombre de paquets et y est le logarithme des effectifs correspondants. Les points vers le bas ne sont pas très significatifs tels qu'ils sont présentés dans la mesure où les effectifs représentés en dessous de 10 sont trop faibles. Malheureusement, les points qui correspondent aux fichiers de moins de 10 paquets ne sont pas particulièrement alignés. On ne peut donc rien conclure de cette courbe et surtout pas que les fichiers seraient distribués selon une loi de puissance (en loglog, le moindre écart signifie une très grosse erreur). En regardant de nouveau après retraitement les données dans le même échelle, on obtient la figure (12.3), où on voit qu'entre 100 et 1000 paquets il y a une sorte d'alignement. Cependant si on fait la courbe d'importance de ces points, on voit qu'on s'est trompé comme on peut le constater sur la figure (12.4). Les petits fichiers ont effectivement une faible importance, par contre les fichiers entre 3ko et 10Mo ont une importance comparable quand on les groupe en tailles logarithmiques, ceci n'est

²⁸Le mot log signifie enregistrement en anglais.

²⁹Cette affirmation est très controversée, nous allons voir que les mesures dont nous avons pu disposer indiqueraient plutôt une distribution de type log-normale.

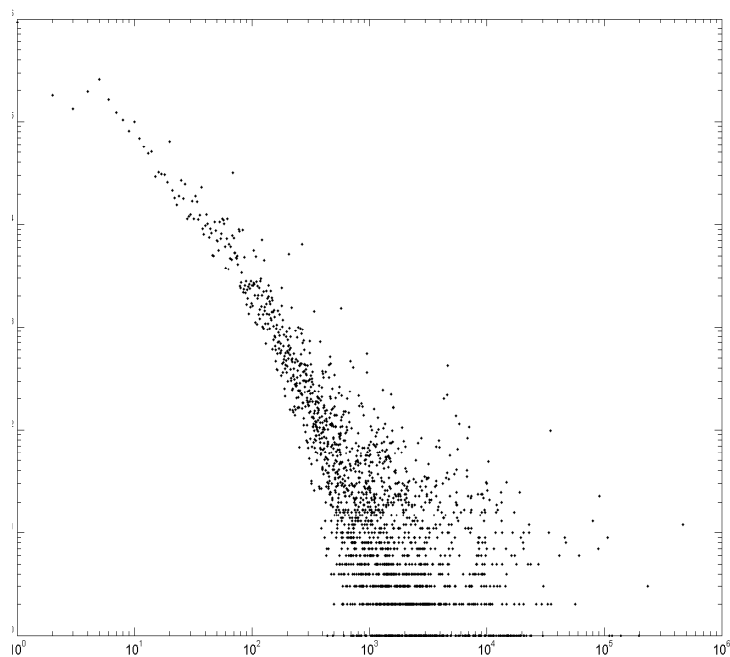


FIG. 12.2 – Taille des fichiers téléchargés sur le serveur WEB de l'ENS du 27 juin au 31 juillet 2005 en échelle log-log (HTTP 1.1 uniquement). En x, la taille de fichier arrondie au paquet, en y la fréquence des fichiers de cette taille

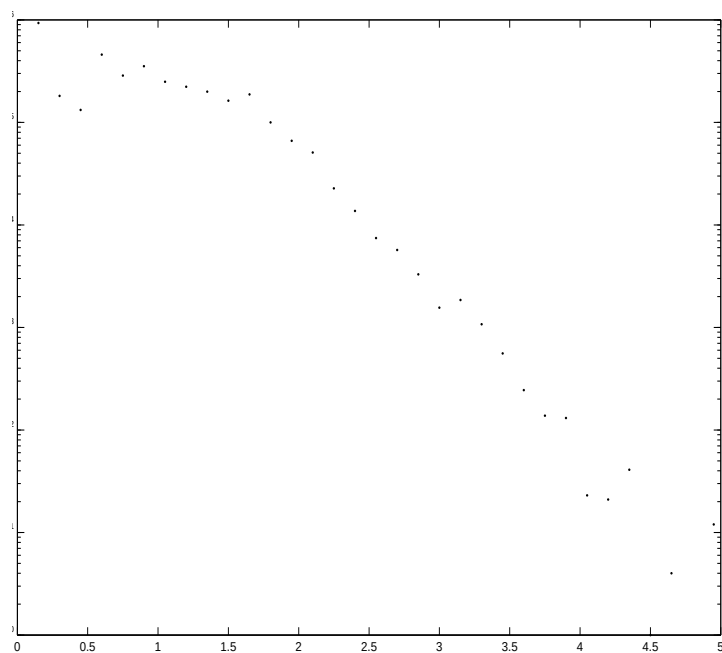


FIG. 12.3 – Trace des tailles de fichiers retraitée en regroupant les fichiers par taille $x \in [1.5^i, 1.5^{i+1}]$ paquets en échelle log-log

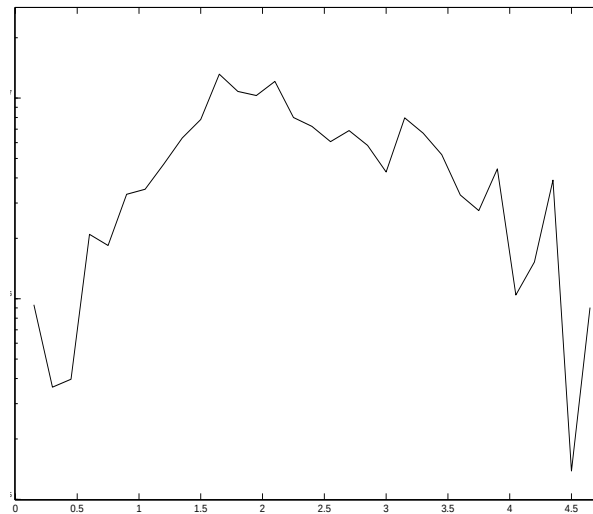


FIG. 12.4 – Quantité de données téléchargée en nombre de paquets correspondant à chaque taille de fichier en échelle log-log à base 10 pendant la période du 27 juin au 31 juillet 2005 sur le serveur WEB de l'ENS. Les masses de téléchargement sont ensuite regroupées par $[1.5^i, 1.5^{i+1}]$. Pour comparaison une loi de puissance stricte serait strictement croissante pour une valeur du paramètres $1 < \alpha < 2$ constante pour $\alpha = 2$ et strictement décroissante pour $\alpha > 2$

pas vraiment compatible avec une loi de puissance. Le fait de croire à un alignement dans la courbe log-log s'explique bien si on va lire [40].

Données de la littérature Un bon point de départ est le début de l'excellente thèse d'Alexandre Proutière [126]. Essentiellement on voit que la notion de session et de temps entre session est quelque chose qui n'est pas évident à définir. Pourtant c'est ce qui fait le lien entre le nombre d'utilisateurs supposés et le nombre de personnes en train de télécharger à un instant donné. Un utilisateur télécharge une page de HTML puis les quelques images contenues dedans à la volée. C'est ce qui rend le trafic compliqué. Les arrivées ne sont alors pas Poisson, on procède par rafales, comme [123] l'a mis en évidence.

Des résultats assez précis existent sur la distribution statistique des tailles de fichier. Il est généralement admis que les tailles de fichier ont une forte variance, on parle de souris et d'éléphants (il y a les pages d'une part et des gros fichiers d'autre part). En fait l'énorme majorité des fichiers transférés fait une faible part du débit alors qu'une faible partie de fichiers très gros crée la partie majoritaire du débit.

On voit d'après la littérature qu'il y a deux types de fichiers, les petits, qu'on télécharge très souvent, et des gros qui plus sont rares, mais qui représentent une masse de trafic très supérieure aux premiers.

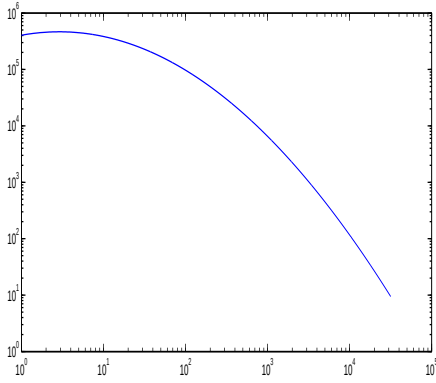


FIG. 12.5 – Distribution lognormale de tailles de fichiers où on regroupe les fichiers par taille (correspond à $f\mathcal{D}(f)$)

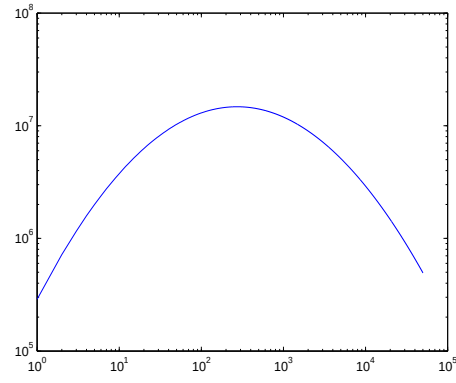


FIG. 12.6 – Quantité de données téléchargées regroupée par groupe de tailles (correspond à $f^2\mathcal{D}(f)$)

Données que nous devrions retenir En fait les données que nous avons pu observer en provenance de logs de serveurs WEB sont en désaccord avec la littérature (les souris et les éléphants) puisque toutes les tailles standards de fichiers sont à peu près équitablement représentées du point de vue de la consommation de débit. Par contre la répartition est très étalée comme on s’y attendait. Dans le cas de la figure (12.4), on notera que par comparaison, une loi de puissance devrait donner quelque chose d’au moins monotone, avec le cas critique pour le paramètre de puissance 2. Aussi nous retiendrons une loi lognormale pour les tailles de fichiers. Pour ce qui nous concerne, comme nous ne regardons pas des utilisateurs individuels on se rend rapidement compte que les résultats ne dépendent que de la moyenne du temps d’attente dans la mesure où comme nous le supposons jusqu’ici sans l’avoir écrit ce temps est indépendant de la taille du fichier. La distribution précise ne modifie pas l’état stationnaire s’il est stable (ce qui est une trivialité) et a assez peu d’impact sur les états transitoires (on le voit en essayant différentes distributions). Pour rendre les simulations le moins stable possible, on peut prendre un temps d’attente déterministe, le cas le plus stable est l’exponentielle.

On prendra une moyenne aux alentours de 35 paquets par fichier, avec la loi lognormale de paramètres $\sigma = 2$ et $f_0 = 5^{30}$ dont la densité est donnée par :

$$\mathcal{D}(f) = \frac{1}{\sigma f \sqrt{2\pi}} \exp\left(-\frac{1}{2\sigma^2} \log\left(\frac{f}{f_0}\right)^2\right).$$

Ceci nous donne une moyenne de taille qui vaut $e^{\log(f_0) + \sigma^2/2} \approx 37$. De plus la distribution de $f\mathcal{D}(f)$ et de $f^2\mathcal{D}(f)$ correspond à ce que nous attendions aux figures (12.3) et (12.4), comme on peut le voir dans les figures (12.5) et (12.6).

Enfin le temps d’attente sera pris avec une moyenne de 20 seconde. On a vu que ce temps résulte d’une douzaine de téléchargements de petits fichiers suivis d’une attente

³⁰Ces chiffres ont été choisis pour ajuster à peu près les courbes en utilisant des valeurs simple. Compte tenu de l’imprécision des mesures, un ajustement précis sur la courbe lognormale n’aurait pas beaucoup de sens.

de 100 s. Aussi on pourra prendre une distribution moyenne pondérée avec 10 fois une exponentielle de moyenne 1,5 secondes et une fois une exponentielle de moyenne 100 secondes.

Les simulations que nous verrons dans la suite n'ont malheureusement pas pu bénéficier de l'adaptation précise sur les tailles de fichiers. En effet, une distribution à queue un peu trop lourde rend les simulations sous leur forme actuelle peu fiables et les résultats ne seraient pas présentables. Aussi, nous retiendrons une distribution exponentielle des tailles de fichiers tout en précisant bien que c'est un point qui devait être amélioré, le lecteur qui souhaite jouer avec le simulateur [135] pourra mettre tout type de distribution en entrée dans un tableau excel, il est pour l'instant déconseillé de mettre des distributions trop étalées, mais on pourra choisir des distributions avec des pics par exemple.

12.1.4 Difficultés techniques

Le choix du langage de programmation De nombreux choix étaient possibles. Nous avons finalement choisi Matlab qui est à notre sens la plateforme de développement la plus rapide disponible aujourd'hui pour résoudre numériquement des problèmes mathématiques. Ce langage présente plusieurs avantages comme :

- avoir une communauté de développement très étendue et un support développé et entretenu,
- être un outil tout compris qui contient dans ses bibliothèques standard la plupart des algorithmes qui peuvent aider les mathématiciens,
- permettre le développement en standard de fenêtres et d'afficher facilement tout type de graphique.

Interface graphique La conception d'une bonne interface graphique utilisateur (GUI) est une des choses les plus importantes dans la réalisation d'un programme. Dans le cas de cette simulation, le choix a été fait de donner une représentation de l'évolution de la distribution des fenêtres et de la taille de la file d'attente. Les paramètres peuvent être modifiés au vol ce qui évite l'utilisation de scripts de simulation qu'il faut modifier à la main avant la simulation. Ce GUI nous a permis de bien mieux comprendre les phénomènes mis en jeu et d'acquérir une bonne intuition des problèmes rencontrés. De plus le côté ludique de l'utilisation permet de passer plus de temps à travailler dessus dans des conditions agréables.

On peut voir le GUI du programme dans les figures (12.7) et (12.8).

12.2 Validations du simulateur pour des sources persistantes

12.2.1 Comparaison entre le simulateur et OPNET

Mise en garde La comparaison de notre simulateur avec OPNET n'a qu'un intérêt assez relatif. En effet il s'agit de deux modèles différents de la réalité qu'est Internet. Nul ne sait vraiment si OPNET ou son confrère NS sont de très bons modèles pour

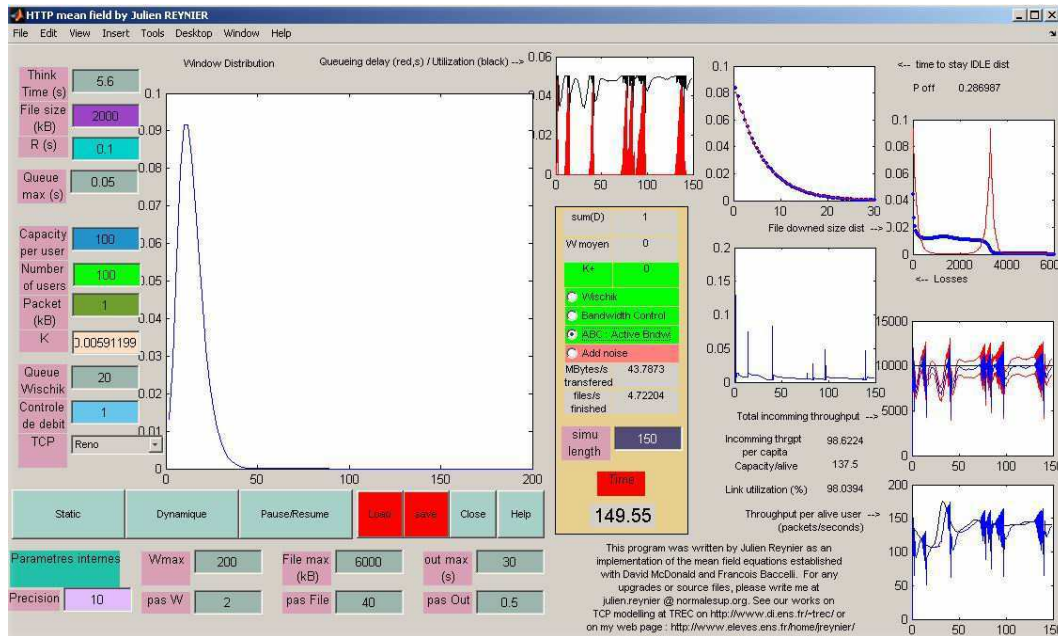


FIG. 12.7 – Interface graphique

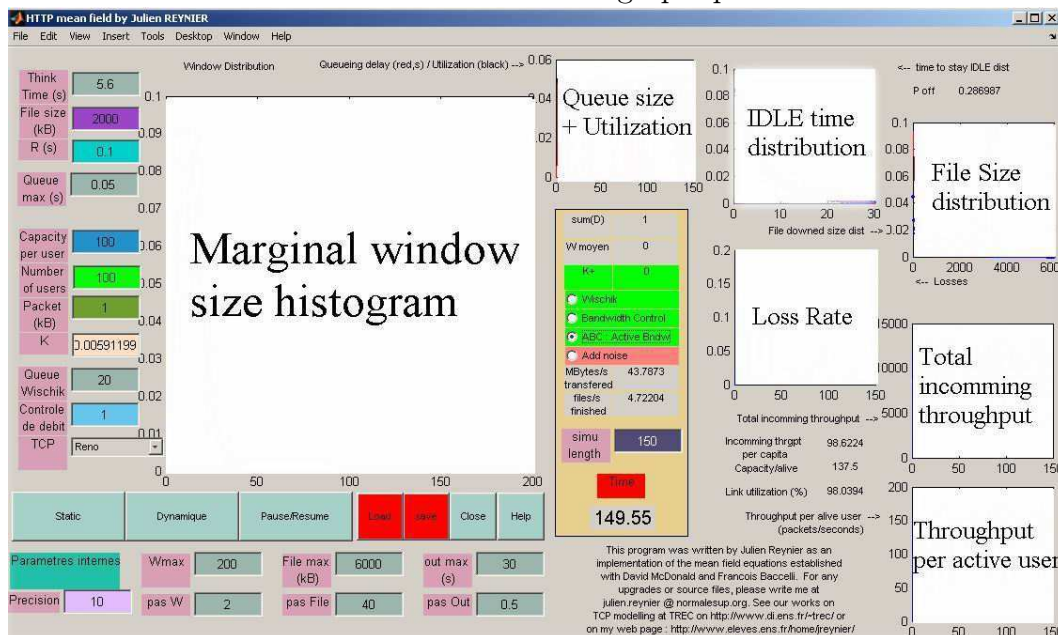


FIG. 12.8 – Description des fonctions

ce qui nous intéresse (une bonne partie de la communauté des réseau est même plutôt convaincue du contraire). Aussi il aurait été plus probant de comparer à de véritables traces obtenues en étudiant un véritable réseau d'ordinateurs qui téléchargent des fichiers par FTP. Malheureusement nous n'avons pas pu nous fournir de telles traces.

Comme nous l'avions indiqué au début du chapitre nous reprenons des résultats d'articles, c'est pourquoi nous pensons que cette section mérite de figurer, car ils illustrent notre démarche d'ingénieur dans le développement d'un nouvel outil ; nous nous assurons d'abord de la consistance avec ce qui existe.

Pour plus de rigueur cependant, la section 12.2.3 compare l'équation simulée de la particule du champ moyen avec des réalisations des systèmes originaux.

Contenu de la section Tout d'abord, dans la mesure où on étudie un modèle en limite des grands nombres, on peut se demander combien d'utilisateurs constituent un grand nombre. Dans la pratique 25 est un bon seuil, même si dès 5 ou 6 utilisateurs on observe des choses intéressantes.

Le système (12.4) et (10.8) se résout donc numériquement. Nous découvrons ainsi de nombreux comportements possibles qui dépendent du choix de la fonction F , de L et de T . Pour simplifier nous allons prendre $F(q) = 0$ si $q < q_{min}$ et $F(q) = 1$ si $q \geq q_{max}$. Ceci couvre le cas usuel de drop tail dans le cas où $q_{max} = B$. Nous regardons d'abord le point fixe dont il est question à la section 10.2.1. Ensuite nous discutons du cas où il y a des cycles limites et finalement des cas avec plusieurs points fixes.

Corrections à apporter dans le cas de stabilisation Nous allons prendre un exemple avec $N' = 200$ et une capacité de 44.736 Mégabits par seconde. Nous supposons que chaque paquet a la taille minimale autorisée par le transport IP de 536 octets. Ceci fait un débit de 10433 paquets par seconde. Nous supposons que nous n'avons pas de timeout pour commencer, ainsi $N = N'$. Donc avec $N = 200$, $L = 20.866$ paquets par utilisateur et par seconde. Le délai de transmission vaut $T = .1$ secondes. $Q_{min} = 0$, $p_{min} = 0$, $Q_{max} = 1000$, $p_{max} = .05$. Par (10.17), le taux de pertes au point fixe vaut $k = 2.6\%$ avec une taille de file d'attente $q = 520.9$ et le RTT 0.15 secondes. Pour les mêmes paramètres, la solution numérique approximée de (10.7) et (10.8) et de (12.4) et (10.8) se stabilisent toutes deux à la même taille de file d'attente, même RTT, et même taux de pertes. Ceci n'est pas surprenant dans la mesure où le taux de pertes reste assez faible.

Une simulation du logiciel OPNET (v6) [74] avec 200 sources donne une taille de file d'attente moyenne de 452, un taux de pertes moyen de 2.5% et un RTT moyen de 0.145 environ. Les résultats sont les mêmes avec 400 puis 800 utilisateurs à normalisation près (voir les figures (12.11)).

Le taux de pertes plus petit pour OPNET et la file d'attente plus petite proviennent des timeouts qui se produisent quand la taille des fenêtres est petite. Si nous tenons compte des timeouts avec $200 - N = TO(k)N$ et (10.9), on trouve $N = 185$ (approximativement) et 15 connexions en attente de timeout ou en slow-start. Le taux de pertes à l'équilibre est alors $k = 0.0240$ ce qui donne une file d'attente de 483, $r = .146$ et $TO = 0.0806$. En réitérant le procédé, il vient $N = 184$, $k = 0.0238$, $q = 481$, $r = .146$ et $TO = 0.0901$. Ceci donne quasiment un point fixe puisque la proportion de timeouts correspond à celle

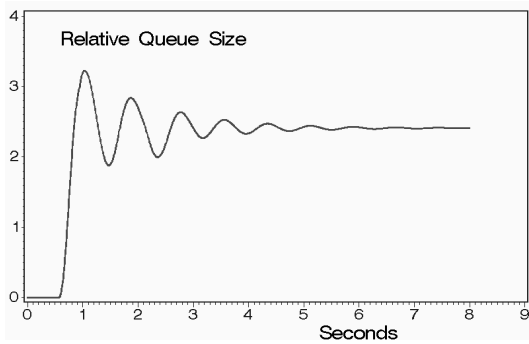


FIG. 12.9 – Solution numérique du système approximé avec $N=184$, normalisé par $N'=200$

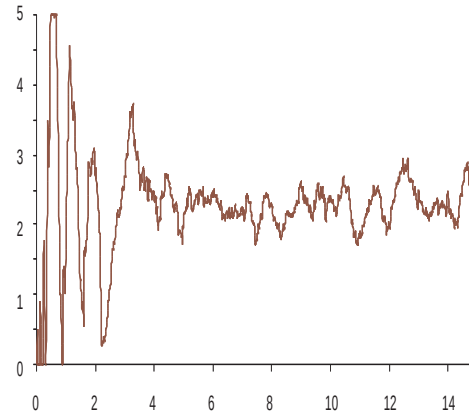


FIG. 12.10 – Simulation Opnet avec $N'=200$: la taille de la file d'attente est divisée par N'

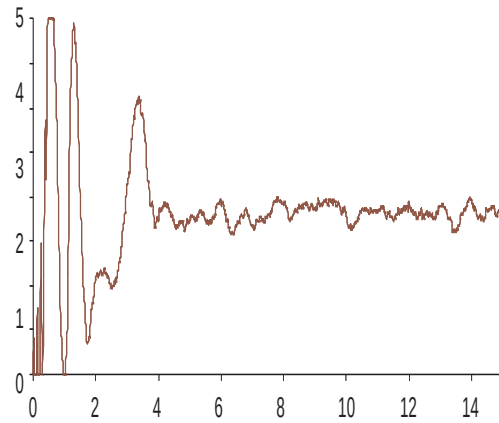
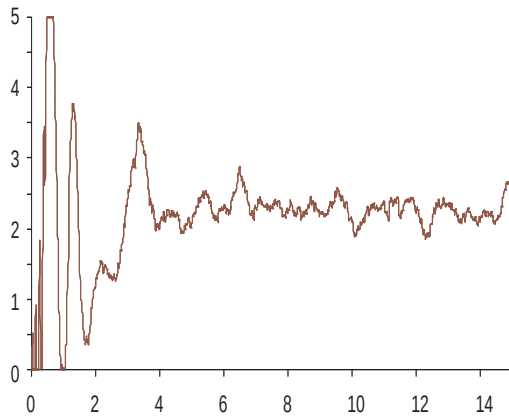


FIG. 12.11 – Simulation Opnet avec $N=400$ (à gauche) $N=800$ (à droite) : la file d'attente est divisée par N

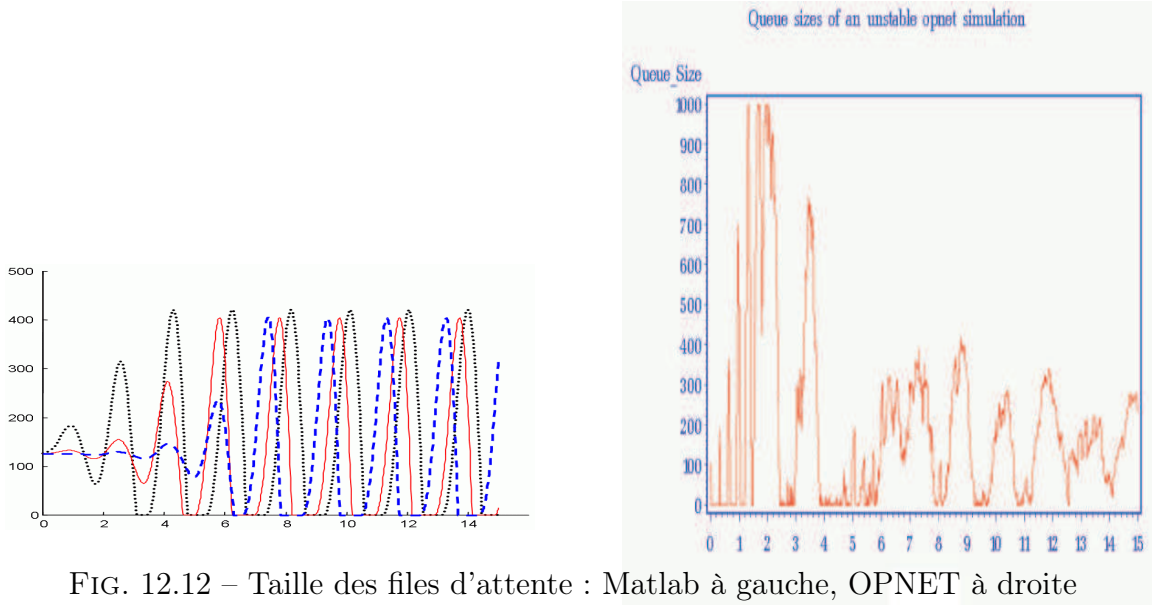


FIG. 12.12 – Taille des files d’attente : Matlab à gauche, OPNET à droite

que nous attendons. Ainsi nous voyons que notre solution devient très proche de celle d’OPNET. D’autre part un œil exercé remarquera qu’au début OPNET remplit tout de suite la file d’attente, ceci est dû au fait que les sources sont totalement synchronisées. Aussi les premiers paquets arrivent exactement en même temps à la file d’attente. Ceci ne se produit pas en réalité, et notre simulateur est meilleur de ce point de vue.

Remarque 12.2. *Nous notons que toutes les simulations OPNET sont semblables à part des oscillations autour de la moyenne qui semble être expérimentalement de l’ordre de $1/\sqrt{N}$. En tout cas nous avons la confirmation expérimentale d’une convergence quand le nombre d’utilisateurs devient grand.*

Périodes : Confrontons la période que nous observons dans le simulateur, environ $0.8s$, avec la période prévue par la formule du $2\pi r = 0.95s$. Nous voyons que l’approximation est consistante mais pas extrêmement précise.

Un cas d’équilibre instable Ensuite nous utilisons les mêmes paramètres mais avec un délai de transmission de $T = .3$ secondes.

Dans la figure de gauche de (12.12), l’approximation (10.7) et (10.8) est indiquée en pointillés, l’approximation (12.4) et (10.8) est indiquée en ligne pleine. L’approximation par les équations de [67] est figurée par la ligne brisée. Toutes ces courbes semblent osciller autour du point fixe mais avec des oscillations qui vont en s’amplifiant jusqu’à ce que la queue touche 0. La simulation OPNET avec 200 sources est à droite de la figure (12.12), elle est aussi instable (mais nous n’avons pas réussi à faire partir OPNET d’une distribution d’équilibre à $t = 0$). Nous n’avons pas corrigé la proportion des timeouts que nous ne savons pas calculer dans ce cas. Il faut noter que l’amplitude des périodes d’OPNET est bien approximée par les trois niveaux d’approximation. Nous reviendrons plus tard sur la précision des simulations faites par OPNET, mais elles ne sont pas si fiables qu’on voudrait le croire.

- Max p = 5 %
- Qmax=1000 paquets

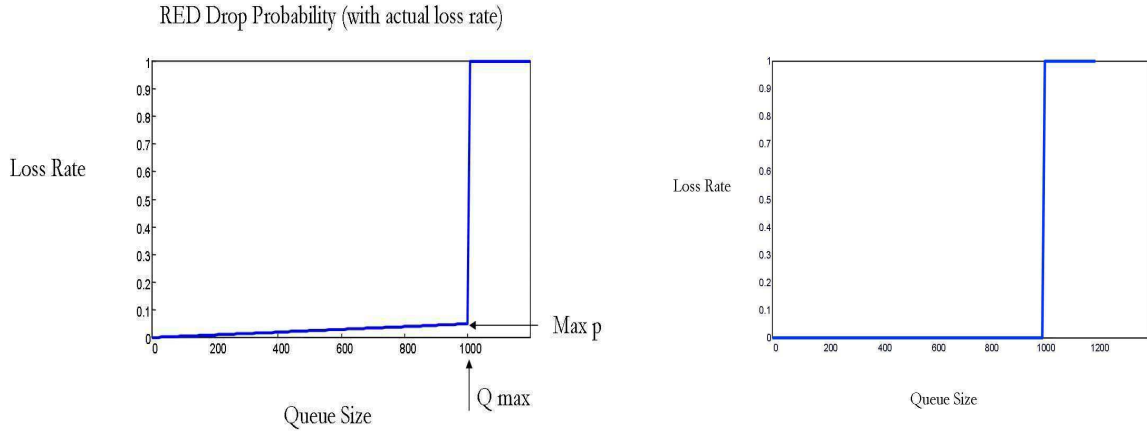


FIG. 12.13 – Fonctions de pertes pour RED et Tail Drop

12.2.2 Analyse de la qualité de service

L'attente des clients du modèle Une question qu'on peut légitimement se poser est pourquoi chercher à obtenir un équilibre stable. Pour cela il faut se placer du point de vue des attentes des utilisateurs et de l'opérateur. L'opérateur va souhaiter maximiser l'utilisation de sa bande passante, sinon il aura l'impression de perdre de l'argent, et il se demandera pourquoi il paye des mathématiciens pour résoudre ses problèmes plutôt que d'acheter des fibres optiques supplémentaires. Le client quant à lui est intéressé par la qualité du service qui lui est rendu, c'est à dire un bon débit, ceci dépend plus de la taille des tuyaux que l'opérateur lui loue que des optimisations que l'on peut réaliser. Mais bien plus qu'un débit de crête élevé, il sera intéressé s'il utilise par exemple de la téléphonie sur IP par un débit moyen à peu près stable et un délai le plus faible possible. C'est pour cela qu'il est intéressant qu'il y ait un contrôle stable à la file d'attente : on a un débit un délai de transmission à peu près constants.

L'indicateur de QoS On note l'indicateur de QoS (Quality of Service) en débit :

$$QoS(X\%) = \sup_{T \in \mathbb{R}} \left[\mathbb{P}_{flots}(\text{débit} > T) > X \right]. \quad (12.7)$$

Toujours dans notre cas du routeur T3 à 10000 paquets par seconde, nous étudions les indicateurs de QoS en débit pour $Minth = 100$ paquets et un $Maxth$ variable. Nous pouvons voir dans 12.14 et 12.15 que l'instabilité pénalise les QoS, en particulier une des plus importantes qui est la QoS à 95% qui correspond à donner un débit correct 95% du temps ou/et à 95% des utilisateurs en un instant donné dans les cas stationnaires (sinon on a des instants où tous les utilisateurs ont un mauvais débit en même temps). De plus nous voyons que pour les faibles taux de pertes nous avons une très bonne correspondance avec les résultats d'OPNET, en particulier pour la rentrée dans la région stable, le résultat

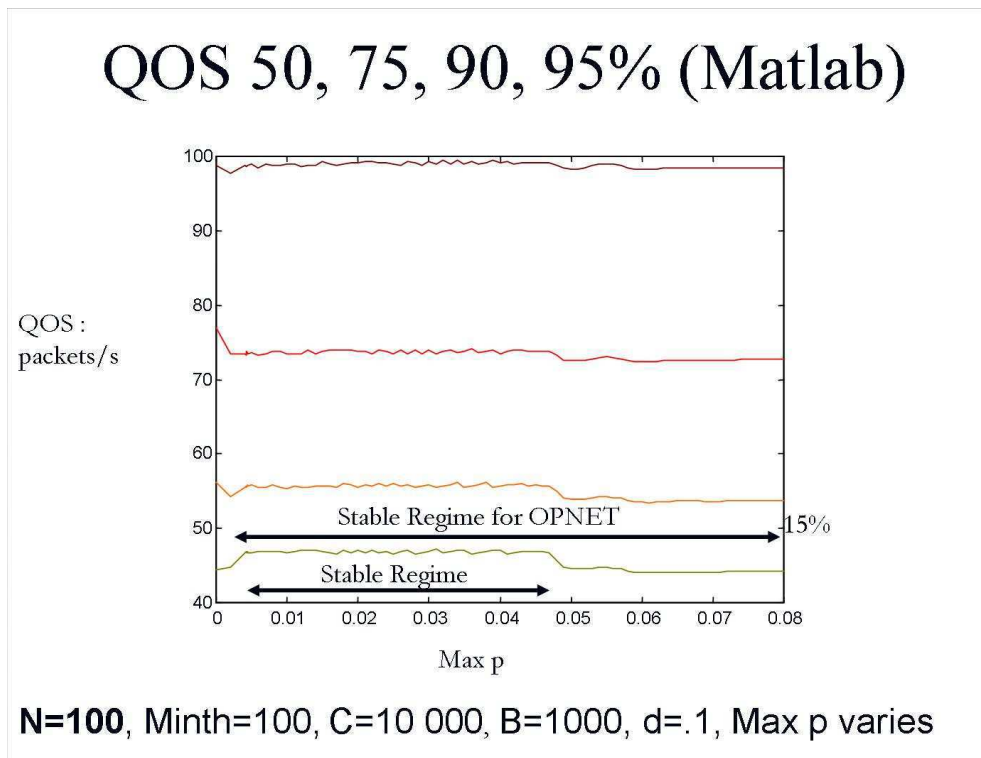


FIG. 12.14 – Indicateurs de QoS pour N=100 utilisateurs

est quasiment exact. En ce qui concerne les domaines de stabilité, ils sont en réalité aussi semblables, mais il est difficile de faire figurer avec OPNET la différence entre ne pas vider la file d'attente et être stable. Dans le cas de Matlab, on voit bien que pour $N = 200$ par exemple entre 9% et 12% on a une dégradation de la QoS. Cette dégradation est due non pas à une utilisation qui n'est pas maximale du lien, mais à un régime périodique autour de $Minth$ qui ne parvient pas à vider la file d'attente, mais qui produit un résultat presque aussi mauvais que Tail Drop. Ensuite, OPNET fournit une meilleure stabilisation, ceci provient essentiellement de «l'hypothèse de Little» sur les débits. En effet les résultats de stabilité seraient semblables en améliorant le modèle de débits comme nous allons le voir dans la section 12.3.

Si on compare la stabilité avec ce qui est donné par la formule (10.39), on voit qu'il y a un rapport 5 entre la borne que nous prévoyons et celle qui est observée (on devrait avoir une pente qui correspond à des valeurs de $Maxth$ de 3.8% et 7.7% environ pour N valant 100 et 200). 10.39 est donc un bon indicateur de stabilité.

Les trois régimes de RED On observe donc quatre régimes lorsqu'on augmente le paramètre $Maxth$:

- Le drop tail à grand buffer avec dégradation de la performance par rapport au Tail Drop pur (le pire cas).
- La stabilisation à buffer partiellement rempli avec performances optimales (le meilleur cas).
- Des cycles périodiques autour de $Minth$ avec utilisation maximale, mais QoS dégra-

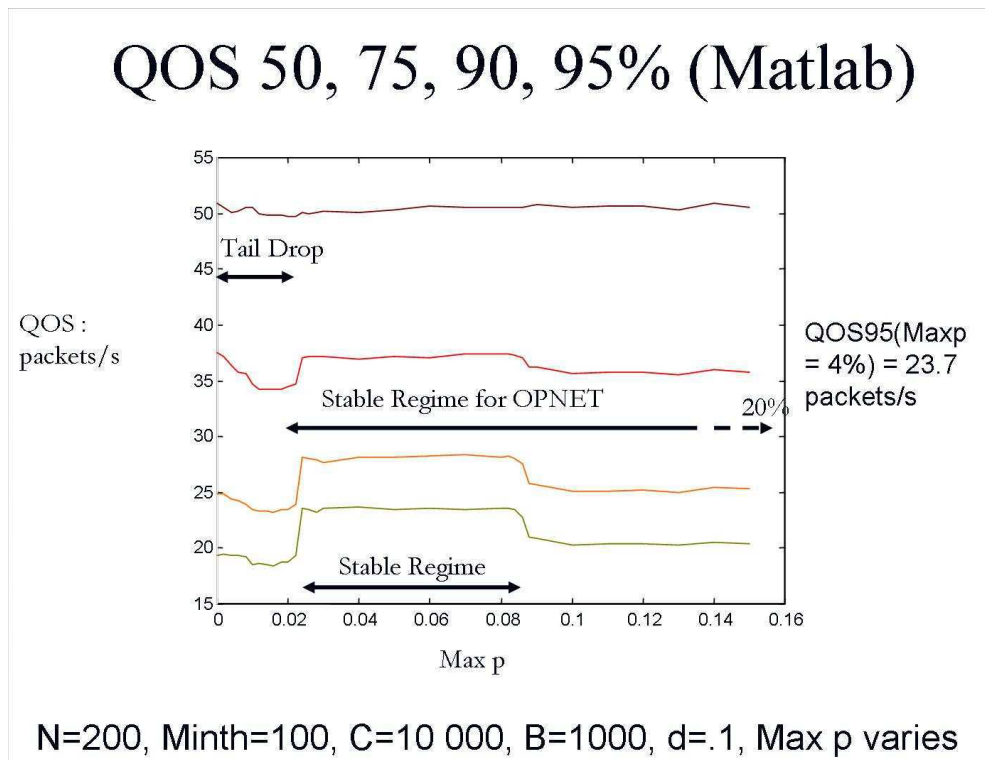


FIG. 12.15 – Indicateurs de QoS pour $N=200$ utilisateurs

dée (cas intermédiaire du point de vue utilisateur, bon du point de vue opérateur qui n'aime pas ses clients).

- Un drop tail avec buffer maximal de $Minth$ et des performances à peu près identiques au cas précédent.

On voit donc que Tail Drop avec un buffer relativement petit se comporte plutôt mieux du point de vue QoS qu'un Tail Drop avec un très grand buffer. Pour des considérations plus précises sur la taille du buffer à choisir pour Tail Drop, on peut se reporter à [129] (il faut prendre garde en particulier à ne pas avoir un buffer trop petit pour que l'approximation fluide tienne; voir l'annexe A à ce sujet).

Ce que nous n'avons jamais observé par contre ce sont des cycles limites dans la région où la pente du taux de pertes par rapport à la taille de la file d'attente est constante, même en prenant des paramètres extrêmes mais qui correspondent à la réalité (notons cependant que des phénomènes étranges apparaissent dans la zone où $W \ll 1$).

12.2.3 Comparaison entre le simulateur et des réalisations du système à N particules

Un simulateur du système à N particules Dans cette partie nous avons écrit un petit simulateur du système à N particules avec lequel nous testons les équations du champ moyen. Ceci va nous permettre de discuter des différentes hypothèses et simplifications que nous avons réalisé après avoir écrit les équations du modèle. Le système original à N particules n'étant pas déterministe nous étudions une réalisation particulière des processus

de Poisson dans le plan. Les équations sont ($I_i(t)$ est l'intensité de λ_i) :

$$dW_i(t) = \frac{dt}{R(t)} - \frac{W_i(t)}{2} d\lambda_i(t), \quad (12.8)$$

$$D(t) = \frac{1}{R(t)} \sum_i W_i(t), \quad (12.9)$$

$$I_i(t) = K(t - R(t)) \frac{W_i(t)}{R(t)}, \quad (12.10)$$

$$K(t) = f(Q(t)) \text{ ou } \max(0, 1 - C/D(t)), \quad (12.11)$$

$$R(t) = T + \frac{Q(t - R(t))}{NL}, \quad (12.12)$$

$$\frac{dQ}{dt}(t) = D(t)(1 - K(t)) - C. \quad (12.13)$$

Nous avons implémenté une classe de fonctions «CADLAG», ce qui permet une résolution exacte lorsque le RTT est fixe. Le principe est de résoudre le système d'équations de RTT en RTT comme dans la démonstration de l'existence du système auxiliaire.

Convergence du théorème de la limite centrale Dans le cas où on impose un taux de pertes fixé indépendant du temps avec un RTT fixe (et sans limite de débit), les particules sont indépendantes, et la convergence a lieu en théorème central limite (si la condition initiale converge suffisamment bien). Ce type de résultat une fois notre normalisation appliquée donne une convergence en $\frac{1}{\sqrt{N}}$ en ce qui concerne les débits. Cette convergence est déjà assez rapide, comme nous l'illustrons ici : avec $N = 10, 100$ et 1000 utilisateurs et un taux de pertes et un RTT constants $k = 1\%$ et $r = .1s$ (figure (12.16)). La condition initiale est $W_i(0) = 1$ pour tout i .

De plus deux phénomènes peuvent rendre la convergence plus rapide. D'une part en présence de synchronisation, la convergence va (visiblement) plus vite que le théorème central limite, mais surtout nous étudions la file d'attente qui est une intégrale par rapport aux débits. Par ailleurs, rappelons que le champ moyen qui sort du simulateur correspond à une équation un peu simplifiée.

Convergence avec RED Rajoutons une file d'attente qui implémente RED et une limite sur le débit. La file est susceptible de se remplir, les pertes et les délais ne sont plus constants. Nous obtenons alors la figure (12.17).

Dans la figure (12.18), nous faisons l'hypothèse $W(t - r) = W(t)$ comme dans le simulateur du champ moyen. Le résultat est assez semblable, les fluctuations sont semblables-elles uniquement dues au bruit. Ceci constitue une validation importante : la simplification que nous avons utilisée dans le simulateur des équations du champ moyen ne modifie pas sensiblement les résultats des simulations.

Convergence avec la drop tail On voit dans la figure (12.19) la convergence très rapide dans le cas de la file Drop Tail. Comme on pouvait s'y attendre les corrélations dues à la synchronisation augmentent la vitesse de convergence par rapport à un cas

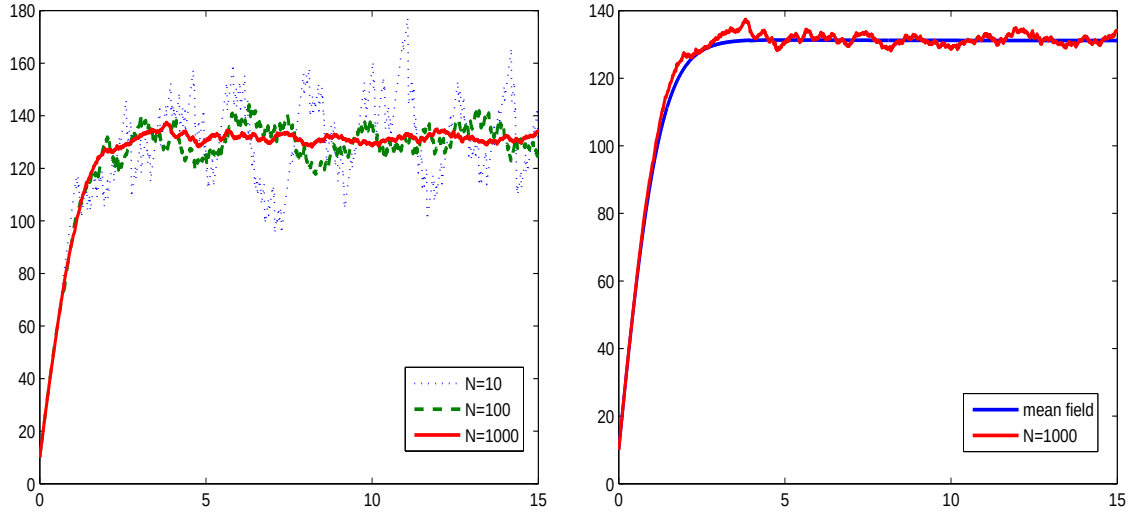


FIG. 12.16 – Convergence des débits normalisés $\bar{x}^N(t)$ et débit du champ moyen $k = 1\%$ et $r = 100ms$.

indépendant. Dans ce cas, la loi des grands nombres fonctionne déjà expérimentalement à $N = 10$ utilisateurs !

Ceci signifie que le système avec 10 utilisateurs peut être utilisé à la place du champ moyen pour ce type de simulations, et que quelque soit le nombre d'utilisateurs, le comportement est bien décrit par le système à 10 utilisateurs et par les équations du champ moyen.

12.3 Amélioration du modèle de débit

12.3.1 Diagnostic d'un problème dans la prise en compte des débits

Le fait de ne pas bien prendre en compte les débits donne lieu à certaines incohérences dans les simulations en présence d'une file d'attente en particulier lorsque la taille de cette dernière est grande. On peut le voir dans les figures (12.20) et (12.21) où on a pris $T = .1s$, $Q_{max} = .5s$ et des flots longs. Les résultats obtenus ne sont pas bien cohérents avec l'approximation puisque après une vague de pertes, les fenêtres sont réduites et la file reste pleine. Ensuite, la file se vide ce qui provoque un accroissement retardé du débit (puisque le RTT diminue).

Étudions les courbes (12.20) et (12.21). Regardons ce qui s'est passé dans notre cas après 20 secondes (la plus grosse oscillation). La fenêtre moyenne est passée de 30 à 15, et la file d'attente passe de .5 secondes à moins de .05 secondes. Avant de ressentir les pertes, on a peu près $W_{max}/RTT_{max} = L$. Ensuite, la formule de Little (débit = W/RTT) fait passer le débit par un maximum à une valeur supérieure à $W_{min}/RTT_{min} = L \times 1/2 \cdot \frac{.6s}{.15s} = 2L$.

Ceci n'est pas correct car le débit ne peut pas excéder significativement le débit d'ACK

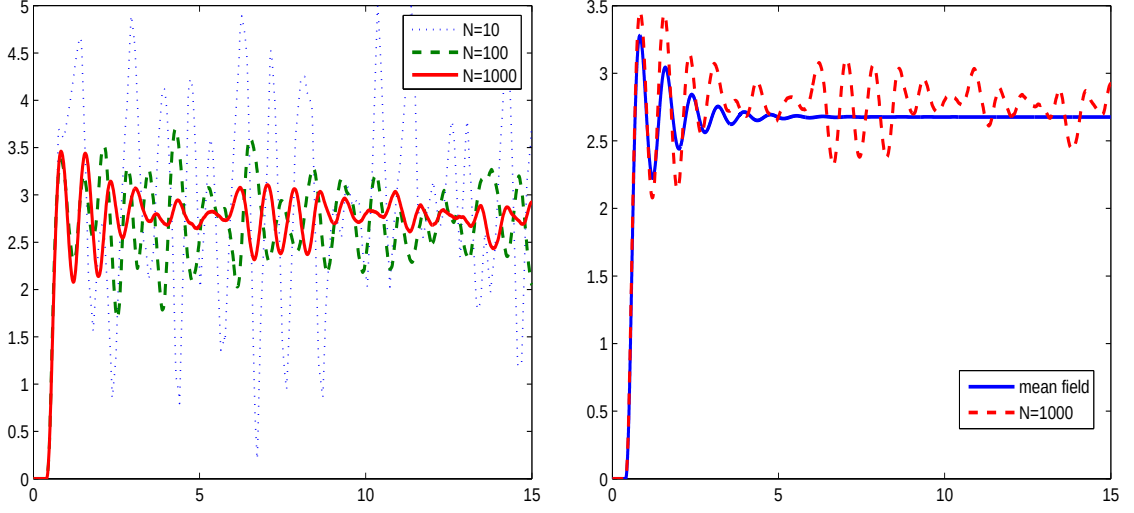


FIG. 12.17 – Convergence des files d’attentes normalisées $q^N(t)$ et $q(t)$ du champ moyen avec RED $p_{max} = 5\%$ et $r = 100ms$, $q_{max} =$, $q_{min} = 0$ et $L = 50$.

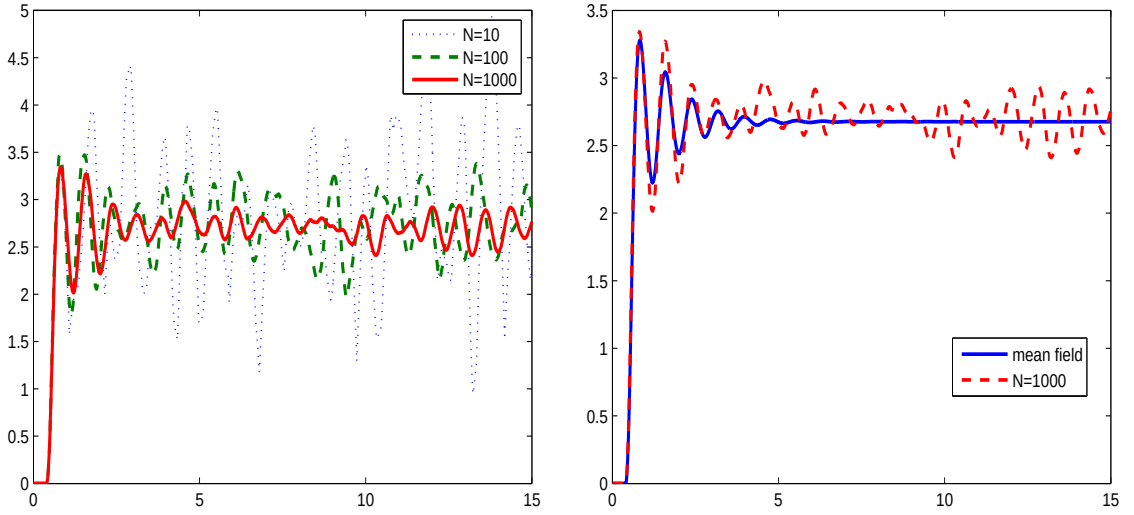


FIG. 12.18 – Convergence des files d’attentes normalisées avec simplification de la dynamique ($W(t-r)$ remplacé par $W(t)$) $q^N(t)$ et $q(t)$ du champ moyen avec RED $p_{max} = 5\%$ et $r = 100ms$, $q_{max} = 5$, $q_{min} = 0$ et $L = 50$.

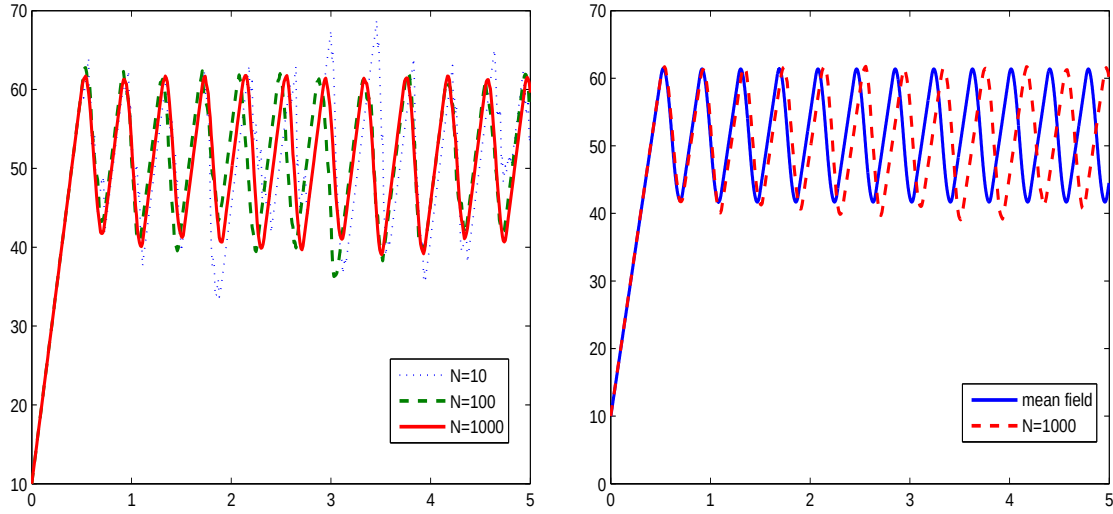


FIG. 12.19 – Convergence des débits normalisés avec une file de type drop tail sans buffer de capacité $C = 50$ paquets par utilisateur.

reçus. Ce débit d'ACK est majoré par le débit du routeur. Ceci n'est pas une erreur du simulateur, mais bien une incohérence due à l'utilisation d'une formule approximative.

12.3.2 Changement du modèle de lien et de file d'attente

Nouvelles notations La solution donnée ici est la modélisation fluide précise, un grand classique de la physique des fluides. On peut également la retrouver dans le simulateur N2N de Dohy Hong [68] (on lira la note [69] à ce sujet). Notons $D_i(t)$ le débit entrant dans le routeur, une proportion $K_i(t)$ de ces entrées est perdue ("i" comme «in», "D" «dimension of the link» ou débit). Notons $D_o(t)$ le débit sortant de la file d'attente ("o" comme «out»). La figure (12.22) permet de visualiser les notations. L'introduction du taux de pertes en sortie K_o nous permet d'évaluer quel est le débit virtuel d'ACK négatifs que la source va recevoir. En effet comme nous l'avons déjà vu plusieurs fois, le fait que les paquets n'arrivent pas est une information.

Notons de plus que dans notre système il y a en réalité trois étapes uniquement, le traitement par la source, le traitement par la file d'attente et le traitement par le lien après la file d'attente. Les passages dans le lien avant la file d'attente et par le récepteur des paquets qui les transforme en ACK sont totalement transparents de notre point de vue.

Relation entre D_i et D_o , K_i et K_o au passage de la file d'attente On a un délai et une limitation du débit disponible. Nous savons déjà calculer le délai grâce à la grandeur $Q(t)$ que nous avons introduite. Ainsi

$$K_o(t + Q(t)) = K_i(t) \quad (12.14)$$

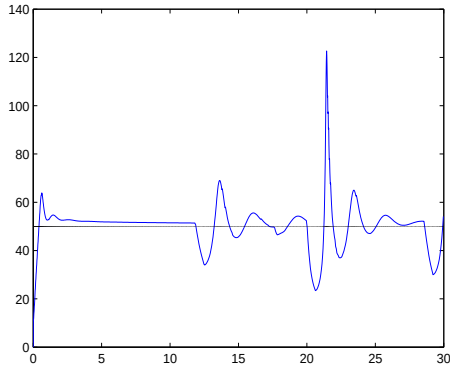


FIG. 12.20 – Débit et capacité disponible (en pkts/s/utilisateur) avec l'approximation Little

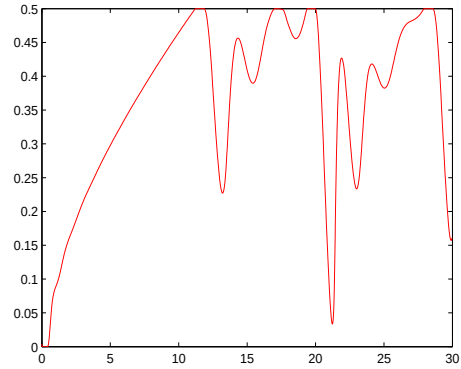


FIG. 12.21 – Taille de la file d'attente en secondes d'attente correspondante

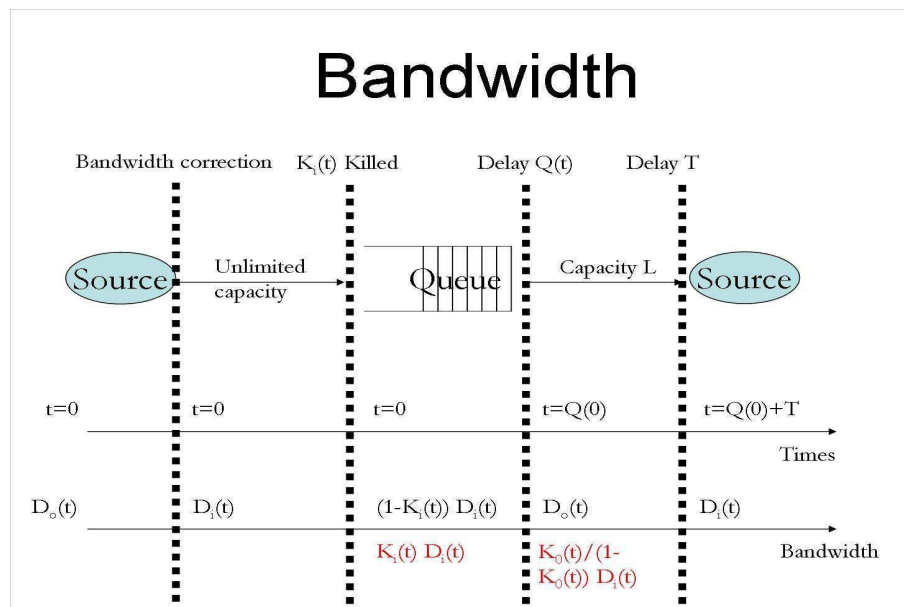


FIG. 12.22 – Description des notations

et

$$D_o(t + Q(t)) = \begin{cases} \min((1 - K_i(t))D_i(t), L) & \text{si } Q(t+Q(t))=0 \\ L & \text{sinon} \end{cases}$$

il revient au même d'écrire plus simplement :

$$D_o(t) = \begin{cases} \min(L, D_i(t)) & \text{si } Q(t) = 0 \\ L & \text{sinon} \end{cases} . \quad (12.15)$$

Relation entre D_i et D_o , K_i et K_o au passage du lien de sortie et de la source

Le lien sortant de la file d'attente impose un délai fixe de longueur T , ce qui ne pose pas de problème. La condition de passage par la source est que le nombre total de paquets en vol (morts ou vifs) vaut exactement la fenêtre (c'est sa définition). Si on note $V(t)$ la quantité moyenne de paquets en vol par source (en limite des grands nombre en supposant qu'on a toujours un champ moyen qui converge), alors $V(t) = M_1 = \int wp(t, w)$. La seule opération qui fasse varier la masse de $V(t)$ est le passage par la source. Ainsi

$$\frac{dM_1}{dt}(t) = \frac{dV}{dt}(t) = D_i(t) - \frac{D_o(t - R)}{1 - K_o(t - R)}. \quad (12.16)$$

12.3.3 Implémentation du modèle de débits réel

Astuce d'implémentation de cette amélioration Nous allons implémenter directement une file d'attente et le lien sortant dans nos données. Dans la mesure où nous manipulons déjà une matrice carrée pour calculer la distribution (W, F) des utilisateurs actifs, ceci ne représentera pas une grosse contrainte en termes de mémoire et de calculs. L'idée est de faire ainsi : on met dans une file des coordonnées $(t, D_i(t)\Delta t, K_i(t))$, et on les défile au débit maximal L et uniquement sur la première coordonnée. Ainsi on défile exactement à l'instant où les paquets reviennent à la source sous forme d'ACK, et on résout implicitement le problème d'effet Doppler rencontré à cause des distortions de temps dues aux variations de la longueur de la file d'attente. Ceci se nomme Ligne de retard ou «Delay line» en anglais. Ce modèle était utilisé dans le modèle décrit au chapitre 2 sous le nom Kelly-Kelly-Bain.

Pour le choix de la structure de donnée, voici les principaux objectifs à réaliser :

- des mises à jour faciles pour les opérations de défilage et d'enfilage dans le transporteur (queue et lien ensemble),
- un accès facile à la taille de la file d'attente,
- une bonne gestion de l'invariant : le nombre de paquets dans le système est donné par la fenêtre.

Pour réaliser cela nous avons choisi un tableau déroulant avec 3 positions repérées, l'entrée dans la file d'attente depuis la source, la sortie de la file d'attente vers le lien et la sortie du lien. Ceci résout les deux premières difficultés. Ensuite pour ajuster la masse comme nous ne pouvons pas nous permettre d'enfiler des paquets négatifs, on garde une variable d'ajustement courant du nombre de paquets à renvoyer par la source.

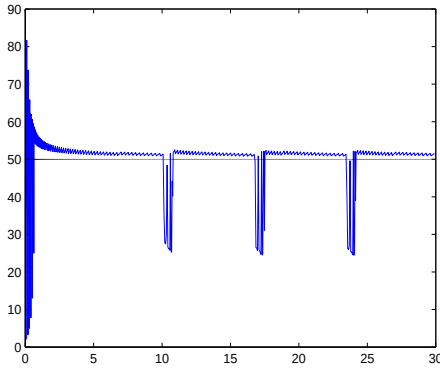


FIG. 12.23 – Débit et capacité disponibles (en pkts/s/utilisateur) avec le modèle précis

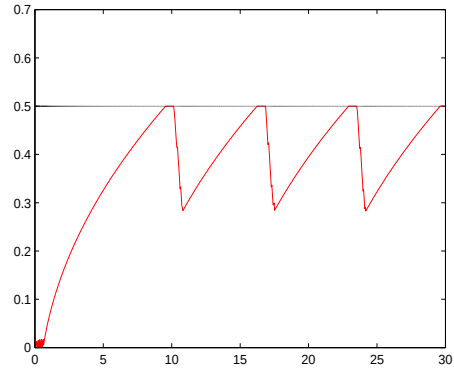


FIG. 12.24 – Taille de la file d'attente en secondes d'attente correspondante

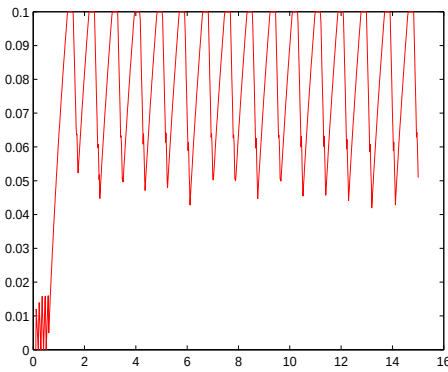


FIG. 12.25 – Temps d'attente dû à la file par l'EDP du champ moyen modifiée.

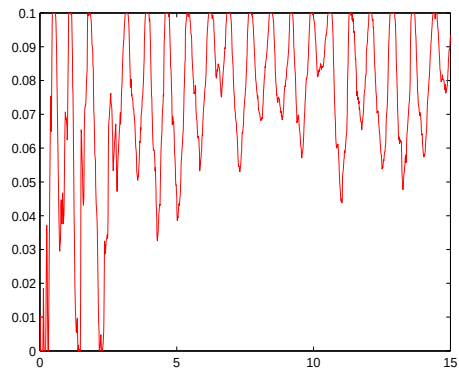


FIG. 12.26 – Temps dans la queue par le simulateur OPNET.

Le problème initial est corrigé Les résultats sont spectaculaires, on peut voir que le problème initial a été résolu sur les figures (12.23) et (12.24). De plus cette meilleure prise en compte des débits améliore aussi considérablement les résultats concernant la stabilité. Un petit défaut que l'on peut voir ici cependant est que le pas de calcul du modèle défluidifie la file d'attente et les débits : dans la première seconde alors que la file devrait être vide, on voit qu'elle se remplit et se vide, cela est lié au pas de calcul. Dans la version avec l'approximation de Little, ce problème n'apparaissait pas.

Un autre motif de satisfaction (assez relatif celui-là) de cette nouvelle modélisation des débits est l'accord que l'on trouve dans la simulation de Tail Drop dans le cas des connexions persistantes. Avec les paramètres TCP long, 200 utilisateurs se partagent une bande passante de 10 000 paquets/s, un temps de transmission $T = 100ms$ et une file d'attente de longueur correspondant à une attente de $100ms$. On peut voir les résultats dans les figures (12.25) et (12.26).

On voit que la période d'OPNET est plus petite que celle du simulateur Matlab, mais

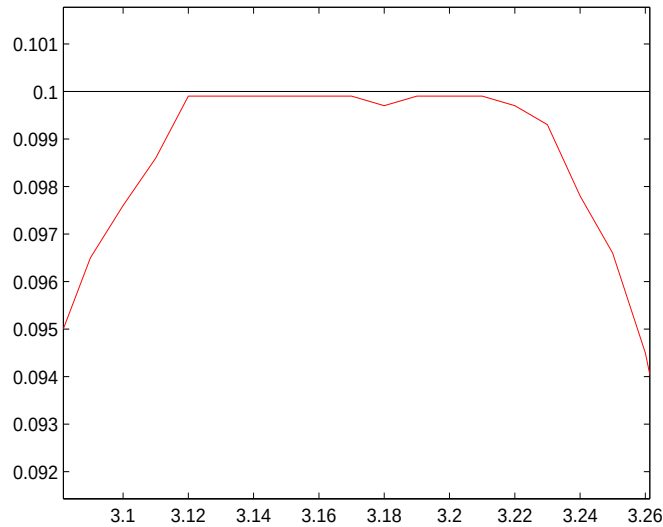


FIG. 12.27 – OPNET donne des résultats incohérents : zoom sur une partie où la réaction à l'atteinte de q_{max} devrait durer au moins 200 ms.

que les formes sont maintenant identiques (on sait même si ce n'est pas formellement prouvé que les modèles à N utilisateurs convergent extrêmement rapidement dans le cas très synchronisé que nous étudions). On observe chez OPNET un comportement doublement périodique qui provient sans doute des timeouts qui ne sont pas intégrés dans la simulation Matlab. On remarque aussi un problème de programmation d'OPNET, il réagit plus vite aux pertes qu'il ne devrait le faire (voir la figure (12.27)). Ce problème explique pourquoi nous allons compléter dans la suite les simulations avec NS.

12.4 Intégration des Timeouts

Nous décrivons ici des améliorations que nous pourrions apporter à notre simulateur.

Retour sur la simulation de Tail Drop avec NS-2, description des problèmes

La raison principale qui nous avait mené à utiliser OPNET plutôt que NS-2 est que NS est extrêmement technique à faire fonctionner et surtout bénéficie de très peu de support.

OPNET pour se part est propriétaire, un peu cher, et les versions 6, 7 et 8 testées avaient le bug que nous avons décrit précédemment ; en outre les compétences disponibles en décryptage de code C n'ont pas permis de régler le problème. Peut-être ce défaut a-t-il été corrigé depuis. Aussi il a fallu refaire les simulations avec NS-2 (voir [34]). Pour le prendre en main on fera bien de se reporter à l'excellent travail d'Eitan Altman et Tania Jimenez ([3, 4]).

La file d'attente obtenue dans une configuration habituelle ($Q_{max} = .1s$, $N = 200$, $T = .1s$, $L = 50pkt/s/user$, $C = 10000pkt/s$, TCP long lived, Drop Tail) est représentée dans la figure (12.28). Remarquons que la période plus faible de NS par rapport à l'équation

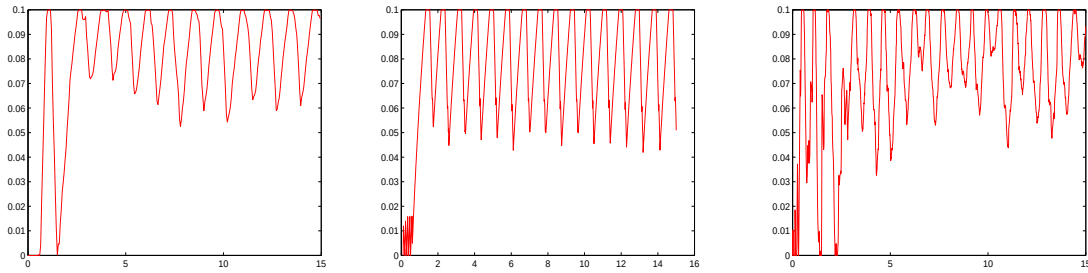


FIG. 12.28 – Comparaison de NS (gauche), la PDE sous Matlab (milieu) et OPNET (à droite)

différentielle s'explique bien en rajoutant les timeouts : le taux d'accroissement de la file d'attente est très légèrement plus faible que celui de Matlab, et vaut 1 paquet par connexion active et par RTT. La différence provient sans doute des connexions en timeout.

Si nous reprenons la modélisation du chapitre 10, un timeout se produit essentiellement lorsqu'un utilisateur avec une fenêtre inférieure à 4 reçoit une perte, car le dispositif du duplicate ACK ne peut pas fonctionner. À ce titre l'exemple que nous avons choisi est particulièrement difficile pour notre simulation puisqu'on a une fenêtre moyenne très petite, et donc de relativement nombreux timeouts. Il semble que dans la pratique le timeout typique des simulateurs soit de l'ordre de 1 seconde, ce qui ne suit pas les recommandations du RFC, mais est plus simple à programmer : c'est donc la valeur que nous retiendrons.

Nous ne poursuivrons pas plus les validations par rapport aux simulateurs existants, dans la mesure où nous avons sans doute atteint les limites de leur validité, et il faudrait comparer maintenant les résultats à des traces réelles. En effet on observe de nombreuses incohérences : dans la simulation de la figure (12.28) par exemple, si on veut croire que la période est toujours comme dans le cas stable pour RED de l'ordre de $T = 2\pi r$ (voir section 10.38), on voit que pour NS, $r = 0.20s$ environ (période de $1.25s$), pour Matlab $r = 0,15$ (période $0.9s$) et pour OPNET $r = .11s$ (période $0.7s$), ce qui est assez consistant avec le fait qu'OPNET ne prend pas en compte correctement le délai d'attente dans la file et que NS majore franchement le délai.

Solution finalement adoptée Devant les difficultés rencontrées pour trouver un simulateur plus fiable que les équations que nous simulons, nous avons finalement ajouté des timeouts dans la version précisée précédemment sans réussir à faire une optimisation. Nous avons choisi une valeur dans une fenêtre autour de 1 seconde pour la valeur de RTO. Globalement l'intérêt de la modélisation des timeouts est de savoir dans quels cas le modèle sans timeouts est loin de la réalité. De plus une modélisation des timeouts devrait s'accompagner d'une prise en compte du mécanisme de slow start au moins en supposant que les utilisateurs commencent à une fenêtre plus grande que 1 à défaut de décrire l'évolution exacte.

Et slow start ? Les équations de slow start ne présenteraient pas de difficultés particulières d'implémentation, il s'agit juste d'une limitation de temps qui n'a pas permis d'en tenir compte. Par contre cette intégration soulève une question que l'on pourrait se poser, est-il normal de modéliser par une dérive une augmentation de même nature que la diminution que l'on modélise par des sauts ? Une réponse pourrait être que la division par deux ne modélise pas la fenêtre mais plutôt le débit sortant de la source. Nous laissons ces questions en suspens pour nous lancer dans l'exploitation du simulateur.

Chapitre 13

Exemples d'utilisation du simulateur

Dans ce chapitre nous supposons que notre simulateur donne une certaine description de la réalité. Nous allons dresser un certain nombre d'observations et de conclusions que l'on peut tirer sur le comportement de certains routeurs d'Internet. Nous commençons en analysant l'effet de la congestion dans des réglages typiques sur des utilisateurs de HTTP ; la dégradation est pire que ce qu'on aurait pu attendre, ceci à cause d'un effet de levier sur le taux d'activité des utilisateurs. Dans un second temps nous tentons d'expliquer les turbulences observées dans l'article [9]. Il s'agit de montrer qu'un routeur devant un site WEB peut se trouver dans une situation où à même comportement des utilisateurs, deux situations sont possibles, l'une stable, l'autre oscillatoire (on dira encore turbulente).

Sommaire

13.1 Effet de la congestion pour un routeur sous ses réglages habituels	232
13.1.1 Données numériques	232
13.1.2 Premiers résultats	232
13.1.3 Détails sur la congestion	232
13.1.4 Conclusion sur la congestion avec HTTP	233
13.2 Turbulences	233
13.2.1 Première approche	234
13.2.2 Facteurs de turbulence	235
13.2.3 Des cas réalistes de régime transitoire long et de turbulences pour TCP-AIMD (valable pour SACK, Reno et New-Reno)	236
13.2.4 Conclusions sur le phénomène de turbulence	236

13.1 Effet de la congestion pour un routeur sous ses réglages habituels

13.1.1 Données numériques

Nous allons étudier l'exemple du chapitre précédent du site WEB de l'École Normale Supérieure. Ce routeur a un débit de 150 Mbits/s environ. Nous avons noté que la taille moyenne des fichiers présents était de 40 paquets de 1ko environ. De manière classique le routeur de l'ÉNS est réglé en fonction d'un RTT de 100 ms, c'est-à-dire que le nombre de paquets dans la file d'attente est 1/10 du nombre de paquets envoyés par le routeur en 1 seconde. La taille du buffer présumé est donc de 2000 paquets environ, ce qui nous permet de négliger l'effet de pertes anticipée dû aux fluctuations décrit dans [129] (voir l'annexe A pour une discussion sur ce sujet). Enfin on prend des temps de réflexion de 20 secondes. Dans les conditions optimales de fonctionnement (régime permanent et capacité disponible suffisante), un utilisateur a besoin de 1.9 paquets par seconde en moyenne et 3.65% des utilisateurs sont en phase de transmission à un instant donné, c'est-à-dire que la capacité disponible par utilisateur effectivement en train de télécharger est de 360 kbits/s.

13.1.2 Premiers résultats

Nous commençons par illustrer le phénomène de congestion que le modèle explique particulièrement bien. Dans la figure (13.1), on peut voir que nous avons fait baisser le débit total disponible de 25% au temps $t = 5s$, ceci se traduit pour les utilisateurs par une baisse ressentie de plus de 80%, à tel point qu'on sort des conditions de fonctionnement de TCP avec une fenêtre moyenne de l'ordre de 4 et des taux de pertes bien supérieurs à 10% (on considère que TCP fonctionne correctement jusqu'à 5% de pertes). Aussi, le mécanisme de timeout va rentrer en jeu massivement.

Ce qui arriverait avec un afflux supplémentaire de clients dans la même proportion (+33% de clients) provoquerait les mêmes conséquences (division par plus de 5 du débit disponible) ; c'est ce que l'on ressent à certaines heures de la journée quand un site semble ne plus fonctionner correctement et qu'on doit faire recharger plusieurs fois la page avant de péniblement obtenir quelque chose.

13.1.3 Détails sur la congestion

Les figures (13.2) et (13.3) permettent d'interpréter simplement le phénomène que nous avons observé. Le débit diminuant, il faut plus de temps pour télécharger une page. Ce faisant la proportion du temps que l'on passe à télécharger par rapport au temps pendant lequel on lit devient de plus en plus grande. Ceci explique le phénomène d'auto-amplification de la congestion. D'autre part la file d'attente ne sert à rien dans cet exemple à part à augmenter le temps de transmission, ce qui permet des fenêtres plus grandes pour le même débit. Dans tous les cas, on est ici face à un problème de congestion sévère, la seule solution pour l'éviter est d'augmenter le débit disponible ou de diminuer le nombre d'utilisateurs.

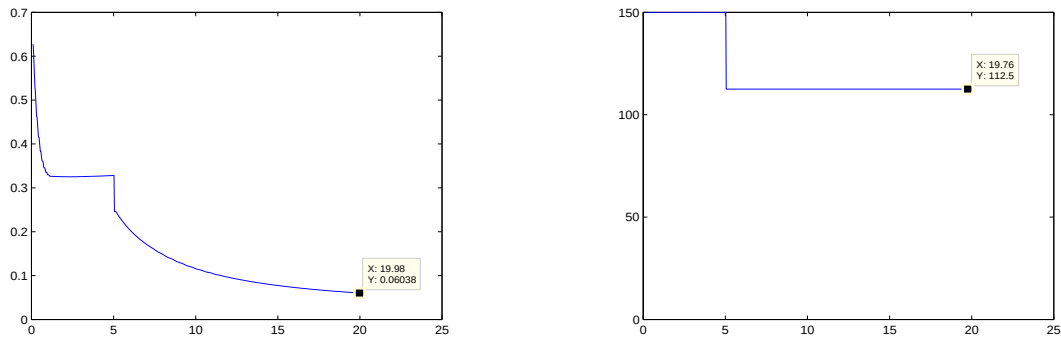


FIG. 13.1 – Débits disponibles en Mbits/s par utilisateur en cours de transmission à gauche et capacité du routeur à droite. Au temps $t = 5s$ on impose une réduction du débit disponible

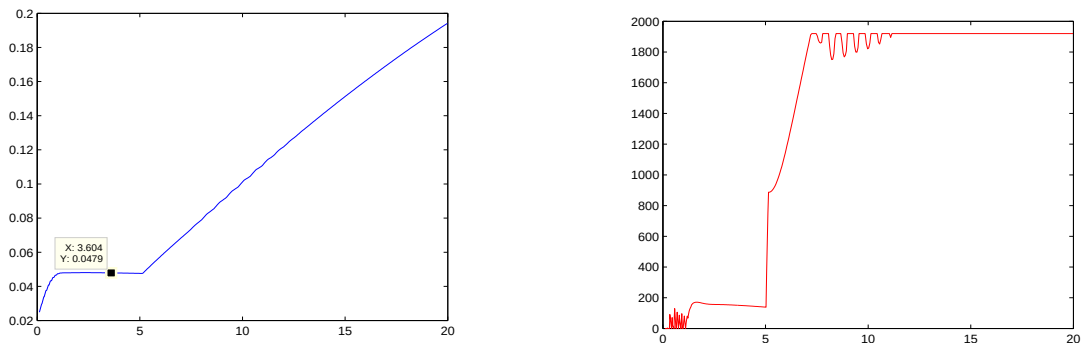


FIG. 13.2 – Proportion des utilisateurs actif à un instant donné en fonction du temps (réduction de débit à $t = 5s$)

FIG. 13.3 – Taille de la file d'attente en paquets qui correspond à une attente maximale de 100 ms

13.1.4 Conclusion sur la congestion avec HTTP

Nous venons de mettre à jour un effet de levier qui dégrade les performance ressenties par des utilisateurs de HTTP. Dans le cas de téléchargement persistants, une hausse de 33% du trafic occasionne une baisse de 25% du débit ressenti par les utilisateurs. Dans le cas de HTTP, un nouvel effet rentre en jeu, les utilisateurs mettent plus longtemps à télécharger leurs pages, et passent proportionnellement moins de temps à lire qu'à télécharger. Ceci représente un effet de levier qui provoque une hausse du nombre des utilisateurs actifs à un instant donné et finalement une baisse de 80% du débit. Ceci peut mener à faire sortir du domaine de fonctionnement habituel de TCP.

13.2 Turbulences

Ce que nous entendons par le mot «turbulences» Avant de commencer, il convient de clarifier ce point. En effet le mot turbulence est apparu à la suite des travaux de Lixia

Zhang, Scott Shenker et David Clark[152] où on voit dans un cas particulier que le trafic dans Internet a parfois un aspect oscillatoire. Ici ce n'est pas à ce genre de chose que nous faisons allusion. Dans le cas de HTTP, nous avons vu qu'il existait un équilibre stable quand le débit offert est suffisant. En effet la dynamique que nous avons retenue dans la modélisation fait alterner les phases actives et de sommeil ; le sommeil n'est pas plus court si la phase active a été courte. Ce qui est marquant dans les observations de [9] est l'existence d'un régime turbulent (c'est-à-dire oscillatoire) qui côtoie le régime stable pour les mêmes paramètres. L'existence de ces deux régime est démontré dans le cas de TCP Tahoe pour un modèle du type HTTP-AIMD tel que nous avons vu au chapitre 7, et il est observé dans le cas de TCP Reno pour le simulateur N2N de Dohy Hong et François Baccelli. Le principe de turbulences dont il est question ici surtout une curiosité mathématique portant sur l'étude des équations des modèles. En effet, il ne fait aucun doute que dans la réalité, le régime oscillatoire est privilégié à cause du bruit ambiant.

13.2.1 Première approche

Rappelons que la turbulence correspond à un état où selon des conditions de départ on peut avoir deux comportements limites différents. Dans nos cas il s'agit d'une situation où le débit est constant et d'une autre où le débit oscille. Ce phénomène s'observe très facilement dans le cas de TCP Tahoe (lorsqu'au lieu de diviser la fenêtre par deux, on la réinitialise à 1), et l'effet a été démontré pour le modèle HTTP-AIMD dans [9]. Par contre cet effet est plus délicat à observer pour TCP Reno, et si on n'a pas de mal à trouver des régimes transitoires extrêmement longs, il est plus délicat de tomber sur un vrai cas turbulent.

Le phénomène de turbulence s'explique bien par l'effet de congestion dont nous venons de parler : lorsque des pertes sont engendrées en dépassant la taille de la file d'attente, on peut provoquer une sous-utilisation du lien. Ceci a deux conséquences, une augmentation de la proportion d'utilisateurs actifs et une baisse du débit global partagé. Une chose intéressante à regarder est donc l'influence d'un taux de pertes sur le débit global réalisé par des sources dans HTTP-RED en régime permanent lorsqu'il n'y a pas de limite de débit. En effet ce débit est la résultante de l'augmentation du nombre d'utilisateurs et de la baisse du débit par utilisateur. C'est ce que nous pouvons voir dans la figure (13.4).

Il est vrai qu'on ne peut en général que difficilement comparer les résultats d'un régime à pertes constantes à un régime oscillatoire. Cependant dans notre cas, les situations sont assez semblables. Nous pourrions vérifier dans le cas de TCP Reno (ou de TCP AIMD³¹) que le régime oscillatoire est à peine un peu meilleur à taux de pertes moyen identiques. Passons sur ce point et essayons de voir quelles indications tirer de la figure (13.4). On peut voir premièrement que si on impose un taux de pertes identique pour les deux, la version Tahoe utilisera moins de bande passante. C'est précisément ce point qui est crucial dans le phénomène de turbulence, le régime turbulent doit saturer le routeur tout en n'utilisant pas non plus toute la bande passante disponible. En fait le régime turbulent doit gaspiller plus de bande passante que le régime constant pour continuer à osciller ;

³¹Rappelons que AIMD se réfère à Additive Increase Multiplicative Decrease, c'est-à-dire la manière dont le protocole gère l'évolution des fenêtres en fonction de la réponse du réseau : il croît linéairement quand tout se passe bien, quand on découvre un problème, il divise par une constante la fenêtre d'envoi.

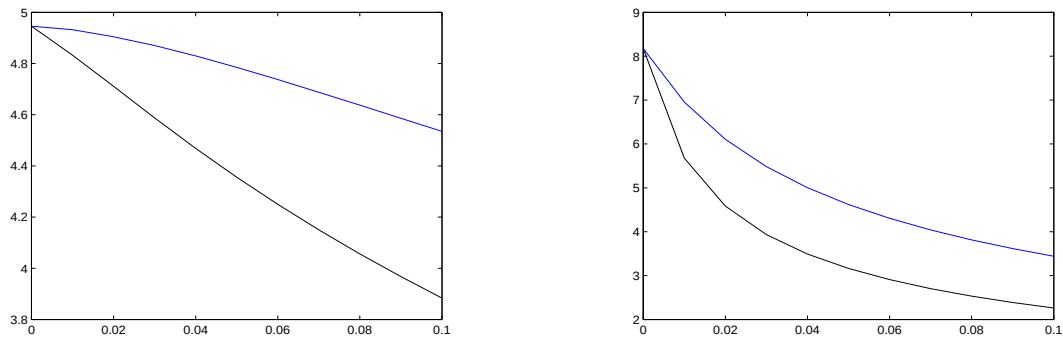


FIG. 13.4 – Débits moyen total (à gauche) et fenêtres moyenne des utilisateurs actifs (à droite) HTTP en paquets par utilisateur et par RTT pour des taux de pertes constants entre 0 et 10%. Deux versions de TCP, Reno en bleu, Tahoe en noir. Dans cet exemple le think time moyen est de 20 s et la taille moyenne des fichiers est de 100 (distribution exponentielle).

sinon la proportion d'utilisateurs actifs va baisser et finalement on va sortir du régime turbulent.

Remarque 13.1. *Un point important pour la compréhension des turbulences est que le régime turbulent ne peut pas offrir plus de bande passante en moyenne que le régime constant.*

Ainsi nous voyons dans la figure de gauche de (13.4) que la version TCP AIMD utilise beaucoup mieux la bande passante en présence de pertes, ce qui rend moins probable le gaspillage de bande passante lorsque le routeur est saturé.

13.2.2 Facteurs de turbulence

Le gaspillage de bande passante Les premiers facteurs de turbulence sont assez simples à décrire, en fait il s'agit de tout ce qui va provoquer un gaspillage de bande passante lorsque le routeur partagé est congestionné. Il s'agit :

1. du délai important par rapport à la taille de la file d'attente,
2. d'un mécanisme de réponse trop brutal à la congestion (comme dans le cas de Tahoe par rapport à AIMD),

Ensuite il est délicat de dire d'emblée si c'est en ayant de petites ou de grandes fenêtres que le cas est le pire, en effet les petites fenêtres vont provoquer plus de pertes, mais elles y réagiront moins que les grandes. Par expérience nous dirons que sans doute les très grandes fenêtres subiront un choc plus important que les autres.

L'augmentation de l'activité Le second facteur important est l'effet de congestion dont nous avons parlé dans la section précédente, le taux d'activité augmente si on utilise mal la bande passante. Ce qui va favoriser cet effet sera essentiellement le faible taux

d'activité à la normale, c'est-à-dire une taille de fichier à télécharger en moyenne (donnée en temps optimal de téléchargement par exemple qui est de l'ordre de $R\sqrt{f}$) assez petite par rapport au temps moyen d'attente entre deux téléchargements.

Les turbulences dans les cas réalistes Les cas que nous avons décrits comme typiques satisfont bien les facteurs que nous avons relevés pour la turbulence. Le fait essentiel étant le taux d'activité relativement bas. Une des difficultés qui pourrait empêcher de découvrir les turbulences est que pour une même configuration, deux taux d'activités différents au départ peuvent pour l'un donner un régime transitoire très long et pour l'autre donner le régime oscillatoire. Il faut bien penser à faire augmenter de manière importante le taux d'activité par rapport à celui observé en régime constant.

De plus pour se placer dans le cas de petit tampon grand par rapport au considérations de la première section il faut choisir un RTT assez important de l'ordre de 200 ms.

13.2.3 Des cas réalistes de régime transitoire long et de turbulences pour TCP-AIMD (valable pour SACK, Reno et New-Reno)

On prend une taille moyenne des fichiers de 100 paquets, un RTT de 100 ms, pour un temps de réflexion de 20 s. La taille de la file d'attente est prise à 10 ms pour un débit de 5 paquets par utilisateurs, on suppose que le nombre d'utilisateurs qui partagent le routeur est de l'ordre de 40000, ce qui nous donne une taille de file d'attente de 2000 paquets qui nous permet de négliger l'effet de la première section. Le débit global du routeur est de 1.5Gbits/s ce qui est une valeur assez commune aujourd'hui (l'Ethernet moderne fait déjà 1 Gbit/s).

Ces hypothèses sont réalistes, en particulier dans la mesure où on peut sans grand risque prédire que la taille des fichiers va continuer à augmenter à mesure que le prix de la bande passante va baisser.

Pour commencer la figure (13.5) présente la stabilisation au début de notre expérience. La figure (13.6) présente ensuite une période de congestion où le débit est diminué. Le but est d'augmenter le taux d'activité. Enfin dans la figure (13.7), on observe ce qui pourrait être interprété comme une turbulence, mais qui en fait se trouve être un cas transitoire très long.

Ensuite nous réalisons la même expérience dans la figure (13.8)(avec des paramètres qui permettent une stabilisation plus rapide vite au départ), la différence essentielle est qu'à partir de la situation $t = 10$ secondes les sources qui sont encore synchronisées sont laissées avec le débit 5 paquets par utilisateur par secondes. On observe semble-t-il un vrai cas de turbulence, la figure (13.9) montre comment se comporte le débit utilisé par les utilisateurs actifs.

13.2.4 Conclusions sur le phénomène de turbulence

Finalement le problème des turbulences est réglé pour les problèmes pratiques, le cas oscillatoire est un phénomène qui existe bien dans les cas que nous avons indiqué. En fait

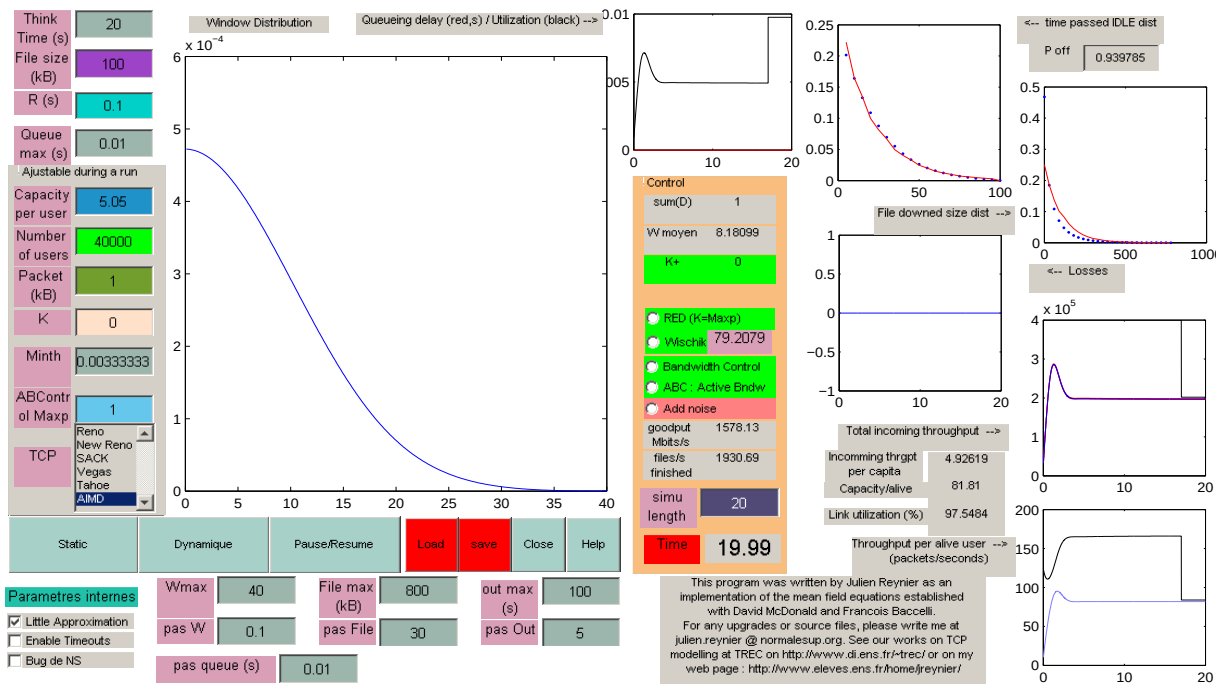
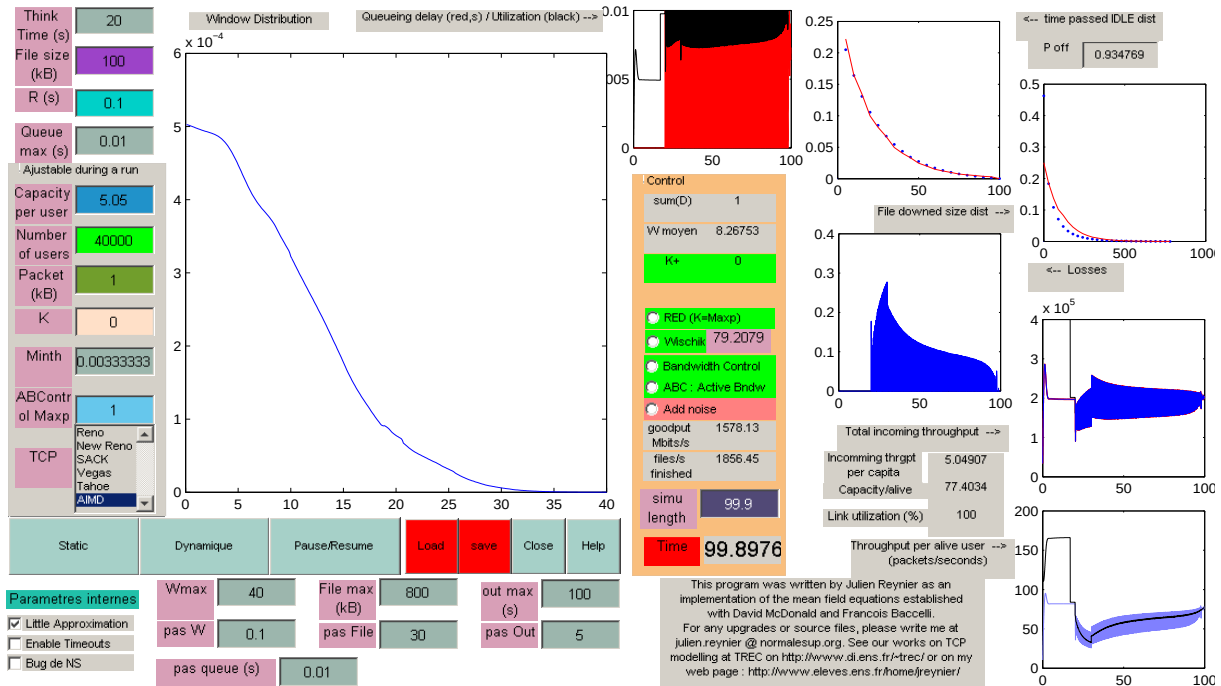
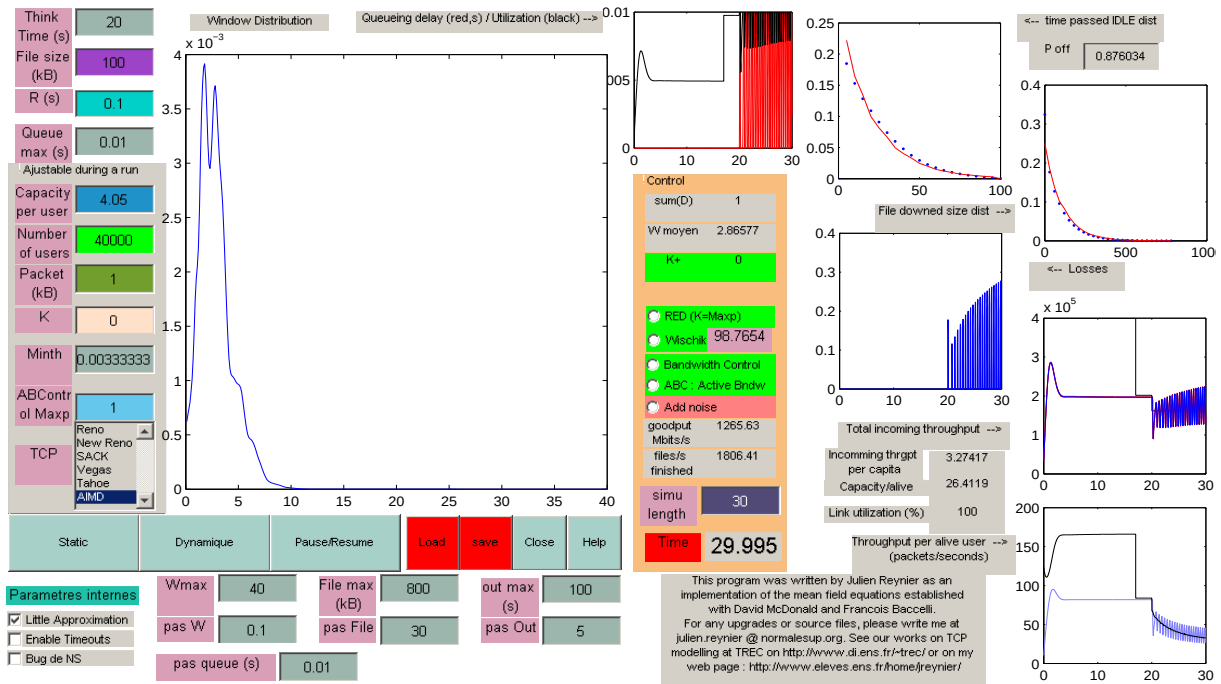


FIG. 13.5 – Jusqu’à $t=17$ s on offre 10 paquets par seconde et par utilisateur. La bande passante utilisée vaut 4.93 paquets par utilisateur et par seconde. On peut estimer que le régime est alors totalement stabilisé. Ensuite on baisse l’offre à 5.05 paquets par utilisateurs jusqu’à $t = 20$ s, ce qui essentiellement ne change rien puisque l’offre reste supérieure à la demande qui est stabilisée. On observe que l’utilisation est alors de 97.5% environ.

le cas difficile à observer dans des cas réels sera le cas constant, et on peut penser qu’on aura des passages d’un état à l’autre au cours du temps.

Cependant il semble qu’une étude théorique plus poussée soit à faire pour savoir si les turbulences que nous observons viennent du bruit du simulateur ou font bien partie du système mathématique (sans bruit rajouté bien sûr).



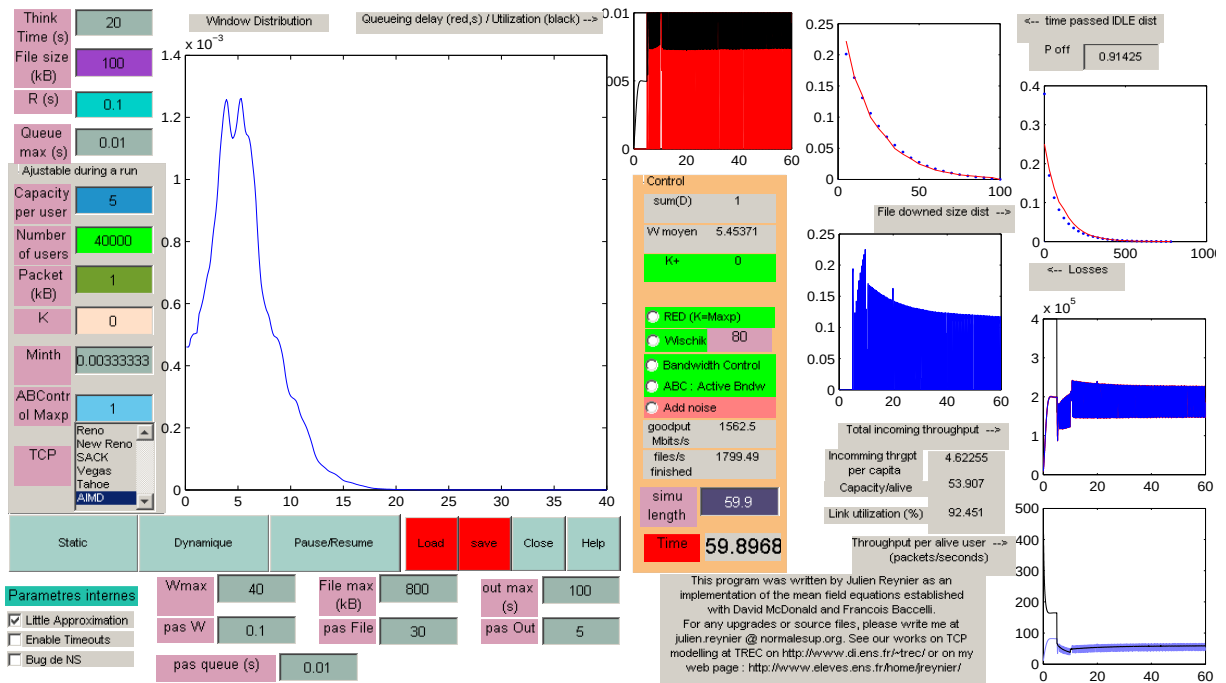
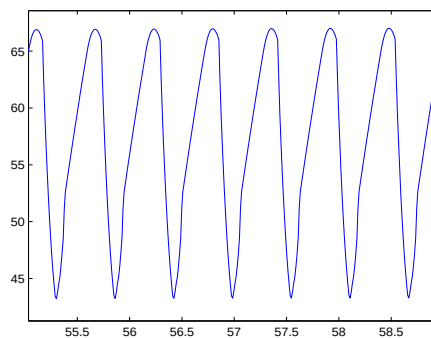
FIG. 13.8 – À partir de $t = 10$ s. On a la turbulence.

FIG. 13.9 – Débit utilisé en moyenne par chaque utilisateur actif en paquets par secondes dans le cas turbulent.

Annexe A

Des files d'attente de petite taille

Le champ d'investigation : Buffer et $M/D/1$ Tout d'abord rappelons que les paquets sont de tailles variables. Cependant il a été remarqué que la modélisation par la $M/D/1$ permettait de décrire une borne dégradée des résultats réels. Nous ne rentrerons pas trop dans les détails de ceci, on consultera l'introduction de la thèse d'Alexandre Proutière ([126]) pour des précisions intéressantes à ce sujet. Le but de cette section est d'étudier les limites possibles à l'effet de turbulences dans HTTP étudié dans [9]. La turbulence du système étant définie par la présence dans certaines configurations de deux états limites :

- l'un stable où le débit reste toujours en dessous de la capacité disponible,
- l'autre oscillatoire où débit varie et utilise en moyenne moins bien la bande passante que l'état stable.

Cette turbulence est étudiée pour un routeur sans buffer. Une question est de savoir si cette double dynamique peut exister dans la réalité. Une première étape est de voir si l'effet de pertes anticipées discuté dans l'article [129] élimine totalement la possibilité d'observer des turbulences. Nous étudions dans cette annexe l'effet de ces pertes (dans le pire des cas) sur la dynamique de TCP. Nous ne discutons pas ici de la validité du modèle de $M/D/1/B$ pour étudier l'arrivée des paquets dans une file d'attente et les fluctuations se reporter à [129] ou [95] pour plus de précisions sur la validité du modèle.

A.1 Pourquoi une file d'attente ?

Sans file d'attente ou bien lorsque la file d'attente est trop petite, le trafic souffre du manque de fluidité de TCP et des arrivées simultanées de paquets au routeur. Notre simulation ne le verra pas mais il y aura effectivement des pertes anticipées. Ce problème est traité par Damon Wischik et Gaurav Raina [129]. Nous allons brièvement reprendre le principe présenté car nous avons eu les pires difficultés à implémenter la solution proposée. En effet l'article de Kuusela & Al [95] qui est mis en référence par [129] fait lui-même appel à une troisième référence qui donne une formule grossièrement fausse sur les probabilités d'occupation de la $M/D/1/B$. En recherchant dans la littérature nous avons d'ailleurs trouvé beaucoup de formules fausses, avec des probabilités supérieures à 1, des divisions par 0 pour ne citer que cela. Finalement nous avons trouvé seulement deux références

avec un résultat sans erreurs [136, 91]. Le principe lorsque la file d'attente est petite est de l'assimiler à une $M/D/1/B$ en régime stationnaire dont on connaît le taux de pertes de la forme $L_B(\frac{\text{Throughput}}{C})$. La fonction L_B dépend du nombre B de paquets que peut contenir la file d'attente et vaut lorsque l'utilisation est plus petite que 1 :

$$L_B(\rho) = 1 - \frac{1}{\frac{\pi_0}{\Pi_B} + 1 - \pi_0} \quad (\text{A.1})$$

$$\pi_0 = e^{-\rho} \quad (\text{A.2})$$

$$\Pi_B = \sum_{n=1}^B \frac{(n\rho)^{n-1}}{n!} e^{-n\rho}. \quad (\text{A.3})$$

Attention la taille du tampon à considérer ici est bien la vraie taille et pas la version normalisée par le nombre d'utilisateurs. Soulignons que cette formule est exacte et pas une approximation.

Nous allons rappeler quelques propriétés relatives aux formules ci-dessus.

Proposition A.1.

- Π_B converge quand $B \rightarrow \infty$ de manière uniforme pour $\rho \in \mathbb{R}^+$,
- $\Pi_B \rightarrow 1$ quand $B \rightarrow \infty$ uniformément sur $[0, 1]$,
- $L_B(\rho) \rightarrow 0$ uniformément quand $B \rightarrow \infty$,
- $\sup_{\rho \in [0,1]} (L_B(\rho)) = L_B(1)$,
- $L_B(1) \sim_{B \rightarrow \infty} \frac{\sqrt{2}}{e\sqrt{\pi}} \frac{1}{\sqrt{B}} = \frac{\alpha_2}{\sqrt{B}}$ où $\alpha_2 \approx 0.2935$ est une constante.

Preuve : Soit $\Pi_\infty(\rho) = \sum_{n=1}^\infty \frac{(n\rho)^{n-1}}{n!} e^{-n\rho}$. Une application de la formule de Stirling permet de dire que cette série est sommable pour tout $\rho \geq 0$, on a l'équivalent du terme général :

$$\frac{(n\rho)^{n-1}}{n!} e^{-n\rho} \sim \frac{(n\rho)^{n-1}}{\sqrt{2\pi n} (n/e)^n} e^{-n\rho} \quad (\text{A.4})$$

$$\sim \frac{e^{1-\rho}}{n\sqrt{2\pi n}} e^{(n-1)(1-\rho+\log(\rho))} \text{ pour } \rho > 0. \quad (\text{A.5})$$

Le terme dans l'exponentielle est toujours négatif ou nul (par concavité du logarithme comparé à sa tangente en 1). Ainsi le terme général de la série appartient à $O(n^{-3/2})$ qui est le terme général d'une série sommable. De plus pour $\rho = 0$ on voit tout de suite que $\Pi_\infty(0) = 1$. Notons au passage qu'en tant que série de fonction on peut prendre le sup sur ρ avant de sommer, ce qui signifie en particulier que la convergence est uniforme par rapport à ρ .

Nous savons que Π_∞ est la somme d'une distribution stationnaire, c'est à dire d'une probabilité pour $\rho \in [0, 1]$ (voir [91]). Comme on a une limite uniforme sur $\mathbb{R}^+$ à fortiori sur $[0, 1]$ de fonctions continues, on voit que la deuxième assertion est aussi vérifiée par continuité de la limite.

Lorsque $x \geq 0$, la convexité de l'hyperbole permet de dire que $1 - \frac{1}{1+x} \leq x$, ainsi :

$$L_B(\rho) \leq e^{-\rho} \left(\frac{1}{\Pi_B} - 1 \right) \leq \frac{1}{\Pi_B} - 1,$$

d'où la convergence uniforme.

Ensuite l'assertion $\sup_{\rho \in [0,1]} (L_B(\rho)) = L_B(1)$ est aussi une conséquence directe du fait qu'on étudie une distribution stationnaire avec utilisation ρ , le pire cas est quand ρ est maximal. De plus il faut toujours utiliser la continuité des quantités considérées pour majorer par la valeur en 1 qui n'a pas d'interprétation à priori.

Enfin le dernier équivalent ne pose pas de problème particulier. ■

Théorème A.1. *En régime stationnaire, pour des utilisateurs de HTTP ou de FTP (TCP long lived) dans l'hypothèse AIMD, on suppose que le routeur considéré est en utilisation maximale. Notons $B \gg 1$, la taille de son buffer, aucun mécanisme de pertes anticipées n'étant implémenté. Alors la fenêtre moyenne maximale que chaque utilisateur peut avoir vaut environ :*

$$W_{max}(B) := \alpha \sqrt{\frac{\sqrt{B}}{\alpha_2}} \approx 2.418 B^{1/4}$$

où $\alpha = 1.310$ est la constante de la distribution stationnaire des fenêtres du chapitre 10 et $\alpha_2 = 0.2935$ est définie dans la proposition précédente.

Par conséquent si son RTT a pour valeur r , le débit moyen maximal auquel chaque utilisateur peut prétendre en congestion avoidance vaut :

$$D_{max} := \frac{W_{max}(B)}{r} = \frac{\alpha}{r} \sqrt{\frac{\sqrt{B}}{\alpha_2}}.$$

Ce théorème est une conséquence directe de notre "square-root formula" du chapitre 10³² et de la propriété précédente. Le "vaut environ" du théorème doit être entendu comme "est équivalent lorsque B tend vers l'infini à".

A.2 Limite de fonctionnement des files d'attente face à TCP

Nous allons étudier les conséquences pratiques du théorème précédent. Dans le chapitre 11 nous avons vu que sans perte, HTTP pouvait se stabiliser. Cependant, cette stabilisation ne peut se faire dans les conditions prévues que si la file d'attente est grande ou l'utilisation assez faible.

Application numérique On suppose que le temps de propagation est $T = .25s$, sur un routeur congestionné par des utilisateurs de HTTP, son débit est de $45Mbit/s$ soit pour fixer les idées de $10000paquets/s$. Alors la règle habituelle du "Bandwidth Delay product" dit que son buffer doit avoir pour taille $2500paquets$.

On se demande si le routeur fera bien son travail d'accorder une bande passante maximale à 50 utilisateurs, ce qui donnerait une bande passante de $0.9Mbit$ par utilisateur,

³²rappelons que la formule est valable à fortiori pour des utilisateurs de HTTP quand on se place uniquement sous la phase d'évitement de congestion

soit 200 paquets par seconde et une fenêtre d'au moins 50 paquets (qui peut aller jusqu'à 100 si la file d'attente se remplit). Le calcul précédent indique que le taux de pertes dû aux effets probabilistes de dépassement de la file ne permettra pas une fenêtre moyenne supérieure à $W_{max} = 17$ avec utilisation maximale, ce qui est incompatible avec la fenêtre recherchée.

Conséquence pratique Ceci indique que dans le cas précédent, il n'existe pas de régime stable qui utilise toutes les capacités disponibles dans le cas où les seules pertes existantes proviennent des fluctuations. Par contre il ne faut pas conclure que le débit sera aussi faible qu'on le dit ici, puisque le taux de perte calculé est donné en supposant une utilisation maximale. Ce taux de pertes est quasiment nul si l'utilisation est plus faible, et on aura un régime oscillatoire pour HTTP.

Par contre comme l'expérience le confirme sur certains cas on peut espérer stabiliser les versions actuelles de TCP avec une utilisation maximale des routeurs congestionnés en utilisant des mécanismes d'AQM. En effet le calcul précédent ne donne rien sur l'évolution dynamique. De plus les fluctuations sont entendues cumulées et à une échelle de temps relativement longue par rapport au RTT et tiennent compte du fait que les seules pertes sont dues au dépassement de la taille du buffer.

Le cas des fluctuations Nous allons tirer de ceci une règle de dimensionnement que les routeurs devraient satisfaire pour ne pas subir ces pertes anticipées. Sur la figure A.1, on peut voir qu'un buffer trop petit de $B = 100$ paquets peut provoquer des pertes importantes dès 80% d'utilisation. Comme [129] le fait remarquer ceci peut mener à une stabilisation de TCP, mais avec une utilisation sous-optimale, ce qui n'est pas nécessairement mauvais si la bande passante n'a pas beaucoup de valeur et que le temps de transmission est ce qu'on veut optimiser. On peut par exemple souhaiter configurer son routeur d'accès à Internet pour assurer de bonnes connexions pour les jeux ou la voix sur IP qui réclament de faibles temps de latence en choisissant un buffer petit quitte à avoir des pertes et à pénaliser les téléchargements de fichiers dès que l'utilisation commence à être importante.

Dans la figure A.2, on voit que le taux de pertes reste extrêmement faible tant que l'utilisation ne se rapproche pas très près de 1.

A.3 Dimensionnement minimal du buffer d'un routeur

Première application numérique Voici une application numérique qui peut donner une idée de la méthode à utiliser pour dimensionner un routeur. Le scénario est assez peu réaliste, mais ce n'est pas le point important. Un utilisateur en voyage à Paris (en France) veut regarder depuis son bureau une vidéo en qualité DVD (720×480) en encodage MPEG4/MP3 qui se trouve sur le site d'un de ses amis hébergé à Sidney en Australie. L'ami en question a rendu la vidéo accessible sur le port 443 en sFTP, ce qui permet de la faire passer sur le port HTTP sécurisé à travers n'importe quel firewall (ceci explique pourquoi notre utilisateur utilise TCP plutôt qu'UDP). Notre utilisateur espère bien réussir à regarder sa vidéo en même temps qu'il la télécharge, en effet il veut en

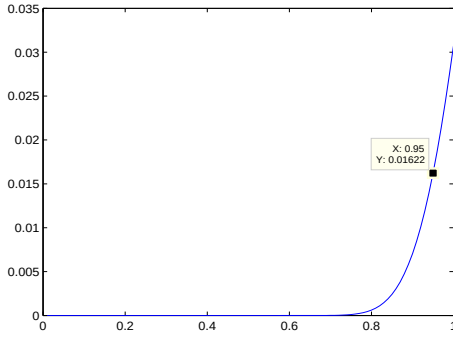


FIG. A.1 – Pertes engendrées par $B = 100$ paquets en fonction de ρ dans un régime stationnaire de M/D/1

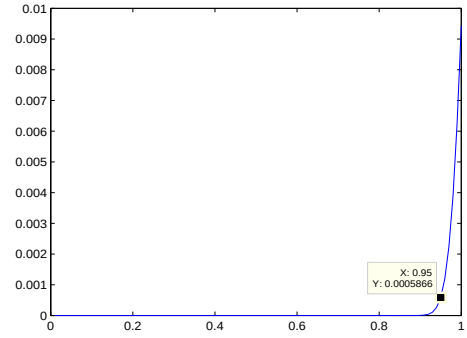


FIG. A.2 – Pertes engendrées par $B = 1000$ paquets en fonction de ρ dans un régime stationnaire de M/D/1

parler avec d'autres amis avec qui il a rendez-vous dès la sortie du bureau. Il est déjà tard et la vidéo dure 45 minutes.

Le RTT entre l'utilisateur et son serveur est d'environ $250ms$ en tenant compte de temps d'attentes divers et du temps de transmission à la vitesse de la lumière. Le débit à assurer pour la qualité demandée est de l'ordre de 1 Mbits/s, soit de l'ordre de 240 paquets par seconde, donc une fenêtre de l'ordre de 60 paquets. Information supplémentaire : cet après-midi le serveur est très populaire à Paris à cause d'un mail envoyé à une large liste de diffusion. Aussi on peut estimer que 18 utilisateurs se partagent les 20 Mbits/s d'upload (débit ascendant) disponible au niveau du serveur. On estime que chaque utilisateur va télécharger à vitesse maximale et mettre en cache ce qu'il a téléchargé d'avance.

Par la formule de la racine carrée, on sait que cette fenêtre ne peut être atteinte en moyenne que pour un taux de pertes inférieur à 5×10^{-4} pour pouvoir regarder la vidéo en même temps qu'on la télécharge. D'après les données, on doit assurer un débit moyen de 18 Mbits/s sur 20 disponibles, soit une utilisation cible supérieure à 90%. On regarde donc la courbe donnant le taux de pertes en fonction de la taille du buffer pour les utilisations $\rho = 92,5\%$ et $\rho = 95\%$ dans les figures (A.3) et (A.4). On voit qu'une faible augmentation dans l'objectif d'utilisation (ou bien un utilisateur supplémentaire) a des conséquences importantes sur la taille du buffer que l'on devrait choisir pour espérer se stabiliser, on passe de moins de 600 paquets à plus de 1000 pour un faible gain d'utilisation.

Une limite basse pour la taille des buffers Pour avoir une idée plus générale regardons ce que nous pouvons conclure du graphe 3D (figure A.5) lorsque à la fois B et ρ varient. Dans la pratique on est intéressé par un choix intelligent de B . Aussi on constate une forte inflexion aux alentours de $B = 300$ paquets. Aussi on peut conclure du point de vue général (indépendamment de l'usage d'une version particulière de TCP) qu'un choix de taille de buffer inférieur à la barrière 300-400 paquets ne sera pas très judicieux à moins d'avoir un objectif de délai très serré.

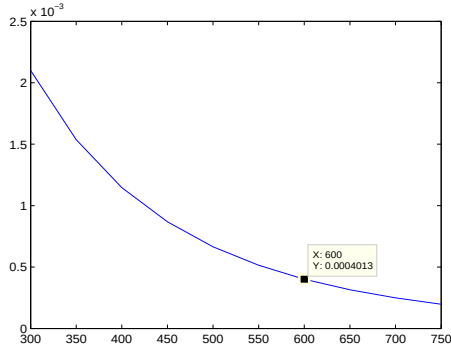


FIG. A.3 – Pertes engendrées par $\rho = 92.5\%$ fonction de B dans un régime stationnaire de $M/D/1$

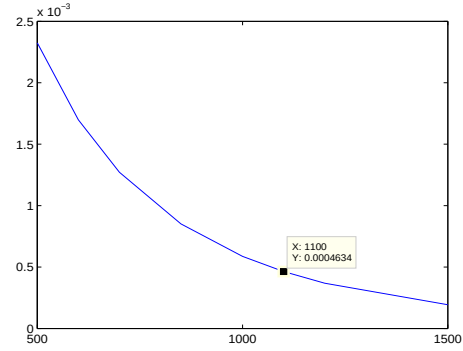


FIG. A.4 – Pertes engendrées par $\rho = 95\%$ en fonction de B dans un régime stationnaire de $M/D/1$

Que se passe-t-il pour une utilisation plus grande que 1 ? Dans les fluctuations de TCP en tail drop pour un routeur congestionné, nous avons déjà constaté un mouvement assez périodique où le débit augmente jusqu'à dépasser la capacité disponible, puis la taille de la file d'attente se met à augmenter, enfin on touche la taille maximale pendant un temps et les sources réagissent, en envoyant un débit bien moindre. Pendant cette durée, la file se vide progressivement jusqu'à 0 dans certains cas et on reprend. Pendant une certaine durée on a une utilisation plus grande que 1. Notons bien qu'on n'est précisément pas en régime permanent et que ce que nous avons dit précédemment n'a de sens que pour les fluctuations de durée propre très rapide devant l'augmentation du débit. Jusqu'à présent nous avons dit que les pertes avaient lieu comme l'indiquent les figures (A.6,A.7) avec un buffer très petit par rapport au nombre d'utilisateurs : lorsqu'un débit D_i entre avec une capacité totale C , on tue les paquets en plus de la capacité, c'est à dire $D_i - C \vee 0$, sur le total de D_i , soit un taux de pertes en notant $\rho = \frac{D_i}{C}$ qui vaut

$$0 \vee 1 - \frac{1}{\rho}. \quad (\text{A.6})$$

On voit ainsi que le recollement n'est pas très difficile, on peut prendre le max de $\rho = 1$ pour $M/D/1$ et le taux de la formule (A.6) sans faire une trop grande discontinuité. En plus il faut tenir compte du fait que la $M/D/1$ devient un modèle particulièrement mauvais lorsque l'utilisation est proche de 1. On peut donc avoir certains doutes sur la validité de la prédiction de pertes anticipées dans ce cas. Cependant nous restons dans le pire des cas pour être certains de nous placer dans un cas qui ne pose pas de problèmes du point de vue de notre modélisation.

Les pertes $M/D/1$ et TCP Comme il est démontré dans l'article [129], nous retiendrons que si on prend un buffer d'une taille de l'ordre de 1000 paquets, l'effet d'époque de congestion sera parfaitement vérifié dès lors que les fenêtres considérées seront pour la plupart plus petites que 20. En effet si on regarde la figure (A.7), et en tenant compte du fait que les pertes petites devant 0.5% n'ont pas d'effet sur des fenêtres de cette taille, on

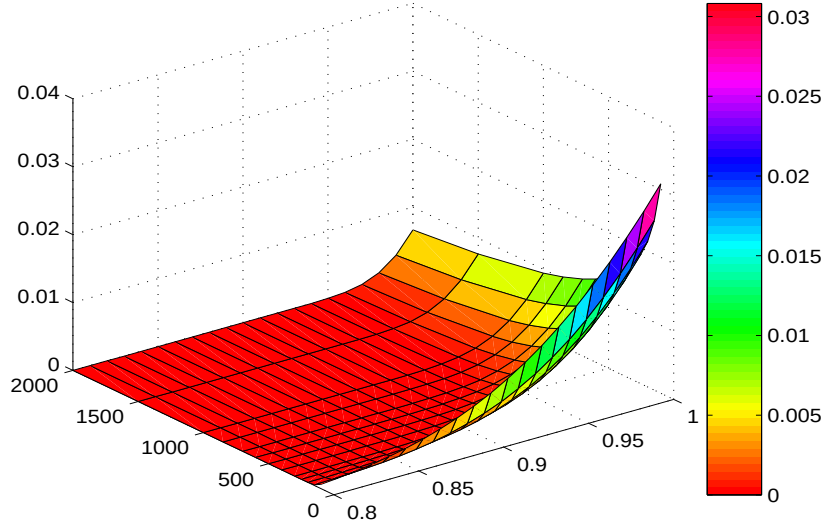


FIG. A.5 – Pertes engendrées $\rho = 80..100\%$ en fonction de $B = 100..2000$ dans un régime stationnaire de M/D/1. Les tailles de buffer qui correspondent aux lignes valent [100, 150, 200, 250, 300, 400, 500, 600, 700, 850, 1000, 1500, 2000] paquets.

voit qu'on doit s'inquiéter à partir de 98% d'utilisation environ. Comparons cela à l'augmentation de débit pendant un RTT en l'absence de pertes ; nous avons avant période de congestion :

$$N \frac{W}{R} \approx C,$$

et un RTT plus tard le débit entrant vaut :

$$N \frac{W+1}{R} \approx C(1 + \frac{1}{W}).$$

En un RTT, l'augmentation d'utilisation est donc de l'ordre de 5% pour $W = 20$, certes dans le cas de HTTP, la croissance peut être moindre compte tenu des fins de téléchargement ; cependant nous noterons que dans les cas de turbulences le point important est que le nombre d'utilisateurs actifs augmente dans le cas instable, et que le débit est mal utilisé à cause des fluctuations, ce qui donne deux raisons d'une augmentation assez dynamique du débit en phase ascendante.

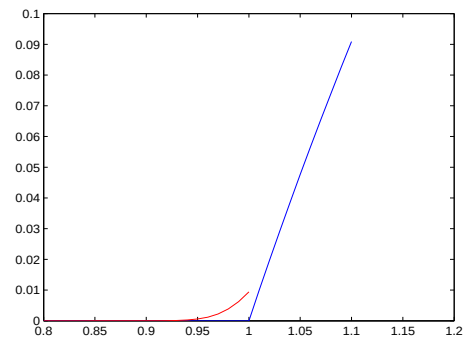
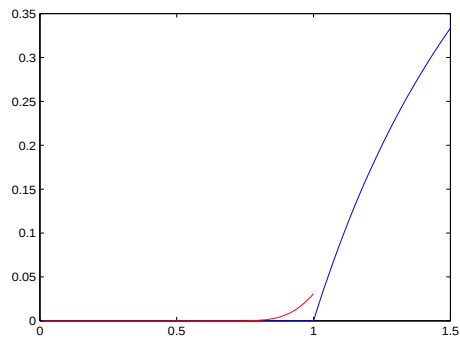


FIG. A.6 – Pertes engendrées par $B = 100$ paquets en fonction de B dans un régime stationnaire de M/D/1 (rouge) et prolongement sans pertes anticipées (bleu).

FIG. A.7 – Pertes engendrées par $B = 1000$ paquets en fonction de B dans un régime stationnaire de M/D/1 (rouge) et prolongement sans pertes anticipées (bleu).

Annexe B

Quelques théorèmes de mathématiques générales

B.1 Lemmes de Gronwall

La présentation ici est inspirée de [144]. Il n'y a pas un mais des lemmes de Grönwall. Dans ce cas on a confirmation qu'un bon lemme est souvent beaucoup plus important que la plupart des théorèmes. La première version du lemme date de 1918 est due au mathématicien suédois Thomas Hakon Grönwall. Les applications à la théorie des équation différentielles dont il est aujourd'hui un des résultats fondateurs après les théorèmes d'existence et d'unicité de solutions sont dues à Richard Bellman.

Le résultat en soi est très simple puisqu'il s'agit de tirer d'une inégalité différentielle dont le cas d'égalité correspondrait à une exponentielle une inégalité sur la solution. Par contre l'utilisation de ce lemme est fort utile et évite bien souvent des raisonnements de connexité avec les histoires d'ensembles ouverts et fermés sur lesquels une certaine égalité est vérifiée. En gros intuitivement on doit voir ce lemme comme une borne sur l'erreur de la solution d'une équation différentielle connaissant une borne sur l'erreur.

Théorème B.2. *Pour une fonction f réelle, supposons que sur $[0, a]$ on ait l'estimé :*

$$f(t) \leq A + B \int_0^t f(s)g(s)ds$$

avec

- A, B positifs (ou nuls)
 - f et g continues à valeurs positives
- alors pour tout t dans l'intervalle $[0, a]$

$$0 \leq f(t) \leq Ae^{B \int_0^t g(s)ds} \tag{B.1}$$

Preuve : Dans le cas où $A > 0$,

Pour $0 \leq r, t \leq a$, si l'on divise les deux membres de l'estimé par la valeur strictement positive $A + B \int_0^r f(s)g(s)ds$, il vient

$$\frac{f(r)}{A + B \int_0^r f(s)g(s)ds} \leq 1$$

en multipliant par la valeur positive ou nulle $Bg(r)$ et en intégrant entre 0 et t , on obtient :

$$\int_0^t \frac{Bf(r)g(r)}{A + B \int_0^r f(s)g(s)ds} dr \leq B \int_0^t g(r)dr$$

mais comme les fonctions considérées sont continues, la fonction du dénominateur de gauche est continûment dérivable en r de dérivée le numérateur de gauche. Il vient donc :

$$\log \left(A + B \int_0^t f(s)g(s)ds \right) - \log A \leq B \int_0^t g(r)dr$$

Il ne reste plus qu'à prendre l'exponentielle de cette inégalité en rappelant que la fonction exponentielle est croissante, et à identifier la fonction majorante de f dans l'estimé pour obtenir

$$f(t) \leq Ae^{B \int_0^t g(s)ds}$$

qui est le résultat que nous recherchions.

Dans le cas $A = 0$ on a la majoration qui est valable pour tous les $\tilde{A} > 0$ ce qui fait que

$$f(t) \leq \inf_{\tilde{A} \in \mathbb{R}_*^+} \tilde{A} e^{B \int_0^t g(s)ds} = 0$$

■

B.2 Autour des lemmes de Dini et son corollaire probabiliste

Un premier lemme de Dini amélioré

Les lemmes de Dini assurent une convergence uniforme à partir d'une convergence point à point par un argument de compacité et une hypothèse sur les fonctions, par exemple suite croissante de fonctions, suite de fonctions chacune croissante, suite de fonctions lipschitziennes. Nous allons rappeler son argument dans le cas de fonctions uniformément semi-lipschitziennes supérieurement, c'est à dire dont le taux de croissance reste uniformément borné supérieurement et dont la limite est continue. Cette hypothèse signifie que les fonctions ont une croissance contrôlée uniformément, mais qu'on peut décroître aussi vite qu'on le veut y compris en faisant des sauts. Nous allons présenter le théorème dans $\mathbb{R}$, pour avoir le droit de faire le taux d'accroissement sans norme, ce qui simplifie l'énoncé et la démonstration. Une version un peu plus générale suit, elle donne les idées pour toutes sortes de généralisations ad hoc.

Théorème B.3. *Soit une suite f_n de fonctions d'un connexe compact de $\mathbb{R}$ (ie : un segment) dans $\mathbb{R}$ qui convergent ps vers une fonction f qui est continue. On suppose de plus que le taux de croissance $\frac{f_p(x) - f_p(y)}{x - y}$ reste uniformément borné supérieurement par M par rapport à x , y et p .*

Alors la suite f_n converge vers f ps uniformément (par rapport à la mesure de Lebesgue).

Preuve : Pour commencer remarquons que l'inégalité de majoration du taux d'accroissement sur les f_n se propage à f .

Si on se fixe une $\epsilon > 0$, soit $\epsilon' = \dots$ (on le choisit après), on peut recouvrir le compact par un nombre fini boules de rayon ϵ' . Donc on peut dériver un recouvrement par des boules de rayon $2\epsilon'$ dont le centre est un point de convergence de la suite (on a aucun mal puisque toutes les boules de rayon $2\epsilon'$ dont le centre est dans celle de rayon ϵ' conviennent). Doublons encore le rayon de nos boules et on a de plus la propriété que le centre de chaque boule est dans les boules précédente et suivante.

Notons $x_1 \dots x_{k-1}$ les centres de ces boules dans l'ordre croissant on rajoute les boules extrêmes de centres x_0 et x_k

Soit x un point de convergence de f_n vers f . Alors il existe i tel que $x_i \leq x \leq x_{i+1}$, ainsi $f(x) - f(x_i) \leq M(x - x_i)$ et $f(x_{i+1}) - f(x) \leq M(x_{i+1} - x)$ (de même pour les f_n). De plus $|x - x_i| \leq 4\epsilon'$ et $|x_{i+1} - x| < 4\epsilon'$ ce qui encadre $f(x)$ entre $f(x_{i+1}) - 4\epsilon'M$ et $f(x_i) + 4\epsilon'M$ (de même pour f_n). De même $f(x_{i+1}) \leq f(x_i) + 4\epsilon'M$ et

Ainsi en découpant correctement on prouve que

$$|f(x) - f_n(x)| \leq 8\epsilon'M + |f(x_i) - f_n(x_i)| + |f(x_{i+1}) - f_n(x_{i+1})| + |f(x_i) - f(x_{i+1})|$$

comme on a la convergence point a point, sur tous les x_i , il existe un N tel que pour tout $n \geq N$ on ait la convergence à ces points à ϵ' près. De plus le théorème de Heine pour la fonction f qui est continue sur un compact donc uniformément continue donne : $\exists \eta_\epsilon, |x - y| < \eta_\epsilon \Rightarrow |f(x) - f(y)| < \epsilon/2$. D'où il vient en choisissant $\epsilon' = \min(\frac{\epsilon}{2 \times (8M+2)}, \eta_\epsilon)$ la propriété de convergence uniforme recherchée.

■

Corollaire B.1. NOTATIONS : Soit X_t^N une suite processus d'espace de temps un segment I de $\mathbb{R}$ à valeurs dans $\mathbb{R}$. Soit $f(t)$, une fonction (déterministe) de $\mathbb{R}$ dans $\mathbb{R}$

HYPOTHÈSES : pour tout t , $X_t^N \rightarrow f(t)$ en probabilité, f est continue. On suppose de plus que le taux de croissance $\frac{X_x^p - X_y^p}{x-y}$ reste uniformément borné supérieurement par M par rapport à x, y, p et ω .

CONCLUSION :

$$\sup_{t \in I} |X^N(t) - f(t)| \rightarrow 0 \text{ en probabilité.}$$

Preuve : Nous sommes dans un cas plus simple que le théorème précédent pour le début car la convergence a lieu point à point et pas presque sûrement, ce qui permet d'extraire la suite de points x_i qui donne (à fortiori) pour tout temps $x \in I$:

$$|f(x) - X_x^n| \leq 8\epsilon'M + |f(x_i) - X_{x_i}^n| + |f(x_{i+1}) - X_{x_{i+1}}^n| + |f(x_i) - f(x_{i+1})|.$$

À l'argument d'uniforme continuité de f qui reste valable, il faut ajouter que la séquence finie de variables aléatoires $X_{x_i}^N$ converge en probabilité. Ce qui signifie que pour le même choix de ϵ' , avec une probabilité aussi grande que l'on veut (quand $N \rightarrow \infty$), $|f(x) - X_x^n| < \epsilon$. C'est-à-dire exactement la conclusion souhaitée.

■

B.2.1 Deuxième généralisation : sur les fonctions

Les généralisations qui suivent ne sont pas utilisées dans le manuscrit mais sont des curiosités mathématiques intéressantes. Nous n'écrivons pas les généralisations probabilistes, mais le raisonnement serait le même que dans le corollaire B.1.

Théorème B.4. *HYPOTHÈSES Soit une suite f_n de fonctions d'un compact convexe K de $\mathbb{R}^N$ dans un espace de Hilbert H qui convergent ps vers une fonction continue f .*

On se fixe un vecteur non nul R de $\mathbb{R}^N$. Soit $C = \{x \in \mathbb{R}^N, \langle R, x \rangle \geq \|x\|\}$, le cône de révolution convexe ouvert de sommet 0 de direction R et d'angle au sommet $2\alpha = 2\text{Arccos}(\frac{1}{\|R\|}) > 0$.

On fixe dans H un vecteur non nul S , et on définit de même le cône fermé de sommet 0 d'angle $2\beta < \pi$: $D = \{y \in H, \langle S, y \rangle \leq \|y\|\}$

On suppose que les variations de f vérifient $\{f_p(x) - f_p(x'), x' - x \in C, p \in \mathbb{N}\} \subset D$.

CONCLUSION *Alors la suite f_n converge vers f ps uniformément (par rapport à Lebesgue) sur tout compact K_{int} de l'intérieur K .*

Remarque B.2. *La continuité de la limite et la définition sur un compact connexe sont les ingrédients indispensables. Ensuite on doit supposer que le taux d'accroissement est majoré uniformément à partir de tout point x dans une partie $x + A$, et que l'union des intérieurs des ensembles $x + A$ représente un recouvrement ouvert du compact. Une hypothèse simple sur A peut être de dire qu'il s'agit d'un cône ouvert.*

Remarque B.3. *Fondamentalement ce qui fait fonctionner la démonstration est constitué des arguments suivants :*

- l'ensemble de départ des compact
- quand on part de deux points "proches" x et y , on peut déduire un ouvert O de l'ensemble de départ sur lequel l'ensemble $\{f_p(t), p \geq \phi(n), t \in O\}$ est dans un ensemble de dimension bornée par une fonctionnelle des quatre points $f(x), f(y), f_n(x)$ et $f_n(y)$.
- on a moyen de s'assurer que la borne à l'arrivée devienne petite

Preuve : Nous allons procéder en trois étapes

- expliciter l'uniforme continuité de la limite
- recouvrir l'ensemble de départ par des cônes tronqués bien agencés
- remarquer que les deux premiers arguments permettent de confiner l'image de f dans des ensembles de petits diamètres

PREMIER POINT Soit $\epsilon > 0$ par uniforme continuité de la limite il existe une constante η_ϵ pour laquelle $\|x - x'\| < \eta_\epsilon \Rightarrow |f(x) - f(x')| < \frac{\epsilon}{2(1+\max(\tan \beta, 1))}$. De plus la distance (continue car 1-lipschitzienne) entre ∂K et K_0 est atteinte, elle vaut donc $\eta > 0$.

DEUXIÈME POINT Soit $r_0 = \min(\eta, \frac{\eta_\epsilon}{1+\cos(2\alpha)\tan(\alpha)/4})$, il faut trouver un recouvrement fini de notre compact K par des $C_{x_i, r} := B(x_i, r) \cap (x_i + C)$ qui ait la propriété selon laquelle pour tout x de K_{int} il existe un sommet de cône tronqué du recouvrement dans C_{x, r_0} et que x soit dans un cône tronqué dont le sommet est dans K . On réalise ce recouvrement ainsi : le diamètre du cône tronqué $C_{\cdot, r}$ est inférieur à $2r$ et il est linéaire par rapport à r (et strictement positif). Donc pour $r < \eta/2$, $x \in K$. Pour le deuxième point remarquons

que C_{x,r_0} est contenu dans K_{int} et contient en particulier l'intérieur la boule inscrite du tronc de cône de révolution, cette boule est $B(a := x + \frac{R}{\|R\|} r_0 \frac{\cos(2\alpha)}{\cos \alpha}, r_1 := r_0 \cos(2\alpha) \tan \alpha)$. Ainsi tout cône tronqué de rayon $r_1/2$ qui recouvre un point de la boule $B(a, r_1/4)$ est entièrement contenu dans $B(a, r_1)$.

Résumons si nous recouvrons K avec des cône tronqués avec $r = \min(r_0/4 \cos(2\alpha) \tan \alpha, \eta)$ dont nous extrayons un recouvrement fini. Le lieu sommets des cônes tronqués au rayon $2r$ qui contiennent un des cônes de rayon r et qui touche K_{int} est de mesure non nulle et entièrement contenu dans K , on peut donc en choisir un où la convergence de la suite est assurée. Notons x_i , les sommets de ces cônes de rayon $2r$. De plus nous savons que pour tout point x de K_{int} il existe deux indices i et j tels que $x \in C(x_i, 2r)$ et $x_j \in C(x, r_0)$.

TROISIÈME POINT on sait que $f_n(x) - f_n(x_i) \in D$ et que $f_n(x_j) - f_n(x) \in D$, en particulier $\|f_n(x) - f_n(x_i)\| \leq \|f_n(x_j) - f_n(x_i)\| \max(1, \tan(\beta))$,

Ainsi

$$\begin{aligned} \|f_n(x) - f(x)\| &\leq \|f_n(x) - f_n(x_i)\| + \|f_n(x_i) - f(x_i)\| + \|f(x_i) - f(x)\| \\ &\leq \|f_n(x_j) - f_n(x_i)\| \max(\tan(\beta), 1) + \|f_n(x_i) - f(x_i)\| + \|f(x_i) - f(x)\| \\ &\leq \|f_n(x_j) - f(x_j)\| (\max(1, \tan(\beta)) + \|f_n(x_i) - f(x_i)\| (1 + \max(1, \tan(\beta)))) \\ &\quad + \|f(x_i) - f(x)\| + \|f(x_i) - f(x)\| \max(\tan(\beta), 1) \\ &\leq \epsilon \end{aligned}$$

la dernière inégalité tient pour n grand tel que les $f_n(x_i)$ valent tous $f(x_i)$ à $\frac{\epsilon}{2+4\max(1, \tan \beta)}$ et en remarquant que les rayons des boules sont choisis pour que $\|x_i - x_j\| < \eta_\epsilon$ ainsi l'inégalité d'uniforme continuité sur f s'applique et permet de conclure.

■

Troisième généralisation : de la convergence en probabilité vers la convergence uniforme

La suite de fonctions que nous définissons a un commun un principe de compacité qui permet de transmettre la convergence en un nombre fini de points à tous les points de manière uniforme. Mais en fait on peut faire mieux, avec le même jeu d'hypothèses exactement, on démontre directement la transmission de la convergence en probabilité vers la convergence uniforme. Ceci est vrai pour toute suite de fonctions dont la classe permet de dire que la convergence presque sûre implique la convergence uniforme par un argument de compacité du type de celui employé ci-dessus. Nous nous bornerons à indiquer une démonstration

On le sait moralement, la convergence en probabilités est presque identique à la convergence presque sûre à part qu'elle est parfois un peu trop lente (on peut toujours accélérer cette convergence en effectuant une extraction). Le seul type de contre-exemples est une suite du type $Y_n = X_n + \chi_{[S_n, S_{n+1}]} \bmod 1$ avec $X_n \rightarrow X$ ps et S_n une suite qui tend vers ∞ tout en vérifiant $S_{n+1} - S_n \rightarrow 0$. On ne doit donc pas être surpris par le fait qu'une hypothèse de régularité ne permette pas de se trouver dans des cas où la convergence est ralentie.

Théorème B.5. *HYPOTHÈSES Soit une suite f_n de fonctions d'un compact convexe K de $\mathbb{R}^N$ dans un espace de Hilbert H qui convergent en probabilité vers une fonction continue f pour la mesure de Lebesgue de $\mathbb{R}^N$.*

On se fixe un vecteur non nul R de $\mathbb{R}^N$. Soit $C = \{x \in \mathbb{R}^N, \langle R, x \rangle > \|x\|\}$, le cône de révolution convexe ouvert de sommet 0 de direction R et d'angle au sommet $2\alpha = 2\text{Arccos}(\frac{1}{\|R\|}) > 0$.

On fixe dans H un vecteur non nul S , et on définit de même le cône fermé de sommet 0 d'angle $2\beta < \pi$: $D = \{y \in H, \langle S, y \rangle \leq \|y\|\}$

On suppose que les variations de f vérifient ps $\{f_p(x) - f_p(x'), x' - x \in C, p \in \mathbb{N}\} \subset D$ (ps pour la mesure de Lebesgue sur l'espace produit).

CONCLUSION *Alors la suite f_n converge vers f ps uniformément (par rapport à Lebesgue) sur tout compact K_{int} de l'intérieur K .*

Preuve : Si par contraposée on a une convergence qui n'est pas ps sur K_{int} , il existe une ensemble de mesure non nulle A inclus dans K_{int} de points où la convergence n'a pas lieu. Comme on a beaucoup de points il en existe au moins un, $x \in A$ qui vérifie l'encadrement pour tous les n . Il existe donc une valeur associée η_x et une extraction ϕ_x telle que $\|f_{\phi_x(n)}(x) - f(x)\| > \eta_x$. Mais alors l'uniforme continuité de f et le fait que pour presque tous les x' qui voient x dans leur cône (ie : $x - x' \in C$) on ait $f_{\phi_x(n)}(x) - f_{\phi_x(n)}(x') \in D$ pour tout n permet de déduire un ensemble B_x de mesure non nulle (l'intersection du miroir de C en x et d'une boule suffisamment petite) telle que la différence entre $f_{\phi_x(n)}(x')$ et $f(x')$ soit au moins $\eta_x/2$ pour presque tout x' dans B_x . Ce qui est en opposition avec la convergence en probabilité puisque $\mathbb{P}(f_{\phi_x(n)} - f > \eta/2) \geq \lambda(B_x) > 0$. Nous ne reviendrons pas sur la généralisation ps dans la convergence uniforme, le raisonnement est évident à mettre en place et ne présente pas d'intérêt particulier dans une démonstration qui est déjà suffisamment complexe.

■

Bibliographie

- [1] C. Adjih, P. Jacquet, and N. Vvedenskaya, *Performance evaluation of a single queue under multi-user tcp/ip connections*, Tech. report, INRIA Research Report number 4478, 6 2002.
- [2] M. Allman, V. Paxson, and W. Stevens, *RFC 2581 - tcp congestion control*, 1999.
- [3] E. Altman and Jimenez T., *Ns simulator course for beginners*, 2002, <http://www-sop.inria.fr/mistral/personnel/Eitan.Altman/COURS-NS/n3.pdf>.
- [4] ———, *Ns simulator course for beginners*, 2002, <http://www-sop.inria.fr/mistral/personnel/Eitan.Altman/ns.htm> consulté en août 2005.
- [5] G. Appenzeller, McKeown N., Sommers J., and Barford P., *Recent results on sizing router buffers*, Network Systems Design Conference (2004).
- [6] J. Aweya, M. Ouellette, Y. M. Delfin, and A. Chapman, *A load adaptive mechanism for buffer management.*, 2000, <http://www.commentcamarche.net> consulté en mai 2005.
- [7] F. Baccelli and T. Bonald, *Window flow control in FIFO networks with cross traffic*, Queuing Syst. Theory Appl. **32** (1999), no. 1-3, 195–231.
- [8] F. Baccelli and P. Bremaud, *Elements of queuing theory*, second ed., Springer-Verlag, 2003.
- [9] F. Baccelli, A. Chaintreau, D. De Vleeschauwer, and D. R. McDonald, *A mean-field analysis of short lived interacting tcp flows*, SIGMETRICS 2004/PERFORMANCE 2004 : Proceedings of the joint international conference on Measurement and modeling of computer systems, ACM Press, 2004, pp. 343–354.
- [10] F. Baccelli and D. Hong, *Tcp is max-plus linear and what it tells us on its throughput*, SIGCOMM '00 : Proceedings of the conference on Applications, Technologies, Architectures, and Protocols for Computer Communication, ACM Press, 2000, pp. 219–230.
- [11] ———, *Aimd, fairness and fractal scaling of tcp traffic.*, in Proceedings of the Conference of the IEEE Computer and Communications Societies, 2002.
- [12] ———, *Flow level simulation of large ip networks.*, in Proceedings of the Conference of the IEEE Computer and Communications Societies, 2003.
- [13] ———, *Interaction of tcp flows as billiards.*, in Proceedings of the Conference of the IEEE Computer and Communications Societies, 2003.
- [14] F. Baccelli and K. B. Kim, *TCP throughput analysis under transmission error and congestion losses*, Proc. of IEEE Infocom (2004).

- [15] F. Baccelli, K. B. Kim, and D. De Vleeschauwer, *Analysis of the competition between wired, DSL and wireless users in an access network*, Proc. of IEEE Infocom (2005).
- [16] F. Baccelli, K. B. Kim, and D. R. McDonald, *Equilibria of a class of transport equations arising in congestion control*, Rapport de Recherche INRIA (2005).
- [17] F. Baccelli and D. McDonald, *A square root formula for the rate of non-persistent tcp flows*, Tech. report, INRIA Research Report number 5301, 8 2004.
- [18] F. Baccelli, D. R. McDonald, and J. Reynier, *A mean-field model for multiple tcp connections through a buffer implementing RED*, Perform. Eval. **49** (2002), no. 1-4, 77–97, <http://www.eleves.ens.fr/home/jreynier/Recherche/Performance2002-final.pdf>.
- [19] P. Bachy, *Histoire du web*, <http://users.skynet.be/pierre.bachy/histoireweb.html> consulté en juin 2005.
- [20] A. Bain, *Fluid limits for congestion control in networks*, Ph.D. thesis, UL : Order in Manuscripts Room (Not borrowable), Classmark : PhD.27674, 2004.
- [21] H. Balakrishnan, S. Padmanabhan, V. N. Seshan, and R. H. Katz, *A comparison of mechanisms for improving TCP performance over wireless links*, IEEE/ACM Transactions on Networking (1997), no. 5, 756–769.
- [22] S. Ben Fredj, T. Bonald, A. Proutière, G. Régnié, and J. W. Roberts, *Statistical bandwidth sharing : a study of congestion at flow level*, SIGCOMM '01 : Proceedings of the 2001 conference on Applications, technologies, architectures, and protocols for computer communications, ACM Press, 2001, pp. 111–122.
- [23] P. Billingsley, *Convergence of probability measures.*, Wiley, 1968.
- [24] T. Bonald and L. Massoulié, *Impact of fairness on internet performance*, SIGMETRICS '01 : Proceedings of the 2001 ACM SIGMETRICS international conference on Measurement and modeling of computer systems, ACM Press, 2001, pp. 82–91.
- [25] R. Braden, RFC 1122 - requirements for internet hosts – communication layers, 1989.
- [26] L. S. Brakmo and L. L. Peterson, *TCP vegas : end to end congestion avoidance on a global internet*, IEEE Journal on Selected Areas in Communications (1995), no. 8, <http://netweb.usc.edu/yaxu/Vegas/Reference/brakmo.ps>.
- [27] P. Brémaud, *Point processes and queues : Martingale dynamics.*, Springer Verlag, 1981.
- [28] Bursenberg and Martelli, *Delay differential equations and dynamical systems*, Springer.
- [29] A. Chaintreau and D. De Vleeschauwer, *A closed form formula for long-lived tcp connections throughput*, Perform. Eval. **49** (2002), no. 1-4, 57–76.
- [30] C-S. Chang and Z. Liu, *A bandwidth sharing theory for a large number of HTTP-like connections*, IEEE/ACM Trans. Netw. **12** (2004), no. 5, 952–962.
- [31] M. Christiansen, K. Jeffay, D. Ott, and F. D. Smith, *Tuning RED for web traffic*, Proc. of ACM/SIGCOMM (2000).

-
- [32] M. Crovella and A. Bestavros, *Self-similarity in www traffic : Evidence and possible cause*, Proc. of ACM Sigmetrics (1996).
- [33] D. Dacunha-Castelle and M. Duflo, *Probabilités et statistiques.*, 2nd ed., Masson, Paris, 1993.
- [34] DARPA, *Comment ça marche ? encyclopédie informatique*, <http://www.isi.edu/nsnam/ns/>.
- [35] D. A. Dawson, *Critical dynamics and fluctuations for a mean-field model of cooperative behavior.*, J. Statistical Phys. (1983), no. 31, 29–85.
- [36] ———, *Measure valued markov processes.*, no. Vol. XXI, Lecture Notes in Math., **1541**, P. L. Hennequin (ed.), Springer Verlag, Berlin., 1991.
- [37] S. Deb and R. Srikant, *Rate-based versus queue-based models of congestion control*, 2004.
- [38] Zeitouni Dembo, *Large deviations techniques*, Jones and Bartlett Publishers, 1993.
- [39] X. Deng, S. Yi, G. Kesidis, and Das C., *A control theoretic approach in designing adaptive aqm schemes.*, Globecom 2003.
- [40] A. Desolneux, L. Moisan, and J. M. Morel, *Seeing, thinking and knowing, chapter gestalt theory and computer vision*, Arturo Carseti ed., 2004.
- [41] P. Donnelly and T. Kurtz, *Particle representations for measure-valued population models.*, Ann. Probab. (1999), no. 27, 166–205.
- [42] R. M. Dudley, *Real analysis and probability*, 2nd ed., Wadsworth and Brooks, Pacific Grove, 2002.
- [43] V. Dumas, F. Guillemin, and P. Robert, *A markovian analysis of additive- increase multiplicative-decrease*, Advances in Applied Probability **34** (2002), no. 1, 85–111.
- [44] S. N. Ethier and T. G. Kurtz, *Markov processes : Characterization and convergence.*, Wiley, 1986.
- [45] W. Feller, *An introduction to probability theory and its applications*.
- [46] W. Feng, D. Kandlur, D. Saha, and K. Shin, *A self-configuring RED gateway*, Proc. of IEEE Infocom (1999).
- [47] S. Floyd, *Connections with multiple congested gateways in packet switched networks, part 1 : One way traffic.*, ACM Computer Communications Review **21** (1991), no. 5, 30–47.
- [48] ———, *Tcp and explicit congestion notification*, SIGCOMM Comput. Commun. Rev. **24** (1994), no. 5, 8–23.
- [49] S. Floyd and T. Henderson, *RFC 2582 - the newreno modification to tcp's fast recovery algorithm*, 1999.
- [50] S. Floyd, T. Henderson, and A. Gurtov, *RFC 3782 - the newreno modification to tcp's fast recovery algorithm*, 2004.
- [51] S. Floyd and V. Jacobson, *Random early detection gateways for congestion avoidance*, IEEE/ACM Trans. Netw. **1** (1993), no. 4, 397–413.

- [52] S. Floyd, J. Mahdavi, M. Mathis, and M. Podolsky, *RFC 2883 - an extension to the selective acknowledgement (sack) option for tcp.*, 2000.
- [53] S. Floyd, S. Ratnasamy, and S. Shenker, *Modifying tcp's congestion control for high speeds*, (2002), <http://www.icir.org/floyd/hstcp.html>.
- [54] Dictionnaire Français, *Le dictionnaire hachette en ligne*, <http://www.dictionnaire.com/hachette/> consulté en juin 2005.
- [55] R. Gibbens, P. Hunt, and F. Kelly, *Bistability in communications networks*, In (Eds. G.R. Grimmett and D.J.A. Welsh) *Disorder in Physical Systems* (1990), 113–127.
- [56] R. Gibbens and F. Kelly, *Resource pricing and the evolution of congestion control*, *Automatica* **35** (1999), 1969–1985.
- [57] C. M. Goldie and C. Klüppelberg, *Subexponential distributions survey*, (1997).
- [58] Graham, Kurtz, Méléard, Protter, Pulvirenti, and Talay, *Probabilistic models for nonlinear pde*, Springer.
- [59] C. Graham, *Chaoticity on path space for a queuing network with selection of the shortest queue among several*, *Journal of Applied Probabilities* (2000), no. 37, 198–211.
- [60] ———, *Kinetic limits for large communication networks*, *Modelling in Applied Sciences : a Kinetic Theory Approach*, Bellomo and Puvirenti (ed.), Birkhauser, Boston, 2000.
- [61] ———, *Functional central limit theorems for a large network in which customers join the shortest of several queues*, *Probability Theory and Related Fields* 131 (2005), 97–120.
- [62] C. Graham and S. Méléard, *Chaos hypothesis for a system interacting through shared ressources*, *Probability Theory and Related fields* 100 (1994), 157–173.
- [63] ———, *Fluctuations for a fully connected loss network with alternate routing*, *Stochastic Processes and their Applications* (1994), 97–115.
- [64] H. C. Gromoll, A. L. Puha, and R. J. Williams, *The fluid limit of a heavily loaded processor sharing queue.*, *Ann. Applied Probab.* (2001), no. 12, 797–859.
- [65] F. Guillemin, P. Robert, and B. Zwart, *Aimd algorithms and exponential functionals*, *Annals of Applied Probability* **14** (2004), no. 1, 90–117.
- [66] ———, *Aimd algorithms and exponential functionals*, *Annals of Applied Probability* (February 2004), 90–117.
- [67] C.V. Hollot, V. Misra, D. Towsley, and W-B. Gong, *A control theoretic analysis of red.*, *Proceedings of IEEE INFOCOM* (2001), 10pp.
- [68] D. Hong, *N2nsoft*, 2003, <http://www.n2nsoft.com/>.
- [69] ———, *A note on the tcp fluid model*, Tech. Report RR-4703, INRIA - Rocquencourt, January 2003.
- [70] D. Hong and D. Lebedev, *Many tcp user asymptotic analysis of the aimd model*, Tech. Report 3971, INRIA Research Report number 3971.

-
- [71] P. Hurley, Le Boudec, and P. J. Y., Thiran, *A note on the fairness of additive increase and multiplicative decrease*, Proceedings of ITC-16 (Edinburgh), June 1999.
 - [72] Ikeda, Watanabe, Fukushima, and Kunita, *Itô's stochastic calculus and probability theory*, Springer.
 - [73] N. Ikeda and S. Watanabe, *Stochastic differential equations and diffusion processes*, North Holland and Kodansha, Amsterdam, Oxford, New York, 1981.
 - [74] OPNET INC, *Opnet technologies inc*, <http://www.opnet.com/> consulté en juin 2005.
 - [75] C. Itzykson and J.-M. Drouffe, *Théorie statistique des champs*, CNRS Editions, 2000.
 - [76] V. Jacobson, *Congestion avoidance and control*, SIGCOMM '88 : Symposium proceedings on Communications architectures and protocols, ACM Press, 1988, pp. 314–329.
 - [77] J. Jacod and A. Shiryaev, *Limit theorems for stochastic processes.*, Springer-Verlag, 1987.
 - [78] C. Jin, D. X. Wei, and S. H. Low, *Fast tcp : motivation, architecture, algorithms, performance*, Proc. of IEEE Infocom (2004).
 - [79] I. Karatzas and S. E. Shreve, *Brownian motion and stochastic calculus*, Springer-Verlag, New York, 1991.
 - [80] F. Kelly, *Charging and rate control for elastic traffic*, European Transactions on Telecommunications. **8** (1997), 33–37.
 - [81] F. Kelly, A. Maulloo, and D. Tan, *Rate control in communication networks : shadow prices, proportional fairness and stability*, Journal of the Operational Research Society **49** (1998), 237–252.
 - [82] F. Kelly and R. Williams, *Fluid model for a network operating under a fair bandwidth-sharing policy*, Annals of Applied Probabilities **14** (2004), 1055–1083.
 - [83] T. Kelly, *Scalable tcp : Improving performance in highspeed wide area networks*, (2002).
 - [84] ———, *Engineering flow controls for the internet*, Ph.D. thesis, University of Cambridge, February 2004.
 - [85] S. Keshav, *A control-theoretic approach to flow control*, SIGCOMM '91 : Proceedings of the conference on Communications architecture & protocols (New York, NY, USA), ACM Press, 1991, pp. 3–15.
 - [86] ———, *An engineering approach to computer networking : Atm networks, the internet, and the telephone network*, Addison-Wesley Longman Publishing Co., Inc., Boston, MA, USA, 1997.
 - [87] A. Kherani and A. Kumar, *Stochastic models for throughput analysis of randomly arriving elastic flows in the internet*, in Proceedings of IEEE INFOCOM, New York City, NY, July 2002., 2002.
 - [88] ———, *Closed loop analysis of the bottleneck buffer under adaptive window controlled transfer of HTTP-like traffic.*, Infocom, 2003.

- [89] K. B. Kim, *Design of feedback controls supporting TCP based on modern control theories*, IEEE Transactions on Automatic Control (2004).
- [90] K. B. Kim and S. H. Low, *Design of receding horizon AQM in stabilizing TCP with multiple links and heterogeneous delays*, Proc. of 4th Asian Control Conference (2002).
- [91] L. Kleinrock, *Queueing systems, volume i et ii*, Wiley Interscience, New York, 1975.
- [92] D. E. Knuth, *The art of computer programming, volume 3 : (2nd ed.) sorting and searching*, Addison Wesley Longman Publishing Co., Inc., Redwood City, CA, USA, 1998.
- [93] H. Kunita, *Stochastic flows and stochastic differential equations*, Cambridge University Press.
- [94] S. Kunniyur and R. Srikant, *End-to-end congestion control schemes : utility functions, random losses and ECN marks*, Proc. of IEEE Infocom (2000).
- [95] P. Kuusela, P. Lassila, J. Virtamo, and P. Key, *Modeling red with idealized tcp sources.*, IFIP Conference on performance modelling and evaluation of ATM and IP networks, Budapest (2001), no. 9.
- [96] T. V. Lakshman and U. Madhow, *The performance of tcp/ip for networks with high bandwidth-delay products and random loss*, IEEE/ACM Trans. Netw. **5** (1997), no. 3, 336–350.
- [97] J-Y. Le Boudec, *Rate adaptation, congestion control and fairness : A tutorial*.
- [98] A. Legout, *Contrôle de congestion multipoint pour les réseaux best effort*, Ph.D. thesis, Iniversité de Nice-Sophia Antipolis, Octobre 2000.
- [99] J. Liu, I. Matta, and M. Crovella, *End-to-end interface of loss nature in a hybrid wired/wireless environment*, Proc. of the workshop (WiOpt'03) (2003).
- [100] Y. Liu, Lo F., P. Misra, and D. Towsley, *Fluid models and solutions for large-scale ip networks*, 2003.
- [101] S. H. Low and D. E. Lapsley, *Optimization flow control, i : basic algorithm and convergence*, IEEE/ACM Transactions on Networking (1999), no. 7, 861–874, <http://netlab.caltech.edu>.
- [102] L. Massoulié and J. Roberts, *Bandwidth sharing : objectives and algorithms*, IEEE/ACM Trans. Netw. **10** (2002), no. 3, 320–328.
- [103] F. Mathieu, *Graphes du web, mesures d'importance à la pagerank*, Ph.D. thesis, Université Montpellier 2 - LIRMM, Décembre 2004.
- [104] M. Mathis, Mahdavi, Floyd J., and A. S., Romanow, RFC 2018 - *tcp selective acknowledgment options*, 1996.
- [105] M. Mathis, J. Semke, J. Mahdavi, and T. Ott, *The macroscopic behavior of the tcp congestion avoidance algorithm*, SIGCOMM Comput. Commun. Rev. **27** (1997), no. 3, 67–82.
- [106] M. May, T. Bonald, and J.-C. Bolot, *Analytic evaluation of RED performance*, Proc. of IEEE Infocom (2000).

-
- [107] D. R. McDonald and J. Reynier, *Slides for performance "mean-field model for multiple tcp connections through a buffer implementing RED"*, 2002, <http://www.eleves.ens.fr/home/jreynier/Recherche/PresentationPerf2002>.
 - [108] ———, *Slides for aap "mean-field model for multiple tcp connections through a buffer implementing RED"*, 2005, <http://www.eleves.ens.fr/home/jreynier/Recherche/Janvier2005-MeanFieldForTCP.ppt>.
 - [109] ———, *Mean field convergence of a model of multiple tcp connections through a buffer implementing red*, Ann. Appl. Probab. **16** (2006), no. 1, 244–294, <http://www.eleves.ens.fr/home/jreynier/Recherche/AAP.pdf>.
 - [110] P. A. Meyer, *Probabilités et potentiel*, Hermann, Paris, 1966.
 - [111] B. Miller, K. Avrachenkov, K. Stepanyan, and G. Miller, *Flow control as stochastic optimal control problem with incomplete information*, Proc. of IEEE Infocom (2004).
 - [112] A. Misra and T. J. Ott, *Performance sensitivity and fairness of ecn-aware "modified tcp"*, International IFIP-TC6 Networking Conference Proceedings (2002), no. 2.
 - [113] V. Misra, W. B. Gong, and D. Towsley, *Fluid-based analysis of a network of AQM routers supporting TCP flows with an application to RED*, Proc. of ACM/SIGCOMM (2000).
 - [114] J. Mo, R. La, V. Anantharam, and J. Walrand, *Analysis and comparison of tcp reno and vegas*, Proc. of IEEE Infocom (1999).
 - [115] Diekmann. O., *The cell size distribution and semigroups of linear operators.*, Lecture notes in biomathematics : The dynamics of physiologically structured populations ;.
 - [116] T. Ott, Kemperman, and M. J., Mathis, *The stationary behavior of ideal tcp congestion avoidance*.
 - [117] T. J. Ott, T. V. Lakshman, and L. Wong, *Sred : Stabilized red*, Proc. of IEEE Infocom (1999).
 - [118] Tinnakornsrisuphap P. and A. Makowski, *Many flow asymptotics for tcp with ecn/-red*.
 - [119] ———, *Queue dynamics of red gateways under a large number of tcp flows.*, Globecom (2001).
 - [120] ———, *Limit behavior of ecn/red gateways under a large number of tcp flows*, 2003.
 - [121] J. Padhye, , V. Firoiu, D. Towsley, and J. F. Kurose, *Modeling tcp reno performance : a simple model and its empirical validation*, IEEE/ACM Trans. Netw. **8** (2000), no. 2, 133–145.
 - [122] J. Padhye, Firoiu, Towsley V., and J. D., Kurose, *Modeling tcp throughput : a simple model and its empirical validation*, SIGCOMM '98 : Proceedings of the ACM SIGCOMM '98 conference on Applications, technologies, architectures, and protocols for computer communication, ACM Press, 1998, pp. 303–314.
 - [123] V. Paxson and S. Floyd, *Wide-area traffic : The failure of poisson modelling*, IEEE/ACM Trans. Networking (1995), 226–255.
 - [124] S. Pilosof, R. Ramjee, D. Raz, Y. Shavitt, and P. Sinha, *Understanding TCP fairness over wireless lan*, Proc. of IEEE Infocom (2003).

- [125] M. Porter, *Competitive strategy : Techniques for analyzing industries and competitors*, 1 ed., Free Press, 1980.
- [126] A. Proutière, *Insensibilité et bornes stochastiques dans les réseaux de files d'attente. application à la modélisation des réseaux de télécommunications au niveau flot*, Ph.D. thesis, École Polytechnique, January 2004.
- [127] B. Radunovic and J. Y. Le Boudec, *A unified framework for max-min and min-max fairness with applications*, Proceedings of 40th Annual Allerton Conference on Communication, Control, and Computing (Allerton, IL), October 2002.
- [128] ———, *Rate performance objectives of multi-hop wireless networks*, Proceedings of INFOCOM 04 (Hong Kong), March 2004.
- [129] G. Raina and D. Wischik, *How good are deterministic fluid models of internet congestion control ?*, INFOCOM (2002).
- [130] ———, *Buffer sizes for large multiplexers : Tcp queueing theory and instability analysis.*, IFIP Conference on performance modelling and evaluation of ATM and IP networks, Budapest (2004), no. 9.
- [131] K. Ramakrishnan, S. Floyd, and D. Black, *RFC 3168 - the addition of explicit congestion notification (ecn) to ip*, 2001.
- [132] K. K. Ramakrishnan and R. Jain, *A binary feedback scheme for congestion avoidance in computer networks with a connectionless network layer*, SIGCOMM '88 : Symposium proceedings on Communications architectures and protocols (New York, NY, USA), ACM Press, 1988, pp. 303–313.
- [133] ———, *A binary feedback scheme for congestion avoidance in computer networks*, ACM Trans. Comput. Syst. **8** (1990), no. 2, 158–181.
- [134] P. Ranjan, E. H. Abed, and R. J. La, *Nonlinear instabilities in TCP-RED*, Proc. of IEEE Infocom (2002).
- [135] J. Reynier, *Simulation des équations du champ moyen pour http*, <http://www.eleves.ens.fr/home/jreynier/Recherche/MatlabHTTP.rar>.
- [136] J. Roberts, U. Mocchi, and J. T. (Editors) Virtamo, *Broadband network traffic : Performance evaluation and design of broadband multiservice networks, final report of action cost 242 (lecture notes in computer science)*.
- [137] V. Rosolen, O. Bonaventure, and G. Leduc, *A red discard strategy for atm networks and its performance evaluation with tcp/ip traffic*, ACM Computer Communication Review (1999).
- [138] S. Schenker, L. Zhang, and D. D. Clark, *Some observations on the dynamics of a congestion control algorithm*, SIGCOMM Comput. Commun. Rev. **20** (1990), no. 5, 30–39.
- [139] W. Stevens, *Tcp/ip illustrated : the protocols*, no. 1, Addison–Wesley, 1999, 15th printing.
- [140] D. Stoyan, W. S. Kendall, and J. Mecke, *Stochastic geometry and its applications.*, Wiley, 1995.

-
- [141] A. S. Sznitman, *Topics in propagation of chaos*, École d'Été de Probabilités de Saint-Flour (1989), no. Vol. XIX., 165–251.
 - [142] J. W. A. Tang and S. H. Low, *Is fair allocation always inefficient*, Proc. of IEEE Infocom (2004).
 - [143] E.C. Titchmarsh, *The theory of functions*, Oxford University Press, 1932.
 - [144] F. Verhulst, *Nonlinear differential equations and dynamical systems*, Springer.
 - [145] C. Villamizar and Song C., *High performance tcp in ansnet*, ACM Computer Communication Review **24** (1994), no. 5, 45–60.
 - [146] G. Vinnicombe, *On the stability of networks operating tcp-like congestion control*, Proc. of 15th IFAC World Congress on Automatic Control (2002).
 - [147] M. Vojnovic, J. Y. Le Boudec, and C. Boutremans, *Global fairness of additive-increase and multiplicative-decrease with heterogeneous round-trip times*, in Proceedings of the Conference of the IEEE Computer and Communications Societies (Tel Aviv, Israel), March 2000, pp. 1303–1312.
 - [148] Wiki, *Comment ça marche ? encyclopédie informatique*, <http://www.commentcamarche.net> consulté en juin 2005.
 - [149] ———, *Technical history of the internet and other network protocols*, <http://www.cs.utexas.edu/users/chris/think/> consulté en juin 2005.
 - [150] D. Wischik and N. McKeown, *Part i : Buffer sizes for core routers.*, ACM/SIGCOMM Computer Communication Review **35**, no. 3.
 - [151] L. Xu, K. Harfoush, and I. Rhee, *Binary increase congestion control for fast long-distance networks*, Proc. of IEEE Infocom (2004).
 - [152] L. Zhang, S. Shenker, and D. D. Clark, *Observations on the dynamics of a congestion control algorithm : The effects of two-way traffic*, SIGCOMM, 1991, pp. 133–147.