



AVERTISSEMENT

Ce document est le fruit d'un long travail approuvé par le jury de soutenance et mis à disposition de l'ensemble de la communauté universitaire élargie.

Il est soumis à la propriété intellectuelle de l'auteur. Ceci implique une obligation de citation et de référencement lors de l'utilisation de ce document.

D'autre part, toute contrefaçon, plagiat, reproduction illicite encourt une poursuite pénale.

Contact : ddoc-theses-contact@univ-lorraine.fr

LIENS

Code de la Propriété Intellectuelle. articles L 122. 4

Code de la Propriété Intellectuelle. articles L 335.2- L 335.10

http://www.cfcopies.com/V2/leg/leg_droi.php

<http://www.culture.gouv.fr/culture/infos-pratiques/droits/protection.htm>

THÈSE

Présentée à

L'UNIVERSITÉ PAUL VERLAINE-METZ

UFR Mathématique, Informatique, Mécanique et Automatique

Pour obtenir le grade de

Docteur de l'Université de Paul Verlaine de Metz

Spécialité : **Automatique**

par

Jie LI

**Evaluation et optimisation des performances
des systèmes de production distribution**

Soutenue le 26 Septembre 2006 devant le jury composé de :

M. Enver YUCESAN

Professeur à INSEAD Fontainebleau (Rapporteur)

M. Michel GOURGAND

Professeur à Université Blaise Pascal (Rapporteur)

M. Patrick CHARPENTIER

Professeur à l'ENSTIB Vandoeuvre (Examinateur)

M. Xiaolan XIE

Professeur à l'ENSM Saint-Etienne (Directeur de thèse)

M. Alexandre SAVA

Maître de Conférences à ENIM Metz (Co-directeur)

献给我的父母, 我的妻子

Remerciements

Le travail de thèse présenté dans ce mémoire a été effectué à l'INRIA-LORRAINE (Unité Lorraine de l'Institut National de Recherche en Informatique et en Automatique) dans l'équipe MACSI (Modélisation, analyse et conduite des systèmes industriels) et au LGIPM (Laboratoire de Génie Industriel et Production Mécanique) dans l'équipe SdP (Systèmes de Production) à l'Université Paul Verlaine-Metz.

Je tiens à traduire ma respectueuse reconnaissance à mon directeur de thèse, Monsieur le Professeur Xiaolan XIE, de m'avoir accueilli au sein de son équipe de recherche et d'avoir accepté de diriger ce travail. Qu'il soit remercié en premier pour sa grande disponibilité, son suivi continu, ses conseils constructifs et la qualité de ses idées qui m'ont permis d'atteindre la finalité de ce travail.

J'exprime toute ma gratitude à mon Co-encadrant de thèse M. Alexandre SAVA, pour son aide, sa compréhension, sa gentillesse et sa compétence qu'il m'a témoigné tout au long de ces années.

Ma gratitude et mes remerciements vont ensuite aux membres du jury qui ont bien voulu me faire l'honneur de participer à ce jury :

- **M. Enver YUCESAN**, Professeur, INSEAD, Boulevard de Constance, F-77305 Fontainebleau, France
- **M. Michel GOURGAND** Professeur, Université Blaise Pascal, ISIMA Campus des Cézeaux, 63173, Aubière cedex, France
- **M. Patrick CHARPENTIER**, Professeur, Directeur Adjoint de l'ENSTIB, CRAN UMR CNRS 7039 OR ENSTIB, BP n° 239, Vandoeuvre France

J'adresse mes remerciements le plus chaleureux aux membres de l'équipe de MACSI et SDP pour leur sympathie et leur convivialité. Particulièrement, merci M. Lyes BENYUCEF de son aide dans le travail et la vie.

Que Mme Christel Wiemert reçoive l'expression de reconnaissance et d'amitié pour son aide et ses services durant toutes ces années.

C'est le moment aussi de dire un grand merci à tous ceux qui m'ont aidé pour que je puisse réaliser mes études supérieures par leurs encouragements. Je cite surtout mon père, ma mère, mes sœurs et tous mes amis. J'exprime mes plus sincères gratitudes à toutes ces personnes : vous étiez toujours présents ici, malgré la distance, avec votre humour, votre charme et votre chaleur.

Enfin un remerciement très particulier à celui donné le courage pour surmonter les moments difficiles durant cette thèse, pour son soutien et sa confiance à Mme. HENNEQUIN Sophie, M. HU Heng, M. DING Hongwei, M. MOURANI Iyad, M.ACHOUR Zied et M. TANONKOU Guy Aimé.

Encore une fois, merci à toutes et à tous.

Résumé : Nous considérons un système de production-distribution composé d'un ensemble d'unités de production et d'entrepôts disposées en série et connectés par des facilités de transport. Chaque entrepôt est contrôlé par une politique donnée de gestion de stock et il est approvisionnée par une unité de production amont de capacité finie. On considère un seul type de produit fini. Le problème à résoudre est de déterminer le meilleur paramétrage de la politique de gestion de stock sur l'ensemble du système tout en prenant en compte les contraintes sur la capacité de production, le taux de service client souhaité et les délais de transport. Nos travaux de recherche ont permis le développement d'une méthodologie d'optimisation des paramètres de gestion du stock basée sur la simulation, qui permet de trouver les paramètres de la politique de gestion du stock qui minimisent le coût sur l'ensemble du système de production-distribution. L'efficacité de l'approche proposée a été vérifiée par des expériences numériques.

Mots clés : optimisation, politiques de gestion de stock, multi-niveaux, système de production-distribution

Abstract: This thesis considers a production-distribution system made up of a set of production sites and distribution centers connected by the transport facilities. Each distribution center is managed by a given stock management policy. The aim is to find the best parameters setting of the stock management policy through the whole network in order to optimize the overall performances of the production-distribution system while taking into account the finite production capacity, the customer service requirement, the transport time, and the random customer demand. The results obtained during this PhD thesis allowed to develop a methodology of optimization for the parameter setting of inventory control policies management in the production-distribution systems. Thus, we proposed a simulation based optimization approach which computes the setting that minimizes the overall inventory cost of the production-distribution system, while taking into account fill rate constraints. The approach is validated by a large number of numerical experiments.

Keywords: simulation, optimization, production distribution system, inventory control policy.

Table des matières

Liste des figures.....	5
Liste des tableaux.....	7
Notations.....	9
Introduction générale.....	11

Chapitre 1

L'état de l'art sur l'évaluation et l'optimisation des chaînes logistiques

1.1 Introduction à la chaîne logistique.....	19
1.2 Contexte et motivation de cette recherche.....	24
1.3 Etat de l'art sur l'optimisation des chaînes logistiques.....	26
1.3.1 Politiques de gestion du stock d'un entrepôt.....	27
1.3.2 Etat de l'art sur les systèmes de production-distribution multi-niveau.....	30
1.3.3 Etat de l'art sur la simulation des chaînes logistiques.....	34
1.3.4 Conclusions sur l'optimisation des chaînes logistiques.....	36
1.4 Etat de l'art sur l'optimisation discrète par la simulation.....	37
1.5 Positionnement des contributions de cette thèse.....	41
1.5.1 Analyse et optimisation d'un système de production-distribution.....	42
1.5.2 Optimisation par simulation des systèmes de production-distribution.....	43

Chapitre 2

Une approche analytique pour évaluation et optimisation des performances des systèmes de production-distribution à deux niveaux

2.1 Introduction.....	49
2.2 Topologie du système et approche de modélisation.....	50
2.3 Evaluation des performances.....	52
2.3.1. Le modèle de l'unité de production.....	53
2.3.2 Le modèle de transport.....	56
2.3.3 Délai d'approvisionnement et stock en commande.....	57
2.3.4 Estimation du coût global de gestion du stock.....	58

2.3.5 Estimation du taux de remplissage.....	59
2.4. Optimisation du coût global de gestion du stock.....	60
2.5 Résultats numériques.....	60
2.6 Conclusions.....	63

Chapitre 3

Une approche d'optimisation basée sur la simulation pour les systèmes dynamiques à événements discrets

3.1 Introduction.....	67
3.1.1 Préliminaires.....	67
3.1.2 L'optimisation par simulation.....	68
3.1.3 Méthodologies d'optimisation par simulation.....	72
3.2 Optimisation pas simulation discrète avec contraintes analytiques.....	74
3.2.1 Algorithme compass*- cas d'un espace d'état borné.....	74
3.2.2 Algorithme compass*- cas d'un espace d'état non borné.....	82
3.2.3 Implémentation et résultats numériques.....	87
3.3 Optimisation pas simulation discrète avec contraintes simulées (non-closed form)	92
3.3.1 Algorithme compass* avec contrainte- cas d'un espace d'état borné.....	93
3.3.2 Algorithme compass* avec contrainte- cas d'un espace d'état non borné.....	98
3.3.3 Implémentation et résultats numériques.....	102
3.4 Conclusions	106

Chapitre 4.

Optimisation par simulation des systèmes de production-distribution

4.1. Modélisation des systèmes de production distribution.....	111
4.1.1. Clients.....	112
4.1.2. Entrepôts.....	112
4.1.3. Unités de production.....	113
4.2. Systèmes de production-distribution multi-niveau avec politique base stock.....	114
4.2.1. Problème d'optimisation.....	114
4.2.2. Optimisation des politiques d'approvisionnement.....	115
4.2.3. Analyse de l'influence du taux moyen d'arrivée.....	119
4.2.4. Analyse de l'influence de la taille des commandes.....	122

4.2.5. Prise en compte du taux de service client.	124
4.3. Systèmes de production-distribution multi-niveau avec politique (R, nQ)	126
4.3.1. Politique de gestion du stock (R, nQ)	126
4.3.2. Prise en compte du coût de transport.....	131
4.4 Optimisation d'un système de production-distribution avec politique (s, S)	135
4.4.1 Politique de gestion du stock (s, S)	135
4.4.2. Optimisation sans prise en compte du coût de transport.....	135
4.4.3 Optimisation avec prise en compte du coût de transport.....	137
4.5 Conclusion.....	138
Conclusion générale.....	143
Bibliographie.....	149

Liste des figures

Figure.1.1 Exemple de chaîne logistique.....	20
Figure.1.2 Traitement séquentiel des différents niveaux de décision.....	23
Figure 2.1 La topologie du système.....	50
Figure 2.2 Principe de l'approche de modélisation.....	52
Figure 3.1 Principe de l'optimisation par simulation.....	71
Figure 3.2 Classification des méthodologies d'optimisation par simulation (DOvS)	73
Figure 3.3 la définition de la région la plus prometteuse.....	75
Figure 3.4 Un système de production-distribution	87
Figure 3.5 Cas d'un espace d'état borné et un système de production-distribution à 8.....	90
Figure 3.6 Cas d'un espace d'état borné et un système de production-distribution à 10 niveaux	91
Figure 3.7 Cas d'un espace d'état non borné et un système de production-distribution.....	91
Figure 3.8 Cas d'un espace d'état non borné et un système de production-distribution à 10 niveaux.....	91
Figure 3.9 Cas d'un espace d'état borné et un système de production-distribution à 12 niveaux.....	92
Figure 3.10 Système de production-distribution borné avec 5 niveaux et espace d'état borné.....	104
Figure 3.11 Système de production-distribution borné avec 10 niveaux et espace d'état borné.....	104
Figure 3.12 Système de production-distribution borné avec 5 niveaux et espace d'état non borné.....	105
Figure 3.13 Système de production-distribution borné avec 8 niveaux et espace d'état borné.....	105
Figure 3.14 Système de production-distribution borné avec 12 niveaux et espace d'état borné.....	106
Figure 4.1 Topologie du système.....	112
Figure 4.2 Optimisation d'un système de production à 5 niveaux.....	117

Figure 4.3 Optimisation d'un système de production-distribution à 8 niveaux.....	118
Figure 4.4 Optimisation d'un système de production-distribution à 10 niveaux.....	118
Figure 4.5 Optimisation d'un système de production-distribution à 12 niveaux.....	119
Figure 4.6 Analyse de l'influence du taux moyen d'arrivée.....	122
Figure 4.7 L'influence de l'écart type de la taille des commandes sur le coût global de stockage.....	124
Figure 4.8 Niveau de service et le coût de stockage.....	126
Figure 4.9 Convergence des variables R et Q pour $\lambda=1.0$	129
Figure 4.10 Optimisation d'un système de production distribution multi- niveaux.....	129

Liste des tableaux

Tableau 2.1	Paramètres d'entrée.....	61
Tableau 2.2.	Résultats de simulation et de l'approche analytique pour le délai de fabrication L_s et le délai de re-approvisionnement L	62
Tableau 2.3	Résultats de simulation et de l'approche analytique pour le stock sous commande et coût de gestion du stock	62
Tableau 2.4.	Résultats de simulation et de l'approche analytique pour le taux de service et le coût de gestion du stock pour un niveau de base-stock optimal R^*	63
Tableau 4.1	paramètres de simulation (unité de temps: seconde)	116
Tableau 4.2	solutions calculées lors des 9 dernières itérations pour un système de production-distribution avec 5 niveaux $\lambda=1.0$	120
Tableau 4.3	solutions calculées lors des 9 dernières itérations pour un système de production-distribution avec 12 niveaux $\lambda=1.5$	120
Tableau 4.4	Solutions pour différents niveaux et différents taux moyens d'arrivée de la demande.....	121
Tableau 4.5	paramètres de simulation	123
Tableau 4.6	Solution pour le cas de 5, 8, 10 et 12 nœuds et taille de commandes différentes.....	123
Tableau 4.7	jeux de paramètres (unités de temps: la seconde)	125
Tableau 4.8	solutions pour 5, 8, 10 et 12 niveaux et différents taux de service client.....	125
Tableau 4.9	optimisation d'un système de production distribution à deux niveaux.....	128
Tableau 4.10	solution pour un système de production distribution à 5 niveaux.....	130
Tableau 4.11	solution localement optimale pour un système de production distribution à 5, 8 et 10 niveaux.....	131
Tableau 4.12	paramètres de simulation (unité de temps: seconde)	133
Tableau 4.13	solutions optimums local pour 5, 8 et 10 nœuds.....	133
Tableau 4.14	solution localement optimale pour 5, 8 et 10 nœuds.....	134
Tableau 4.15	solution optimale locale pour 5, 8 et 10 nœuds, $\lambda=1.3$	136
Tableau 4.16	solution pour 5, 8 and 10 nœuds, $\lambda=1.3$	137
Tableau 4.17	solutions locale optimales pour différentes politiques de gestion du stock (5 nœuds)	139

Notations

λ	taux d'arrivée
X	ordre de quantité
m_X	la moyenne de X
σ_X	l'écart type de X
τ	Le temps de fabrication d'une unité de produit
m_τ	la moyenne de τ
σ_τ	l'écart type de τ
L_s	délai de fabrication d'une unité de produit (pas d'un lot)
L_t	Le temps nécessaire pour le transport d'un lot à l'entrepôt est appelé délai de transport
C_h	Le stockage d'une unité de produit engendre un coût C_h par unité de temps
C_b	un produit qui ne peut pas être livré au client en raison d'une rupture de stock engendre un coût C_b par unité de temps.
N_s	nombre d'unités de produit des ordres d'approvisionnement en attente ou en cours de fabrication,
N_t	nombre total de produits en transit vers l'entrepôt,
L	délai de approvisionnement de l'entrepôt,
N	stock en commande, i.e. nombre total de produits commandés par l'entrepôt et pas encore reçues.
T_B	le temps de service d'un ordre de approvisionnement par le serveur
N_B	Le nombre de lots
ρ	l'intensité du trafic de la file d'attente
R	le niveau de base-stock
I	le stock disponible
B	la quantité de demandes en attente
N	le stock en commande
IN	l'état du stock.
θ	l'ensemble des paramètres contrôlables
$f(\theta)$	une mesure de performance choisie du système étudié
$g(\theta; \omega)$	La performance déterminée en exécutant une fois le modèle de simulation pour un paramétrage donné
$\hat{f}(\theta)$	une estimation de la performance mesurée $f(\theta)$

Θ	l'espace des solutions
V_k	l'ensemble de solutions visitées jusqu'à l'itération k
$N_k(x)$	le nombre total d'observations alloués à la solution x jusqu'à l'itération k
$\bar{G}_k(x)$	la moyenne des résultats $G(x, w_i)$
x_k^*	la solution optimale observée durant l'itération k
C_k	une région appelée la région la plus prometteuse
$a_k(x)$	d'observations à allouer à la solution x durant l'itération k
T	la capacité de transport d'un équipement
C_T	coût unitaire de transport

Introduction générale

Introduction

Les travaux de recherche présentés dans ce mémoire concernent l'évaluation de performances et l'optimisation des systèmes de production-distribution. Les décisions prises lors de la conception d'un tel système ont une incidence majeure sur les performances à long terme. D'un côté, le choix de l'emplacement des différentes composants du réseau est une décision de niveau stratégique. D'un autre côté, les décisions opérationnelles concernent plutôt la planification de la production et la gestion du stock. La plupart des systèmes de production-distribution sont stochastiques à cause de l'aspect aléatoire des événements qui déterminent leur dynamique. Dans ce cas, les approches analytiques sont difficiles à utiliser pour des systèmes de taille réelle et la simulation est le seul outil qui permet d'évaluer les solutions proposées. Dans ce mémoire, nous utilisons à la fois des méthodes d'optimisation analytique et des techniques d'optimisation par la simulation pour optimiser les performances d'un système de production-distribution.

Plus précisément, nous considérons un système de production-distribution composé d'un ensemble d'unités de production et d'entrepôts disposées en série et connectés par des facilités de transport. Chaque entrepôt est contrôlé par une politique donnée de gestion de stock et il est approvisionnée par une unité de production amont de capacité finie. On considère un seul type de produit fini. Le problème à résoudre est de déterminer le meilleur paramétrage de la politique de gestion de stock sur l'ensemble du système tout en prenant en compte les contraintes sur la capacité de production, le taux de service client souhaité et les délais de transport. L'objectif de nos travaux de recherche est le développement d'une méthodologie d'optimisation des paramètres de gestion du stock afin de minimiser le coût sur l'ensemble du système de production-distribution.

Le premier chapitre du mémoire débute par une introduction sur les systèmes de production-distribution. Ensuite, nous donnons les objectifs de cette thèse de doctorat. Le premier objectif concerne le développement d'une approche analytique pour l'évaluation de performances et l'optimisation d'un système de production-distribution simple, composé d'un entrepôt alimenté par une unité de production. Le deuxième

objectif est la mise en place d'une méthode d'optimisation par simulation des systèmes de production-distribution complexes. Nous présentons également un état de l'art des principales techniques d'évaluation de performances et d'optimisation des systèmes de production-distribution.

Le premier objectif est traité dans le deuxième chapitre. Nous proposons une approche analytique pour l'évaluation des performances et l'optimisation d'un système composé d'un entrepôt alimenté par une unité de production qui a une capacité de fabrication finie. L'arrivée des clients est aléatoire, ainsi que la quantité commandée. Le stock est contrôlé selon une politique de base-stock. Le transport des produits entre l'unité de fabrication et l'entrepôt est effectué selon la quantité demandée par l'ordre d'approvisionnement. Le délai de transport est constant et non négligeable. L'approche analytique que nous proposons pour évaluer les performances et minimiser le coût total de gestion du stock est basée sur des approximations qui ont prouvé leur efficacité pour la plupart de systèmes de production : i) caractérisation des variables aléatoires par les deux premiers moments, ii) approximation markovienne des processus généraux, iii) approximation de processus dépendants par des processus indépendants. Plus précisément, l'approche proposée utilise les deux premiers moments des variables aléatoires pour calculer le délai d'approvisionnement de l'entrepôt, le stock disponible et la quantité en rupture. Le coût total de gestion du stock, défini par la somme de coûts de possession et de rupture, est minimisé à l'aide d'une technique de gradient, tout en respectant un taux de service client minimal imposée. Des expérimentations numériques ont permis de valider l'approche proposée et de montrer son efficacité.

Malgré les très bons résultats obtenus pour un système de production-distribution à deux niveaux, une approche analytique semble difficile à mettre en place pour un système plus compliqué. Ainsi, dans le chapitre 3 nous proposons une méthode d'optimisation par simulation. L'approche que nous proposons, appelé COMPASS*, est basée sur la méthode COMPASS existante dans la littérature. Notre contribution consiste à améliorer la convergence de cette algorithmes, à optimiser l'utilisation du budget de simulation et à intégrer la possibilité de prendre en compte des contraintes évaluées par la simulation. Les résultats présentés dans ce chapitre nous ont permis d'atteindre notre deuxième objectif de recherche.

Dans le chapitre 4, nous analysons les performances de différentes politiques de gestion de stock dans le cadre du système de production-distribution multi-niveaux considéré, avec prise en compte du taux de service client minimal imposé et du coût de transport. Nous testons les politiques de gestion de stock suivantes : i) la politique base-stock, ii) la politique (R, nQ) et iii) la politique (s, S) . L'objectif est de déterminer la politique ainsi que le paramétrage associé, qui minimise le coût global de gestion du stock dans le système.

Nous terminons par quelques conclusions et nous indiquons quelques directions de recherche future pour la poursuite des nos travaux de recherche.

Chapitre 1

L'état de l'art sur l'évaluation et l'optimisation des chaînes logistiques

Dans ce chapitre, nous présentons un état de l'art sur le problème d'évaluation de performances et d'optimisation des systèmes de production-distribution. La section 1.1 introduit le problème de la conception des chaînes logistiques. La section 1.2 présente les objectifs de nos travaux de recherche. Par la suite, nous illustrons les travaux existantes dans la littérature qui sont les plus en relation avec nos objectifs de recherche. D'abord, dans la section 1.3 nous concentrons notre attention sur les politiques de gestion du stock ainsi que sur les approches d'optimisation analytiques. Ensuite, dans la section 1.4 nous présentons des techniques d'optimisation par la simulation. Nous concluons ce chapitre par une mise en évidence des contributions de cette thèse.

1.1 Introduction à la chaîne logistique

Il y a plusieurs définitions du "Supply Chain Management" ou du management de la chaîne logistique dans la littérature. La définition la plus répandue présente une chaîne logistique comme un système de fournisseurs, des unités de production, de distributeurs, des détaillants et des clients ([Houlihan (1985)], [Stevens (1989)], [Lee et Billington (1993)], et [Lamming (1996)]). Le flux de matériaux part des fournisseurs vers les clients, alors que le flux d'informations va dans les deux directions. Dans les années 1990 le management de la chaîne logistique s'est imposé en tant que nouveau concept de management. Il s'appuie sur une vision globale et prend en compte les interactions à différents niveaux au sein d'une chaîne logistique.

Dans cette thèse, nous considérons une partie de la chaîne logistique, notamment l'interface entre le processus de production et celui de distribution. Dans ce contexte, deux éléments ont une importance considérable : le coût de gestion et les attentes des clients.

Une chaîne logistique est également un réseau d'unités de production et de distribution qui assure l'approvisionnement des matières, la transformation de ces matières en sous-ensembles et produits finis et la distribution de ces produits aux clients (figure 1.1). Les chaînes logistiques concernent les organisations industrielles ainsi que les organisations des services. Cependant, la complexité varie selon le type d'industrie et d'une compagnie à l'autre. En pratique, une chaîne logistique concerne plusieurs produits finis qui partagent les composants, les équipements de fabrication et les capacités. Le flux de matières n'est pas toujours organisé d'une manière arborescente. Différentes modes de transport peuvent être envisagés. De plus, la nomenclature des produits finis peut être constituée de plusieurs niveaux et peut concerner un nombre important d'articles.

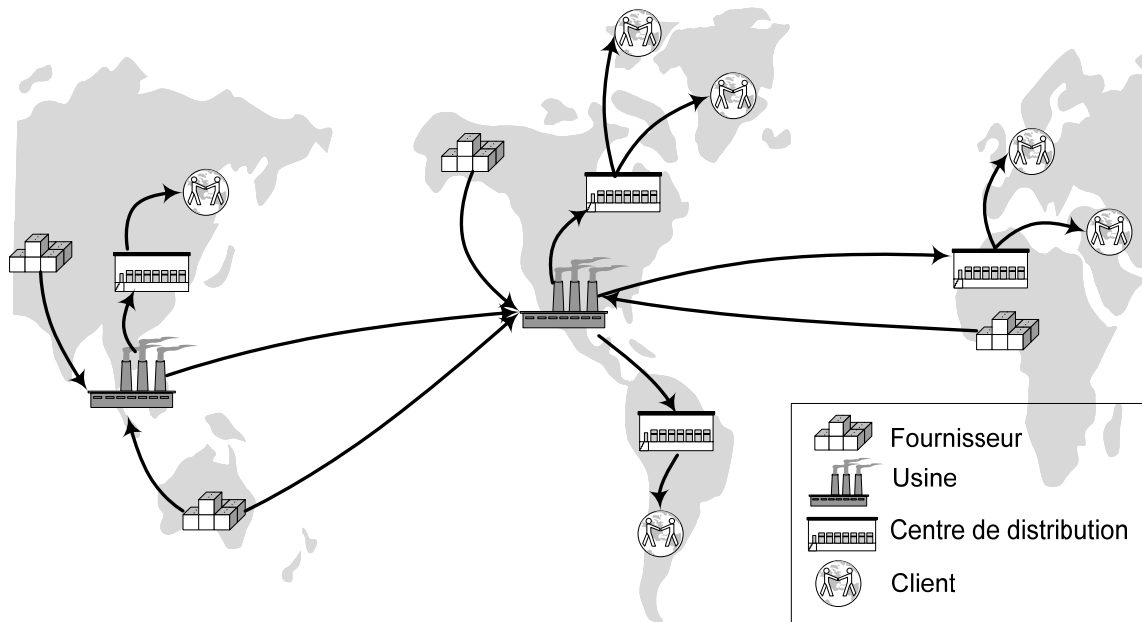


Fig. 1.1 Exemple de chaîne logistique

En utilisant le concept de flux à partir de la matière première vers le client, certains auteurs ont présenté une structure générale de chaîne logistique et l'ensemble d'éléments qui la constituent [Mabert et Venkataramanan (1998)]. Cette structure présente cinq éléments principaux qui modélisent les phases principales dans l'évolution des flux :

- l'**approvisionnement** implique non seulement l'acquisition de la matière première par le biais d'un réseau d'acheteurs, mais il participe également au développement des produits par la conception des outils et des sous-ensembles ;
- la **logistique amont** gère le mouvement et le stockage efficace et efficient des matières afin de satisfaire le programme de fabrication ;
- la **fabrication** doit fournir des produits de qualité à un prix compétitif ;
- la **logistique aval** gère le mouvement des produits finis dans le réseau de distribution, jusqu'au clients finaux.
- le **service après-vente** répond à la nécessité de fournir le support technique pour les produits commercialisés.

Traditionnellement, dans une entreprise, le service marketing, la planification, la production, les achats et la distribution opèrent d'une manière indépendante. Dans une telle organisation chaque service a ses propres objectifs. De plus, souvent les objectifs de

différents services sont contradictoires. Par exemple, les objectifs du marketing de maximiser le taux de service client sont en contradiction avec les objectifs de la production et de la distribution. Souvent, l'activité de production est organisée pour maximiser la productivité tout en minimisant le coût de revient sans se soucier de l'impact sur le volume de stock et la capacité de la distribution. Les contrats d'achat sont souvent négociés sans prendre en compte de nouvelles informations sur les ventes. La conséquence de ces observations est qu'il n'y a pas un plan unique et intégré pour l'ensemble d'organisation. Donc, il est nécessaire de disposer d'un mécanisme qui permet à ces fonctions de se coordonner et de travailler ensemble. Le management de la chaîne logistique fournit le cadre nécessaire pour obtenir cette intégration.

Une chaîne logistique peut être définie par un processus intégré au sein duquel différentes entités (fournisseurs, fabricants, distributeurs et détaillants) travaillent ensemble afin de : (i) acquérir de la matière première ; (ii) la transformer en produits finis et (iii) délivrer les produits finis aux détaillants et aux clients finaux. Cette chaîne est traditionnellement caractérisée par un flux de matières « en avant » et un flux d'information « en arrière ». Récemment, le concept traditionnel de management de la chaîne logistique a été étendu avec la prise en compte du recyclage, de la re-fabrication et de la re-utilisation.

Trois niveaux de la prise de décision dans un chaîne logistique ont été identifiés selon l'horizon de temps impliqué [Ballou (1992)] : le niveau stratégique, le niveau tactique et le niveau opérationnel. Le niveau stratégique prend en compte un horizon de temps de plus d'un an. A l'opposé, les décisions opérationnelles concernent des périodes de temps courtes, souvent inférieures à un jour ou à une heure. Le niveau tactique se situe entre les deux.

Niveau stratégique : Les décisions stratégiques concernent typiquement l'ouverture, la fermeture ou la localisation de sites de production et des entrepôts. On souhaite par exemple déterminer l'emplacement des nouveaux entrepôts ou leur capacité. Un autre exemple de décision stratégique concerne le système de distribution. Est-ce que la compagnie doit posséder ses propres camions ou sous-traiter l'activité de transport? Certaines décisions stratégiques peuvent également être affectées au niveau tactique suivant la flexibilité de la production ou de la distribution. Par exemple, lorsque les activités de transport et d'entreposage sont sous-traités, la flexibilité de la distribution augmente et il est

possible d'augmenter rapidement l'espace utilisé pour l'entreposage ou le nombre de camions utilisés.

Niveau tactique : Les décisions au niveau tactique concernent principalement le contrôle des processus opérationnels : le stockage, la fabrication, le transport. D'un côté, le rôle de la gestion du stock est de répondre aux questions suivantes : quoi, quand et combien faut-il approvisionner. Est-ce qu'on doit émettre des ordres d'approvisionnement périodiquement ou lorsque le niveau du stock baisse au-dessous d'une certaine valeur ? Faut-il privilégier les achats groupés ou acheter les produits individuellement ? Comment doit-on conditionner les produits ? D'un autre côté, le contrôle de la production s'intéresse surtout à l'ordonnancement des opérations dans un atelier, à la définition des priorités et au suivi du programme de fabrication.

Niveau opérationnel : Les décisions opérationnelles sont liées eux aussi aux activités suivantes : stockage, production et transport. Un exemple de décision opérationnelle sur la gestion d'un stock est de commander une quantité supplémentaire pour faire face à une augmentation de la demande dans un future très proche. Une décision opérationnelle dans un système de transport est le choix de la route à suivre, ou de l'ordre de passage chez différentes destinations. Le fait de fabriquer un lot de produits en urgence représente une décision opérationnelle pour la production.

Différents modèles sont utilisés pour différents niveaux de décision. Les modèles utilisés au niveau stratégique sont conçus pour analyser et comparer un nombre important de réseaux logistiques différentes. Les modèles linéaires sont souvent utilisés afin d'obtenir une solution dans un temps de calcul raisonnable. Si le modèle est linéaire, des techniques efficaces telles que la programmation linéaire peuvent être utilisées pour trouver rapidement une solution optimale au problème considéré. Les modèles tactiques sont généralement caractérisés par un niveau de détail plus élevé et peuvent inclure des non-linéarités. Enfin, les modèles opérationnels sont généralement linéaires parce que l'horizon de temps est court et il y a des contraintes fortes sur les temps de calcul.

A cause de la complexité du problème, rares sont les méthodes qui adressent en même temps des décisions stratégiques, tactiques et opérationnelles ainsi que les risques liés à la

dynamique de la chaîne logistique (fluctuations dans les commandes clients, ainsi que dans les délais de fabrication et de transport).

Durant des années, les chercheurs et les praticiens ont étudié séparément les différents processus d'une chaîne logistique. Ces trois types de problèmes de décision sont généralement traités dans l'ordre suivant : d'abord les décisions stratégiques, ensuite les décisions tactiques et finalement opérationnelles. Cette approche est illustrée dans la figure 1.2.

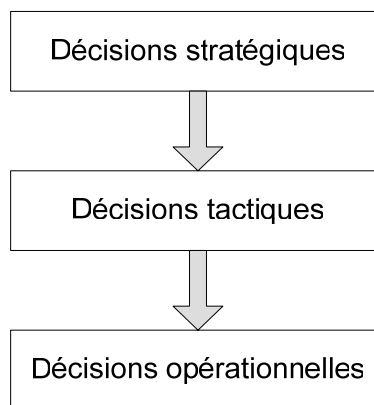


fig. 1.2 Traitement séquentiel des différents niveaux de décision

Une question importante apparaît lors de la conception d'une chaîne logistique : comment peut-on assurer que les décisions stratégiques sont effectives aux niveaux tactique et opérationnel ? Des modèles déterministes et stochastiques ont été proposés dans la littérature pour l'optimisation de la chaîne logistique.

Récemment, il y a un intérêt accru pour la conception, l'évaluation de performances et l'optimisation d'une chaîne logistique dans sa globalité. Des travaux apparaissent sur l'interaction des deux catégories de décisions rencontrées dans la configuration d'une chaîne logistique :

- (1) Décisions structurales (à long terme et au niveau stratégique), par exemple :
 - ❑ localisation d'une unité de production, d'un entrepôt
 - ❑ capacité d'une unité de production,
 - ❑ mode de transport.

(2) Coordination de décisions (plus court terme et aux niveaux tactique et opérationnel), par exemple :

- ❑ répartition des stocks : où et combien,
- ❑ approvisionnement centralisé ou décentralisé,
- ❑ produire à la demande ou produire pour stock (make-to-stock, make-to-order).

La configuration d'une chaîne logistique détermine le niveau du flux de production et de stockage de la matière première et des sous-ensembles associées aux différents niveaux de la nomenclature. Les produits finis et le flux d'information interagissent à travers un réseau d'unités de production et de centres de distribution afin de satisfaire une demande généralement fluctuante.

La reconfiguration périodique d'une chaîne logistique est essentielle pour préserver la compétitivité d'une entreprise. L'optimisation de la configuration d'une chaîne logistique consiste, entre autres à :

- ❑ dimensionner le stock de sécurité et le point de commande,
- ❑ fixer la localisation des entrepôts,
- ❑ choisir la politique de production (make-to-stock, make-to-order),
- ❑ définir la capacité de production en terme de quantité et de flexibilité,
- ❑ affecter les ressources de distribution et les modes de transport,
- ❑ imposer des objectifs de performance à atteindre pour les unités opérationnelles.

Par conséquent, l'objectif de l'optimisation de la configuration d'une chaîne logistique est de trouver la meilleure ou une suffisamment bonne configuration telle que la chaîne logistique atteint le niveau de performance demandé.

1.2 Contexte et motivation de cette recherche

Cette thèse s'inscrit dans une recherche de long terme du projet MACSI de l'INRIA-Lorraine, démarrée dans le cadre du projet européen GROWTH-ONE (2001-2004),

sur l'optimisation des entreprises en réseau. L'objectif est de développer des modèles réalistes pour la conception et la gestion des chaînes logistiques en prenant compte des coûts, des niveaux de service des clients, des impacts sociaux et environnementaux. Les cas d'étude proposés par les partenaires industriels ont montré clairement la nécessité de tenir compte des phénomènes aléatoires, des risques et des stratégies de pilotage.

Ces exigences ont été prises en compte dans le cadre de la thèse de H. Ding [ding, 2003a,b,2004a~d]. Les auteurs ont développé une approche qui couple la simulation et l'optimisation pour la conception des chaînes logistiques. La difficulté majeure d'une telle approche vient du fait que le modèle de simulation dépend de la structure de la chaîne logistique. Pour cela, ils ont proposé dans Ding [ding,2004a] un modèle de simulation similaire au modèle de référence SCORE composé des entités de base tels que les fournisseurs et les entrepôts. La particularité de cette modèle est l'utilisation d'une règle de pilotage paramétrée pour chaque problème de décision rencontré. Ces paramètres de contrôle représentent le point de commande d'un dépôt, les ratios d'allocation entre les fournisseurs, etc. Ils dépendent de la structure du réseau choisi et doivent être optimisés.

De même, dans Ding [ding, 2004d] les auteurs ont proposé une approche d'optimisation multicritères basée sur la simulation et les pour la conception optimale des chaînes logistiques dans un contexte dynamique et incertain. Cette approche comprend un module d'optimisation basé sur des algorithmes génétiques multicritères et un module de simulation permettant l'évaluation de la chaîne étudiée au cours de l'optimisation. L'optimisation concerne à la fois la configuration de la chaîne logistique et le système de pilotage caractérisé par les règles de pilotage et leurs paramètres de contrôle pour la prise de décision durant la simulation. L'algorithme génétique explore les solutions candidates intéressantes (configuration de la chaîne + règles de pilotage de la chaîne) dont les performances sont évaluées par le simulateur. Dès réception d'une solution candidate, le simulateur génère automatiquement le modèle de simulation correspondant et évalue un ensemble d'indicateurs de performances tels que les coûts, les délais et la consommation moyenne en carburant. En tenant compte des résultats fournis par le simulateur, une évaluation globale est associée à cette solution candidate. Toutes les solutions candidates, proposées par le module d'optimisation, seront évaluées au fil de l'eau par le simulateur. La technique des algorithmes génétiques multi-objectif (MOGA), utilisée dans le module

d'optimisation, a pour objectif de guider la recherche (dans un espace de solutions possibles) vers une solution proche de l'optimum. La paramétrisation du module de simulation prend en compte les incertitudes liées à la demande, à la production, au stockage et à la distribution. L'approche a été validée sur deux cas d'étude : pour la configuration d'un réseau de distribution automobile (FIAT) et pour le choix des fournisseurs en textile (HI-TEC).

L'avantage de l'approche précédente est l'optimisation des décisions à la fois stratégiques (configuration du réseau) et tactiques (systèmes de pilotage) en tenant compte de tous les mesures de performance importants et en s'appuyant sur un modèle d'évaluation des performances réaliste. L'inconvénient de cette approche est le temps de calcul excessif à cause du temps de simulation d'un modèle réaliste et de la difficulté de l'approche génétique à capter les propriétés du problème.

Cette thèse poursuit la recherche démarrée dans la thèse de Ding. Nous nous restreignons à l'optimisation du système de pilotage et développons des méthodes d'optimisation efficaces. Plus précisément, cette thèse considère un système de production-distribution composé de différentes sites de production et des centres de distribution reliés par des moyens de transport. Chaque site est gérée par une politique de gestion de stock donnée. Le problème consiste à déterminer le meilleur paramétrage des politiques de gestion de stock de l'ensemble des sites afin d'optimiser les performances globales du système de chaînes logistiques tout en tenant compte des contraintes de production, des aléas de la demande, des aléas, des délais de transport, etc. L'objectif de cette thèse est de développer une méthodologie d'optimisation pour le paramétrage des politiques de gestion de stock fréquemment rencontrés dans les systèmes de production-distribution. Cette thèse exploite les techniques développées dans le cadre des systèmes à événements discrets pour l'évaluation et l'optimisation des performances, en particulier les techniques d'optimisation fondés sur la simulation.

1.3 Etat de l'art sur l'optimisation des chaînes logistiques

Dans la suite de ce chapitre, nous passons en revue la littérature sur les décisions tactiques

dans les chaînes logistiques et nous décrivons la manière dont les travaux de recherche existants déterminent les politiques de gestion du stock. Nous commençons par les études des politiques de gestion de stock d'un entrepôt avec un seul type de produit. Nous abordons ensuite la très riche littérature sur l'évaluation et l'optimisation des performances des systèmes de production-distribution à plusieurs niveaux. Les travaux sur la simulation des chaînes logistiques sont ensuite passés en revue. Nous terminons par un survol des travaux sur l'optimisation discrète par la simulation.

1.3.1 Politiques de gestion du stock d'un entrepôt

Considérons, pour commencer, le modèle logistique le plus simple : un seul produit et un seul point de stockage. Les deux décisions importantes à prendre pour la gestion de ce type de système sont quand et combien approvisionner.

Par rapport à la date de re-approvisionnement, la littérature considère deux modes de contrôle du niveau du stock : continu et périodique. Les politiques de gestion du stock basées sur un contrôle continu du niveau du stock lancent un ordre d'approvisionnement dès que le niveau du stock descend au-dessous d'un seuil. Le niveau du stock est égal à la quantité physiquement existante en stock plus la quantité en commande moins la demande en attente. Nous notons le niveau du stock avec P . Le seuil qui déclenche l'approvisionnement est noté R ou s . Les politiques de gestion du stock périodiques lancent un ordre d'approvisionnement chaque fois que le niveau du stock est au-dessous du seuil R ou s au moment du contrôle. La période de contrôle du niveau du stock est noté r .

Par rapport à la quantité à approvisionner, la littérature présente deux types de politiques : approvisionnement par quantité fixe et approvisionnement par reapprovisionnement. Dans le cas d'approvisionnement par reapprovisionnement, un ordre est lancé afin d'amener le niveau du stock à un certain niveau S . Selon les politiques d'approvisionnement par quantité fixe, un multiple de la quantité fixe est commandé à chaque fois. Cette quantité fixe est notée Q . La quantité commandée est calculée de telle manière que le niveau du stock augmente à une valeur comprise entre S et $S+Q$.

Les combinaisons des méthodes de décision concernant la date et la quantité à

approvisionner donnent naissance à trois politiques importantes de gestion du stock.

La première est la politique (R, nQ) . Cette politique de gestion du stock est caractérisée par un suivi continu ou périodique du niveau du stock P . Lorsque P descend au-dessous de R , alors une quantité nQ est commandée afin que la valeur de du stock augmente à une valeur comprise entre R et $R+nQ$. Dans la littérature, on rencontre souvent la politique (R, Q) , où la quantité approvisionnée est toujours égale à Q . Cette politique est utilisable seulement lorsque la probabilité que la demande soit supérieure à la taille du lot est négligeable.

Le deuxième type de politique est (s, S) . Cette politique considère que le niveau du stock P est contrôlé d'une manière continue et lorsque P diminue au-dessous s , un ordre d'approvisionnement est placé afin de remonter la valeur de P à S .

Le troisième type de politique est appelée base stock (R) et peut être utilisé avec un suivi continu ou périodique. Cette politique place un ordre d'approvisionnement chaque fois que le niveau du stock P diminue au-dessous d'une valeur R . La quantité commandée est calculée tel que la valeur de P est remontée à R . Dans le chapitre 3 de ce mémoire nous considérons la politique base stock, tandis que dans le chapitre 4 nous étudions les politiques (R, nQ) et (s, S) .

Des états de l'art excellents sur la gestion du stock d'un seul type de produit peuvent être trouvés dans [Lee et Nahmias, 1993], [Silver et al., 1998] et [Scarf, 2002]. Les principales contributions sont brièvement présentées par la suite.

La gestion de stock a été étudié pour la première fois dans [Harris 1913]. L'auteur considère le cas d'une demande stationnaire et déterministe. Le stock est contrôlé par une politique (R, Q) et le délai d'approvisionnement est supposé nul. Le taux de service client est une contrainte. On considère que la demande est toujours satisfaite à partir du stock. Le point de commande R est toujours égal à zéro, vu que le délai d'approvisionnement est nul. L'auteur détermine la taille de lot optimale en fonction du coût. Cette valeur de Q est dénommée quantité économique (EOQ). Elle est assez robuste et peut être utilisée à des applications plus larges.

Dans [Wagner et Whitin 1958], les auteurs considèrent que la demande est dynamique et déterministe. Ils construisent un modèle de programmation dynamique qui permet de déterminer la taille du lot qui minimise le coût.

[Scarf et Karlin 1958] considère le cas d'une demande par période stationnaire mais stochastique. Dans cet article, les auteurs considèrent que la durée du cycle de vie du produit est d'une seule période de planification. C'est le fameux modèle de marchands de journaux (News boy). Ils considèrent des coûts de possession, de commande et de pénalité ainsi que la quantité commandée. Dans ce modèle, le taux de service client exigé n'est pas présenté comme une contrainte mais il est pris en compte implicitement dans le calcul des pénalités.

L'approche présentée dans [Hadley et Whitin 1963] fournit un traitement rigoureux du modèle (R, nQ) . La demande par unité de temps est aléatoirement distribuée. Ils supposent que le point de commande, la taille des lots et la demande sont des variables continues. Ils prennent en considération les coûts de possession, de passation d'une commande et de pénalité et ils déterminent R et Q qui minimisent ces coûts. Le résultat principal obtenu est que le niveau du stock après un re-approvisionnement est uniformément distribué dans l'intervalle $(R, R+Q)$.

L'approche présentée dans [Tijms et Groeneveld, 1984] utilise des simples approximations pour le point de commande dans le cadre d'une politique de gestion de stock (s, S) . Les coûts de possession et de passation d'une commande sont pris en compte. Le point de commande optimal est calculé pour un S donné. De plus, ils supposent que la demande des clients arrive selon une loi compound Poisson et que les délais de réapprovisionnement sont aléatoires. Les temps entre deux arrivées consécutives des événements sont exponentiellement distribués et la taille de chaque commande client est une variable aléatoire de distribution quelconque.

Dans [De Kok, 1991], l'auteur utilise des approximations déduites de résultats asymptotiques pour des processus de naissance et de mort afin de déterminer le taux de service client par rapport aux paramètres de la politique de gestion du stock. Il traite les politiques de gestion du stock suivantes: (s, S) , (R, s, nQ) et (R, s, S) . Il suppose que le

temps entre deux commandes consécutives et la taille de chaque commande sont des variables aléatoires de distribution quelconque ce qui correspond à un "renewal compound process".

1.3.2 Etat de l'art sur les systèmes de production-distribution multi-niveau

Les décisions tactiques dans un réseau de production-distribution multi-niveau à un seul type de produits concernent la date et la quantité d'approvisionnement pour chaque point de stockage du système. La différence par rapport à un système de production-distribution à un seul niveau est que, dans un système multi-niveaux, le délai d'approvisionnement est une variable endogène. La composante endogène du délai d'approvisionnement est le temps d'attente à cause d'un stock insuffisant à un entrepôt intermédiaire. Des états de l'art sur ce type de modèle peuvent être trouvés dans [Federgruen, 1993], [Axsater, 1993], [Diks et al., 1996], [Houtum et al., 1996], [Zipkin, 2000], [Ettl, *et al.*, 2000] et [Liu, *et al.*, 2004].

Dans [Axsater, 1993], l'auteur propose une revue des modèles et des algorithmes pour l'analyse des politiques de gestion de stock avec un suivi continu des stocks. Le modèle classique d'une telle politique est le système METRIC [Sherbrooke 1968], développé pour les USA Air Force pour contrôler le stock de pièces détachées. Ce modèle est composé de deux échelons. L'échelon inférieur est composé de n points de stockage identiques. L'échelon supérieur est représenté par un seul entrepôt qui alimente les entrepôts de l'échelon inférieur. La modèle initial est caractérisé par un seul type de pièces, des délais de livraison indépendants et identiquement distribués, une politique d'approvisionnement unité par unité, capacité illimitée et une demande stationnaire suivant une loi de distribution Poisson composée à chaque entrepôt de l'échelon inférieur. L'objectif de cette approche est de minimiser la probabilité de rupture de stock en respectant une contrainte sur le budget. Ce modèle a été étendu de différentes manières. Par exemple, le modèle MOD-METRIC [Muckstadt 1973] étend le modèle initial avec la prise en compte de plusieurs types de produits. [Sherbrooke, 1986] et [Graves, 1985] proposent des méthodes pour améliorer la précision du modèle. [Lee et Moynadeh, 1987a, 1987b] prend en compte l'approvisionnement et la livraison par lots. [Deuermeyer et Schwartz, 1981] et [Svoronos et Zipkin, 1988] développent des approches simplifiées basées sur la décomposition. [Svoronos et Zipkin, 1991] considèrent un modèle dans lequel les ordres

d'approvisionnement sont livrés dans l'ordre d'émission. Ceci contraste les autres modèles dans lesquels les délais d'approvisionnement sont des variables aléatoires indépendantes et donc un ordre d'approvisionnement peut être livré avant un autre ordre émis plus tôt.

[Clark et Scarf, 1960], [Federgruen et Zipkin, 1984], et [Rosling, 1989] utilisent une approche de gestion centralisée. Le coût total de gestion d'un système avec un contrôle centralisé est généralement inférieur à celui d'un système décentralisé [Axsater et Rosling 1993]. Cependant, les politiques centralisées peuvent être très compliquées pour la plupart des systèmes réels. Dans la pratique, la plupart des chaînes logistiques sont pilotées d'une manière décentralisée.

[Graves, 1988] utilise un mécanisme décentralisé pour contrôler les réseaux de production. Une politique base-stock avec un suivi périodique est utilisée et la demande est stationnaire et normalement distribuée. L'approche proposée est basée sur l'agrégation des résultats obtenus pour chaque entrepôt afin d'évaluer la performance du réseau. Ce modèle est limitée par l'hypothèse de flexibilité totale du taux de production à chaque site.

Dans [Cohen et Lee, 1988] les auteurs modélisent un réseau décentralisé plus général où chaque site de production a des entrées multiples et les sorties qui alimentent un réseau divergent de distribution. Ils décomposent ce réseau en plusieurs modèles : contrôle du flux de matières, de production et de distribution. Le modèle de contrôle du flux de matières alimente les unités de production. Les sorties des unités de production alimentent un réseau divergent de distribution. La liaison principale entre la production et le système de distribution est le délai de fabrication qui devient délai d'approvisionnement pour le modèle de distribution. Le modèle de production utilise une politique de type (R, nQ) , alors que le modèle de distribution utilise une politique (s, S) au niveau de chaque site. La demande de produits finis respecte une distribution de Poisson stationnaire. La limitation du modèle vient du fait qu'une seule unité de production est admise dans le modèle.

Le modèle présenté dans [Lee et Billington 1993], dénoté par la suite LB, suppose que chaque site opère selon une politique base stock avec un suivi périodique. La demande de produits finis respecte une distribution normale. Une hypothèse importante du modèle est que le délai d'approvisionnement d'un produit à un nœud du réseau est composé de délai de

production et d'un retard. De la même manière que [Graves 1988], LB évalue les performances du réseau par l'agrégation des solutions de plusieurs problèmes caractérisés par un site et un seul type de produits. Un des avantages de cette approche est la possibilité de calculer les deux premiers moments de chaque indicateur de performance.

Dans [Garg, 1999], l'auteur développe un outil de modélisation et analyse d'une chaîne de distribution décentralisée (SCMAT) avec application pour le management de la chaîne logistique d'un grand fabricant d'équipements électroniques (LEM). SCMAT contient deux sous-modèles : (1) un modèle de files d'attente, utilisé pour calculer le délai de fabrication à chaque site et (2) le sous-modèle de gestion du stock qui calcule le niveau base stock pour chaque SKU (stock keeping unit) dans chaque entrepôt. La sortie du sous-modèle de production (le délai de fabrication) est une des entrées du sous-modèle de gestion du stock. A partir du niveau de base stock et de la valeur de la demande, du délai d'approvisionnement et du taux de service client on peut calculer différents indicateurs de performance. Ces mesures de performance incluent la valeur moyenne et la variance du stock physiquement disponible, la valeur moyenne et la variance du temps de réponse et l'utilisation de la capacité de chaque site de production.

Le modèle SCMAT étend les travaux existants à travers l'utilisation des files d'attente pour modéliser des phénomènes de congestion au niveau de sites de production. De la même manière que ses prédécesseurs, le modèle SCMAT utilise quelques approximations afin de limiter la complexité du modèle analytique. Ainsi, la demande de produits finis est stationnaire et normalement distribuée.

Dans la pratique, plusieurs produits partagent des ressources communes : lignes de fabrication, espace de stockage, etc. Il est très difficile de pré-déterminer la capacité à allouer pour chaque produit qui passe par une unité de production. De plus, les managers des LEM sont également intéressés par la capacité nécessaire à chaque site ainsi que par l'influence de cette capacité sur le délai de livraison et les stocks. Ce type d'analyse ne peut pas être effectué par des modèles de chaîne logistique avec un suivi périodique. SCMAT considère explicitement l'effet de congestion occasionné par la capacité limitée de chaque site de production, alors que les modèles développés par [Graves 1988], [Cohen et Lee, 1988] et LB ne considèrent pas cet aspect.

Les travaux de recherche présentés dans [Ettl et al., 2000] visent à modéliser des réseaux logistiques larges, comme par exemple dans l'industrie de PCs. Les auteurs développent un modèle de chaîne logistique dont les entrées sont la nomenclature du produit, les délais de livraison, la demande, les coûts et le taux de service client. La sortie du modèle est le niveau de base stock de chaque entrepôt du réseau. Ce modèle considère un mécanisme de contrôle distribué où chaque site est contrôlé par une politique base stock. En effet, il s'agit d'un mécanisme de contrôle continu, connu également en tant que politique (S-1, S). L'utilisation de cette politique de gestion de stock au niveau de chaque entrepôt élimine la nécessité de prendre en considération la taille des lots. De plus, les auteurs supposent que, vu l'absence de rupture au niveau des entrepôts amont, le délai d'approvisionnement de n'importe quel entrepôt ne dépend pas du nombre des ordres en attente. Dans ce sens, les entrepôts et la production ont une capacité illimitée. L'optimisation est effectuée à travers un modèle non-linéaire par contraintes, qui minimise le coût total de gestion du stock tout en respectant le taux de service client demandé. La méthode d'optimisation utilisée est celle du gradient conjugué, où l'estimation des gradients est calculée analytiquement. Le principal désavantage de cette approche est qu'il ne prend pas en compte la capacité de production, qui est une contrainte majeure dans un système de production.

La caractéristique clé de l'approche proposée par [Ettl et al. 2000] est l'analyse des délais de livraison actuels au niveau de chaque entrepôt et de la demande durant ce délai. Le fonctionnement de chaque site de production est modélisé par une file d'attente. La nature dynamique de la chaîne logistique et plus particulièrement l'évolution non linéaire de la demande est prise en compte en adoptant une approche d'horizon glissant. Ceci veut dire que lorsque le temps évolue et une nouvelle information est disponible, les variables de décision et les mesures de performance changent également.

De la même manière que dans le modèle LB, Ettl et al. (2000), utilise l'approximation normale pour caractériser l'évolution de la demande et pour calculer les indicateurs de performance. Cependant, contrairement au modèle LB, ils déterminent les différents délais à l'aide de l'outil file d'attente et considèrent également le problème d'optimisation. Au lieu de considérer une demande stationnaire, ils traitent une demande non-stationnaire à travers une approche à horizon glissant.

Dans [Liu, *et al.*, 2004] les auteurs introduisent un système de gestion de stock multi-niveau basé sur les files d'attente. Une approche de décomposition est utilisée pour évaluer les performances d'un système de production-distribution serial avec une politique de base stock à chaque niveau. Une approche analytique est proposée pour minimiser le coût de stockage. Le modèle est basé sur une file d'attente GI/G/1, où les produits arrivent un par un. La taille des lots ainsi que les délais de transport ne sont pas prises en compte.

1.3.3 Etat de l'art sur la simulation des chaînes logistiques

La performance d'une chaîne logistique peut être améliorée par la réduction des incertitudes. Il est clair qu'il est nécessaire de disposer d'un mécanisme de coordination des activités et des processus au sein et entre les parties d'une chaîne logistique afin de réduire les incertitudes et d'augmenter la valeur ajoutée aux clients. Ceci nécessite d'établir des relations d'interdépendance entre les variables de décision et les différentes composantes d'une chaîne logistique. Ces relations peuvent varier avec le temps et il est très difficile de les étudier à travers un modèle analytique. La simulation fournit des moyens bien plus flexibles pour modéliser la dynamique de réseaux complexes. Elle est considérée la méthode la plus fiable pour l'analyse des performances dynamiques des chaînes logistiques. De même, la simulation fournit un outil effectif pour évaluer les efforts de re-ingénierie relatifs à une chaîne logistique en termes de performances et de risque.

Dans [Towill, 1991] et [Towill *et al.*, 1992] des techniques de simulation sont utilisées afin d'évaluer les effets des différentes stratégies de pilotage d'une chaîne logistique face à une augmentation de la demande. Les stratégies identifiées sont les suivantes [Towill *et al.* 1992]:

1. éliminer le niveau de distribution de la chaîne logistique en intégrant cette fonction au niveau de production ;
2. intégrer le flux d'information à travers la chaîne logistique ;
3. implémenter une politique de gestion du stock à flux tiré pour réduire les délais ;
4. améliorer le mouvement des produits intermédiaires en modifiant la taille des lots ;
5. modifier les paramètres des procédures de lancement d'ordres de fabrication existantes.

L'objectif d'un modèle de simulation est de déterminer quelles stratégies sont les plus intéressantes pour réduire la variation de la demande. Les stratégies 3 et 1 ont fourni les meilleurs résultats dans le cas étudié.

[Bhaskaran,1998] montre l'importance d'un processus de re-ingénierie d'un projet pour des opérations de fabrication chez General Motors en utilisant la simulation en tant qu'outil principal. L'article décrit le niveau de détail nécessaire pour comprendre les flux de matières et de l'information. Il évalue différentes configurations du système afin d'identifier les améliorations. Cependant, les interactions au niveau d'une chaîne logistique impliquent généralement des mécanismes de contrôle plus sophistiqués. Par exemple, lorsqu'un ordre important arrive, il doit être traité avec priorité par rapport aux autres ordres. De même, le traitement d'une unité de produit représente plus qu'attendre à un centre de traitement pendant un certain temps. Par exemple, lorsqu'un ordre est traité, des composants peuvent être assemblés ce qui, en revanche, engendre certains événements selon la quantité disponible en stock. Des règles de décision doivent être utilisées à chaque occurrence des événements.

Malgré les avantages des modèles de simulation, il y a deux problèmes majeurs liés à la construction des modèles de simulation « sur mesure » : (1) durée de développement importante et (2) ils sont très spécifiques, donc on ne peut pas les réutiliser. Dans [Swaminathan, Smith et Sadeh, 1998], les auteurs ont proposé un environnement qui permet le développement rapide des outils pour aide à la décision pour le management de la chaîne logistique. L'objectif est de re-configurer rapidement la chaîne logistique basée sur des études menées sur différents composants de chaîne logistique mises à disposition à travers une librairie. Ils utilisent un paradigme multi-agent pour modéliser et analyser une chaîne logistique. Les environnements de calcul multi-agents sont adaptés pour l'étude d'un large ensemble de problèmes de coordination incluant plusieurs agents de résolution de problèmes autonomes ou semi-autonomes. Dans [Swaminathan, Smith et Sadeh, 1998] on identifie plusieurs agents au sein d'une chaîne logistique et on munit chaque agent de la capacité à utiliser un sous-ensemble d'éléments de contrôle. Les éléments de contrôle aident à la prise de décision au niveau agent à l'aide de différentes politiques pour la demande, l'approvisionnement, l'information et le contrôle du flux de matières à travers la chaîne logistique. Leur analyse est basée sur la simulation à événements discrets des

différentes alternatives et politiques de contrôle. Par conséquent, la combinaison de modèles analytiques avec des modèles de simulation offre un cadre intéressant pour l'analyse des aspects à la fois statiques et dynamiques d'un problème.

Rappelons que, comme précisé dans la section 1.2, un modèle générique de simulation a été développé dans [Ding *et al.* 2004a] pour la simulation des chaînes logistiques dans une perspective de l'optimisation de la conception des chaînes logistiques.

1.3.4 Conclusions sur l'optimisation des chaînes logistiques

D'un coté, les problèmes rencontrés dans l'optimisation des chaînes logistiques sont souvent difficiles à transposer en modèles mathématiques d'optimisation. Selon la taille d'un système de production-distribution réel, souvent il y a des dizaines de milliers de contraintes et des variables dans le cas déterministe. Par conséquent, les méthodes d'optimisation déterministe traditionnelles n'arrivent pas à capter avec une précision suffisante la dynamique de la plupart des systèmes de production-distribution réels. La raison principale est que ces applications sont caractérisées par des incertitudes au niveau de leur évolution. Ce phénomène est généré par le fait qu'une partie de l'information nécessaire à la prise de décision n'est pas disponible au moment où la décision doit être prise [Beamon, 1998]. De même, très souvent le problème à résoudre doit être simplifié afin de permettre des solutions analytiques. Même si ces méthodes analytiques aide à trouver des bonnes solutions, elles sont applicables seulement sous certaines hypothèses particulières.

D'un autre coté, la simulation est une technique de plus en plus utilisée dans la conception des chaînes logistiques à cause de sa capacité à analyser la variabilité et les interdépendances dans l'évolution d'un système. La simulation permet à un décideur d'évaluer les changements dans une partie de la chaîne logistique et de visualiser l'effet de ces changements sur d'autres composants du système et sur la performance globale du système.

Les modèles de simulation peuvent fournir une série de résultats et permettent aux utilisateurs d'améliorer la précision de ces résultats afin de construire un système robuste et

prédictif. En effet, la simulation est le seul outil qui peut suivre tous les indicateurs de performance du système et de prévoir les interactions entre ces indicateurs dus au changement des entrées du système. De plus, la simulation permet d'inclure les variations dans le système dans le modèle. Parmi les facteurs qui génèrent une variation importante dans le système sont : la demande du client, les dates de livraison de la matière première ainsi que des produits finis, l'arrivée des pannes, la capacité de stockage. Chacune de ces variables peut causer des variations importantes dans les performances de la chaîne logistique et doit être surveiller de près afin de comprendre son influence.

L'association des techniques d'optimisation avec des techniques de simulation est donc naturelle et fait l'objet d'un important développement ce dernier temps. L'optimisation des modèles de simulation s'adresse à la situation où le décideur souhaite savoir lequel des jeux de valeurs des paramètres permet d'obtenir une performance maximale. Dans le domaine de plans d'expériences, les paramètres d'entrée, ainsi que les paramètres structuraux associés à un modèle de simulation sont appelés des facteurs. Une mesure de la performance obtenue est appelée réponse. Par exemple, un modèle de simulation d'une unité de fabrication peut inclure des facteurs tels que le nombre de machines de chaque type, le réglage des machines le nombre d'opérateurs ayant un certain niveau de compétence. Des exemples de réponse peuvent être le temps de cycle, le niveau d'en-cours et l'utilisation des ressources.

Dans le domaine de l'optimisation les facteurs deviennent des variables de décision. Les réponses sont utilisées pour modéliser les fonctions objectif et les contraintes. L'objectif des plans d'expériences est de trouver quels sont les facteurs qui ont l'effet le plus important sur la réponse. Par contre, l'objectif de l'optimisation est de trouver la combinaison de valeurs de facteurs qui minimise ou maximise la réponse compte tenue des contraintes imposées.

1.4 Etat de l'art sur l'optimisation discrète par la simulation

La simulation stochastique à événements discrets est un outil largement utilisé pour l'analyse dynamique des incertitudes dans l'évolution des systèmes. Une caractéristique

importante des expériences de simulation est que les utilisateurs peuvent utiliser différents jeux de valeurs pour les paramètres du système afin d'essayer d'améliorer ses performances. Par conséquent, il est naturel de chercher des jeux de valeurs pour les paramètres qui optimisent les performances. Cette technique, appelée optimisation par la simulation (OvS), fournit une approche structurée pour déterminer une valeur optimale pour les variables de décision, où l'optimalité est mesurée par rapport aux variables de sortie du modèle de simulation. Ce dernier temps, l'intérêt des problèmes d'OvS s'est penché plus vers des problèmes avec des variables de décision discrètes (DOvS).

Plusieurs techniques ont été développées pour les problèmes DOvS [James, Paul etc.2004]. Si la taille de l'espace de solutions est finie et petite, la méthode de tri et sélection ou rank and selection [Goldsman et Nelson,1998b] et celle des comparaisons multiples sont appropriées. Par contre, si l'espace de solutions est très large, voir infini, alors des méthodes telles que l'optimisation ordinale [Ho, Sreeniva, et Vakili,1992], [Ho,1994] [Ho,2002] [Ho, et al.2000] et [Jia,et al.2005], le recuit simulé [Eglese,1990], [Flescher,1995], la méthode de la comparaison stochastique [Gong et al. 1999], la recherche tabu [Glover et Laguna,1997] et les algorithmes génétiques [Liepins et Hilliard,1989], [Muhlenbein,1997] ont été adaptés. Des états de l'art sur les techniques d'optimisation par la simulation peuvent être trouvés dans [Fu, 2002], [Andradottir, 1998a,b] et [Olafsson et Kim, 2002].

Optimisation ordinale : L'objectif de l'optimisation ordinale [Ho, Sreeniva, et Vakili,1992] [Ho,1994] est de trouver des solutions suffisamment bonnes, plutôt que de chercher la meilleure solution et évaluer précisément les mesures de performance (i.e., lissage de l'objectif) [Lee *et al.*,1999]. Pour cela, on s'intéresse au classement des solutions au lieu des valeurs exactes de leur performances. Par conséquent, l'optimisation ordinale réduit la recherche d'une solution optimale d'un ensemble important de solutions et la valeur de critère correspondant à une analyse sur un espace plus petit de bonnes solutions. La propriété de convergence exponentielle de l'optimisation ordinale a été analysée dans [Dai, 1995] et [Xie, 1997]. Cette dernière montre que dans un processus régénératif la probabilité d'obtenir la solution désirée par l'optimisation ordinale converge exponentiellement alors que la variance du mesure de performance converge avec un taux de $O(1/t^2)$, où t est le temps de simulation.

Recuit simulé : Le recuit simulé (SA) [Eglese,1990], [Flescher,1995] est une méthode de recherche inspirée du processus de recuit où un alliage est refroidit graduellement afin d'augmenter ses performances. Il accepte ou refuse une direction d'avancement selon une probabilité donnée. Le recuit simulé commence avec une température haute, où la probabilité d'accepter une direction apparemment mauvaise est grande. Ensuite la température diminue au bout de quelques itérations. La probabilité d'accepter une solution plus mauvaise diminue avec la température pour avoisiner zéro.

[Baretto et al., 1999] les auteurs proposent un algorithme basé sur le principe du recuit simulé et l'appliquent dans le cas de la production de l'acier. [Yucesan et Jacobson, 1996] présente un algorithme de recuit simulé pour résoudre le problème ACCESSIBILITY. [Andradottir, 1996] utilise une approche de marche aléatoire pour développer un algorithme d'optimisation par simulation pour un ensemble large de paramètres. Dans [Alrefaei et Andradottir, 1999] les auteurs adaptent l'algorithme de recuit simulé aux problèmes DOvS. Ils définissent également les conditions pour que l'algorithme converge le plus sûr vers un optimum global, avec un programme de refroidissement fixé.

Il y a plusieurs autres variantes de recuit simulé. Cependant, peu d'approches basées sur cette méthode ont été utilisées pour l'optimisation par la simulation [Bulgak et Sanders, 1988], [Haddock et Mittenthal,1992]. La faible utilisation dans la pratique du recuit simulé est due à la sensibilité du résultat au bruit dans les mesures.

Algorithmes génétiques : Les algorithmes génétiques (GA) [Liepins et Hilliard,1989], [Muhlenbein,1997] et leurs variantes sont similaires au processus de sélection naturelle et au processus d'évolution biologique. Différentes configurations sont codées par des chaînes de valeurs similaires aux chromosomes. GA garde une population d'individus qui peut évoluer au fur et à mesure que les chromosomes sont créés et modifiés. Chaque chromosome est composé des gènes qui décrivent les paramètres dans les jeux de paramètres de la simulation. Une génération est composée des individus de la génération antérieure qui ont survécu et des nouveaux individus qui sont générés par des processus de mutation, de croisement et de sélection.

Les GA sont très intéressants et génériques. Le calcul des certains paramètres nécessaires

aux GA (taille de la population, probabilité de croisement et de mutation,...) peut être difficile et il est lié au problème à résoudre. Le temps de calcul risque d'être l'obstacle principale de cette méthode.

Nested partitions : La méthode d'optimisation appelée Nested Partition [Shi et Olafsson 2000a,b] et [Pichitlamken et Nelson 2003] est dédiée à résoudre des problèmes d'optimisation globale. Cette méthode partage d'une manière systématique l'espace des solutions et concentre sa recherche sur la partition qui semble la plus prometteuse. Une région la plus prometteuse est sélectionnée à chaque itération en fonction de l'information obtenue à travers de l'analyse de quelques solutions choisies aléatoirement parmi les solutions admissibles et d'une recherche locale. Cette région est ensuite partagée plus finement et les autres partitions sont regroupées en une seule. Par conséquent, cette méthode combine à la fois la recherche locale que la recherche globale. Elle converge avec la probabilité 1 vers un optimum global en temps fini. Dans [Olafsson et Shi 2002] on montre que la combinaison avec une optimisation ordinale offre des nouvelles perspectives pour améliorer la convergence de la méthode nested partitions. Cette approche est également utilisée pour montrer que la convergence globale nécessite un nombre de simulations important et propose des nouvelles variantes plus efficaces pour l'algorithme.

COMPASS : La méthode appelée Convergent Optimization via Most-Promising-Area Stochastic Search (COMPASS) est proposée par [Hong et Nelson 2006]. Elle est destinée aux problèmes d'optimisation où les variables de décision sont des nombres entiers et les performances sont estimées par simulation stochastique à événements discrets. Cette méthode est basée sur une recherche aléatoire dans une nouvelle structure de voisinage appelée région la plus prometteuse. Les auteurs montrent que la méthode COMPASS permet de résoudre une large variété de problèmes d'optimisation, y compris lorsque l'espace des solutions admissibles est très large. La supériorité de cette méthode par rapport à d'autres algorithmes est illustrée à travers des expériences numériques. Ils montrent également que l'algorithme converge avec la probabilité 1 vers un ensemble de solutions localement optimales à condition que chaque solution visitée soit simulée un nombre de fois qui tend vers l'infini. Nous jugeons cette condition comme étant trop restrictive. Elle engendre la consommation inutile d'une partie du budget de simulation.

OCBA : Différentes méthodes ont été proposées pour optimiser l'utilisation du budget de simulation pour sélectionner l'optimum parmi un nombre important de solutions candidates. Ainsi, dans [Chen et Shi, 2000], [Chen,Dai,Chen et Yücesan,1998] et [Chen, Lin, Yücesan, et Chick S. E, 2000] les auteurs développent une approche pour allouer d'une manière intelligente le budget de simulation aux différentes solutions candidates. Cette approche, appelée OCBA (Optimal Computing Budget Allocation), permet de déterminer d'une manière intelligente la meilleure longueur de simulation pour chaque solution candidate afin de réduire le coût de calcul pour obtenir le même niveau de confiance pour la solution optimale. Ils comparent également cette approche avec la procédure Rinott à deux niveaux [Dudewicz et Dalal, 1975] issue du plan d'expériences. Les résultats numériques du test ont montré que leur approche est dix fois plus rapide.

Dans [Andradóttir, 1999], l'auteur fournit un schéma pour accélérer la convergence des algorithmes DOvS en accumulant les observations antérieures. Elle discute l'estimation de la solution optimale lorsque des méthodes de recherche aléatoires sont appliquées pour résoudre des problèmes DOvS. Les méthodes d'optimisation existantes choisissent en tant que solution optimale soit la solution admissible courante soit celle qui a été visité le plus souvent. Elle propose l'utilisation de toutes les valeurs de la fonction objectif observées durant l'exploration du domaine admissible par la méthode de recherche. L'objectif est d'obtenir des estimations de plus en plus précises de la fonction objectif à des points différents du domaine admissible. A chaque itération, la meilleure solution rencontrée est utilisée en tant qu'estimation de la solution optimale. Par conséquent, cette approche utilise pleinement l'historique de la simulation.

1.5 Positionnement des contributions de cette thèse

Un problème central dans la conception des systèmes de production-distribution est de quantifier et d'optimiser le compromis entre le niveau de service des clients et le coût de stockage pour supporter ces spécifications. Le problème est rendu encore plus difficile à cause de la dynamique d'un système de production-distribution : variété de produits finis, cycle de vie des produits de plus en plus court, lancements fréquents de nouveaux produits,

changements fréquentes des spécifications client. La nature dynamique des systèmes de production-distribution complexes implique que le compromis entre le niveau de service client et le coût global de stockage doit être re-analysé et optimisé périodiquement. Par conséquent, un modèle formel est une aide précieuse pour l'analyse et l'optimisation des systèmes de production-distribution. De plus, afin de répondre en temps utile à diverses questions de type « What if » ces modèles formels doivent être efficaces du point de vue du temps de calcul.

L'activité de recherche que nous menons dans le cadre de cette thèse de doctorat concerne deux aspects de l'optimisation des chaînes logistiques : (i) l'analyse et l'optimisation d'un réseau de production-distribution à travers une approche analytique ; et (ii) l'optimisation par la simulation des réseaux de production-distribution intégrés. L'objectif est de montrer que, sous certaines conditions restrictives et pour des systèmes simples, les solutions analytiques permettent d'apporter rapidement des réponses à des questions sur la conception et l'optimisation des chaînes logistiques. Pour le cas général, la simulation est incontournable mais une utilisation intelligente de la simulation permet de donner une réponse dans un temps raisonnable à des problèmes de conception et d'optimisation.

1.5.1 Analyse et optimisation d'un système de production-distribution

La motivation de notre travail de recherche est de développer une méthode analytique pour l'évaluation de performances et l'optimisation des systèmes de production-distribution réels, caractérisés par production par lots de taille aléatoires, capacité de production finie et délai de transport.

L'approche que nous proposons dans le chapitre 2 concerne le problème d'évaluation de performance et l'optimisation du coût de gestion du stock sous contraintes du niveau de service client dans un système de production-distribution à deux niveaux, composé d'un entrepôt alimenté par une unité de production. L'arrivée des clients à l'entrepôt est aléatoire. Une commande client est satisfaite si la quantité nécessaire est disponible en stock. Elle est mise en attente si le stock est vide. L'entrepôt est géré selon une politique base stock. Les ordres de re-approvisionnement sont envoyés à l'unité de fabrication selon l'arrivée des clients. La capacité de production de l'usine est finie. De plus, le délai de transport d'un lot

entre l'unité de production et l'entrepôt est non-négligeable et constante.

Une approche analytique est développée pour l'évaluation des performances de ce système. Cette approche est basée sur plusieurs approximations techniques, qui ont prouvé leur efficacité pour la plupart des systèmes de production : (i) approximation par les deux premiers moments ; (ii) approximation markovienne des processus généraux ; (iii) approximation des processus aléatoires dépendantes par des processus indépendantes. Plus précisément, cette approche analytique utilise seulement les deux premiers moments des variables aléatoires pour évaluer le délai de re-approvisionnement de l'entrepôt, le stock en commande et le coût de possession et de rupture du stock. L'unité de production est modélisée par une file d'attente $M^X/G/1$ et le transport par une file d'attente $M^X/D/\infty$. Des solutions analytiques approximatives sont obtenues pour le premier et le deuxième moment du délai de re-approvisionnement de l'entrepôt et du stock en commande. Une distribution log-normale est utilisée pour approximer la distribution de probabilité associée au stock en commande et pour calculer le coût total de possession et de rupture de stock à l'entrepôt. Le coût total de gestion du stock est minimisé à l'aide d'une technique de gradient tout en prenant en compte des contraintes de taux de service client. Les résultats numériques montrent l'efficacité de l'approche.

Malgré les résultats excellents obtenus pour un système de production distribution à deux niveaux, l'extension de cette approche pour l'analyse et l'optimisation des systèmes de production-distribution multi-niveaux conduit à des erreurs trop importantes pour être utiles. De plus l'extension à d'autres politiques de pilotage semble très difficile. Par conséquent, nous avons focalisé notre attention sur des méthodes d'évaluation de performances et optimisation par la simulation pour les systèmes généraux.

1.5.2 Optimisation par simulation des systèmes de production-distribution

Dans cette deuxième partie de la thèse, nous nous intéressons à l'optimisation du système de pilotage et développons des méthodes d'optimisation efficaces pour un système de production-distribution multi-niveau. Plus précisément, nous considérons un système de production-distribution composé des différentes sites de production et des centres de distribution reliés par des moyens de transport. Chaque site est gérée par une politique de

gestion de stock donnée. Le problème consiste à déterminer le meilleur paramétrage des politiques de gestion du stock de l'ensemble des sites afin d'optimiser les performances globales du système de production-distribution tout en tenant compte des contraintes de production, des aléas de la demande et des aléas de transport, etc.

L'évaluation analytique d'un système à plusieurs niveaux avec contraintes de capacité de production, des délais de transport, des politiques de gestion de stock variées et des tailles de lot des approvisionnements et de transport est particulièrement difficile. Plusieurs facteurs contribuent à rendre l'analyse analytique difficile : (i) l'interactions entre les décisions prises pour les différents sites. L'émission d'un ordre d'approvisionnement peut déclencher des approvisionnements à d'autres niveaux; (ii) propagation des erreurs d'évaluation à un niveau à d'autres niveaux; (iii) difficulté de tenir en compte d'autres politiques de gestion de stock. L'approche analytique du chapitre deux s'appuie fortement sur le fait que, sous la politique base stock, l'arrivée des ordres d'approvisionnement à chaque niveau suit le même processus que la demande extérieure. Cette propriété n'est plus vraie sous d'autres politiques de gestion de stock; (iv) l'arrivée aléatoire des demandes, la fabrication et le transport par lot compliquent également l'analyse analytique. La plupart des méthodes analytiques se limite à la production et au transport par unité.

Compte tenu de la difficulté d'une analyse analytique, nous nous orientons vers l'association de la simulation et de l'optimisation pour évaluer les performances et optimiser un système de production-distribution multi-niveaux. Le problème se pose de deux manières différentes selon la prise en compte du niveau de service client en tant que un critère à optimiser ou en tant que une contrainte à satisfaire.

Lorsque le critère à optimiser inclut les coûts et le niveau de service des clients, l'optimisation des politiques de gestion de stock peut être décrite par un problème d'optimisation discrète :

$$\min_{x \in \Theta} f(x)$$

dans lequel (i) les variables de décisions x sont des entières et représentent les paramètres des politiques de gestion de stock à optimiser, et (ii) la fonction objective $f(x)$ est inconnue mais peut être estimée pour chaque solution x à l'aide de la simulation par réplifications, i.e.

$$f(x) = E[F(x, w_i)].$$

Nous développons une méthode d'optimisation appelée COMPASS* qui alloue de manière dynamique le temps de simulation (i.e. nombre de répliques) à des solutions jugées intéressantes. Nous prouvons que la méthode converge vers des vraies solutions optimales locales.

Pour résoudre des problèmes d'optimisation des chaînes logistiques avec contraintes de niveaux de service, nous nous intéressons ensuite à une extension du problème d'optimisation précédente avec contrainte $g(x) > a$ où, de la même manière que $f(x)$, la contrainte $g(x)$ est inconnue mais peut être estimée pour chaque solution x à l'aide de la simulation. Nous avons étendu la méthode COMPASS* pour allouer de manière dynamique le temps de simulation à des solutions jugées intéressantes pour l'estimation de $f(x)$ et de $g(x)$. Encore une fois, la méthode converge vers des vraies solutions optimales locales.

La méthode COMPASS* est ensuite appliquée pour étudier le problème de paramétrage optimal des différentes politiques pour la gestion des différents stocks dans un système de production-distribution. Elle s'est révélée très efficace et permet de trouver de très bonnes solutions dans un temps de calcul faible. De plus, nous montrons comment des propriétés importantes des solutions optimales peuvent être prise en compte pour une optimisation plus efficace et plus rapide.

Chapitre 2

Une approche analytique pour évaluation et optimisation des performances des systèmes de production-distribution à deux niveaux

Ce chapitre traite le problème d'évaluation de performances et d'optimisation d'un système de production-distribution composé d'un entrepôt alimenté par une unité de production à capacité finie. L'arrivée des commandes clients à l'entrepôt, ainsi que la quantité commandée à chaque fois est aléatoire. Nous proposons une approche analytique qui utilise seulement les deux premiers moments des variables aléatoires pour évaluer le coût total de gestion du stock et le taux de service client. Une technique de gradient est utilisée pour minimiser le coût total de gestion du stock par rapport à un taux de service donné. Les expériences numériques ont montré l'efficacité de l'approche proposé.

Publications: [LI,SAVA et XIE, 2004,2005a,b,2006a]

2.1 Introduction

La conduite des systèmes de production-distribution nécessite un compromis approprié entre le coût et les performances. L'objectif de ce chapitre est d'introduire une approche analytique pour évaluer et optimiser ces grandeurs dans le cas d'un système de production-distribution composé d'un entrepôt alimenté par un système de production.

Dans le système considéré dans ce chapitre, l'arrivée des ordres des clients et la quantité commandée sont aléatoires. Un client est servi si la quantité de produits en stock est suffisante. Il est mis en attente si le stock est vide. Le niveau du stock est géré selon une politique base-stock. Les ordres d'approvisionnement du stock sont envoyés à l'unité de production au fur et à mesure que les clients arrivent. La capacité de production de l'unité de production est finie. Le transport des produits de l'unité de production à l'entrepôt respecte les ordres d'approvisionnement et le délai de transport est constant.

L'approche analytique proposée pour l'évaluation et l'optimisation des performances est basée sur des approximations techniques dont l'efficacité a été prouvée pour la plupart des systèmes de production : (i) approximation par les deux premiers moments des variables aléatoires ; (ii) approximation des processus généraux par des processus markoviens ; (iii) approximation des processus aléatoires dépendants par des processus aléatoires indépendants.

Concrètement, cette approche utilise seulement les deux premiers moments des variables aléatoires pour évaluer le délai d'approvisionnement de l'entrepôt, le niveau du stock et les coûts de stockage et de rupture. L'unité de production est modélisée par une file d'attente $M^X/G/1$ et le processus de transport entre l'unité de production et l'entrepôt par une file d'attente $M^X/G/\infty$. Le délai d'approvisionnement de l'entrepôt et le niveau du stock sont évalués par des expressions analytiques déduites sur la base des approximations annoncées. Ensuite, la loi de distribution log-normale est utilisée pour approximer le niveau du stock ainsi que pour évaluer les coûts de stockage et de rupture du stock. Le coût global de gestion du stock est minimisé à l'aide d'une méthode d'optimisation basée sur le calcul des gradients, tout en respectant le niveau de service demandé. La comparaison avec les valeurs numériques obtenues par la simulation montre l'efficacité de l'approche ainsi que son

insensibilité à la charge de travail du système de production.

A notre connaissance il n'existe pas dans la littérature aucune approche d'évaluation des performances et d'optimisation de politiques de gestion de stock qui prenne en compte simultanément (i) une demande aléatoire ; (ii) une capacité de production finie ; (iii) demande, fabrication et livraison par lots de taille aléatoire ; (iv) temps nécessaire au transport des produits.

2.2 Topologie du système et approche de modélisation

Le système de production-distribution considéré est composé d'une unité de production qui alimente un entrepôt (figure 2.1). On considère un seul type de produits. L'entrepôt doit faire face à une demande aléatoire par lots et gérer l'approvisionnement à l'unité de production. L'unité de production traite les ordres reçues selon une politique FIFO en fonction de sa capacité de fabrication. Le traitement d'un ordre correspond à la fabrication du nombre des produits demandé. Les produits associés à un ordre d'approvisionnement sont envoyés à l'entrepôt. La durée du transport est suffisamment importante pour qu'on ne puisse pas la négliger. Une description détaillée de ce système est donnée par la suite.

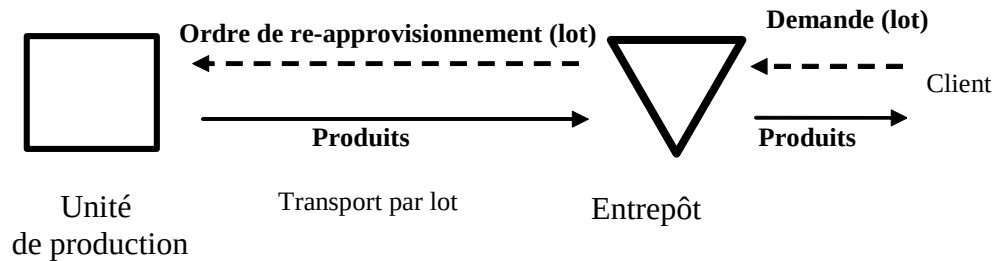


Figure 2.1. La topologie du système

Les clients arrivent à l'entrepôt aléatoirement, selon un processus de Poisson avec un taux d'arrivée λ . La quantité commandée est une variable aléatoire positive X . Les quantités associées à des différentes demandes des clients sont indépendantes. La quantité X associée à un ordre est caractérisée par ses deux premiers moments : la moyenne m_X et l'écart type σ_X . Lorsqu'un ordre de quantité X arrive à l'entrepôt, si le stock disponible est suffisant,

alors la quantité demandée est immédiatement livrée au client. Sinon, on livre la quantité disponible et le client est mis en attente, le temps que l'entrepôt soit re-approvisionné par l'unité de production.

Le stock de l'entrepôt est géré selon une politique base-stock. Cette politique est basée sur le concept de position du stock qui est égal au stock disponible, moins les demandes en attente, plus la quantité commandée à l'unité de production et pas encore reçue. L'objectif de cette politique de gestion de stock est de maintenir la position du stock à une valeur constante appelée base-stock R . Par conséquent, chaque fois qu'un ordre de taille X arrive, l'entrepôt envoie immédiatement un ordre d'approvisionnement de taille X à l'unité de production. Ainsi l'arrivée des ordres d'approvisionnement à l'unité de production est un processus stochastique identique à l'arrivée des clients à l'entrepôt (i.e. un processus de Poisson).

L'unité de production est modélisée par une file d'attente avec un seul serveur. Les ordres d'approvisionnement de l'entrepôt arrivent selon un processus de Poisson et attendent dans une file FIFO pour être prises en charge par la production. Chaque ordre désigne un lot composé de plusieurs unités de produit correspondant à un ordre d'approvisionnement. Le serveur traite chaque ordre unité par unité. Le temps de fabrication d'une unité de produit est une variable aléatoire τ qui suit une distribution générale de moyenne m_τ et écart type σ_τ . La durée de fabrication d'un ordre d'approvisionnement X dépend de sa taille. Plus précisément, l'unité de production peut être modélisée par une file d'attente $M^X/G/1$. Le temps écoulé entre l'arrivée d'un ordre d'approvisionnement à l'unité de production et l'instant où les produits commandés sont prêts à être envoyés à l'entrepôt est appelé délai de fabrication. Il est noté L_s .

Les produits fabriqués sont livrés à l'entrepôt par lots sur la base des ordres d'approvisionnement. Dès que le traitement d'un ordre de taille X est achevé, un lot de X produits est transporté à l'entrepôt. Le temps nécessaire pour le transport d'un lot à l'entrepôt est appelé délai de transport. Il est noté L_t . Nous supposons que le délai de transport est donné et constant.

Le délai d'approvisionnement de l'entrepôt, noté L , représente le temps écoulé depuis

l'envoi d'un ordre d'approvisionnement à l'unité de production et la réception des produits commandés par l'entrepôt. Il est défini par la somme du délai de fabrication et du délai de transport, i.e. $L = L_s + L_t$ (figure 2.2).

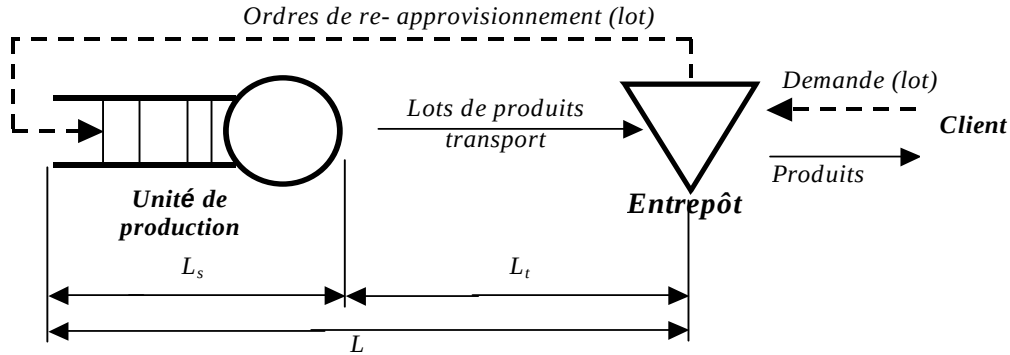


Figure 2.2. Principe de l'approche de modélisation

Un critère important dans la gestion d'un stock est le coût total représenté par les coûts de stockage et de rupture. Le stockage d'une unité de produit engendre un coût C_h par unité de temps. De même, un produit qui ne peut pas être livré au client en raison d'une rupture de stock engendre un coût C_b par unité de temps.

Un autre indicateur pour la gestion de stock est le niveau de service des clients que nous mesurons par le taux de remplissage des demandes des clients. Dans cette thèse, le taux de remplissage est le pourcentage des quantités demandées par les clients satisfaites immédiatement, sans attente.

Par la suite, nous proposons une technique analytique pour évaluer le premier et le deuxième moment des indicateurs de performance décrits précédemment.

2.3 Evaluation des performances

Cette section porte sur l'évaluation des performances et du coût d'un système de production-distribution. Nous évaluons d'abord le premier et/ou le deuxième moment des indicateurs de performance suivants :

- L_s : délai de fabrication d'une unité de produit (pas d'un lot),
- N_s : nombre d'unités de produit associée aux ordres d'approvisionnement en attente ou

en cours de fabrication,

- N_t : nombre total de produits en transit vers l'entrepôt,
- L : délai de approvisionnement de l'entrepôt,
- $N = N_s + N_t$: stock en commande, i.e. nombre total de produits commandés par l'entrepôt et pas encore reçus.

Ces indicateurs de performance seront ensuite utilisés pour l'évaluation des coûts de gestion de stock et le taux de service client.

Nous insistons sur le fait que la taille de lots définies par les ordres d'approvisionnement sont préservées jusqu'à la réception des produits à l'entrepôt. Par conséquent, par abus de langage on désigne un ordre d'approvisionnement par un lot, lorsque aucune confusion n'est possible.

2.3.1. Le modèle de l'unité de production

L'unité de production est modélisée par une file d'attente $M^X/G/1$. Le stock à l'entrepôt étant géré par une politique base-stock, l'arrivée des ordres d'approvisionnement suit le même processus de Poisson que la demande des clients.

Soit T_B le temps de service d'un ordre de approvisionnement par le serveur. Par définition,

$$T_B = \sum_{i=1}^X \tau_i \quad (2.1)$$

où X est la taille de l'ordre et τ_i est le temps du service de la i -ème unité de l'ordre. Puisque que X et τ_i sont des variables aléatoires indépendantes,

$$m_{TB} = E[T_B] = m_\tau m_x \quad (2.2)$$

où $m_\tau := E[\tau]$ et $m_x := E[X]$. Conditionnant en X , l'indépendance entre X et τ_i implique :

$$\begin{aligned}
 E[T_B^2] &= E[E[(\sum_{i=1}^X \tau_i)^2 | X]] \\
 &= E[X \cdot E[\tau^2] + X(X-1)(E[\tau])^2] \\
 &= m_x \sigma_\tau^2 + m_x^2 (\sigma_x^2 + m_x^2)
 \end{aligned} \tag{2.3}$$

où, $\sigma_\tau^2 := \text{Var}[\tau]$ et $\sigma_x^2 := \text{Var}[X]$.

A partir des relations (2.2) et (2.3),

$$\begin{aligned}
 \sigma_{TB}^2 &:= \text{Var}(T_B) \\
 &= E[T_B^2] - E^2[T_B] \\
 &= m_x \sigma_\tau^2 + m_x^2 \sigma_x^2
 \end{aligned} \tag{2.4}$$

Considérons les ordres d'approvisionnement comme les unités de base prises en compte par la file d'attente. Alors, la file d'attente $M^X/G/1$ peut être vue comme une file d'attente $M/G/1$ avec le temps de service T_B . Le nombre moyen N_B de lots présents dans le système de production peut être évalué par la formule Pollazek-Khinchin ([Kleinrock,1975],[Cassandras,1993]):

$$E(N_B) = \frac{\rho}{1-\rho} - \frac{\rho^2}{2(1-\rho)} (1 - \sigma_{TB}^2 / m_{TB}^2) \tag{2.5}$$

où ρ est l'intensité du trafic de la file d'attente, définie par :

$$\rho = \lambda m_{TB} \tag{2.6}$$

A partir de (2.2)-(2.6) on obtient:

$$E(N_B) = \lambda \frac{\lambda(2m_x m_\tau - m_x^2 m_\tau^2 \lambda + \lambda m_x \sigma_\tau^2 + \lambda \sigma_x^2 m_\tau^2)}{2(1 - m_x m_\tau \lambda)} \tag{2.7}$$

Selon la loi de Little, la durée moyenne L_B qu'un ordre d'approvisionnement passé dans l'unité de production est donnée par l'expression suivante:

$$E[L_B] = E[N_B] / \lambda \tag{2.8}$$

Nous attirons l'attention sur les concepts suivantes: (i) la durée moyenne $E[L_B]$ qu'un ordre passe dans l'unité de production et (ii) la durée moyenne $E[L_s]$ qu'une unité de produit passe dans l'unité de production. Ces deux concepts sont différents. Cependant, le temps qu'une unité de produit d'un ordre de taille X passe dans le système de production est

exactement le même que le temps que l'ordre correspondant passe dans le système de production, i.e. $L_s = L_B$. Cependant, $E[L_B]$ et $E[L_s]$ diffèrent selon la distribution de la taille X de l'ordre. Des expériences numériques nous ont permis de conclure que $E[L_B]$ et $E[L_s]$ sont très proches pour les distributions couramment utilisées de X . Par conséquent, on utilise l'approximation :

$$E[L_s] \approx E[L_B] \quad (2.9)$$

La loi de Little, appliquée aux unités de produits, fournit les relations suivantes :

$$E[N_s] = \lambda E[X].E[L_s] \approx \lambda m_x E[L_B] = m_x E[N_B] \quad (2.10)$$

Par la suite, nous évaluons le deuxième moment des indicateurs de performance. D'abord,

$$\sigma_s^2 := \text{Var}[N_s] = E[N_s^2] - E[N_s]^2 \quad (2.11)$$

Le deuxième moment du nombre de produits dans le système de production est :

$$E[N_s^2] = E[E[(\sum_{i=1}^{N_B} X_i)^2 | N_B]] = E[E[(\sum_{i=1}^{N_B} X_i^2 + \sum_{\substack{i,j \\ i \neq j}} X_i X_j) | N_B]]$$

où X_i est la taille de l'ordre d'approvisionnement i . En général, N_B dépend de la taille X_1 de l'ordre en cours de traitement par le serveur et $X_1 | N_B$ n'a pas la même distribution que X_i pour $i \neq 1$. En supposant que N_B et X_1 sont indépendantes, et que $X_1 | N_B$ et X_i sont des variables aléatoires iid, on a l'approximation suivante :

$$E[N_s^2] \approx E[N_B.E[X^2] + N_B(N_B - 1).E^2[X]] = E[N_B].E[X^2] + (E[N_B^2] - E[N_B]).E^2[X] \quad (2.12)$$

Ensuite, il nous reste à estimer le deuxième moment de la longueur N_B de la file d'attente $M/G/1$. Etant donné la distribution du temps de service de la file d'attente $M/G/1$, on peut utiliser la fonction génératrice de N_B pour évaluer le deuxième moment de N_B à l'aide du troisième moment du temps de service ([Cooper, *et al.*, 1981],[Buzacott, *et al.*, 1993]).

$$E[N_B^2] = \frac{\frac{1}{3} \lambda^3 E[T_B^3](1 - \rho) + \lambda^2 E[T_B^2] - \lambda^3 E[T_B]E[T_B^2] + \frac{1}{2} \lambda^4 E^2[T_B^2]}{E[N_{B,M/M/1}]} \quad (2.13)$$

où $E[T_B^3]$ représente le troisième moment du temps de service d'un lot.

Cette méthode implique l'utilisation du troisième moment du temps de service pour un lot de produits. Ainsi, nous proposons d'approximer la distribution du temps de service d'un lot par la distribution log-normale [Saporta, 1990] avec la moyenne m_{TB} et la variance σ_{TB}^2 .

Ensuite, nous pouvons évaluer le troisième moment du temps de service d'un lot $E[T_B^3]$ à l'aide de sa moyenne et de sa variance.

Selon la définition de la distribution log-normale, la variable aléatoire $\log(T_B)$ suit une distribution normale avec la moyenne M et l'écart type S :

$$m_{TB} = e^{M+S^2/2} \quad (2.14)$$

$$\sigma_{TB}^2 = e^{S^2+2M} (e^{S^2} - 1) \quad (2.15)$$

$$E[T_B^3] = \sqrt{e^{S^2}-1}(2+e^{S^2})\sigma_\tau^3 + 3E[\tau^2]E[\tau] - 2E^3[\tau] \quad (2.16)$$

On obtient $S^2 = \ln(\sigma_\tau^2 / m_\tau^2 + 1)$ et $M = \ln m_{TB} - S^2/2$ à partir des formules 2.14 et 2.15.

Ensuite on calcule $E[T_B^3]$ par la formule 2.16.

2.3.2 Le modèle de transport

L'objet de cette sous-section est de déterminer la distribution de probabilité du nombre N_t de produits transportés de l'unité de production vers l'entrepôt. La durée de transport L_t est constante et connue. Le processus d'arrivée pour le processus de transport est en même temps le processus de départ des ordres d'approvisionnement de l'unité de production et il n'est généralement pas un processus de Poisson. Cependant, pour simplifier le calcul nous approximations l'arrivée des ordres au système de transport par un processus de Poisson. Ainsi, le système de transport peut être modélisé par une file d'attente $M^X/D/\infty$ ([Liu, et al.,1990];[Masuyama, and Takine,2002]).

Sachant que le taux moyen d'arrivée des unités de produits au système de transport est $\lambda E[X]$ et que le temps de transport est L_t , en appliquant la loi de Little, on obtient :

$$m_t := E[N_t] = \lambda E[X].L_t = \lambda m_x L_t \quad (2.17)$$

Le deuxième moment de N_t est:

$$E[N_t^2] = E\left[\left(\sum_{i=1}^{N_{Bt}} X_i\right)^2\right]$$

où N_{Bt} est le nombre de lots en cours de transport. En utilisant l'approximation $M^X/D/\infty$ qui implique l'indépendance de N_{Bt} et de X_i , on obtient:

$$\begin{aligned} E[N_t^2] &= E\left[E\left[\left(\sum_{i=1}^{N_{Bt}} X_i\right)^2 \mid N_{Bt}\right]\right] \\ &= E\left[N_{Bt}E[X^2] + (N_{Bt}^2 - N_{Bt})E^2[X]\right] \\ &= E[N_{Bt}]E[X^2] + (E[N_{Bt}^2] - E[N_{Bt}])E^2[X] \\ &= E[N_{Bt}]Var(X) + E[N_{Bt}^2]E^2[X] \end{aligned}$$

où N_{Bt} est le nombre de lots arrivés dans un intervalle de temps de longueur L_t , où $E[N_{Bt}] = \lambda L_t$ et $E[N_{Bt}^2] = (\lambda L_t)^2 + \lambda L_t$. En introduisant ces termes dans l'équation ci-dessus, on obtient:

$$E[N_t^2] = \lambda L_t Var(X) + (\lambda L_t + (\lambda L_t)^2)E^2[X]$$

et

$$\sigma_t^2 := Var(N_t) = E[N_t^2] - E^2[N_t] = \lambda L_t (\sigma_X^2 + m_X^2) \quad (2.18)$$

2.3.3 Délai d'approvisionnement et stock en commande

Considérons le délai d'approvisionnement de l'entrepôt $L = L_s + L_t$ et le stock en commande $N = N_s + N_t$, i.e. quantité d'approvisionnement en cours. Les résultats suivants sont immédiats :

$$E[L] = E[L_s] + L_t \quad (2.19)$$

$$m_N := E[N] = E[N_s] + E[N_t] \quad (2.20)$$

En général, N_s et N_t sont dépendants. Une approximation du deuxième moment de N peut être obtenue en supposant l'indépendance des variables aléatoires N_s et N_t .

$$\sigma_N^2 := Var(N) = Var(N_s + N_t) \approx Var(N_s) + Var(N_t) = \sigma_s^2 + \sigma_t^2 \quad (2.21)$$

2.3.4 Estimation du coût global de gestion du stock

La relation suivante est une conséquence immédiate de la définition de la politique base-stock :

$$R = I - B + N \quad (2.22)$$

où R est le niveau de base-stock, I est le stock disponible, B la quantité de demandes en attente et N est le stock en commande. Sachant qu'un client peut être livré partiellement si le stock disponible est insuffisant, si $I = 0$ alors $B \neq 0$ et inversement. Par conséquent,

$$I = (R - N)^+ \quad (2.23)$$

$$B = (N - R)^+ \quad (2.24)$$

où $(x)^+ = \max\{x, 0\}$. L'état du stock est $IN = I - B = R - N$.

Le coût de gestion du stock comprend deux parties : le coût de stockage et le coût de rupture. Le coût de possession d'une unité de produit pour une unité de temps est C_h . Le coût de rupture par unité de temps engendré par chaque unité de produit que l'entrepôt ne peut pas livrer au client est C_b . Par conséquent, le coût de gestion du stock est :

$$C(R) = C_h E[I] + C_b E[B] \quad (2.25)$$

A partir des équations (2.23) – (2.25) on obtient,

$$C(R) = \int_0^\infty g(R - x) f_N(x) dx \quad (2.26)$$

où $f_N(x)$ est la fonction de densité de probabilité du stock en commande N et

$$g(x) = \begin{cases} C_h x, & \text{si } x \geq 0 \\ -C_b x, & \text{sinon.} \end{cases}$$

Le calcul exact de la fonction de densité de probabilité de N est très difficile. Par conséquent, nous approximons cette fonction par une distribution log-normale avec la moyenne $E[N]$ et la variance $Var(N)$.

$$f_N(x) \approx \frac{1}{Sx\sqrt{2\pi}} e^{-\frac{(\ln x - M)^2}{2S^2}} \quad (2.27)$$

où $S^2 = \ln\left(\frac{\sigma_N^2}{m_N^2} + 1\right)$ et $M = \ln m_N - S^2 / 2$.

Le choix de la distribution log-normale au lieu de la distribution normale est expliqué dans la section suivante et il est déterminé par la variation importante de N par rapport à sa moyenne. A partir des relations (2.26) et (2.27) et de la définition de $g(x)$,

$$C(R) = \int_{-\infty}^{\infty} g(R - e^{M+S.z}) \frac{1}{\sqrt{2\pi}} e^{-z^2/2} dz \quad (2.28)$$

$$C(R) = C_h \int_{-\infty}^{Z^*} (R - e^{M+S.z}) \frac{1}{\sqrt{2\pi}} e^{-z^2/2} dz \\ + C_b \int_{Z^*}^{\infty} (e^{M+S.z} - R) \frac{1}{\sqrt{2\pi}} e^{-z^2/2} dz \quad (2.29)$$

où z^* est la solution de l'équation $R - \exp(M + Sz) = 0$.

2.3.5 Estimation du taux de remplissage

Le taux de service des clients est un indicateur important pour l'entrepôt final. Il dépend fortement du niveau de base-stock. Dans nos travaux, nous définissons le taux de service client par le pourcentage de commandes honorées sans délai à partir du stock.

Soit X la quantité d'une commande reçue de la part d'un certain client. Selon la propriété PASTA (Poisson Arrival See Time Average) du processus de Poisson, l'état de stock IN vu par ce client suit la distribution stationnaire de IN . Par conséquent, le taux de service client est défini par la relation suivante:

$$P(X \leq IN) \quad (2.30)$$

où IN est l'état du stock. Vu que $IN = R - N$,

$$P(X \leq IN) = P(X \leq R - N) = \sum_{i=1}^{\infty} P(N \leq R-i) \cdot P_X(X=i) . \quad (2.31)$$

où P_X est la fonction de densité de probabilité de la quantité par commande.

A partir de la relation (2.27):

$$P(N \leq R-i) = \int_0^{R-i} f_N(N) dN = \int_0^{R-i} \frac{1}{Sx\sqrt{2\pi}} e^{-(\ln x - M)^2 / (2S^2)} dx = \int_0^{z'} \frac{1}{\sqrt{2\pi}} e^{-z^2/2} dz \quad (2.32)$$

où z' est la solution de l'équation suivante:

$$(R-i) - \exp(M + Sz) = 0.$$

2.4. Optimisation du coût global de gestion du stock

L'objectif de notre modèle d'optimisation est de minimiser le coût global de gestion du stock tout en respectant le taux de service client demandé, i.e. $P(X \leq IN) \geq \alpha$.

Examinons d'abord les propriétés de du taux de remplissage et la fonction objectif. Selon les relations (2.31)-(2.32), puisque le stock en commande N est indépendante de R , le taux de service client est croissant en R . Considérons la fonction objectif donnée par la relation (2.28) ou de manière équivalente (2.29). Le terme $R - e^{M+S.z}$ est linéaire et croissant en R . Soit $Y = R - e^{M+S.z}$. Selon la relation (2.28), $g(Y)$ est une fonction convexe et linéaire par morceaux en R . Ainsi, $g(R - e^{M+S.z})$ et $\int_{-\infty}^{\infty} g(R - e^{M+S.z}) \frac{1}{\sqrt{2\pi}} e^{-z^2/2} dz$ sont toutes les deux convexes en R .

Grâce à la monotonie du taux de service client et à la convexité de la fonction objectif, la solution du problème est la suivante $R^* = \text{MAX}\{R_1, R_2\}$ où R_1 est le niveau de base stock tel que le taux de service client est égale exactement à α , i.e. $P(X \leq IN) = \alpha$ et R_2 est le niveau de base stock minimisant la fonction objectif. Nous déterminons R_1 par une méthode de dichotomie. Pour R_2 , nous utilisons une méthode de recherche par dichotomie basée sur le gradient afin de résoudre ce problème d'optimisation avec :

$$\begin{aligned}
 & d \frac{C(R)}{dR} \\
 &= C_h (R - e^{M+S.Z^*}) \frac{1}{\sqrt{2\pi}} e^{-Z^{*2}/2} \frac{dz}{dR} + C_h \int_{-\infty}^{Z^*} \frac{1}{\sqrt{2\pi}} e^{-z^2/2} dz \\
 &\quad + C_b (e^{M+S.Z^*} - R) \frac{1}{\sqrt{2\pi}} e^{-Z^{*2}/2} \frac{dz}{dR} - C_b \int_{Z^*}^{\infty} \frac{1}{\sqrt{2\pi}} e^{-z^2/2} dz \\
 &= C_h \int_{-\infty}^{Z^*} f(z) dz - C_b \int_{Z^*}^{\infty} f(z) dz \\
 &= C_h \Phi(Z^*) - C_b (1 - \Phi(Z^*))
 \end{aligned} \tag{2.33}$$

où $f(z)$ est la fonction de densité et $\Phi(z)$ est la fonction de distribution de la loi normale standard.

2.5 Résultats numériques

L'objectif de cette section est de valider par la simulation les résultats analytiques présentés dans la section précédente. Les expériences numériques sont basées sur l'exemple suivant. Le taux moyen d'arrivée λ d'un client à l'entrepôt varie de 1 à 1.6. La taille d'une commande client est une variable aléatoire uniformément distribuée dans l'intervalle [3, 9]. Le temps de service d'une unité de produit est exponentiellement distribué avec une moyenne de 0.1. La durée de transport est $L_t = 3$. Le coût de stockage d'un produit pour une unité de temps, C_h ainsi que le coût de rupture par unité de temps engendré par chaque unité de produit que l'entrepôt ne peut pas livrer au client, C_b , sont tous les deux égaux à 1. Les paramètres numériques nécessaires à l'approche analytique sont rappelés dans le tableau 1.

Tableau 2.1. Paramètres d'entrée

	λ	Taille d'une commande client	Temps du service de l'usine	Temps de transport	C_h	C_b	R
Moyenne	1 – 1.6	6	0.1	3.0	1.0	1.0	50
Écart-type		2	0.1	0			

Les résultats analytiques sont comparés à ceux obtenus par la simulation. Afin d'obtenir des résultats de simulation avec une précision suffisante, une durée de simulation très importante est utilisée. Le temps de simulation est fixé à 10,000,000 unités de temps.

Dans un premier temps, nous évaluons la précision des résultats analytiques pour l'évaluation des performances. Pour cela, nous fixons le niveau de base-stock à $R=50$.

Les résultats de l'approche analytique ainsi que de la simulation sont représentés dans les tableaux 2.2 et 2.3.

Malgré le fait que l'approche analytique que nous avons proposé prend en compte seulement les deux premiers moments de chaque variable aléatoire, il fournit une estimation très précise du délai d'approvisionnement du stock ainsi que du stock en commande. De plus, malgré les nombreuses approximations faites, le deuxième moment du stock en commande, N , est calculé avec une erreur de seulement 3,5 %. De même,

l'utilisation de la distribution log-normale fournit une bonne approximation du coût global de la gestion du stock.

Une autre remarque intéressante est que la précision du résultat semble être peu sensible à la valeur de l'intensité du trafic ρ .

Nous avons également essayé d'approximer le stock en commande à l'aide de la distribution normale. Les résultats sont mauvais et ils se dégradent avec l'augmentation de l'intensité du trafic. Ce phénomène est principalement due à la variance importante de N par rapport à sa moyenne, ce qui conduit à une probabilité non négligeable du stock en commande N négatif.

Table 2.2. Résultats de la simulation de l'approche analytique pour le délai de production L_s et d'approvisionnement L

Taux d'arrivée	ρ	E[L_s] (sim)	E[L_s] (analyt)	Erreur (%)	E[L] (sim)	E[L] (analyt)	Erreur (%)
1.0	0.60	1.1730	1.1750	0.1706	4.1730	4.1750	0.0480
1.3	0.78	1.9542	1.9591	0.2506	4.9542	4.9591	0.0991
1.5	0.90	4.0304	4.0500	0.4842	7.0304	7.0500	0.2790
1.6	0.96	9.6031	9.8000	2.0092	12.603	12.800	1.5623

Table2. 3. Résultats de la simulation de l'approche analytique pour le stock en commande le coût global de gestion du stock

Taux d'arrivée	ρ	m_N (sim)	m_N (analyt)	Erreur (%)	σ_N (sim)	σ_N (analyt)	Error (%)	Coût (sim)	Coût (analyt)	Erreur (%)
1.0	0.60	25.4324	25.0500	1.5266	14.1190	13.9651	1.1023	25.5746	26.4905	3.5811
1.3	0.78	39.1543	38.6809	1.2237	20.7695	20.6483	0.5873	19.4611	19.7240	1.3510
1.5	0.90	63.8623	63.4500	0.6498	39.3581	39.5504	0.4867	28.6332	27.6893	3.2966
1.6	0.96	121.597	122.880	1.0442	92.2559	95.6122	3.5103	77.7748	77.5252	0.3208

La méthode analytique étant précise pour l'évaluation des performance, les résultats analytiques peuvent alors être utilisés pour l'optimisation. Pour cela, nous considérons le problème de minimisation du coût de gestion de stock avec contrainte d'un taux de service client minimal de 90%. Les résultats numériques sont illustrés dans le tableau 2.4. Ces résultats montrent que l'expression analytique du taux de remplissage est suffisamment

précise et la méthode d'optimisation est capable d'identifier les solutions optimales.

Table 2.4. Résultats de la simulation de l'approche analytique pour le taux de service client le coût optimal de gestion du stock obtenu pour un niveau de base stock R^*

<i>Arrival rate</i>	ρ	<i>Taux Service client (%)</i>	<i>optimal R*</i>	<i>Taux Serv client (sim)</i>	<i>Taux Serv client (analyt)</i>	<i>Erreur (%)</i>	<i>Coût (sim)</i>	<i>Coût (analyt)</i>	<i>Erreur (%)</i>
1.0	0.60	90	50	0.9014	0.9092	0.8572	25.5746	26.4905	3.5811
1.3	0.78	90	71	0.8954	0.9034	0.5448	34.2482	35.1262	2.5635
1.5	0.90	90	119	0.8962	0.9020	0.7371	61.6398	62.2328	0.9621
1.6	0.96	90	241	0.8902	0.9008	1.1813	137.756	137.901	0.1056

2.6 Conclusions

Dans ce chapitre nous avons proposé une méthode analytique pour estimer les performances d'un système de production-distribution à deux niveaux caractérisé par : i) une demande aléatoire, ii) une capacité de fabrication finie, iii) la taille aléatoire des ordres de re-approvisionnement, iv) le délai de transport constant entre l'unité de production et l'entrepôt. Le stock est contrôlé par une politique base-stock. Les principales performances qui ont retenu notre attention sont le délai de re-approvisionnement et le coût total de gestion du stock. De même, une technique de gradient a été utilisée pour minimiser le coût total de gestion du stock. Cette approche a été validée par des expériences numériques.

Cependant, lorsque la taille du système augmente, une approche analytique adaptée est difficile à mettre en place. Par conséquent, dans le chapitre suivant nous proposons une technique d'optimisation par simulation des systèmes de production-distribution à plusieurs niveaux.

Chapitre 3

Optimisation à l'aide de la simulation pour les systèmes dynamiques à événements discrets

Dans ce chapitre, nous proposons une approche d'optimisation par simulation basée sur une technique de recherche dans la zone la plus prometteuse, appelé COMPASS. Cette approche s'adresse aux problèmes d'optimisation caractérisés par des variables de décision discrètes et des mesures de performance estimées par la simulation. Une attention particulière est accordée à l'optimisation du budget de simulation. Nous montrons que les algorithmes proposées dans le cadre de cette approche convergent avec la probabilité 1 vers un optimum local. Cette propriété est respectée également lorsque les contraintes imposées sont estimées par la simulation.*

Publications : [LI, SAVA et XIE, 2006b,c,d,e]

3.1 Introduction

3.1.1 Préliminaires

Les algorithmes d'optimisation pour la résolution de différents types de problèmes existent depuis plusieurs décades. Cependant, les problèmes à résoudre sont devenus de plus en plus complexes due à la complexité croissante des systèmes et à la nécessité d'optimiser des systèmes réels dans des domaines différents tels que l'ingénierie, l'économie et les sciences. Une variété importante de paramètres influent sur les performances de ces systèmes. Cette propriété rend difficile la compréhension même du problème, sans parler de l'optimisation. Considérons, par exemple, une manufacture de semi-conducteurs. Pour ce type d'application un objectif possible est de maximiser le bénéfice (nombre d'unités vendues) tout en minimisant le temps de cycle moyen (durée moyenne nécessaire pour la fabrication d'une unité). Ainsi, on peut viser à trouver la meilleure configuration du processus de fabrication (nombre et type de machines, séquences de fabrication, d'autres politiques) qui va permettre d'atteindre la meilleure combinaison de bénéfice et temps de cycle.

Cependant, vu la complexité du système, il n'existe pas de modèle analytique approprié. De plus, les systèmes réels sont stochastiques par leur nature car il existe des paramètres inconnus qui affectent les performances du système. Ces paramètres inconnus sont d'habitude définis par des perturbations aléatoires ou le bruit. Les pannes des machines représentent un exemple de perturbation aléatoire d'un système de production. Ceci complique encore plus le problème d'optimisation.

Les algorithmes classiques d'optimisation sont principalement conçus pour résoudre des problèmes déterministes. Ils constituent le point de départ des travaux de recherche présentés dans ce mémoire et qui s'appliquent aux systèmes stochastiques complexes.

Dans le cas où un modèle analytique bien défini n'est pas disponible, la simulation par ordinateur est souvent utilisée afin d'approximer le comportement du système tel que les performances du système puissent être estimées à travers des résultats de simulation.

Par conséquent, on a besoin de déterminer un paramétrage pour le modèle de simulation qui conduit à une performance optimale du système. Dans la littérature, le concept qui consiste à combiner les algorithmes d'optimisation et la simulation est dénommé optimisation par simulation ou optimisation via la simulation. Ce concept sera présenté dans la Section 3.1.2. La Section 3.1.3. fournit la taxonomie des méthodologies d'optimisation par la simulation ainsi que quelques détails sur des approches connues d'optimisation par simulation.

3.1.2 L'optimisation par simulation

Dans la pratique, l'optimisation d'un système réel correspond généralement à l'optimisation d'une performance mesurée d'un système stochastique complexe. Si le système est suffisamment complexe, il est fort probable que le seul outil universel disponible pour estimer ses performances soit la simulation par ordinateur. Par conséquent, il y a un besoin croissant de développer des méthodes qui permettent de déterminer le paramétrage optimal de la simulation qui conduit à la meilleure performance du système. De plus, lorsque le nombre de jeux de paramètres est important, l'évaluation de chacun est pratiquement impossible à cause de la durée et des ressources nécessaires. Ainsi, les travaux de recherche dans ce domaine ont essayé de développer des méthodes pour déterminer le meilleur jeu de paramètres sans considérer toutes les possibilités. Le processus de recherche du meilleur jeu de paramètres de simulation d'un système stochastique complexe, dont les performances sont évaluées à partir des résultats d'un modèle de simulation est dénommé optimisation par simulation ([Andradottir,1998], [Olafsson et J. Kim, 2002]). L'objectif de l'optimisation par simulation est d'optimiser les performances choisies du système tout en minimisant le nombre de jeux de paramètres évalués.

La simulation par ordinateur peut fournir une bonne représentation d'un système réel qui permet d'effectuer des expérimentations afin de déterminer le jeu de paramètres qui génère une performance optimale du système. Généralement, afin d'éviter les expériences sur le système réel, on cherche à construire un modèle analytique.

Cependant, lorsque le système est suffisamment complexe, on a recours à la simulation par ordinateur. Une simulation par ordinateur mise en place pour mimer le comportement d'un système réel est dénommée modèle de simulation. Il représente le système réel dont on simule l'évolution dans le temps et il est généralement défini par une collection de variables d'état. Cette collection de variables dépend des objectifs du développement du modèle de simulation. Par exemple, lorsqu'on étudie une banque, les variables d'état peuvent être le nombre de conseillers disponibles/occupés, le nombre de clients, la date d'arrivée/départ de chaque client. Lorsque la valeur des variables d'état change instantanément à des instants de temps séparés, le modèle de simulation est à événements discrets. Le modèle de simulation d'une banque peut être un exemple d'un modèle de simulation à événements discrets car une variable d'état, i.e. le nombre de clients dans la banque change seulement lorsqu'un client arrive ou lorsqu'il part. Un modèle de simulation à temps continu est celui où la valeur des variables d'état évolue d'une manière continue avec le temps. Un exemple de modèle de simulation à temps continu est le déplacement d'une voiture sur l'autoroute car sa vitesse est une variable d'état qui évolue de façon continue avec le temps. L'optimisation par simulation est généralement appliquée à des modèles de simulation à temps discret. Ainsi, dans ce mémoire le concept de modèle de simulation désigne un modèle de simulation à événements discrets.

Dans chaque modèle de simulation il y a des paramètres qui influent son comportement et ses performances. L'ensemble des paramètres qui ont cette propriété désigne un paramétrage de la simulation. Certains de ces paramètres peuvent être contrôlés ou fixés dans un système réel, comme par exemple le paramètre qui détermine le type ou le nombre des machines utilisées dans un atelier. D'autres paramètres ne peuvent pas être contrôlés, comme par exemple l'arrivée d'une panne dans la journée. Même si on souhaite contrôler la date d'arrivée d'une panne, on ne peut pas le faire directement.

Dans le contexte de l'optimisation par simulation on considère strictement l'ensemble des paramètres contrôlables, dénoté par θ . Par la suite de ce mémoire, nous considérons qu'un paramétrage d'une simulation est une collection de paramètres contrôlables.

Un problème d'optimisation par simulation peut être défini par le problème d'optimisation suivant :

$$\min_{\theta \in \Theta} f(\theta) = \max_{\theta \in \Theta} (-f(\theta)) \quad (3.1)$$

où $f(\theta)$ définit une mesure de performance choisie pour le système étudié. Nous rappelons que le modèle de simulation considéré est stochastique. Par conséquent, la mesure estimée de la performance obtenue est $f(\theta) = E_{\omega} [g(\theta; \omega)]$ sur plusieurs réalisations du phénomène ω . La performance déterminée en exécutant une fois le modèle de simulation pour un paramétrage donné (appelée observation ou réplication de simulation) est notée $g(\theta; \omega)$. Le paramétrage θ est le vecteur de N paramètres contrôlables du modèle de simulation, alors que ω représente les perturbations aléatoires, i.e. l'effet des paramètres incontrôlables. Les perturbations stochastiques du système représentent l'influence des paramètres inconnus et des interactions entre les paramètres connus et inconnus qui affectent les performances du système. Enfin, Θ dénote l'ensemble de contraintes, le domaine admissible, pour le paramétrage θ . Par la suite, nous considérons le problème d'optimisation défini par l'équation 3.1 avec l'optimum $\theta^* = \arg \min_{\theta \in \Theta} f(\theta)$.

Après avoir défini le problème d'optimisation par l'équation 3.1, nous expliquons l'utilisation de la simulation pour le résoudre. Nous supposons que la fonction objectif $f(\theta)$, peut être estimée par un modèle de simulation. L'objectif est de déterminer un paramétrage θ qui génère la meilleure ou du moins une très bonne performance tout en limitant le nombre de paramétrages à tester.

Le principe de l'optimisation par simulation est illustré graphiquement dans la figure 3.1. La performance du système est affectée par les paramètres contrôlables et incontrôlables. D'un côté, les paramètres contrôlables définissent le paramétrage θ , qu'on peut ajuster afin d'améliorer la performance du système. De l'autre côté, les paramètres incontrôlables, ω , affectent la performance du système, mais on ne peut pas les modifier directement. Si le modèle de simulation est exécuté plusieurs fois avec le même paramétrage θ , alors des valeurs différentes de la performance mesurée $g(\theta; \omega)$ peuvent être obtenues, selon la distribution de ω . Ainsi, pour un paramétrage donné,

l'estimation de la performance mesurée $f(\theta) = E_{\omega}[g(\theta; \omega)]$ est calculée à l'aide de plusieurs simulations. Dans la pratique, une estimation de $f(\theta)$ est généralement calculée par la moyenne de N simulations indépendantes :

$$\hat{f}(\theta_k) = \frac{1}{N} \sum_{n=1}^N g_n(\theta, \omega) \quad (3.2)$$

où $\hat{f}(\theta)$ est une estimation de la performance mesurée $f(\theta)$, N est le nombre de répliques de simulation réalisées pour le paramétrage θ , et $g_n(\theta; \omega)$ est le résultat obtenu par la $n^{\text{ème}}$ réplique de simulation. L'estimation $\hat{f}(\theta)$ est asymptotique non biaisée et converge vers $f(\theta)$ selon la loi des grands nombres. L'information disponible, notamment θ et $\hat{f}(\theta)$, est ensuite utilisée par la méthode de recherche implémentée dans le module d'optimisation pour générer un meilleur paramétrage θ , qui sera la nouvelle entrée du modèle de simulation. Le choix du nombre de simulations N à effectuer pour un même paramétrage fait partie de la méthodologie d'optimisation par simulation. Nous cherchons à améliorer itérativement la performance calculée par le modèle de simulation jusqu'à ce qu'un critère d'arrêt ou d'optimalité soit satisfait.

Dans ce chapitre nous présentons une méthode de recherche à implémenter dans le module d'optimisation.

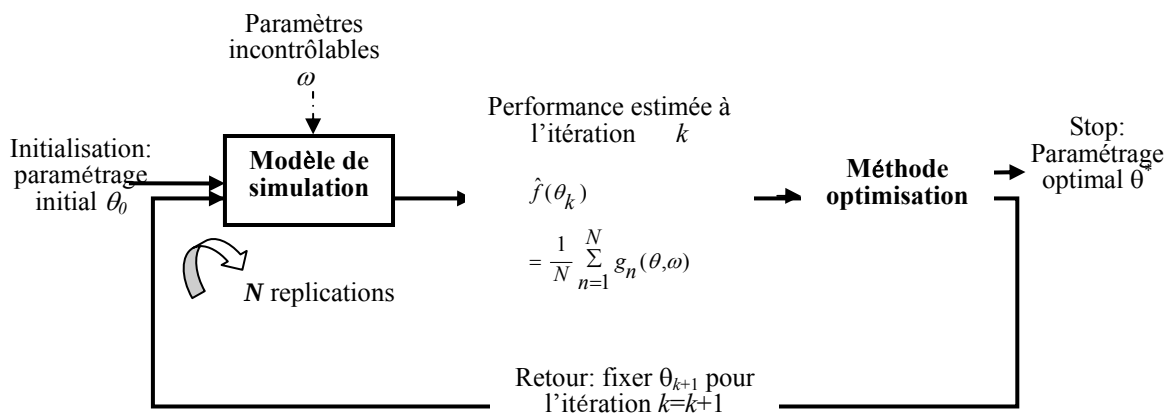


Fig 3.1 Principe de l'optimisation par simulation

Les avantages de l'utilisation d'un modèle de simulation ont été en partie illustrés précédemment. Par la suite, nous donnons quelques arguments qui montrent l'intérêt de

la simulation par ordinateurs. Un modèle de simulation est souvent la seule technique disponible pour évaluer les performances d'un système stochastique complexe parce que la plupart des systèmes réels ne peuvent pas être décrits avec une précision suffisante par un modèle mathématique qui puisse être évalué d'une manière analytique. La simulation par ordinateur permet l'estimation des performances d'un système existant pour des paramétrages différents. Cette technique permet également d'évaluer des systèmes « futures », afin d'affiner leur conception. Un modèle de simulation permet un meilleur contrôle des conditions expérimentales que dans le cas où l'expérimentation s'effectuerait directement sur le système réel. De plus, elle permet d'étudier le comportement du système sur une longue période dans un intervalle de temps comprimé. À l'inverse, elle permet d'analyser en détail l'évolution du système dans un intervalle de temps dilaté.

Cependant, l'optimisation par simulation a également des limites, surtout si la complexité du système à optimiser est importante. Nous rappelons qu'une hypothèse importante est que la performance choisie du système, $f(\theta)$, puisse être estimée à l'aide du modèle de simulation. Ceci signifie que le modèle de simulation est capable d'imiter le comportement dynamique du système réel avec une précision suffisante telle que ce modèle fournit une bonne représentation du système réel. Donc, le modèle de simulation doit être vérifié et validé ce qui n'est généralement pas une tâche triviale. De même, un autre problème qui apparaît lors du développement d'un modèle de simulation est la prise en compte des paramètres incontrôlables. Déterminer la distribution de probabilité de ω , n'est pas une tâche facile. La durée de la simulation doit être bien choisie et cette décision n'est pas triviale non plus. De plus, il est possible d'obtenir seulement une estimation statistique des performances du système. Cependant, malgré ces limitations, l'optimisation par simulation fournit une alternative efficace pour l'optimisation des systèmes stochastiques complexes.

3.1.3 Méthodologies d'optimisation par simulation

Généralement, les méthodes d'optimisation par simulation sont classifiées selon la

nature de l'espace de solution pour le paramétrage du modèle de simulation, i.e. espace continu ou discret. Même si ces approches sont classifiées différemment dans la littérature, une classification généralement acceptée est proposée dans la Figure 3.2. On peut distinguer cinq catégories de méthodologies de recherche: méthodes de gradients, méthodes déterministe, la surface de réponse, la recherche statistique, la recherche stochastique. L'application de ces méthodologies de recherche varie en fonction de la nature de l'espace de solution Θ .

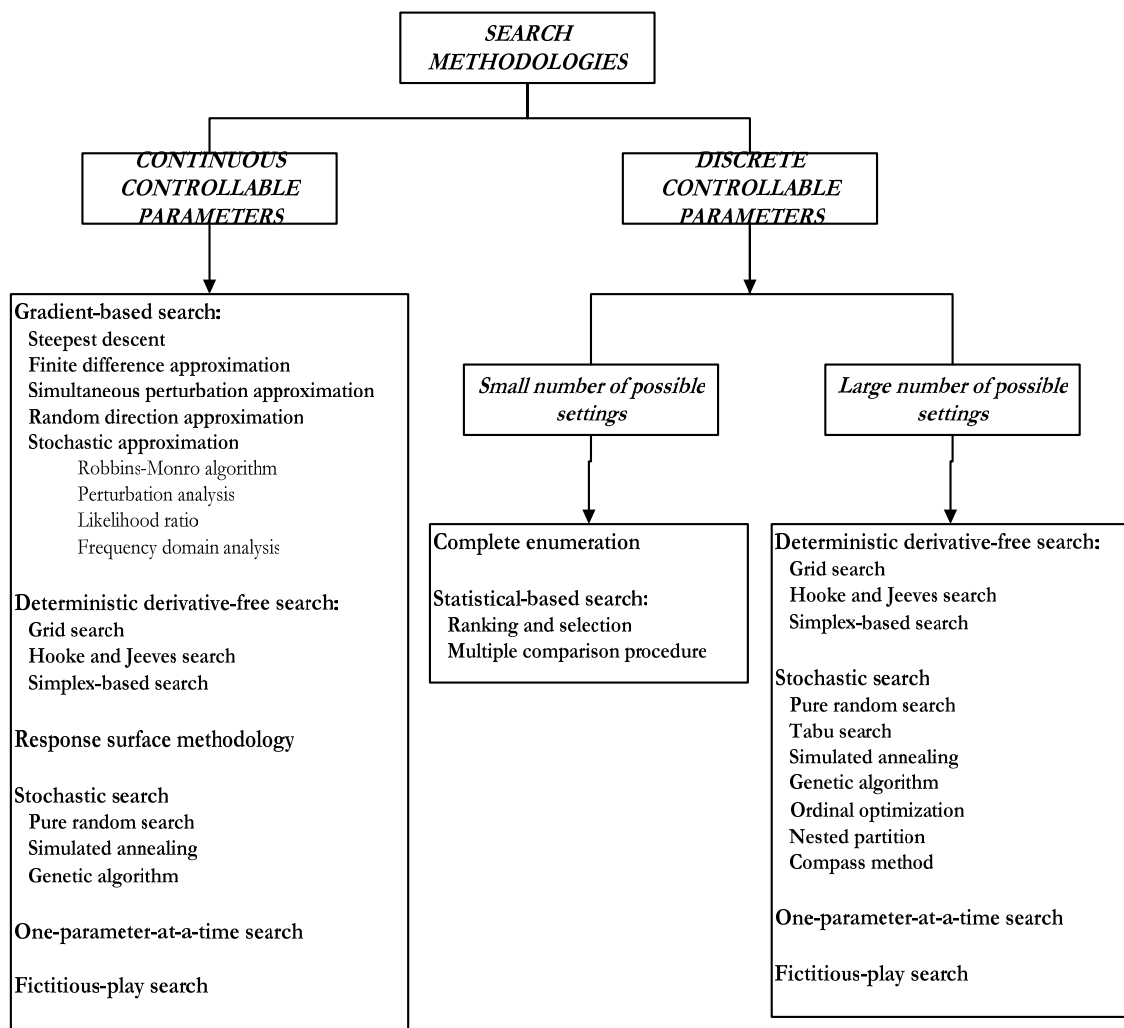


Fig 3.2 Classification des méthodologies d'optimisation par simulation (DOvS)

3.2 Optimisation discrète par simulation avec contraintes

analytiques

Dans cette section, nous considérons le problème suivant d'optimisation discrète par simulation (ou DOvS pour Discret Optimisation via Simulation):

$$\min_{x \in \Theta} E[G(x, w_i)]$$

où $\Theta = \Phi \cap Z^d$ est l'espace des solutions, Φ est un ensemble compact sur $\mathcal{R}^d$, et Z^d est l'ensemble de vecteurs de dimension d , dont les éléments sont des nombres entiers. Nous supposons que $G(x, w_i)$ est la $i^{\text{ème}}$ mesure de performance obtenue pour la solution x et que w_i sont des variables aléatoires i.i.d pour différentes valeurs de i . Soit $g(x) = E[G(x, w_i)]$. Nous supposons que $g(x)$ peut être évaluée à travers des réplifications de simulation de la solution x . Nous supposons également que l'expression analytique de $g(x)$ n'est pas disponible mais que l'ensemble Φ est défini analytiquement.

3.2.1 Optimisation discrète par simulation- cas d'un espace de solution borné

Dans cette section, nous présentons un algorithme itératif appelé COMPASS* qui permet de résoudre le problème d'optimisation discrète par simulation dans le cas où l'ensemble de contraintes Φ est fermé et borné. Dans ce cas, l'espace de solutions $\Theta = \Phi \cap Z^d$ est fini.

L'algorithme COMPASS* n'évalue pas toutes les solutions de l'espace de solutions. Il génère progressivement les solutions qui seront évaluées à chaque itération. Soit V_k l'ensemble de solutions visitées jusqu'à l'itération k . De même, le budget de simulation, i.e. le nombre de réplifications ou d'observations, est alloué dynamiquement à chaque solution. Soit $N_k(x)$ le nombre total d'observations alloués à la solution x jusqu'à l'itération k et $\bar{G}_k(x)$ la moyenne des résultats $G(x, w_i)$ obtenus pour chacune des $N_k(x)$

observations à l'itération k :

$$\bar{G}_k(x) = \frac{1}{N_k(x)} \sum_{i=1}^{N_k(x)} G_i(x, w_i)$$

A chaque itération k , l'algorithme COMPASS* calcule la solution optimale observée x_k^* , qui fournit la plus petite valeur de la fonction objectif parmi toutes les solutions $x \in V_k$. Ensuite il génère aléatoirement d'autres solutions dans une région appelée la région la plus prometteuse C_k définie comme suit (le exemple de 2 dimension est montré par figure 3.3) :

$$C_k = \{x \in \Theta : \|x - x_k^*\| \leq \|x - y\|, \forall y \in V_k \text{ and } y \neq x_k^*\}$$

où $\|x - y\|$ est la distance Euclidienne entre les solutions x et y . L'ensemble C_k contient toutes les solutions admissibles qui sont au moins aussi proches de x_k^* que des autres solutions dans l'ensemble V_k .

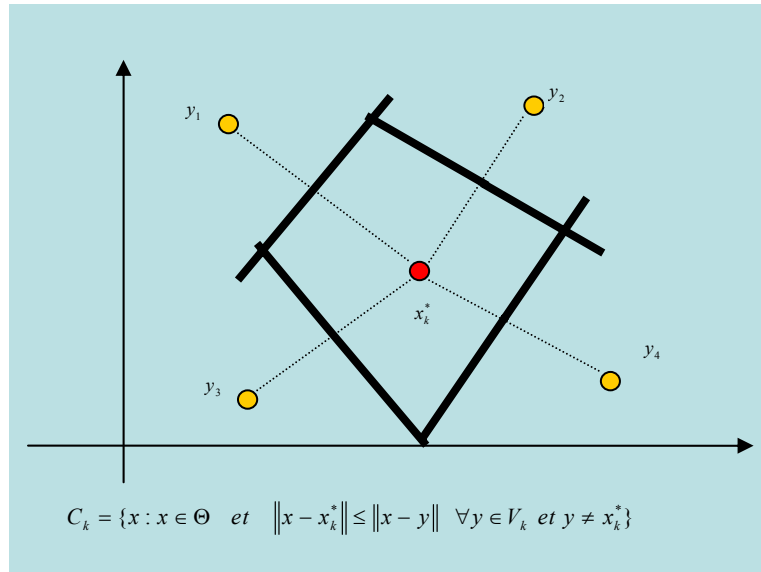


Figure 3.3 la définition de la région la plus prometteuse

A chaque itération k , une règle d'allocation du budget de simulation (SAR) est utilisée afin de calculer le nombre supplémentaire $a_k(\mathbf{x})$ d'observations à allouer à la solution x durant l'itération k . Ainsi, $N_k(x) = \sum_{i=0}^k a_i(x)$ définit le nombre total de fois que la solution x est simulée entre le début de l'algorithme et la fin de l'itération k .

Algorithme 1 (COMPASS* pour DOvS avec espace de solution borné)**Etape 1:** Initialisation

- 1.1 Initialiser le compteur d'itérations $k = 0$. Choisir $x_0 \in \Theta$, faire $V_0 = \{x_0\}$ et $x_0^* = x_0$.
- 1.2 Déterminer $a_0(x_0)$ selon la règle d'allocation du budget de simulation (SAR).
- 1.3 Effectuer $a_0(x_0)$ réplifications de simulation pour la solution x_0 , faire $N_0(x_0) = a_0(x_0)$,
et calculer $\overline{G}_0(x_0)$.
- 1.4 Soit $C_0 = \Theta$.

Etape 2: Soit $k = k + 1$. Choisir aléatoirement des nouvelles solutions et allouer le budget de simulation.

- 2.1. Choisir m solutions $x_{k1}, x_{k2}, \dots, x_{km}$ de manière uniforme et indépendante dans C_{k-1} . Soit $V_k = V_{k-1} \cup \{x_{k1}, x_{k1}, \dots, x_{km}\}$.
- 2.2. Déterminer $a_k(x)$ selon la SAR pour chaque $x \in V_k$.
- 2.3. Pour chaque $x \in V_k$, faire $a_k(x)$ observations, ensuite mettre à jour $N_k(x)$ et $\overline{G}_k(x)$.

Etape 3: Déterminer le nouvel optimum observé et mettre à jour la région la plus prometteuse.

- 3.1. Déterminer le nouvel optimum $x_k^* = \arg \min_{x \in V_k} \overline{G}_k(x)$.
- 3.2. Construire la région la plus prometteuse:

$$C_k = \{x \in \Theta : \|x - x_k^*\| \leq \|x - y\|, \forall y \in V_k \text{ and } y \neq x_k^*\}$$

- 3.3. Aller à l'étape 2.

■

Remarque: A chaque itération k , l'algorithme Compass* choisit d'une manière uniforme des solutions à partir de la région C_k . Ce choix des solutions dans un ensemble de vecteurs d'entiers de dimension d est généralement difficile. L'algorithme que nous utilisons est basé sur la méthode RMD proposée dans [Hong et Nelson ,2006] pour le choix uniforme des nombres entiers à l'intérieur d'un polyèdre.

Par la suite, nous présentons les notations et les hypothèses nécessaires pour la preuve de convergence de l'algorithme COMPASS*.

Hypothèse 1 : Pour chaque $x \in \Theta$, $P[\lim_{N_k(x) \rightarrow \infty} \frac{1}{N_k(x)} \sum_{i=1}^{N_k(x)} G_i(x, w_i) = g(x)] = 1$

■

L'hypothèse 1 implique que la moyenne de $G(x, w_i)$ est une estimation appropriée pour $g(x)$. Si $G(x, w_i)$, $i = 1, 2, \dots$ sont indépendantes et identiquement distribuées, alors l'hypothèse 1 représente la loi des nombres grands. Si $G(x, w_i)$, $i = 1, 2, \dots$ sont ergodiques, alors l'hypothèse 1 représente le théorème d'ergodicité. La plupart de résultats de simulation satisfont cette hypothèse.

Soit $NH(x) = \{y : y \in \Theta \text{ et } \|x - y\| = 1\}$ le voisinage de la solution $x \in \Theta$.

Hypothèse 2 : La SAR est telle que $a_k(x) \geq 1$ si x est une solution visitée pour la première fois à l'itération k , i.e. ($x \in V_k \setminus V_{k-1}$) ou x est dans le voisinage de l'optimum observé à itération $k-1$, i.e. ($x \in NH(x_{k-1}^*)$).

■

Selon cette hypothèse, la SAR garantit que chaque nouvelle solution est simulée au moins une fois et qu'un budget de simulation supplémentaire est alloué aux solutions déjà visitées et qui se trouvent dans le voisinage de l'optimum observé à l'itération précédente. Par conséquent, seulement les solutions qui se trouvent dans le voisinage d'un optimum observé sont simulées un nombre illimité de fois.

Cette hypothèse est bien moins restrictive que celle demandée par l'algorithme COMPASS initial [Hong et Nelson , 2006], qui exige que chacune des solutions visitées soit simulée un nombre illimité de fois. Par rapport à COMPASS, COMPASS* permet de concentrer l'effort de simulation sur des solutions réellement intéressantes alors que COMPASS le répartit entre toutes les solutions. Par conséquent, l'algorithme

COMPASS* est plus efficace et présente de meilleure vitesse de convergence que l'algorithme COMPASS.

La SAR la plus simple qui satisfait l'hypothèse 2 est :

$$a_k(x) = \begin{cases} 1 & \text{if } x \in V_k \setminus V_{k-1} \text{ or } x \in NH(x_{k-1}^*) \\ 0 & \text{otherwise} \end{cases}$$

Plusieurs méthodes d'optimisation des DovS comme par exemple l'optimisation ordinaire [Ho, 1994] et OCBA – Optimal Computation Budget Allocation [Chen et al. 2000] peuvent être utilisées dans l'étape 2 afin de définir des SAR efficaces pour accélérer la convergence de l'algorithme COMPASS*.

Le théorème suivant prouve la convergence de l'algorithme 1 vers un optimum local, défini de la manière suivante. Une solution $x \in \Theta$ est un optimum local si $g(x) \leq g(y)$, $\forall y \in NH(x) \cap \Theta$. Soit M l'ensemble des optimums locaux. Si M est singleton ou si $g(x)$ est uni-modale, alors chaque élément de M est une solution optimale globale.

Théorème 1. Sous les hypothèses 1 et 2, la séquence des optimums observés $\{x_0^*, x_1^*, \dots\}$ générée par l'algorithme 1 converge avec la probabilité 1 vers l'ensemble M dans le sens que $P\{x_k^* \notin M \text{ i.o.}\} = 0$ où *i.o.* signifie infiniment souvent.

■

Notons que non seulement nous établissons la convergence sous des hypothèses bien moins restrictives que celles dans (Hong and Nelson 2006) pour la convergence de COMPASS, notre preuve de convergence est bien plus simple.

Preuve du théorème 1. Quelque soit la séquence infinie $\{V_0, V_1, \dots\}$ générée par l'algorithme 1, vu que $V_k \subseteq V_{k+1} \subseteq \Theta$, $\forall k \geq 0$, $V_\infty = \bigcup_{k=0}^{\infty} V_k$ existe et $V_\infty \subseteq \Theta$. Vu que Θ est un ensemble fini, V_∞ est également un ensemble fini.

Soit U_∞ l'ensemble de solutions qui bénéficient d'un nombre illimité de simulations,

i.e. $U_\infty = \{x : x \in V_\infty \text{ and } \lim_{k \rightarrow \infty} N_k(x) = +\infty\}$. Selon l'hypothèse 2 et vu que Θ est un ensemble

fini, U_∞ est un ensemble fini et non vide. De plus $U_\infty \subseteq V_\infty \subseteq \Theta$.

Par conséquent,

$$P\{x_k^* \notin M \text{ i.o.}\} = \sum_{A \subset \Theta} P\{x_k^* \notin M \text{ i.o.} | U_\infty = A\} P\{U_\infty = A\} \quad (3.3)$$

où A est un sous-ensemble fini de Θ .

Prouver que $P\{x_k^* \notin M \text{ i.o.}\} = 0$ est équivalent à prouver que

$P\{x_k^* \notin M \text{ i.o.} | U_\infty = A\} = 0$ pour tout sous-ensemble $A \subset \Theta$ non vide et fini tel que

$P\{U_\infty = A\} > 0$.

Nous montrons maintenant $P\{x_k^* \notin M \text{ i.o.} | U_\infty = A\} = 0$, $\forall A \subset \Theta$ tel que

$P\{U_\infty = A\} > 0$. Si $x_k^* \notin M \text{ i.o.}$, alors il existe une solution $x \in A$ et $x \notin M$ telle que $x_k^* = x$

i.o. Comme $x \notin M$, il existe $y \in NH(x) \cap \Theta$ tel que $g(y) < g(x)$. Vu que $\|y - x\| = 1 \leq \|y - z\|$

$\forall z \neq y$, $y \in C_{k+1}$ si $x_k^* = x$ et $y \notin V_k$. Donc,

$$P\{y \in V_{k+1} | x_k^* = x \text{ and } y \notin V_k\} \geq \frac{1}{|\Theta|} > 0$$

Si $x_k^* = x \text{ i.o.}$, avec la probabilité 1, il existe $K^* > 0$ tel que $y \in V_k$, $\forall k \geq K^*$. Selon

l'hypothèse 2, $x_k^* = x \text{ i.o.}$ implique $N_k(y) \rightarrow \infty$ lorsque $k \rightarrow \infty$. Donc

$$P\{y \in A | U_\infty = A \text{ and } x_k^* = x \text{ i.o.}\} = 1$$

Comme $g(y) < g(x)$, $\lim_{k \rightarrow \infty} N_k(x) = +\infty$ et $\lim_{k \rightarrow \infty} N_k(y) = +\infty$. Selon l'hypothèse 1, il existe K_2

> 0 tel que $\bar{G}_k(x) > \bar{G}_k(y) \forall k \geq K_2$. Par conséquent, avec la probabilité 1, x_k^* ne peut

être x qu'un nombre fini de fois. Ceci est une contradiction. Donc,

$P\{x_k^* \notin M \text{ i.o.} | U_\infty = A\} = 0$ pour tout sous-ensemble $A \subset \Theta$ non vide et fini tel que $P\{U_\infty = A\} > 0$. Ce résultat appliqué à l'équation (3.3), conclue la théorème.

■

Le prochain théorème montre que l'optimum observé à chaque itération converge vers la solution optimale parmi les solutions qui ont été simulées un nombre illimité de fois. Notons que ce résultat a été utilisé dans la preuve de convergence de (Hong and Nelson 2006) et les auteurs de ce dernier n'étaient pas conscients de l'importance de ce résultat.

Théorème 2. Sous les hypothèses 1 et 2, avec probabilité 1,

$$\lim_{k \rightarrow \infty} g(x_k^*) = \min_{x \in U_\infty} g(x)$$

où U_∞ est l'ensemble des solutions simulées un nombre illimité de fois.

■

Lemme 1 (Hong et Nelson 2006): Soit $a_1, a_2, \dots, a_n$ et $b_1, b_2, \dots, b_n$ des nombres naturels. Alors

$$\left| \min_{i=1, \dots, n} a_i - \min_{i=1, \dots, n} b_i \right| \leq \max_{i=1, \dots, n} |a_i - b_i|$$

■

Preuve du théorème 2.

A partir de la preuve de la théorème 1, $V_k \subseteq V_{k+1} \subseteq \Theta$, $\forall k \geq 0$, il existe un ensemble fini $V_\infty = \bigcup_{k=0}^{\infty} V_k$ tel que $V_\infty \subseteq \Theta$. De plus $U_\infty \subseteq V_\infty \subseteq \Theta$ est un ensemble fini et non vide.

Par conséquent,

$$P\{\lim_{k \rightarrow \infty} g(x_k^*) = \min_{x \in U_\infty} g(x)\} = \sum_{A \subset \Theta} P\{\lim_{k \rightarrow \infty} g(x_k^*) = \min_{x \in A_\infty} g(x) | U_\infty = A\} P\{U_\infty = A\} \quad (3.4)$$

où $A \subset \Theta$, est un sous-ensemble fini de Θ . Le théorème est prouvé si

$$P\{\lim_{k \rightarrow \infty} g(x_k^*) = \min_{x \in U_\infty} g(x) \mid U_\infty = A\} = 1 \quad (3.5)$$

ce qui est équivalent à :

$$P\{|g(x_k^*) - \min_{x \in U_\infty} g(x)| \geq \varepsilon \mid U_\infty = A\} = 0, \forall \varepsilon > 0 \quad (3.6)$$

On peut remarquer que :

$$\begin{aligned} & P\{|g(x_k^*) - \min_{x \in U_\infty} g(x)| \geq \varepsilon \mid U_\infty = A\} \\ & \leq P\{|g(x_k^*) - \bar{G}_k(x_k^*)| \geq \varepsilon/2 \mid U_\infty = A\} \\ & \quad + P\{|\bar{G}_k(x_k^*) - \min_{x \in U_\infty} g(x)| \geq \varepsilon/2 \mid U_\infty = A\} \end{aligned} \quad (3.7)$$

Considérons la séquence croissante $\{V_0, V_1, \dots\}$ générée par l'algorithme 1. comme

$|V_\infty| < \infty, \exists K_0 > 0$ tel que $V_k = V_\infty$ et $x_k^* \in U_\infty \forall k \geq K_0$. Ainsi,

$$\begin{aligned} (3.6) & \leq P\{|g(x) - \bar{G}_k(x)| \geq \varepsilon/2 \mid U_\infty = A\} \\ & \quad + P\{|\min_{x \in U_\infty} \bar{G}_k(x) - \min_{x \in U_\infty} g(x)| \geq \varepsilon/2 \mid U_\infty = A\} \\ & \leq P\{|g(x) - \bar{G}_k(x)| \geq \varepsilon/2 \mid U_\infty = A\} \\ & \quad + P\{\max_{x \in U_\infty} |\bar{G}_k(x) - g(x)| \geq \varepsilon/2 \mid U_\infty = A\} \quad \text{par lemma 1} \\ & \leq 2P\{|g(x) - \bar{G}_k(x)| \geq \varepsilon/2 \mid U_\infty = A\} \\ & \leq 2 \sum_{x \in A} P\{|g(x) - \bar{G}_k(x)| \geq \varepsilon/2 \mid U_\infty = A\} \quad \text{Par l'inégalité Bonferroni} \end{aligned}$$

ce qui mène à:

$$P\{|g(x_k^*) - \min_{x \in U_\infty} g(x)| \geq \varepsilon \mid U_\infty = A\} \leq 2 \sum_{x \in A} P\{|g(x) - \bar{G}_k(x)| \geq \varepsilon/2 \mid x \in U_\infty\} \quad (3.8)$$

La dernière inégalité est vraie car les simulations effectuées pour la solution x lorsque $k \geq K_0$ dépendent seulement de $x \in U_\infty$ et non pas de $U_\infty = A$. Selon les hypothèses 2 et 1,

$\lim_{k \rightarrow \infty} \bar{G}_k = g(x)$ avec la probabilité 1 si $N_k(x) \rightarrow \infty$. Par conséquent

$$P\{|g(x) - \bar{G}_k(x)| \geq \varepsilon/2 \mid x \in U_\infty\} = 0$$

Vu que $|A| < \infty$, la somme de la relation (3.7) a un nombre fini de termes et elle est égale à 0, ce qui prouve les relations (3.5) et (3.4). ■

3.2.2 Optimisation discrète par simulation- cas d'un espace de solutions non borné

Considérons le problème DOvS défini dans la Section 3.2.1. pour un espace de solutions non borné $\Theta = \Phi \cap \mathbb{Z}^d$ où Φ est un ensemble fermé mais non borné dans $\mathbb{R}^d$.

L'application directe de l'algorithme 1 dans ce cas n'est pas possible, car la région la plus prometteuse définie dans la Section 3.2.1 est non bornée. Dans cette section l'algorithme COMPASS* est étendu par l'introduction du concept de la région d'observation à partir de laquelle on choisit les nouvelles solutions. A chaque itération de l'algorithme COMPASS*, la région d'observation D est définie par un cube autour de l'optimum observé x^* . Plus précisément, $D = \prod_{i=1}^d [x^{*(i)} - q, x^{*(i)} + q]$ où $q > 0$ est la largeur de la région d'observation. Nous attirons l'attention sur le fait que la région d'observation est une fenêtre d'observation glissante, mais finie. Ce-ci constitue une différence fondamentale par rapport à l'algorithme COMPASS de (Hong et Nelson, 2006) qui utilise une fenêtre d'observation centrée sur une solution donnée et qui couvre toutes les solutions visitées. L'utilisation d'une fenêtre d'observation finie permet d'assurer la convergence et d'accélérer l'optimisation sous des conditions bien moins restrictives.

Dans ce cas la méthode COMPASS* respecte le même principe que pour un espace de solutions borné, à l'exception que cette fois ci la région la plus prometteuse est limitée à la région d'observation D . Elle contient les solutions dans D qui sont plus proches de l'optimum observé que des autres solutions visitées.

Algorithm 2 (COMPASS* pour DOvS avec espace de solutions non borné)

Etape 1: Initialisation.

- 1.1. Initialiser le compteur des itérations $k=0$. Choisir $x_0 \in \Theta$, faire $V_0 = \{x_0\}$ et $x_0^* = x_0$.
- 1.2. Déterminer $a_0(x_0)$ selon la SAR.
- 1.3. Effectuer $a_0(x_0)$ répliques de simulation pour la solution x_0 , faire $N_0(x_0) = a_0(x_0)$, et calculer $\bar{G}_0(x_0)$.
- 1.4. Faire $D_0 = \Pi_{i=1}^d [x_0^{(i)} - q, x_0^{(i)} + q]$ où $q > 0$ est un nombre entier constant. Faire $C_0 = \Theta$.

Etape 2: Soit $k=k+1$. Choisir des nouvelles solutions et allouer le budget de simulation.

- 2.1. Choisir uniformément m solutions $x_{k1}, x_{k2}, \dots, x_{km}$ indépendantes dans $C_{k-1} \cap D_{k-1}$. Soit $V_k = V_{k-1} \cup \{x_{k1}, x_{k1}, \dots, x_{km}\}$
- 2.2. Déterminer $a_k(x)$ de chaque $x \in V_k$ selon la SAR.
- 2.3. $\forall x \in V_k$, faire $a_k(x)$ répliques de simulation et mettre à jour $N_k(x)$ et $\bar{G}_k(x)$.

Etape 3: Déterminer le nouvel optimum et mettre à jour la région la plus prometteuse

- 3.1. Déterminer $x_k^* = \arg \min_{x \in V_k} \bar{G}_k(x)$,
- 3.2. Construire $D_k = \Pi_{i=1}^d [x_k^{*(i)} - q, x_k^{*(i)} + q]$,
- 3.3. Construire $C_k = \{x \in \Theta : \|x - x_k^*\| \leq \|x - y\|, \forall y \in V_k \text{ and } y \neq x_k^*\}$
- 3.4. Aller à l'étape 2.

Remarque: La région d'observation D_k est un polyèdre et la région la plus prometteuse C_k est un ensemble de points dans ce polyèdre. Par conséquent, la méthode RMD proposée par (Hong et Nelson 2006) peut être utilisée pour le choix uniforme des solutions dans $C_{k-1} \cap D_{k-1}$. De même, par définition, l'ensemble $C_{k-1} \cap D_{k-1}$ contient au moins l'élément x_{k-1}^* . Donc il est non vide.

En plus des hypothèses 1 et 2, d'autres hypothèses sont nécessaires afin de garantir la convergence de l'algorithme 2.

Hypothèse 3. Pour la solution initiale x_0 , il existe un ensemble compact et fini Π et une constante positive $\delta > 0$ tels que $x_0 \in \Pi \cap \Theta$ et $G(x, w_i) \geq g(x_0) + \delta \quad \forall x \in \Pi^c \cap \Theta$.

■

Beaucoup de problèmes d'optimisation par simulation discrets ont une solution étalon, disons x_0 , qui est souvent le paramétrage courant. Toute solution située au-delà d'une certaine distance (inconnue) de la solution étalon est pire que la solution étalon. Par conséquent, il existe une région bornée (mais inconnue) Π telle que $g(x) \geq g(x_0) + \delta, \quad \forall x \in \Pi^c \cap \Theta$. Cette propriété peut être appliquée à la performance mesurée par la simulation $G(x, w_i)$ si sa variance est raisonnable ou si $G(x, w_i)$ contient des éléments de coût non bornés directement liés à x . L'hypothèse 3 assure également que des solutions optimales locales finies existent et se trouvent dans l'ensemble Π .

Hypothèse 4 $\lim_{k \rightarrow \infty} N_k(x_0) = +\infty$.

■

Selon l'hypothèse 4, la solution étalon x_0 est simulée un nombre illimité de fois. Cette hypothèse est nécessaire afin de s'assurer que la région d'observation D_k ne va pas glisser de la solution étalon vers l'infini.

Théorème 3 Sous les hypothèses 1-4, la séquence infinie des optimums observés $\{x_0^*, x_1^*, \dots\}$ générée par l'algorithme 2 converge avec probabilité 1 vers l'ensemble M dans le sens que $P\{x_k^* \notin M \text{ i.o.}\} = 0$.

■

Preuve du théorème 3. Pour toute séquence infinie $\{V_0, V_1, \dots\}$ générée par l'algorithme 2, vu que $V_k \subset V_{k+1}$ et que $V_k \subset \Theta \quad \forall k = 0, 1, \dots$, il existe $V_\infty = \bigcup_{k=0}^{\infty} V_k$ tel que $V_\infty \subset \Theta$.

Selon les hypothèses 1 et 4, il existe $K_1 > 0$ tel que $\forall k \geq K_1, |\bar{G}_k(x_0) - g(x_0)| < \delta$ avec la probabilité 1. A partir de l'hypothèse 3, $\bar{G}_k(x) \geq g(x_0) + \delta \quad \forall x \in \Pi^c \cap \Theta \cap V_k$. Par conséquent, $\bar{G}_k(x) > \bar{G}_k(x_0) \geq \min_{z \in V_k} \bar{G}_k(z) \quad \forall x \in \Pi^c \cap \Theta \cap V_k$ et $\forall k \geq K_1$. Ainsi, $x_k^* \in \Pi$ et $D_k \in \Pi^+ \quad \forall k \geq K_1$, ce qui implique $V_k \subset V_{K_1} \cup \Pi^+$, où $\Pi \subseteq \Pi_{i=1}^d[\underline{b}^{(i)}, \bar{b}^{(i)}]$ et $\Pi^+ = \Pi_{i=1}^d[\underline{b}^{(i)} - q, \bar{b}^{(i)} + q]$. Comme $|V_{K_1}| < \infty$ et $|\Pi^+| < \infty$, l'ensemble V_∞ est fini, i.e. $|V_\infty| < \infty$.

Soit W_∞ l'ensemble de solutions qui sont optimums observés un nombre illimité de fois, i.e. $W_\infty = \{x \in V_\infty : x_k^* = x \text{ i.o.}\}$. L'ensemble W_∞ est différent de l'ensemble U_∞ utilisé dans la preuve du théorème 1.

Comme $W_\infty \subseteq V_\infty$ et $|V_\infty| < \infty$, W_∞ est un ensemble fini. De plus, $x_k^* \in \Pi, \forall k \geq K_1$.

Donc

$$W_\infty \neq \emptyset, W_\infty \subseteq \Pi, \text{ et } P\{W_\infty \subseteq \Pi\} = 1.$$

Par conséquent,

$$P\{x_k^* \notin M \text{ i.o.}\} = \sum_{A \subseteq \Pi} P\{x_k^* \notin M \text{ i.o.} | W_\infty = A\} P\{W_\infty = A\} \quad (3.9)$$

où la somme est effectuée sur tous les sous-ensembles non vides $A \subset \Pi$. Ainsi, prouver $P\{x_k^* \notin M \text{ i.o.}\} = 0$ est équivalent à prouver $P\{x_k^* \notin M \text{ i.o.} | U_\infty = A\} = 0 \quad \forall A \subset \Pi$ non vide et finie tel que $P\{W_\infty = A\} > 0$.

Nous montrons maintenant que $P\{x_k^* \notin M \text{ i.o.} | W_\infty = A\} = 0$ pour tout ensemble non vide $A \subset \Pi$ tel que $P\{W_\infty = A\} > 0$. Si $x_k^* \notin M \text{ i.o.}$, alors il existe une solution $x \in A$ et $x \notin M$ telle que $x_k^* = x \text{ i.o.}$. Comme $x \notin M$, il existe $y \in NH(x) \cap \Theta$ tel que $g(y) < g(x)$.

Puisque $\|y - x\| = 1 \leq \|y - z\| \quad \forall \quad z \neq y$, alors $y \in C_{k+1} \cap D_{k+1}$ si $x_k^* = x$ et $y \notin V_k$.

Comme dans la preuve du théorème 1, on montre que $y \in V_\infty$ et $y \in U_\infty$, i.e. y sera choisi à une certaine itération K_3 et ensuite elle sera simulée un nombre infini de fois. Comme $g(y) < g(x)$, l'hypothèse 1 implique qu'il existe $K_4 > K_3$ tel que $\bar{G}_k(y) < \bar{G}_k(x), \forall k \geq K_4$. Par conséquent, x ne peut être optimum observé qu'un nombre fini de fois, ce qui contredit la supposition $x_k^* = x$ i.o et conclue la preuve. ■

Théorème 4. Sous les hypothèses 1-4, $\lim_{k \rightarrow \infty} g(x_k^*) = \min_{x \in U_\infty} g(x)$ avec probabilité 1, où U_∞ est l'ensemble de solutions qui sont simulées un nombre infini de fois par l'algorithme 2. ■

Preuve du théorème 4.

A partir de la preuve du théorème 3, $V_k \subseteq V_{k+1} \subseteq \Theta$, $\forall k \geq 0$, $V_\infty = \bigcup_{k=0}^\infty V_k$ est fini et $V_\infty \subseteq \Theta$. De plus, $\exists K_1 > 0$ tel que $\forall k \geq K_1$, $x_k^* \in \Pi$, $D_k \in \Pi^+$ et $V_k \subset V_{K_1} \cup \Pi^+ \quad \forall k \geq K_1$ où $\Pi \subseteq \Pi_{i=1}^d [\underline{b}^{(i)}, \bar{b}^{(i)}]$ et $\Pi^+ = \Pi_{i=1}^d [\underline{b}^{(i)} - q, \bar{b}^{(i)} + q]$. Par conséquent, $U_\infty \subseteq V_\infty \subseteq \Theta$ est un ensemble non vide et fini $x_0 \in U_\infty$.

A partir de l'hypothèse 3,

$$\min_{x \in U_\infty} g(x) = \min_{x \in U_\infty \cap \Pi} g(x)$$

ce qui implique

$$P\{\lim_{k \rightarrow \infty} g(x_k^*) = \min_{x \in U_\infty} g(x)\} = \sum_{A \subset \Pi \cap Z^d} P\{\lim_{k \rightarrow \infty} g(x_k^*) = \min_{x \in T_\infty} g(x) \mid T_\infty = A\} P\{T_\infty = A\} \quad (3.4)$$

où $T_\infty = U_\infty \cap \Pi$ et $A \subset \Pi \cap Z^d$.

La suite de la preuve est similaire à celle du théorème 2, en remplaçant U_∞ par T_∞ et en utilisant les propriétés suivantes:

- (i) $x_k^* \in \Pi$, $D_k \in \Pi^+$ et $V_k \subset V_{K_1} \cup \Pi^+ \quad \forall k \geq K_1$
- (ii) selon l'hypothèse 2, il existe $K_2 \geq K_1$ tel que $x_k^* \in T_\infty \quad \forall k \geq K_2$.

■

3.2.3 Implémentation et résultats numériques

Dans cette section nous appliquons les algorithmes COMPASS* 1 et 2 à l'optimisation d'une chaîne logistique et nous comparons les résultats avec ceux obtenus par l'algorithme COMPASS original de (Hong et Nelson, 2006).

La chaîne logistique considérée est un système de production distribution composé d'un entrepôt de produits finis alimenté par une chaîne d'unités de production séparées par des stocks intermédiaires. Ce système est illustré dans la figure 3.4. On considère qu'un seul type de produit est fabriqué et distribué par ce système. Les commandes des clients arrivent à l'entrepôt final selon un processus de Poisson.

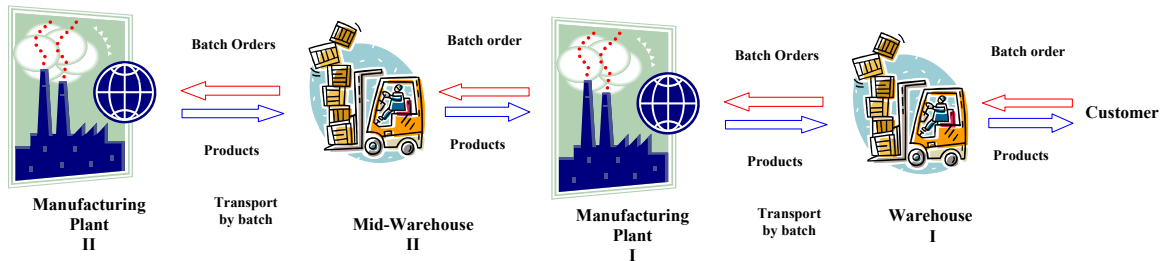


Figure 3.4. Un système de production-distribution

La gestion du stock à chaque entrepôt i est effectuée selon une politique base-stock, avec un niveau base-stock R_i . Cette politique est basée sur le concept de position du stock à chaque entrepôt, qui prend en compte le stock disponible et en transit vers l'entrepôt, plus la quantité commandée en attente d'exécution à l'unité de production en amont, moins les demandes en attente. L'objectif de la gestion du stock d'un entrepôt i par une politique base-stock est de maintenir la position du stock à une valeur constante appelée base-stock R_i . Par conséquent, chaque fois que l'entrepôt final reçoit une

commande d'une quantité X , chacune des unités de production en amont reçoit un ordre de réapprovisionnement d'une quantité X .

Les ordres d'approvisionnement sont traités par chaque unité de production selon une politique FIFO. Dès qu'un ordre est exécuté par l'unité de fabrication, la quantité correspondante est transportée à l'entrepôt en aval. La durée de transport est significative et ne peut pas être négligée.

La performance du système dépend fortement des niveaux de base-stock des différents entrepôts. L'objectif est de déterminer le niveau de base-stock de chaque entrepôt $R = \{R_1, R_2, R_3, \dots\}$ tel que la somme des coûts de stockage sur l'ensemble des entrepôts du système et du coût de rupture à l'entrepôt final (celui qui est en contact avec les clients) soit minimisée. Plus précisément, le problème est de minimiser la fonction objectif suivante :

$$C(R) = \sum_{i=1,2,\dots} C_{hi} E\left[\int_0^T I_i(t) dt\right] + C_b E\left[\int_0^T B(t) dt\right] \quad (3.13)$$

où C_{hi} est le coût de stockage par unité de temps et par unité de produit à l'entrepôt i , C_b est le coût de rupture par unité de temps et par unité de produit à l'entrepôt final, $B(t)$ est la quantité de demande en attente à l'entrepôt final, $I_i(t)$ est le stock disponible de l'entrepôt i et T est l'horizon du problème.

Une approche analytique a été proposée dans le chapitre 2 pour résoudre ce problème dans le cas d'une chaîne logistique à deux niveaux : un entrepôt alimenté par une unité de production. Cependant, lorsque le nombre de niveaux augmente il est très difficile de développer un modèle analytique adapté à ce problème. Ainsi, nous proposons l'utilisation des techniques d'optimisation par simulation pour résoudre ce problème. Nous appliquons l'algorithme COMPASS original (Hong et Nelson, 2006) ainsi que l'algorithme COMPASS* que nous avons introduit dans ce chapitre. L'objectif est de comparer : 1) le budget de simulation nécessaire, i.e. la vitesse de la convergence et 2) le coût global de gestion du stock associé aux solutions optimales fournies par les

algorithmes COMPASS original et COMPASS*.

Les solutions initiales de tous les algorithmes sont générées aléatoirement. Chaque solution est évaluée par des répliques i.i.d. de simulation, leur nombre étant fixé par la règle d'allocation du budget de simulation (SAR).

Pour chacun des algorithmes on utilise une SAR qui satisfait les hypothèses nécessaires pour la convergence. Pour l'algorithme COMPASS original, nous exécutons une simulation supplémentaire pour chaque solution visitée à l'itération k , i.e. $a_k(x)=1$. Pour l'algorithme COMPASS* nous utilisons la règle suivante :

$$a_k(x) = \begin{cases} 1 & \text{if } x \in V_k \setminus V_{k-1} \text{ or } x \in NH(x_{k-1}^*) \\ 0 & \text{otherwise} \end{cases}$$

et $a_k(x_0)=1$, $\forall k$ lorsque l'algorithme COMPASS* est appliqué au cas de l'espace de solutions non borné.

L'expérience numérique est conduite avec le paramétrage suivant. L'arrivée des clients est un processus de Poisson avec le taux moyen d'arrivée $\lambda = 1.5$. La taille de chaque commande est uniformément distribuée dans l'intervalle $[3, 9]$. Le temps de service τ pour une unité de produit dans chaque unité de fabrication est exponentiellement distribué avec une moyenne de 0,1. La durée de transport entre une unité de production et son stock en aval est de 3.0. Le coût unitaire de stockage ainsi que le coût unitaire de rupture sont tous les deux égaux à 1. L'horizon de la simulation est fixé à $T=1000$ unités.

Différentes expérimentations de simulation sont effectuées pour différentes configurations de la chaîne logistique avec nombre différent (5, 8, 10, 12) d'unité de production. Un budget de simulation total (5000, 10000 ou 20000 répliques de simulations) est donné a priori et doit être réparti entre les différentes solutions pour l'optimisation.

Considérons dans un premier temps le cas de l'espace de solutions borné. Ainsi, chaque niveau de base-stock R_i est bornée par l'intervalle $[0, 2000]$. Les résultats pour un

système de production-distribution avec 8 et 10 niveaux sont donnés dans les figures 3.5 et 3.6.

Nous considérons ensuite le cas où l'espace des solutions est non borné. Dans ce cas le niveau de base-stock à chaque entrepôt i est non borné, $R_i \geq 0$. Des expérimentations de simulation sont effectuées pour des systèmes de production-distribution avec 5, 10 et 12 niveaux. Les résultats numériques sont illustrés dans les figures 3.7, 3.8 et 3.9.

Ces résultats confirment que l'algorithme COMPASS* converge vers un optimum local plus rapidement que l'algorithme COMPASS original. Ceci est dû au fait qu'il dépense la plupart du budget de simulation pour les solutions intéressantes. Ainsi, les algorithmes COMPASS* sont capables d'identifier les optimums locaux avec un nombre réduit de simulations.

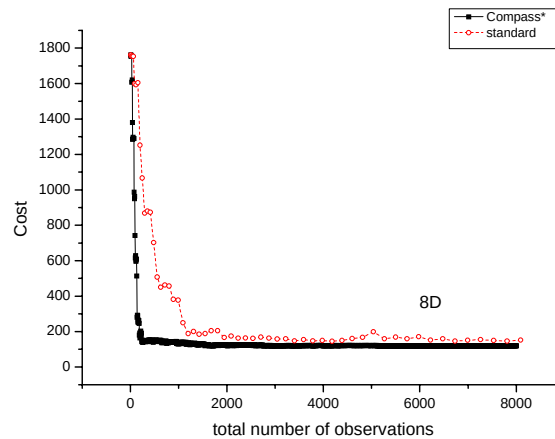


Figure 3.5: Une chaîne logistique à 8 niveaux de l'espace de solutions borné

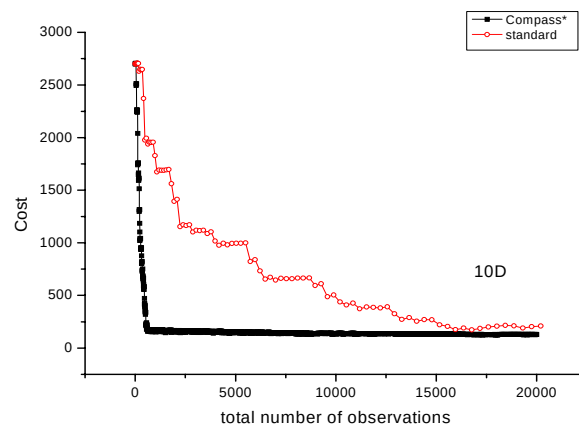


Figure 3.6: Une chaîne logistique à 10 niveaux de l'espace de solutions borné

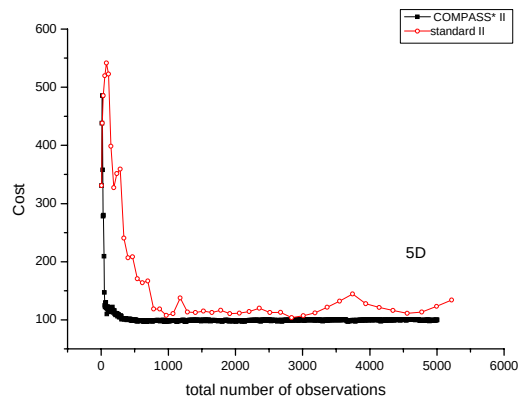


Figure 3.7: Une chaîne logistique à 5 niveaux de l'espace de solutions non borné

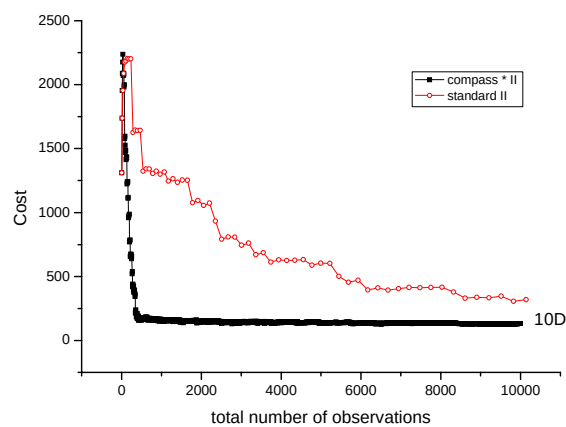


Figure 3.8: Une chaîne logistique à 10 niveaux de l'espace de solutions non borné

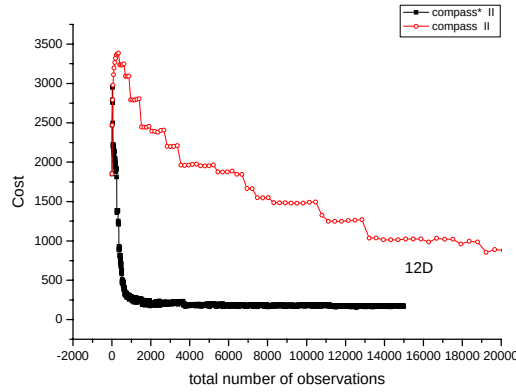


Figure 3.9. Une chaîne logistique à 12 niveaux de l'espace de solutions non borné

3.3 Optimisation discrète par simulation avec contraintes non analytiques

Cette section traite le problème d'optimisation discrète par simulation où la fonction objectif et certaines contraintes n'admet pas de solutions analytiques et doivent être évaluées par la simulation. Plus précisément, nous considérons le problème d'optimisation discrète par simulation suivant :

$$\begin{aligned} \min_{x \in \Theta} E[G(x, w_i)] \\ \text{tel que} \quad E[F(x, w_i)] > \alpha \end{aligned}$$

où l'ensemble de contraintes est décrit par des contraintes analytiques $\Theta = \Phi \cap \mathbb{Z}^d$ ainsi que par une contrainte non analytique $E[F(x, w_i)] > \alpha$. Nous supposons que la fonction objectif $E[G(x, w_i)]$ et la contrainte $E[F(x, w_i)]$ peuvent être estimées à l'aide des répliques de simulations qui sont indépendantes et identiquement distribuées. Soit $G(x, w_i)$ et $F(x, w_i)$ respectivement la $i^{\text{ème}}$ valeur mesurée pour la fonction objectif et pour la contrainte pour la solution x . Notons que $g(x) = E[G(x, w_i)]$ et que $f(x) = E[F(x, w_i)]$.

Par la suite, nous considérons le cas où l'espace de solutions $\Theta = \Phi \cap Z^d$ est borné et le cas où l'espace de solutions $\Theta = \Phi \cap Z^d$ est non borné.

3.3.1 Optimisation discrète par simulation avec contraintes avec l'espace de solutions borné

Dans ce cas, l'algorithme COMPASS* avec des contraintes estimées par simulation est similaire à l'algorithme COMPASS* développé lorsque toutes les contraintes sont définies par une expression analytique, i.e. l'algorithme 1 présenté dans la Section 3.1.1.

La fonction objectif et la contrainte non analytique sont estimées par la moyenne des résultats obtenus suite à plusieurs répliques de simulation. Soit $\bar{G}_k(x)$ et $\bar{F}_k(x)$ respectivement les moyennes sur $N_k(x)$ répliques de simulation de $G(x, w_i)$ et $F(x, w_i)$ suite à l'itération k :

$$\bar{G}_k(x) = \frac{1}{N_k(x)} \sum_{i=1}^{N_k(x)} G_i(x, w_i).$$

$$\bar{F}_k(x) = \frac{1}{N_k(x)} \sum_{i=1}^{N_k(x)} F_i(x, w_i)$$

Afin d'intégrer l'évaluation de la contrainte F , nous modifions la fonction objectif en y ajoutant une fonction de pénalisation. Ainsi, l'optimum observé à chaque itération est déterminé à travers la fonction objectif modifiée suivante :

$$H_k(x) = \bar{G}_k(x) + S_k(x)(\alpha - \bar{F}_k(x))^+$$

où $S_k(x) > 0$ est le facteur de pénalisation pour la violation de la contrainte non analytique et $(x)^+ = \max(0, x)$.

Algorithme 3 (COMPASS* pour l'optimisation discrète par simulation avec contraintes avec l'espace de solutions borné)

Etape 1: Initialisation

- 1.1 Initialiser le compteur d'itérations $k = 0$. Choisir $x_0 \in \Theta$, et faire $V_0 = \{x_0\}$ et $x_0^* = x_0$.
- 1.2 Déterminer $a_0(x_0)$ selon une règle d'allocation du budget de simulation (SAR).
- 1.3 Effectuer $a_0(x_0)$ répliques de simulation pour la solution x_0 , faire $N_0(x_0) = a_0(x_0)$, et calculer $\bar{G}_0(x_0)$ et $\bar{F}_0(x_0)$. Soit $C_0 = \Theta$.

Etape 2 : Mettre à jour $S_k(x)$, où $S_k(x)$ est une fonction croissante non négative en k .

Etape 3: Soit $k = k + 1$. Choisir des nouvelles solutions et allouer le budget de simulation.

- 3.1 Choisir uniformément m solutions $x_{k1}, x_{k2}, \dots, x_{km}$ indépendantes dans C_{k-1} . Soit $V_k = V_{k-1} \cup \{x_{k1}, x_{k2}, \dots, x_{km}\}$.
- 3.2 Déterminer $a_k(x)$ selon la SAR pour chaque solution $x \in V_k$.
- 3.3 Pour chaque $x \in V_k$, effectuer $a_k(x)$ observations, ensuite mettre à jour $N_k(x), \bar{G}_k(x)$ et $\bar{F}_k(x)$.

Etape 4: Déterminer la nouvelle solution localement optimale et mettre à jour la région la plus intéressante.

- 4.1 Soit $x_k^* = \arg \min_{x \in V_k} H_k(x)$ où $H_k(x) = \bar{G}_k(x) + S_k(x)(\alpha - \bar{F}_k(x))^+$.
- 4.2 Construire la région la plus prometteuse:

$$C_k = \{x : x \in \Theta \text{ et } \|x - x_k^*\| \leq \|x - y\| \quad \forall y \in V_k \text{ et } y \neq x_k^*\}$$
- 4.3 Aller à l'étape 2.

Hypothèse 5: Pour chaque $x \in \Theta$, $N_k(x) \rightarrow \infty \Rightarrow S_k(x) \rightarrow \infty$.

Remarques:

1. Dans l'étape 2 de l'algorithme 3, la contrainte est prise en compte dans la fonction objectif sous la forme d'une pénalité. La fonction croissante S_k est une pénalité introduite si le niveau souhaité pour la contrainte non analytique est violé selon la

valeur de $\bar{F}_k(x)$. Cette fonction de pénalisation assure la convergence de l'algorithme vers des solutions admissibles.

2. Nous ne fixons pas de critère d'arrêt pour l'algorithme afin de prouver sa convergence. Dans la pratique l'exécution de l'algorithme peut être arrêtée lorsque x_k^* ne change pas pendant plusieurs itérations et toutes les solutions dans son voisinage ont été visités.

3. Il n'est pas nécessaire de visiter toutes les solutions dans Θ , même si le budget de simulation est infini, ce qui n'est pas le cas dans les algorithmes DOvS globalement convergents.

■

Les hypothèses suivantes sont similaires aux hypothèses nécessaires pour la convergence de l'algorithme 1.

Hypothèse 6: Pour tout $x \in \Theta$ les relations suivantes sont vraies :

$$P\left[\lim_{N_k(x) \rightarrow \infty} \frac{1}{N_k(x)} \sum_{i=1}^{N_k(x)} G_i(x, w_i) = g(x)\right] = 1$$

$$\text{et } P\left[\lim_{N_k(x) \rightarrow \infty} \frac{1}{N_k(x)} \sum_{i=1}^{N_k(x)} F_i(x, w_i) = f(x)\right] = 1$$

■

Hypothèse 7: La SAR garantit que $a_k(x) \geq 1$ si x est une solution visitée pour la première fois à l'itération k , i.e. ($x \in V_k \setminus V_{k-1}$) ou si x se trouve dans le voisinage de l'optimum observé à l'itération précédente, i.e. ($x \in NH(x_{k-1}^*)$).

■

Des hypothèses supplémentaires sont nécessaires pour garantir l'existence d'une solution pour le problème d'optimisation par simulation et pour la convergence de l'algorithme vers un optimum local.

Hypothèse 8 x_0 est une solution admissible, i.e. $f(x_0) > \alpha$.

■

Hypothèse 9 $\lim_{k \rightarrow \infty} N_k(x_0) = +\infty$.

■

La contrainte technique suivante est également nécessaire.

Hypothèse 10 Il n'y a aucune solution $x \in \Phi$ telle que $f(x) = \alpha$.

■

Le théorème suivant caractérise la propriété de convergence vers un optimum local de l'algorithme 3. Soit M l'ensemble de solutions x dites optimums locaux telles que $f(x) > \alpha$ et $g(x) < g(y)$ pour tout $y \in NH(x) \cap \Theta$ et $f(y) > \alpha$. Si M est un singleton, alors l'élément de M est la solution globale optimale.

Théorème 5. Sous les hypothèses 5-10, la séquence infinie des optimums observés $\{x_0^*, x_1^*, \dots\}$ générée par l'algorithme 3 converge avec la probabilité 1 vers l'ensemble M dans le sens que $P\{x_k^* \notin M \text{ i.o.}\} = 0$.

■

Lemme 2. Sous les hypothèses 5, 6 et 10, si $\lim_{k \rightarrow \infty} N_k(x) = \infty$, alors $\lim_{k \rightarrow \infty} H_k(x) = \infty$ si $f(x) \leq \alpha$ et $\lim_{k \rightarrow \infty} H_k(x) = g(x)$ si $f(x) > \alpha$.

■

Preuve. Soit $\varepsilon = |f(x) - \alpha|/2$. Selon l'hypothèse 10, $\varepsilon > 0$. Selon l'hypothèse 6, $\exists K > 0$ tel que $|f(x) - \bar{F}_k(x)| < \varepsilon \forall k > K$. Ainsi, $\bar{F}_k(x) - \alpha > \varepsilon$ si $f(x) > \alpha$ et $\bar{F}_k(x) - \alpha < -\varepsilon$ si $f(x) \leq \alpha$. Par conséquent, si $f(x) > \alpha$, alors

$$\lim_{k \rightarrow \infty} H_k(x) = \lim_{k \rightarrow \infty} \left[\bar{G}_k(x) + S_k(x)(\alpha - \bar{F}_k(x))^+ \right] = \lim_{k \rightarrow \infty} \left[\bar{G}_k(x) \right] = g(x)$$

Si $f(x) \leq \alpha$, l'hypothèse 5 mène à

$$\lim_{k \rightarrow \infty} H_k(x) > \lim_{k \rightarrow \infty} [\bar{G}_k(x) + S_k(x)\varepsilon] = \infty$$

■

Lemme 3. $\forall$ solution x tel que $x_k^* = x$ i.o., avec probabilité 1, x est une solution admissible, i.e. $f(x) > \alpha$.

■

Preuve de la lemme 3 : Supposons qu'il existe une solution x telle que $x_k^* = x$ i.o. et $f(x) \leq \alpha$. Selon l'hypothèse 7, $N_k(x) \rightarrow \infty$ lorsque k augmente. A partir de Lemme 2,

$$\lim_{k \rightarrow \infty} H_k(x) = \infty.$$

De même, selon les hypothèses 8-9 et la Lemme 2,

$$\lim_{k \rightarrow \infty} H_k(x_0) = g(x_0).$$

Ces deux relations impliquent que x ne peut être optimum observé qu'un nombre fini de fois, ce qui contredit le fait que $x_k^* = x$ i.o. et conclue la preuve.

■

Preuve du Théorème 5. Comme dans la preuve du théorème 1, $V_k \subseteq V_{k+1} \subseteq \Theta$, $\forall k \geq 0$, $V_\infty = \bigcup_{k=0}^{\infty} V_k$ est fini et $V_\infty \subseteq \Theta$.

Soit U_∞ l'ensemble de solutions qui sont simulées un nombre illimité de fois par l'algorithme 3 et soit U_∞^* l'ensemble de solutions admissibles qui sont simulées un nombre illimité de fois, i.e. $U_\infty^* = \{x \in U_\infty : f(x) > \alpha\}$. Selon l'hypothèse 8-9, U_∞^* est un ensemble non vide et fini et $x_0 \in U_\infty^*$. Par conséquent,

$$P\{x_k^* \notin M \text{ i.o.}\} = \sum_{A \subset \Theta} P\{x_k^* \notin M \text{ i.o.} | U_\infty^* = A\} P\{U_\infty^* = A\} \quad (3.3)$$

où $A \subset \Theta$, est un sous-ensemble fini de Θ . Selon le Lemme 3, chaque solution x telle que $x_k^* = x$ i.o. est admissible et donc $x \in U_\infty^*$. Comme dans la preuve du théorème 1,

on peut montrer que $P\{x_k^* \notin M \text{ i.o.} | U_\infty^* = A\} = 0 \quad \forall A \subset \Theta$ non vide tel que $P\{U_\infty^* = A\} > 0$. Par conséquent, $P\{x_k^* \notin M \text{ i.o.}\} = 0$ et la preuve est accomplie. ■

Théorème 6. Sous les hypothèses 5-10,

$$\lim_{k \rightarrow \infty} g(x_k^*) = \min_{x \in U_\infty^*} g(x) \text{ avec probabilité 1,}$$

où U_∞^* est l'ensemble de solutions admissibles qui sont simulées un nombre illimité de fois par l'algorithme 3. ■

Preuve du théorème 6.

Selon la preuve du théorème 2, $V_k \subseteq V_{k+1} \subseteq \Theta$, $\forall k \geq 0$, $V_\infty = \bigcup_{k=0}^\infty V_k$ est fini et $V_\infty \subseteq \Theta$. De plus, U_∞^* est un ensemble non vide et fini et $x_0 \in U_\infty^*$. Selon la Lemme 3 et l'hypothèse 7, il existe un nombre entier positif $K > 0$ tel que $x_k^* \in U_\infty^* \quad \forall k > K$. Par conséquent,

$$P\{\lim_{k \rightarrow \infty} g(x_k^*) = \min_{x \in U_\infty^*} g(x)\} = \sum_{A \subset \Theta} P\{\lim_{k \rightarrow \infty} g(x_k^*) = \min_{x \in U_\infty^*} g(x) | U_\infty^* = A\} P\{U_\infty^* = A\}$$

où A est un sous-ensemble fini de Φ . La suite de la démonstration est similaire à celle du théorème 2, en remplaçant U_∞ avec U_∞^* . ■

3.3.2 Optimisation discrète par simulation avec contraintes avec l'espace de solutions non borné

Considérons l'espace de solutions défini par les contraintes analytiques $\Theta = \Phi \cap \mathbb{Z}^d$. Si Φ est un ensemble fermé mais non borné, alors l'espace de solutions est également fermé mais non borné. Dans ce cas nous proposons de résoudre le problème d'optimisation par un algorithme inspiré des algorithmes 2 et 3 déjà présentés.

Algorithme 4 (COMPASS* pour Optimisation discrète par simulation avec contraintes

avec l'espace de solutions non borné)

Etape 1:Initialisation

- 1.1 Initialiser le compteur d'itérations $k = 0$. Choisir $x_0 \in \Theta$, faire $V_0 = \{x_0\}$ et $x_0^* = x_0$.
- 1.2 Déterminer $a_0(x_0)$ selon la SAR.
- 1.3 Effectuer $a_0(x_0)$ répliques de simulation pour la solution x_0 , faire $N_0(x_0) = a_0(x_0)$, et calculer $\overline{G}_0(x_0)$ et $\overline{F}_0(x_0)$. Soit $D_0 = \Pi_{i=1}^d [x_0^{(i)} - q, x_0^{(i)} + q]$ où q est un nombre entier constant.
- 1.4 Soit $C_0 = D_0 \cap \Theta$.

Etape 2: Mettre à jour $S_k(x)$, où $S_k(x)$ est une fonction croissante en k .

Etape 3: Soit $k = k + 1$. Choisir des nouvelles solutions et allouer le budget de simulation.

- 3.1 Choisir uniformément m solutions $x_{k1}, x_{k2}, \dots, x_{km}$ indépendantes à partir de C_{k-1} . Soit $V_k = V_{k-1} \cup \{x_{k1}, x_{k2}, \dots, x_{km}\}$.
- 3.2 Déterminer $a_k(x)$ selon la SAR pour chaque $x \in V_k$.
- 3.3 Pour chaque $x \in V_k$, faire $a_k(x)$ répliques de simulation, et mettre à jour $N_k(x)$, $\overline{G}_k(x)$ et $\overline{F}_k(x)$.

Etape 4: Déterminer le nouvel optimum local et mettre à jour la région la plus prometteuse.

- 4.1 Soit $x_k^* = \arg \min_{x \in V_k} H_k(x)$ où $H_k(x) = \overline{G}_k(x) + S_k(x)(\alpha - \overline{F}_k(x))^+$.

- 4.2 Construire $D_k = \Pi_{i=1}^d [x_k^{*(i)} - q, x_k^{*(i)} + q]$,

$$C_k = D_k \cap \{x : x \in \Theta \text{ et } \|x - x_k^*\| \leq \|x - y\| \quad \forall y \in V_k \text{ et } y \neq x_k^*\}$$

- 4.3 Aller à l'étape 2.

Théorème 7 Sous les hypothèses 3, et 5-10, la séquence infinie des optimums observés $\{x_0^*, x_1^*, \dots\}$ générée par l'algorithme 4 converge avec la probabilité 1 dans l'ensemble M dans le sens que $P\{x_k^* \notin M \text{ i.o.}\} = 0$.

■

Notons que les Lemmes 2 et 3 restent vraies pour l'algorithme 4.

Preuve du théorème 7. Pour toute séquence infinie $\{V_0, V_1, \dots\}$ générée par l'algorithme 4, comme $V_k \subset V_{k+1}$ et $V_k \subset \Theta \quad \forall k = 0, 1, \dots$, il existe un ensemble $V_\infty = \bigcup_{k=0}^{\infty} V_k$ tel que $V_\infty \subset \Theta$.

Selon les hypothèses 6 et 9, la Lemme 1 implique que $\exists K_1 > 0$ tel que $\forall k \geq K_1$, $|H_k(x_0) - g(x_0)| < \delta$ avec la probabilité 1. Selon l'hypothèse 3, $H_k(x) \geq \bar{G}_k(x) \geq g(x_0) + \delta \quad \forall x \in \Pi^c \cap \Theta \cap V_k$. Ainsi, $H_k(x) > H_k(x_0) \geq \min_{z \in V_k} H_k(z)$ $\forall x \in \Pi^c \cap \Theta \cap V_k$ et $\forall k \geq K_1$. Par conséquent, $x_k^* \in \Pi$ et $D_k \in \Pi^+ \quad \forall k \geq K_1$, ce qui implique $V_k \subset V_{K_1} \cup \Pi^+$, où $\Pi \subseteq \Pi_{i=1}^d[\underline{b}^{(i)}, \bar{b}^{(i)}]$ et $\Pi^+ = \Pi_{i=1}^d[\underline{b}^{(i)} - q, \bar{b}^{(i)} + q]$. Vu que $|V_{K_1}| < \infty$ et $|\Pi^+| < \infty$, l'ensemble V_∞ est fini, i.e. $|V_\infty| < \infty$.

Soit W_∞ l'ensemble de solutions qui sont optimums observés infiniment souvent, i.e.

$W_\infty = \{x \in V_\infty : x_k^* = x \text{ i.o.}\}$. Selon Lemme 2, W_∞ est un ensemble de solutions admissibles.

De plus, puisque $x_k^* \in \Pi \quad \forall k \geq K_1$, W_∞ est non vide et fini.

La suite de la preuve est similaire à celle du théorème 3 en considérant les propriétés suivantes:

- (i) chaque solution de W_∞ est admissible,
- (ii) $\forall x, y \in W_\infty$ solutions admissibles telles que $g(y) < g(x)$, $\exists K' > 0$ tel que $H_k(y) < H_k(x), \forall k \geq K'$.

■

Théorème 8. Sous les hypothèses 3 et 5-10,

$$\lim_{k \rightarrow \infty} g(x_k^*) = \min_{x \in U_\infty^*} g(x) \quad \text{avec la probabilité 1,}$$

où U_∞^* est l'ensemble des solutions admissibles qui sont simulées un nombre illimité de fois par l'algorithme 4. ■

Preuve du théorème 8.

A partir de la preuve du théorème 7, $V_k \subseteq V_{k+1} \subseteq \Theta$, $\forall k \geq 0$, $V_\infty = \bigcup_{k=0}^\infty V_k$ est fini et $V_\infty \subseteq \Theta$. De plus, $\exists K_1 > 0$ tel que $\forall k \geq K_1$, $x_k^* \in \Pi$, $D_k \in \Pi^+$ et $V_k \subset V_{K_1} \cup \Pi^+ \quad \forall k \geq K_1$ où $\Pi \subseteq \Pi_{i=1}^d [\underline{b}^{(i)}, \bar{b}^{(i)}]$ et $\Pi^+ = \Pi_{i=1}^d [\underline{b}^{(i)} - q, \bar{b}^{(i)} + q]$. Ainsi U_∞^* est un ensemble non vide et fini et $x_0 \in U_\infty^*$.

A partir de l'hypothèse 3,

$$\min_{x \in U_\infty^*} g(x) = \min_{x \in U_\infty^* \cap \Pi} g(x)$$

ce qui implique

$$P\{\lim_{k \rightarrow \infty} g(x_k^*) = \min_{x \in U_\infty^*} g(x)\} = \sum_{A \subset \Pi \cap Z^d} P\{\lim_{k \rightarrow \infty} g(x_k^*) = \min_{x \in T_\infty} g(x) \mid T_\infty = A\} P\{T_\infty = A\}$$

où $T_\infty = U_\infty^* \cap \Pi$ et A est un sous-ensemble de l'ensemble fini $\Pi \cap Z^d$.

La suite de la preuve est similaire à celle du théorème 2 en remplaçant U_∞ avec T_∞ et en utilisant les propriétés suivantes:

- (i) $x_k^* \in \Pi$, $D_k \in \Pi^+$ et $V_k \subset V_{K_1} \cup \Pi^+ \quad \forall k \geq K_1$,
- (ii) Grâce à l'hypothèse 7 et à Lemme 2, $\exists K_2 \geq K_1$ tel que $x_k^* \in T_\infty \quad \forall k \geq K_2$,
- (iii) Comme Lemme 1, on peut montrer qu'il existe $K' > 0$ tel que $H_k(x) = \bar{G}_k(x), \forall k \geq K', \forall x \in T_\infty$.

■

3.3.3. Implémentation et résultats numériques

Considérons à nouveau le système de production-distribution multi-niveau présenté dans la Section 3.2.3. Dans cette section, nous nous intéressons seulement à minimiser le coût total de stockage dans l'ensemble de la chaîne logistique tout en respectant un taux de service donné pour les clients finaux, i.e. le pourcentage de commandes livrées sans délai. L'objectif est de minimiser le coût total moyen en sélectionnant la meilleure solution $R^* = \{R_1, R_2, R_3, \dots\}$. Le problème d'optimisation est formalisée de la manière suivante :

$$\text{Minimize } C(R) = \sum_{i=1,2,\dots} C_{hi} E\left[\int_0^T I_i(t) dt\right] \quad (3.32)$$

$$\text{tel que } f(R) > \alpha$$

où C_{hi} est le coût unitaire de stockage du $i^{\text{ème}}$ entrepôt, $I_i(t)$ est le stock disponible au $i^{\text{ème}}$ entrepôt, $C(R)$ est le coût total moyen de stockage pour l'horizon de temps $[0, T]$ et $f(R)$ est le taux de service client pour le même horizon de temps.

Les expérimentations numériques sont conduites avec les paramètres suivants. La taille de chaque commande est un nombre entier aléatoire uniformément distribué dans l'intervalle $[3, 9]$. Le temps de service τ de chaque unité de produit est aléatoire et exponentiellement distribué avec la moyenne de 0.1. Le temps nécessaire pour le transport d'un lot entre une unité de production et l'entrepôt aval est de 3.0 unités de temps. Le coût unitaire de stockage, C_{hi} est égal à 1 et le taux de service client est de 90%. L'horizon de temps considéré est $T=500$ unités de temps.

Nous avons implémenté les algorithmes COMPASS* 3 et 4 à l'aide du langage de programmation C++. Le budget de simulation est alloué selon la SAR suivante :

$$a_k(x) = \begin{cases} 1 & \text{if } x \in V_k \setminus V_{k-1} \text{ or } x \in NH(x_{k-1}^*) \text{ or } x = x_0 \\ 0 & \text{otherwise} \end{cases}$$

A chaque itération, 5 nouvelles solutions sont choisies à partir de la région la plus prometteuse.

La solution initiale x_0 est générée aléatoirement. Comme la solution initiale doit être admissible, on peut choisir aléatoirement une solution et exécuter une simulation longue afin de tester si elle est admissible ou pas. Si elle n'est pas admissible, on choisit aléatoirement une autre solution. On répète cette opération jusqu'à ce qu'on trouve une solution admissible.

Le facteur de pénalité $S_k(x)$ est fixé indépendamment de la solution x et il est défini par la relation suivante :

$$S_k = e^k / \min(\alpha - \bar{F}_k(x))^+ \quad \forall x \in \{x : \alpha > \bar{F}_k(x)\}.$$

Plusieurs expérimentations ont été effectuées pour différentes configurations de systèmes de production-distribution avec différents nombres (5, 8, 10 et 12) d'unités de production. Le budget total de simulation a été fixé à 5000 observations.

D'abord, l'algorithme 3 est utilisé pour résoudre le problème d'optimisation pour un espace de solutions borné, plus précisément avec le niveau de base-stock R_i de chaque entrepôt limité dans l'intervalle $[0, 2000]$. Pour ce problème nous avons analysé des systèmes de production-distribution avec 5 et 10 niveaux. Les résultats sont illustrés dans les figures 3.9 et 3.11.

Ensuite, nous avons appliqué l'algorithme 4 pour résoudre le problème d'optimisation lorsque l'espace de solutions est non borné. Le niveau de base-stock à chaque entrepôt i est un entier positif, $R_i \geq 0$. Nous avons effectué des expérimentations pour des systèmes de production-distribution avec 5, 8 et 12 niveaux. Les résultats sont illustrés dans les figures 3.12, 3.13 et 3.14. Chacune des figures montre l'évolution de deux grandeurs : le coût total moyen et le taux de service de l'optimum observé. On peut remarquer que les algorithmes COMPASS* convergent rapidement. Dans la plupart de cas un nombre de quelques centaines d'observations est suffisant.

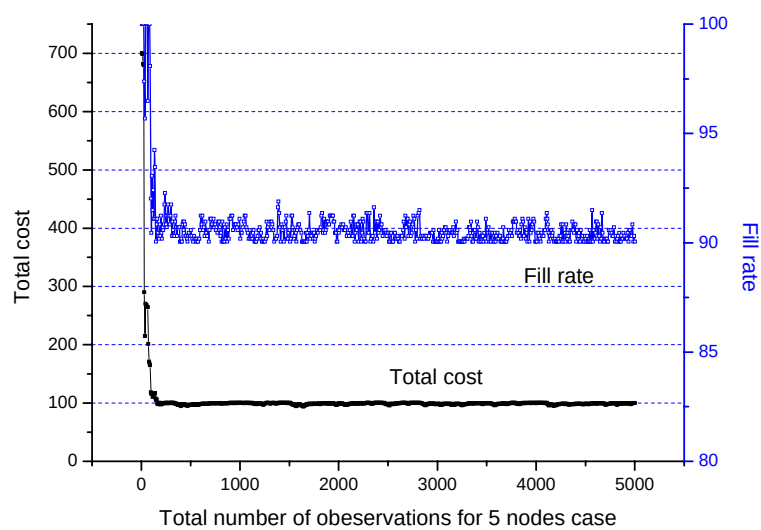


Figure 3.10: Une chaîne logistique à 5 niveaux eu un espace de solutions borné avec contrainte de niveau de service

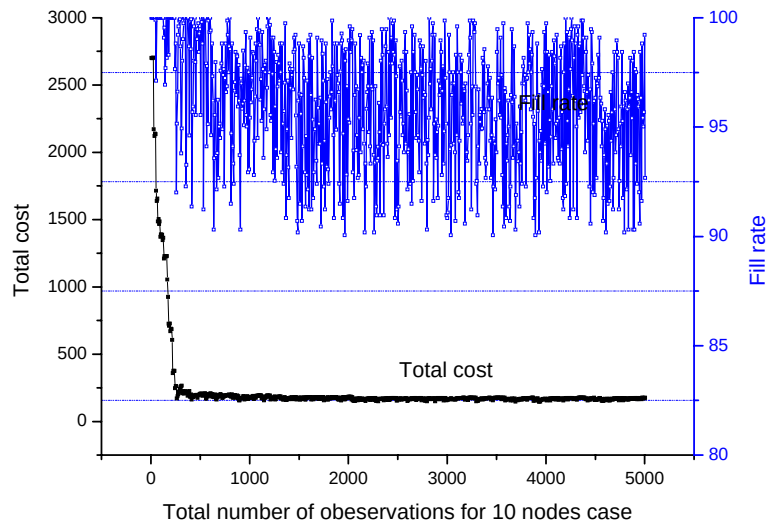


Figure 3.11: Une chaîne logistique à 10 niveaux et un espace de solutions borné avec contrainte de niveau de service

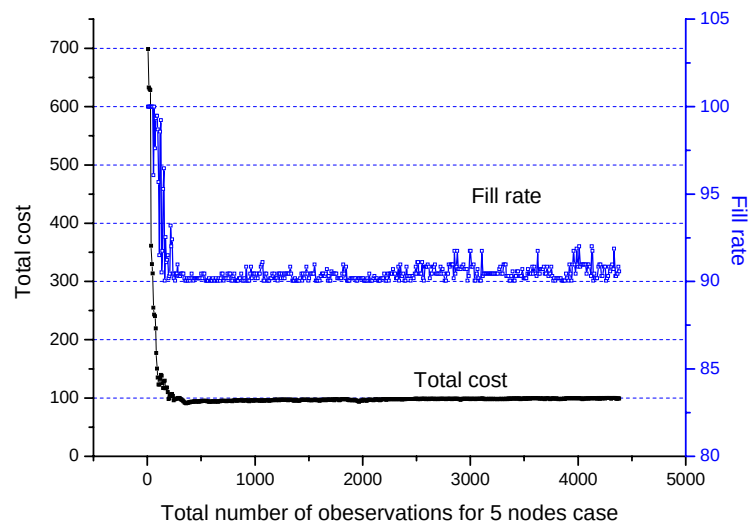


Figure 3.12. Une chaîne logistique à 5 niveaux et un espace de solutions non borné avec contrainte de niveau de service

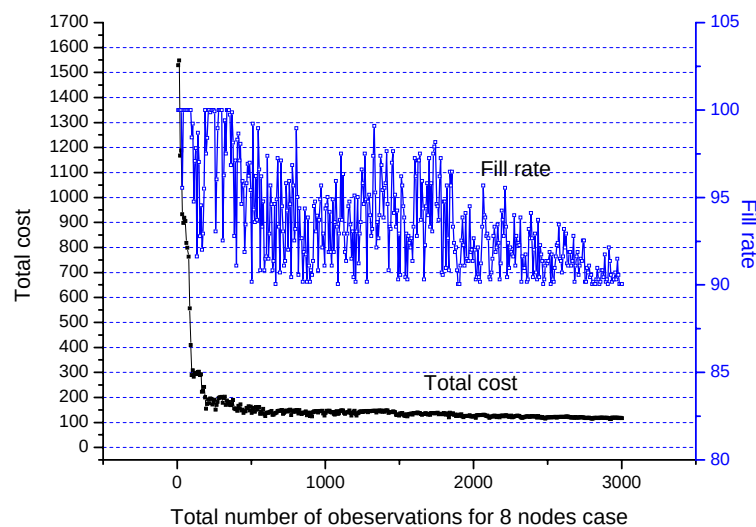


Figure 3.13. Une chaîne logistique à 8 niveaux et un espace de solutions non borné avec contrainte de niveau de service

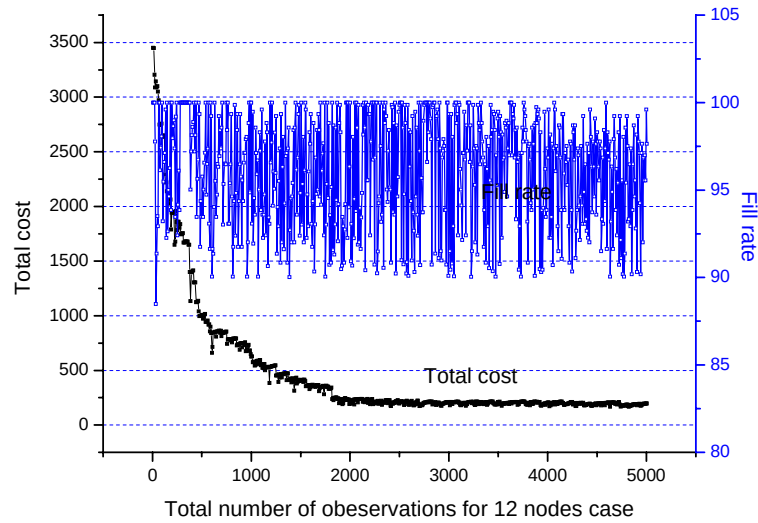


Figure 3.14. Une chaîne logistique à 12 niveaux et un espace de solutions non borné avec contrainte de niveau de service

3.4 Conclusion

Dans la section 2 nous avons présenté une nouvelle approche pour l'optimisation discrète par simulation à événements discrets. Cette approche est simple, robuste et permet de résoudre des problèmes d'optimisation où l'espace de solutions est borné ou non borné. Nous avons montré que l'algorithme COMPASS* converge vers un ensemble des optimums locaux avec probabilité 1 si certaines hypothèses acceptables sont satisfaites. Nous avons également montré que la méthode COMPASS* est plus efficace que la méthode COMPASS originale, proposé par (Hong et Nelson).

Biensûr, les deux algorithmes COMPASS et COMPASS* déterminent seulement des optimums locaux. Plusieurs possibilités d'amélioration existent. D'abord, on peut facilement inclure dans l'algorithme des propriétés connues des solutions optimales. Pour l'optimisation d'une chaîne logistique, on sait que généralement, le niveau de base-stock augmente lorsqu'on remonte une chaîne logistique sérielle, i.e. $R_i \geq R_{i+1} \forall i$. Cette propriété peut être incluse facilement dans le formalisme du problème

d'optimisation. Une deuxième amélioration consiste à introduire le principe d'exploration et exploitation utilisé dans le domaine de l'apprentissage renforcé. Cette idée consiste à combiner le choix des nouvelles solutions dans tout l'espace avec le choix dans la région la plus prometteuse.

Dans la section 3.3 nous introduisons le problème d'optimisation discrète par simulation avec des contraintes qui ne sont pas définies analytiquement. Nous proposons une nouvelle approche pour l'optimisation des variables de décision discrètes avec des contraintes évaluées par la simulation, appelé COMPASS* avec contraintes non analytiques. Cette méthode est simple, robuste, et permet de résoudre des problèmes d'optimisation où l'espace de solutions est borné ou non borné. Nous avons également montré que les algorithmes COMPASS avec contrainte non analytique (algorithmes 3 et 4) convergent vers un ensemble de optimums locaux avec probabilité 1.

Notre recherche future vise à développer des règles d'allocation du budget de simulation qui utilisent les résultats de simulation pour allouer automatiquement des observations supplémentaires à chaque solution afin de minimiser le budget total de simulation. Le principe d'exploration/exploitation est un autre domaine de recherche intéressant.

Chapitre 4

Optimisation par simulation des systèmes de production-distribution

Ce chapitre traite le problème d'optimisation des systèmes de production-distribution à plusieurs niveaux sous contrainte de taux de service. Le système est composé d'un ensemble d'entrepôts et des unités de production disposés en série. L'objectif est de trouver les valeurs des paramètres de contrôle de la politique de gestion de stock de chaque entrepôt afin de minimiser le coût total des stocks dans le système sous contrainte de taux de service client. Dans ce contexte nous analysons trois politiques de gestion de stock : la politique base-stock, la politique (R, nQ) et la politique (s, S) . Nous concluons le chapitre par une analyse des résultats numériques.

4.1. Modélisation des systèmes de production distribution

Un système de production-distribution est composé d'un ensemble de flux de matériaux, d'information, des services et des moyens pour le contrôle de ces flux. Parmi les activités spécifiques à un tel système on retrouve l'approvisionnement en matière première, la gestion des stocks, la production, l'entreposage, le transport et la distribution.

Les principaux composants d'un système de production-distribution sont :

- les fournisseurs de matière première, i.e. le début de la chaîne,
- les unités de production,
- les centres de distribution,
- le système de transport,
- les clients.

Dans ce chapitre nous considérons un système de production–distribution en série. Le système peut être décomposé en un ensemble de mailles en série. Chaque maille du système de production-distribution est composé par un entrepôt associé à une unité de production ou au client final et le système de transport amont. Le système de production-distribution considéré dans ce chapitre est caractérisé par : 1) capacité de production finie ; 2) demande aléatoire en quantité et en date d'arrivée ; 3) délai de transport constant ou aléatoire entre l'unité de production et l'entrepôt aval.

La topologie du système est illustré dans la figure 4.1. Un seul type de produits finis est fabriqué.

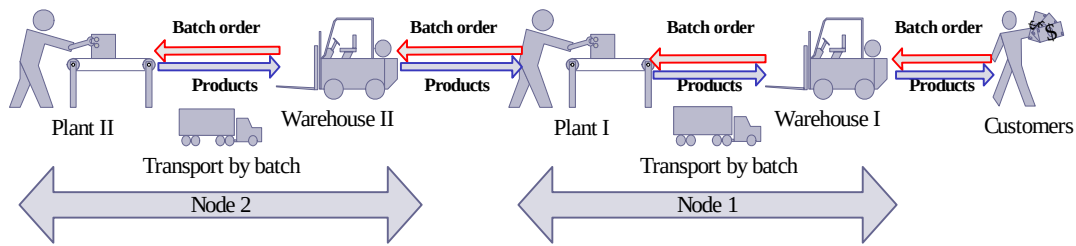


Figure 4.1. Topologie du système

Nous donnons ici les comportements dynamiques des principaux composants du système de production-distribution considéré dans ce chapitre.

4.1.1. Clients

La demande des clients est généralement représentée par une des distributions suivantes : normale, Poisson ou une distribution exponentielle négative. La distribution normale est plutôt utilisée pour modéliser la demande au niveau des unités de production. La loi de Poisson modélise surtout l'arrivée des clients de même que la distribution exponentielle négative. Lorsque la demande est faiblement stochastique, on la modélise généralement par une loi de Poisson. Dans nos travaux de recherche, l'arrivée des demandes des clients est représentée par une loi de compound Poisson avec des demandes arrivant selon la loi de Poisson et la quantité de chaque demande correspondant à une variable aléatoire.

4.1.2. Entrepôts

La principale fonction d'un entrepôt est la gestion de stock afin d'assurer la disponibilité des produits pour satisfaire la demande des clients ou pour production en aval. La gestion de stock s'effectue à travers des politiques d'approvisionnement en fonction de l'état du stock et de demandes en provenance des clients finaux ou des autres mailles. Il est bien évident que les politiques d'approvisionnement dans les différentes mailles ont un impact important sur la performance du système. L'objectif de ce chapitre est d'étudier la coordination de ces différentes politiques d'approvisionnement à l'aide de l'optimisation par simulation.

Le délai de réponse aux demandes des clients et les délais de réapprovisionnement sont aléatoires à cause de la possibilité de rupture de stock au niveau des entrepôts. Chaque entrepôt peut disposer d'une quantité de produits physiquement existante en stock.

Comme les demandes des clients arrivent seulement à l'entrepôt final, on en calcule de pénalités de rupture que pour celui-ci. Les entrepôts intermédiaires sont caractérisés par l'absence de pénalités de rupture. Cependant, une rupture de stock dans un entrepôt intermédiaire peut engendrer un délai supplémentaire d'approvisionnement de l'entrepôt final.

4.1.3. Unités de production

Les flux qui traversent les unités de production sont de deux catégories : 1) le flux de matériaux en provenance de l'entrepôt amont et le flux d'informations représenté par des ordres d'approvisionnement transmis par l'entrepôt aval. La capacité de production est limitée. Le délai de fabrication d'un lot de produits peut être stochastique. Les ordres de fabrication d'une unité de production sont traités selon une politique FIFO (premier arrivé, premier servi). Cette politique n'est pas forcément la plus économique, mais elle permet de simplifier l'analyse.

Les produits fabriqués par une unité de production pour un ordre d'approvisionnement attendent la fin de fabrication de tous les produits de cet ordre d'approvisionnement pour être transportés en ensemble vers l'entrepôt en aval. Le délai de transport entre l'unité de production et son entrepôt aval est considérée significative et elle ne peut pas être négligée. Il peut être constant ou stochastique. Par contre, le délai de transport entre un entrepôt et une unité de production est négligeable.

Dans le cadre de nos travaux de recherche, nous nous sommes appuyés sur l'outil logiciel OMNeT++ pour l'évaluation des performances des systèmes de production-distribution. Cet outil fournit un environnement pour la modélisation et la simulation des systèmes à événements discrets. Il est mis à disposition pour l'utilisation

académique par Omnest Global, Inc. Nous avons utilisé l'outil OMNeT++ afin de créer les composants génériques principaux d'un système de production-distribution.

Par la suite, nous appliquons l'approche d'optimisation par simulations que nous avons développée dans le chapitre 3 afin d'analyser les performances d'un système de production-distribution pour des différentes politiques de gestion du stock. Nous considérons notamment la politique base stock, (R, nQ) et (s, S) .

4.2. Systèmes de production-distribution multi-niveau avec politique base stock

4.2.1. Problème d'optimisation

Dans cette section, nous considérons l'optimisation d'un système de production-distribution multi-niveau présenté dans la section précédente, où le stock de chaque entrepôt est géré par une politique base stock. L'objectif est de déterminer le niveau de base stock R_i de chaque entrepôt i , afin de minimiser le coût global de gestion de stock tout en respectant la contrainte sur le taux de service des clients finaux.

La politique base stock pour l'entrepôt d'une maille est basée sur le concept de la position du stock qui est égale au stock disponible, moins la quantité en rupture des ordres d'approvisionnement ou des demandes des clients finaux, plus la quantité commandée à la maille en amont et pas encore reçue. L'objectif de cette politique de gestion de stock est de préserver la position du stock de chaque maille à une valeur constante appelée niveau base-stock R . Par conséquent, lorsqu'une demande de X produits arrive, l'entrepôt génère immédiatement un ordre d'approvisionnement de la même taille à la destination de la maille en amont.

Lorsqu'une unité de fabrication reçoit un ordre d'approvisionnement de taille X en provenance de l'entrepôt en aval, elle génère immédiatement un ordre de la même taille pour son entrepôt amont. Si l'entrepôt amont dispose d'une quantité suffisante de

produits, alors la quantité commandée est entièrement fabriquée. Sinon, la production est initiée avec le stock disponible dans l'entrepôt amont. La quantité restante est approvisionnée par des livraisons ultérieures. L'unité de production traite les ordres selon une politique FIFO. Si l'entrepôt amont est en rupture de stock, alors l'unité de production doit attendre que les pièces manquantes soient livrées.

La demande client arrive à l'entrepôt selon un processus compound Poisson avec un taux d'arrivée λ . La taille de chaque lot est aléatoire. Lorsqu'une commande client arrive, si l'entrepôt final dispose d'une quantité suffisante, on livre la quantité demandée. Sinon, on livre dans un premier temps la quantité disponible en stock. La quantité restante suite à la livraison de produits de l'unité de production vers l'entrepôt.

Le coût global de gestion du stock est constitué du coût de possession du stock dans chaque entrepôt, ainsi que du coût de rupture à l'entrepôt de produits finis. Le problème à résoudre est de déterminer le niveau base stock de chaque entrepôt, représenté par le vecteur $R = [R_1, R_2, \dots]^T$ pour minimiser le coût global de gestion du stock.

$$\text{Minimize } C(R) = \sum_{i=1,2,\dots} C_{hi} E\left[\int_0^T I_i(t) dt\right] + C_b E\left[\int_0^T B(t) dt\right] \quad (1)$$

où C_{hi} est le coût de possession unitaire de l'entrepôt i , C_b est le coût de rupture unitaire de l'entrepôt de produits finis, $I_i(t)$ est le stock physiquement disponible dans le $i^{\text{ème}}$ entrepôt, alors que $B(t)$ représente la quantité en rupture de l'entrepôt de produits finis. $C(R)$ dénote la somme de coûts moyens de possession et de rupture du stock dans l'horizon de temps $[0, T]$

Ce problème d'optimisation est complexe et il est difficile à résoudre à l'aide d'une approche analytique. Nous avons montré dans le chapitre 3 l'intérêt de l'utilisation des techniques d'optimisation par la simulation pour ce type de problèmes. Par conséquent, dans ce chapitre nous utilisons l'algorithme COMPASS* que nous avons développé dans la section 3.

4.2.2. Optimisation des politiques d'approvisionnement

Les solutions initiales de l'algorithme COMPASS* sont générées aléatoirement. Chaque solution est évaluée par un nombre de réplifications de simulation i.i.d. fixé par la politique SAR (Simulation Allocation Rule) choisie. Cette politique doit satisfaire les hypothèses énoncés pour la convergence de l'algorithme. Un exemple de politique SAR établie pour l'algorithme COMPASS* dans le cas où l'espace de solutions est fermé et borné est la suivante :

$$a_k(x) = \begin{cases} 1 & \text{if } x \in V_k \setminus V_{k-1} \text{ or } x \in NH(x_{k-1}^*) \\ 0 & \text{otherwise} \end{cases}$$

où $a_k(x)$, est le nombre de réplifications de simulation alloués à la solution x lors de l'itération k .

L'expérimentation numérique est conduite avec le paramétrage suivant. La demande respecte une loi de distribution compound Poisson avec un taux d'arrivée $\lambda \in \{1.0, 1.3, 1.5\}$ et la taille de chaque commande est de distribution uniforme dans l'intervalle $[3, 9]$. Le temps de fabrication d'une unité de produit dans chaque unité de production est une variable aléatoire de distribution exponentielle de moyenne 0,1 unité de temps. L'intensité du trafic ρ' est respectivement de 0.6, 0.78 et 0.9. Le délai de transport d'un lot entre une unité de production et l'entrepôt aval est de 3.0 unités de temps. Le coût unitaire de possession est de $C_h = 1$ et le coût unitaire de rupture est de $C_b = 5$. L'horizon de temps considéré pour la simulation est de $T = 1000$ unités de temps. Les valeurs des paramètres sont reprises dans le tableau 4.1.

Tableau 4.1 Paramètres de simulation (unité de temps: seconde)

	Taux d'arrivée client	Taille d'une commande	Temps d'exécution	Délai de transp.	Intensité du trafic ρ'	C_h	C_b
Moyenne	1 ~ 1.5	6	0.1	3.0	0.6~0.9	1.0	5.0
Ecart type		2	0.1	0			

Plusieurs expérimentations numériques ont été effectuées pour de configurations de systèmes de production-distribution avec un nombre de niveaux de 5, 8, 10 et 12. De même, différents budgets de simulation ont été considérés (5000 pour 5, 8 et 10 niveaux

et 10000 pour 12 niveaux). Afin de se retrouver dans le cas d'un espace de solutions borné, nous considérons que le niveau de base-stock $R_i \in [0, 2000]$.

Un aspect très important est le choix de la longueur et du nombre de répliques de simulation à faire pour chaque solution, afin d'obtenir une bonne précision avec un minimum d'effort. D'un côté, si la longueur de la simulation est trop importante, on va perdre trop de temps avec la simulation des mauvaises solutions. D'un autre côté, si la durée de la simulation est trop courte, la précision du résultat n'est pas suffisamment bonne. Suite à un nombre important d'expérimentations numériques, nous avons pu conclure qu'une durée de 1000 unités de temps est suffisante pour chaque réplique de simulation.

Les résultats numériques sont illustrés dans les figures 4.2 à 4.5.

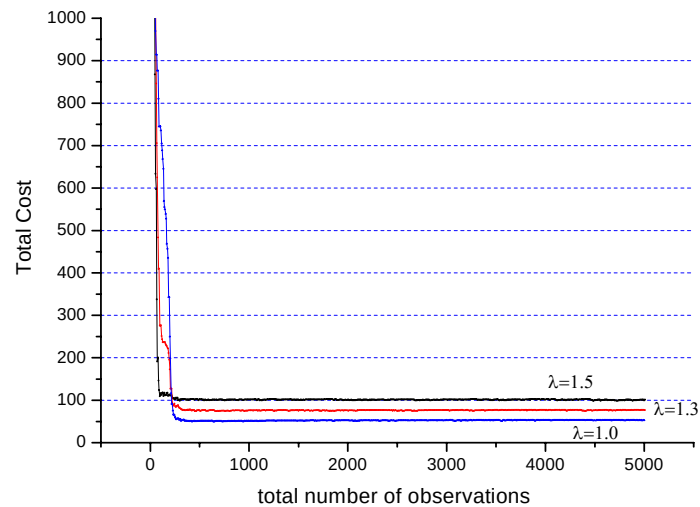


Figure 4.2: Optimisation d'un système de production à 5 niveaux

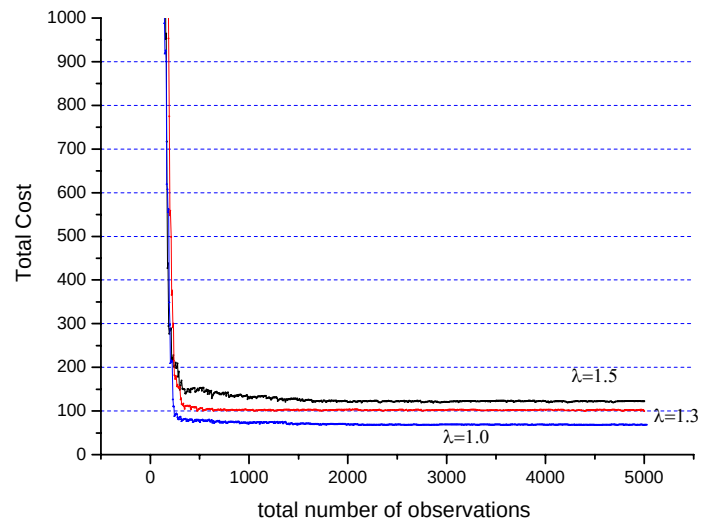


Figure 4.3: Optimisation d'un système de production-distribution à 8 niveaux

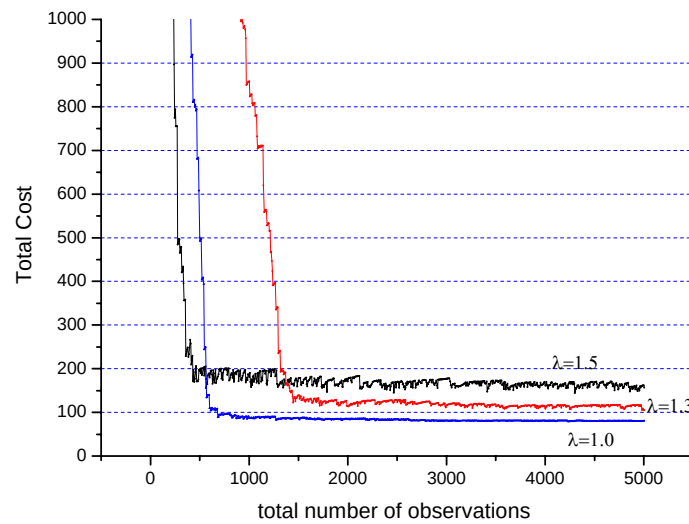


Figure 4.4: Optimisation d'un système de production-distribution à 10 niveaux

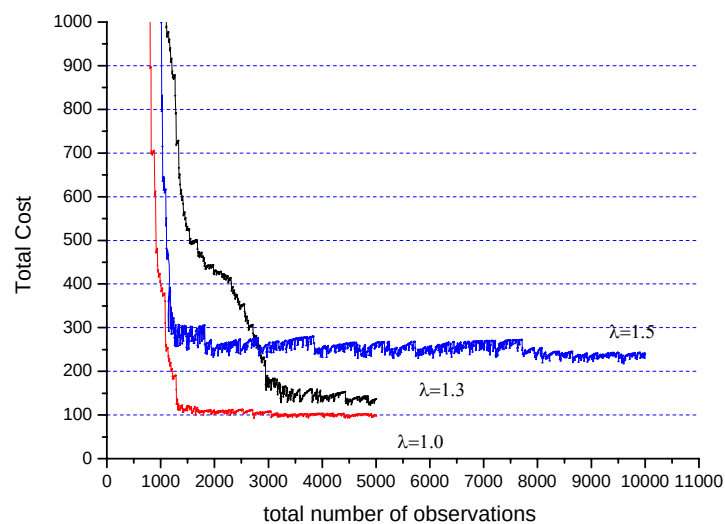


Figure 4.5: Optimisation d'un système de production-distribution à 12 niveaux

4.2.3. Analyse de l'influence du taux moyen d'arrivée

En analysant les graphiques, on peut remarquer que pour le cas d'un système de production avec 5, 8 ou 10 niveaux, le niveau base-stock de chaque entrepôt ainsi que le coût total de gestion de stock calculé convergent d'une manière lisse et rapide. Les neuf dernières solutions localement optimales sont illustrées dans le tableau 4.2. On peut remarquer que le coût total de gestion de stock à chaque itération ainsi que les niveaux de base-stock associés sont stables et raisonnables.

Table 4.2 Solutions calculées lors des 9 dernières itérations pour un système de production-distribution avec 5 niveaux $\lambda=1.0$.

R5	R4	R3	R2	R1	Coût optimal	Nombre de réplifications de chaque solution	Nombre total de réplifications
0	0	0	25	144	59.07	635	4945
0	0	0	25	144	59.07	636	4952
0	0	0	25	144	59.07	637	4959
0	0	0	25	144	59.07	638	4966
0	0	0	25	144	59.07	639	4973
0	0	0	25	144	59.07	640	4980
0	0	0	25	144	59.07	641	4987
0	0	0	25	144	59.07	642	4994
0	0	0	25	144	59.07	643	5001

Cependant, lorsque le nombre de niveaux augmente (par exemple 12) en même temps que le taux d'arrivée ($\lambda=1.5$), le caractère lisse de l'évolution du coût total calculé et sa vitesse de convergence sont diminuées. L'analyse de 9 dernières solutions, dans le tableau 4.3. montre que le coût total calculé converge vers une valeur fixe, alors que la valeur des niveaux de base-stock ne se stabilise même après 10000 observations.

Tableau 4.3 Solutions calculées lors des 9 dernières itérations pour un système de production-distribution avec 12 niveaux $\lambda=1.5$.

R12	R11	R10	R9	8r	R7	R6	R5	R4	R3	R2	R1	Coût optimal
14	11	2	6	21	11	2	57	11	306	97	103	240.376
4	10	1	2	30	7	1	57	10	306	103	103	231.883
5	3	1	10	18	3	5	59	12	309	109	99	236.827
1	5	1	6	16	3	1	69	10	316	115	94	237.749
1	6	22	1	12	3	2	70	21	301	105	107	240.59
1	3	25	11	6	4	12	64	15	297	104	106	240.971
0	5	33	10	1	4	3	71	9	296	100	98	230.825
1	10	28	1	1	6	0	79	8	300	92	91	238.204
0	11	2	1	41	5	1	69	7	310	102	98	241.271

Le nombre important de niveaux ainsi que le taux d'arrivée important engendrent une augmentation du caractère aléatoire de l'évolution du système ainsi que la densité du trafic réelle. Il est bien connu que l'augmentation de la densité du trafic entraîne l'augmentation du budget de simulation nécessaire afin d'obtenir une solution stable.

Par la suite, nous analysons la relation entre le taux moyen d'arrivée de la demande et le coût total de gestion du stock. Les solutions optimales obtenues pour des systèmes de production distribution avec 5, 8, 10 et 12 niveaux ainsi que pour différents taux moyens d'arrivée de la demande sont illustrée dans le tableau 4.4.

Table 4.4 . Solutions pour différents niveaux et différents taux moyens d'arrivée de la demande

	5			8			10			12		
λ	1.0	1.3	1.5	1.0	1.3	1.5	1.0	1.3	1.5	1.0	1.3	1.5
R1	150	190	185	125	72	122	47	67	94	43	74	98
R2	0	1	63	9	125	200	134	301	148	36	61	102
R3	3	1	0	45	9	0	72	9	144	194	272	310
R4	0	0	0	2	9	1	9	3	13	3	13	7
R5	0	18	32	0	9	0	0	0	12	16	25	69
R6				56	93	45	0	0	5	9	7	1
R7				0	0	0	31	0	78	1	3	5
R8				0	0	25	0	0	3	0	1	41
9							0	1	0	0	9	1
R10							0	4	1	33	2	2
R11										8	7	11
R12										0	1	0
Cost	52.67	76.89	101.1	68.93	99.72	122.2	80.48	105.3	156.8	97.58	136.1	241.3

Considérons d'abord la solution optimale R . Les résultats obtenus montrent que le niveau de base-stock est de plus en plus faible lorsqu'on parcourt la processus de fabrication à l'inverse, du produit fini, vers la matière première.

La relation entre le taux moyen d'arrivée et le coût total est illustrée graphiquement dans la figure 4.6. On peut conclure que le coût total de la gestion du stock augmente avec le taux moyen d'arrivée des clients. Cette remarque est raisonnable car le système doit augmenter le stock disponible afin de satisfaire une demande croissante.

Une constatation intéressante est que la pente de la courbe représentant l'influence du taux d'arrivée sur le coût de gestion du stock augmente avec la taille du système. Ceci montre que le coût total de gestion du stock est d'autant plus sensible au taux d'arrivée des clients lorsque la taille du système augmente.

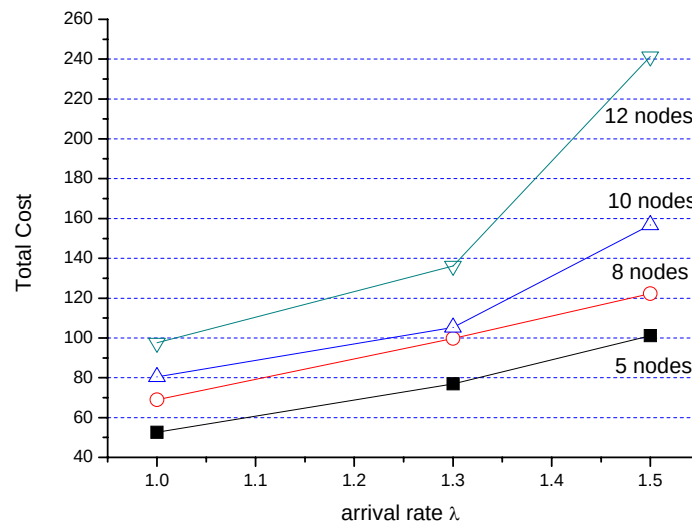


Figure 4.6. Analyse de l'influence du taux moyen d'arrivée

4.2.4. Analyse de l'influence de la taille des commandes

Nous rappelons que la demande arriv  e    l'entrep  t final selon une loi de compound Poisson avec un taux d'arriv  e λ . La taille de chaque commande est une variable al  atoire X et les quantit  s des diff  rentes commandes sont i.i.d. par la suite nous analysons l'effet de la taille des commandes client sur le co  t de gestion du stock au niveau du syst  me de production-distribution. Les jeux de param  tres que nous utilisons pour les exp  rimentations num  riques sont synth  tis  es dans le tableau 4.5.

Tableau 4.5 Paramètres de simulation

	Taux d'arrivée	Taille de chaque commande			Temps de service	Temps de transp.	Intensité du trafic ρ'	C_h	C_b
Moyenne	1.3	6			0.1	3.0	0.78	1.0	5.0
Ecart type		3.2	2.0	0.82	0.1	0			

Plus particulièrement, nous avons fixé le taux d'arrivée $\lambda=1.3$ et la densité du trafic $\rho=0.75$. La taille d'une commande est distribuée uniformément avec la même valeur moyenne mais des écarts types différentes. Nous analysons trois cas, selon que la taille de la commande est comprise dans l'intervalle [1, 11], [3, 9], et [5, 7]. Les résultats sont présentés dans le tableau 4.6.

Tableau 4.6 . Solution pour le cas de 5, 8, 10 et 12 noeuds et taille de commandes différentes

	5			8			10			12		
X	[1,11]	[3,9]	[5,7]	[1,11]	[3,9]	[5,7]	[1,11]	[3,9]	[5,7]	[1,11]	[3,9]	[5,7]
R1	136	190	201	100	72	188	113	67	64	51	74	74
R2	98	1	0	133	125	8	213	301	266	54	61	21
R3	1	1	1	11	9	7	29	9	44	689	272	324
R4	0	0	1	35	9	43	1	3	0	7	13	4
R5	1	18	0	0	9	1	0	0	0	6	25	20
R6				69	93	61	1	0	0	6	7	1
R7				1	0	0	30	0	0	3	3	2
R8				11	0	0	0	0	0	3	1	1
9							53	1	1	0	9	1
R10							0	4	1	2	2	0
R11										0	7	5
R12										0	1	0
Cost	91.97	76.89	73.21	120.4	99.72	87.10	138.9	105.3	98.54	469.2	136.1	125.0

La relation entre l'écart type de la taille des commandes et le coût total de la gestion du stock est présenté dans la figure 4.7.

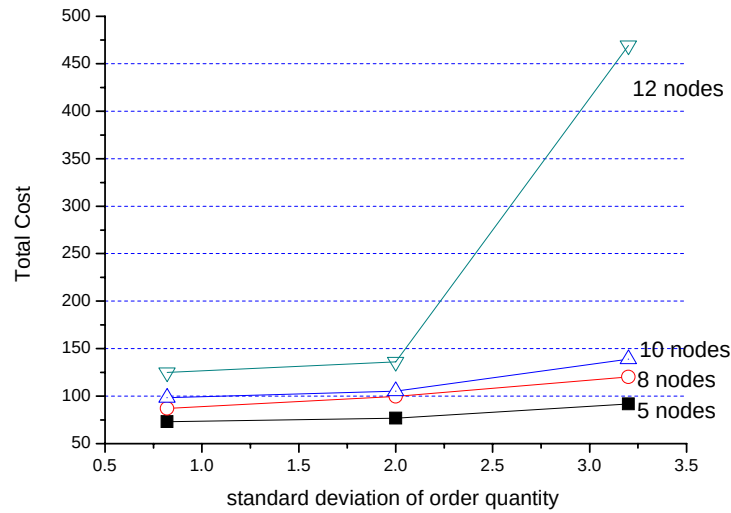


Figure 4.7 L'influence de l'écart type de la taille des commandes sur le coût de stockage

On remarque que le coût global de gestion du stock augmente avec l'écart type de la taille des commandes. De la même manière que pour l'influence du taux d'arrivée des clients, on remarque que la sensibilité du coût total de gestion du stock par rapport à l'écart type de la taille des commandes augmente avec la taille du système.

4.2.5. Prise en compte du taux de service client.

Dans ce cas, on prend en compte seulement le coût de possession du stock au niveau de chaque entrepôt. Les pénalités de rupture sont implicitement prises en compte par le taux de service client imposé. Le problème consiste à choisir le niveau base-stock de chaque entrepôt, représenté par le vecteur $R = [R_1, R_2, R_3, \dots]^T$ afin de minimiser le coût total de gestion du stock tout en garantissant un taux de service client supérieur à α , $0 < \alpha < 1$:

$$\text{Minimize } C(R) = \sum_{i=1,2,\dots} C_{hi} E\left[\int_0^T I_i(t) dt\right]$$

$$\text{t.q.} \quad P(R) > \alpha.$$

où C_{hi} est le coût unitaire de possession et $I_i(t)$ est le stock physiquement disponible de l'entrepôt i . $C(R)$ est le coût total de gestion du stock pour l'horizon de temps $[0, T]$. $P(R)$ dénote le taux de service client pour le même horizon.

Afin de résoudre ce problème, nous utilisons la méthode COMPASS*, algorithme 3, proposée dans le chapitre 3. Les jeux de paramètres utilisés pour les expérimentations numériques sont donnés dans le tableau 4.7. Plus spécialement, nous considérons un taux moyen d'arrivée des clients $\lambda=1.3$ et l'intensité du trafic $\rho=0.75$. Afin d'analyser l'influence du taux de service client sur le coût de gestion du stock, nous considérons les valeurs suivantes pour $P(R)$: 0.85, 0.90, 0.95 et 0.98.

Tableau 4.7 Jeux de paramètres (unités de temps: la seconde)

	Taux d'arrivée	Taille commande	Taux de service usine	Temps transp.	Intensité du trafic ρ'	Taux de service client $P(R)$					C_h
Moyenne	1.3	6	0.1	3.0	0.75	0.85	0.9	0.95	0.98	1.0	
Ecart type		2.0	0.1	0							

Tableau 4.8 Solutions pour 5, 8, 10 et 12 niveaux et différents Taux de service client

	5				8				10				12			
$P(R)$	85	90	95	98	85	90	95	98	85	90	95	98	85	90	95	98
$R1$	194	104	174	151	112	197	120	247	97	79	107	102	71	88	100	95
$R2$	17	126	61	81	142	96	166	53	137	226	139	117	69	117	64	85
$R3$	0	0	13	1	9	2	11	0	106	97	46	94	106	124	58	247
$R4$	0	1	1	2	0	0	0	1	1	3	10	64	2	6	12	34
$R5$	12	0	0	31	48	0	9	0	3	1	0	3	206	138	179	33
$R6$					0	1	45	61	2	0	0	0	3	0	22	5
$R7$					1	35	0	0	48	1	109	69	1	4	19	17
$R8$					17	0	0	0	2	0	0	0	1	6	7	7
9									4	1	22	2	7	1	36	7
$R10$									6	1	4	0	2	0	5	9
$R11$													12	4	20	28
$R12$													0	0	3	12
$C(R)$	69.4	76.7	92.8	109	87.0	89.7	108	119	105	110	134	151	121	128	160	215
Taux serv.	85.5	90.0	95.3	98.2	85.8	90.4	95.7	98.9	85.3	90.7	95.2	98.1	85.1	90.8	95.3	98.9

L'influence du taux de service client souhaité sur le coût de gestion du stock est illustré dans la graphique de la figure 4.8.

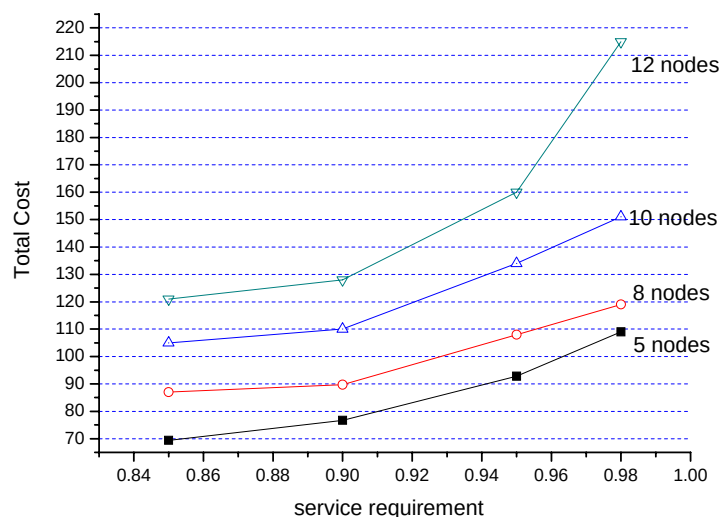


fig 4.8 Niveau de service et le coût de stockage

On remarque à nouveau que la sensibilité du coût total de gestion de stock par rapport au taux de service client augmente avec la taille du système.

4.3. Systèmes de production-distribution multi-niveau avec politique (R, nQ)

Dans cette section nous analysons un système de production-distribution avec la même configuration que dans la section 4.2., sauf que dans ce cas le stock est géré selon une politique (R, nQ) .

4.3.1. Politique de gestion du stock (R, nQ)

La politique (R, nQ) est une politique à révision continue. De même que la politique de base-stock, elle est basée sur le concept de la position de stock. Son objectif est de garantir que la position du stock à un entrepôt i n'est jamais inférieure au niveau de base-stock R_i et les ordres d'approvisionnement correspondent toujours à une quantité

multiple de Q_i . Par conséquent, lorsqu'un client arrive et commande une quantité X , si la position du stock devient inférieur au niveau de base-stock R_i , alors l'entrepôt génère un ordre d'approvisionnement pour une quantité nQ_i telle que la position du stock retrouve une valeur comprise entre R_i et $R_i + Q_i - 1$.

Afin d'optimiser les performances du système de production-distribution, nous prenons en compte le coût unitaire de possession au niveau de chaque entrepôt ainsi que le taux de service client. Le problème à résoudre est de définir les niveaux de base-stock de chaque entrepôt représenté par le vecteur $R = [R_1, R_2, R_3, \dots]^T$ ainsi la taille du lot associée $Q = [Q_1, Q_2, Q_3, \dots]^T$ afin de minimiser le coût total de gestion du stock tout en respectant le taux de service client souhaité α avec $0 < \alpha < 1$.

$$\begin{aligned} \text{Minimize } C(R, Q) &= \sum_{i=1,2,\dots} C_{hi} E\left[\int_0^T I_i(t) dt\right] \\ \text{t.q.} \quad P(R, Q) &> \alpha \end{aligned}$$

où C_{hi} est le coût unitaire de possession du stock, $I_i(t)$ est le stock physiquement disponible au niveau de l'entrepôt i . $C(R, Q)$ est le coût moyen de possession du stock calculé pour l'horizon de temps $[0, T]$ et $P(R, Q)$ le taux de service client pour le même horizon. Les performances $C(R, Q)$ et $P(R, Q)$ peuvent être déterminées facilement par la simulation. Ainsi, on peut appliquer l'algorithme COMPASS*.

Les expérimentations numériques sont effectuées avec le même jeu de paramètres que dans la section 4.2. La taille d'une commande, dénotée X , est un nombre entier, uniforme distribué dans l'intervalle $[3, 9]$. Le temps nécessaire à la fabrication d'une unité de produit τ est exponentiellement distribué avec une moyenne de 0.1. la durée de transport entre une unité de production et l'entrepôt en aval est de 3.0 unités de temps. Le coût unitaire de possession du stock est égal à 1.0 et le taux de service client doit être supérieur à 90%. Le niveau de base-stock et la taille du lot de chaque entrepôt i sont cherchées dans un domaine large : $R_i \in [0, 10000]$ et $Q_i \in [0, 3000]$. L'horizon de temps est $T=5000$ secondes.

Pour résoudre ce problème, nous utilisons la méthode COMPASS*, algorithme 3, que

nous avons présenté dans le chapitre 3. La solution initiale x_0 est générée aléatoirement. Le budget de simulation est alloué selon la SAR suivante :

$$a_k(x) = \begin{cases} 1 & \text{if } x \in V_k \setminus V_{k-1} \text{ or } x \in NH(x_{k-1}^*) \\ 0 & \text{otherwise} \end{cases}$$

A chaque itération, on choisit aléatoirement cinq nouvelles solutions à l'intérieur de la région la plus prometteuse.

Dans un premier temps, nous considérons le cas d'un système de production-distribution à deux niveaux, composé d'un entrepôt approvisionné par une unité de production. Le budget de simulation est de 5000 observations. Comme la méthode COMPASS* fournit une solution localement optimale, on applique la méthode plusieurs fois et on va retenir la meilleure solution. Les valeurs optimales calculés pour le niveau de base-stock R^* et la taille de lot Q^* pour différents taux d'arrivée des clients λ sont donnés dans le tableau 1. Les colonnes Obs. et Temps CPU définissent respectivement le nombre moyen de réplifications de simulation et le temps CPU nécessaire pour obtenir la convergence de l'algorithme vers la solution optimale. On peut remarquer que la taille de lot Q est toujours égale à 1. Ceci est due au fait que le coût du transport n'est pas pris en compte. Dans ce cas, chaque fois qu'un client passe une commande, l'entrepôt de produits finis génère un ordre d'approvisionnement d'une même taille que la commande client [LI, SAVA et XIE,2005]. La convergence des variables de décision R et Q est illustrée dans la figure 4.9.

Tableau 4.9 Optimisation d'un système de production distribution à deux niveaux

Taux moyen d'arrivée	ρ	Taux de service	R^*	Q^*	Coût optimal	Obs	Temps CPU
1.0	0.60	0.9016	49	1	25.0670	100	45
1.1	0.66	0.9053	55	1	27.4983	200	130
1.3	0.78	0.9013	71	1	33.9590	500	400

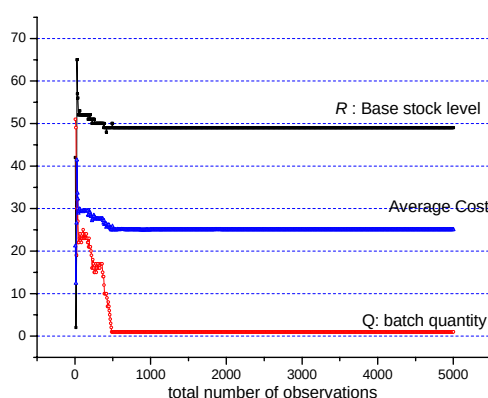


Figure 4.9. Convergence des variables R et Q pour $\lambda=1.0$

Par la suite, nous utilisons la méthode COMPASS*, pour paramétrer la politique de gestion de stock (R, nQ) afin de minimiser le coût total de gestion du stock pour un système de production distribution à 5, 8 et 10 niveaux. Le budget de simulation est limité à 5000 réplifications de simulation. Nous négligeons les coûts de transport, donc la taille du lot, Q , est toujours égale à 1. Les résultats sont illustrés dans la figure 4.10.

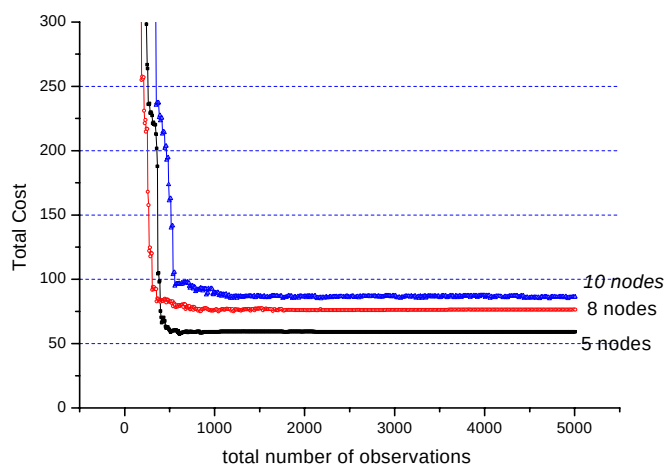


Figure 4.10. Optimisation d'un système de production distribution multi-niveaux

Lorsque le système de production-distribution étudié a un nombre de niveaux inférieur à 5, la méthode COMPASS* converge en moins de 800 réplifications de simulation. La solution obtenue pour le cas d'un système de production distribution à 5 niveaux est détaillée dans le tableau 4.10. Les niveaux de base-stock sont donnés à partir du

l'entrepôt de produits finis, en remontant vers l'entrepôt de matière première.

Tableau 4.10. Solution pour un système de production distribution à 5 niveaux

Nombre niveaux	R^*	Coût optimal	Taux de service	Obs	Temps CPU
5	114				
	55				
	0	59.116	0.9030	800	1100
	0				
	0				

La méthode COMPASS* garantit la convergence vers une solution optimum locale. Par conséquent, le choix de la solution initiale joue sur la performance de la solution proposée. Le niveau de base-stock, le coût total de gestion du stock et le taux de service client pour un système de production distribution à 5, 8 et 10 niveaux sont respectivement illustrés dans le tableau 4.11. On remarque, que le niveau de base-stock décroît en partant de l'entrepôt de produits finis, vers les entrepôts en amont. En effet, l'entrepôt de produits finis a une influence très importante sur le taux de service fournit au client. Cette remarque nous a poussé à utiliser la contrainte suivante dans le modèle d'optimisation :

$$R_i \geq R_{i+1} \text{ for } \forall i=1,2,\dots,$$

A travers nombreux expérimentations numériques, nous avons pu constater que cette contrainte permet d'améliorer la solution trouvée.

Tableau 4.11. *Solution Localement Optimale* pour un système de production distribution à 5, 8 et 10 niveaux

	5	8	10	5_{con}	8_{con}	10_{con}
R1	114	154	140	144	243	156
R2	55	47	65	25	0	69
R3	0	8	9	0	0	36
R4	0	0	54	0	0	34
R5	0	8	0	0	0	8
R6		0	0		0	1
R7		35	0		0	0
R8		0	8		0	0
R9			30			0
R10			0			0
Cost	59.116	76.442	86.422	59.072	70.731	84.746
FR	0.9030	0.9034	0.9016	0.9035	0.9015	0.9018

4.3.2. Prise en compte du coût de transport

Considérons à nouveau le système de production distribution décrit dans la section 4.2. Cette fois-ci, on suppose que le coût de transport ne peut plus être négligé. De même, la capacité de transport d'un équipement T_i entre l'unité de production i et l'entrepôt aval (l'entrepôt i) est limitée. Soit CT_i le coût de transport entre l'unité de production i et l'entrepôt aval. Il est défini de la manière suivante :

$$CT_i^1 = C_{Ti} \left[\frac{Q_i}{T_i} \right]^+$$

où C_{Ti} est en coût unitaire de transport, $[x]^+$ dénote le nombre entier le plus petit supérieur à x . Par exemple, si la capacité de l'équipement de transport est $T_i = 10$ et la taille du lot est $Q_i = 15$, alors le coût de transport est de $2 C_{Ti}$.

Dans ce cas, le coût total de gestion du stock est composé du coût de transport et du coût de possession du stock. Le problème à résoudre consiste à définir le niveau de base-stock $R = [R_1, R_2, R_3, \dots]^T$ ainsi que la taille de lot $Q = [Q_1, Q_2, Q_3, \dots]^T$ au

niveau de chaque entrepôt afin de minimiser le coût total de gestion de stock tout en respectant la contrainte sur le taux de service client α où $0 < \alpha < 1$.

$$\begin{aligned} \text{Minimize } C(R, Q) &= \sum_{i=1,2,\dots} (C_{hi} E[\int_0^T I_i(t) dt] + C_{Ti} E[\int_0^T CT_i(t) dt]) \\ \text{tel que } P(R, Q) &> \alpha \end{aligned}$$

où C_{hi} est le coût unitaire de possession, au niveau du entrepôt i , C_{Ti} est le coût unitaire de transport de l'unité de production i à l'entrepôt en aval, $I_i(t)$ le niveau du stock à l'entrepôt i , $C(R, Q)$ est le coût global de gestion du stock durant l'horizon de temps $[0, T]$ et $P(R, Q)$ est le taux de service client pour le même horizon de temps.

Dans ce cas on doit faire un bon compromis entre la quantité transportée en une fois, qui a une incidence directe sur le coût du transport et le coût total de gestion du stock. Plus la quantité transportée en une fois est importante, plus le coût de transport est faible. En échange, le délai d'approvisionnement augmente. Par conséquent, il faut augmenter le niveau du stock afin de préserver le même taux de service client. Ceci entraîne une augmentation du coût de possession et donc du coût total de gestion de stock. Au contraire, si on transporte par de petites quantités, le délai d'approvisionnement de l'entrepôt est réduit, mais le coût de transport augmente.

Une première idée consiste à choisir la taille de lot Q_i au niveau de chaque entrepôt égale à la capacité du moyen de transport utilisé T_i . Nous testons cette idée à travers des expérimentations numériques.

Soit la capacité de transport $T_i = 7$, le coût unitaire de transport $C_{Ti} = 5$ et le taux d'arrivée des clients $\lambda = 1.3$. Nous analysons deux cas : 1) $Q_i = T_i = 7$ et 2) $Q_i \in [1, 10]$. Le détail des paramètres est illustré dans le tableau 4.12. Nous effectuons des expérimentations numériques pour un système de production distribution 5, 8 et 10 niveaux. Les résultats obtenus sont donnés dans le tableau 4.13, où R dénote le vecteur de niveaux de base-stock et FR est le taux de service client.

Tableau 4.12 Paramètres de simulation (unité de temps: seconde)

	Taux d'arrivée clients	Taille de lot	Taux de service usine	Délai de transp.	Intensité du trafic ρ'	T_i	C_{Ti}	C_h
Moyenne	1.3	6	0.1	3.0	0.78	7	5.0	1.0
Ecart type		2.0	0.1	0				

Considérons d'abord les résultats obtenus pour le système de production-distribution avec 5 nœuds. On remarque que le niveaux de base-stock calculés sont très proches, voir identiques, pour les deux cas (Q_i fixé et $Q_i \in [1, 10]$). Cependant, le coût total de gestion du stock est plus faible lorsque la taille de lot $Q_i \in [1, 10]$. Cette propriété reste valable lorsque le nombre de niveaux augmente. Par la suite, nous analysons l'influence entre les niveaux de base-stock R et la taille de lot Q à travers deux jeux d'expérimentations.

Tableau 4.13. Solutions optimums local pour 5, 8 et 10 nœuds

	5(fixé)	8(fixé)	10(fixé)	5(borné)		8(borné)		10(borné)	
	R	R	R	R	Q	R	Q	R	Q
1	266	381	185	266	10	292	10	187	10
2	0	6	140	0	10	50	4	181	1
3	0	2	91	0	3	18	2	62	1
4	0	1	70	0	1	1	4	8	4
5	0	1	2	0	2	1	2	0	5
6		1	1			0	10	0	2
7		1	1			0	2	0	2
8		0	0			0	1	0	2
9			0					0	10
10			0					0	1
Coût	139.2	198.7	243.6	114.8		155.3		186.5	
FR	90.01	90.20	90.68	91.72		92.40		90.22	

Dans le premier cas, nous considérons seulement le coût total de possession. Le coût de

transport est négligé. Dans le deuxième cas on considère le coût de possession et le coût de transport. Nous conservons les paramètres de simulation donnés dans le tableau 4.12. Les résultats obtenus sont donnés dans le tableau 4.14. On peut remarquer que le niveaux de base-stock ont une même allure. Prenons par exemple le cas où le nombre de nœuds est égal à 5. Dans le cas où on néglige le coût de transport, le vecteur du niveau de base-stock est [167, 61, 1, 0, 0], alors que si on prend en compte le coût de transport, il est [175, 74, 1, 0, 0]. On peut conclure que les changements au niveau du système de transport ont peu d'influence sur le niveau de base-stock. Cette propriété peut être utilisée pour améliorer la convergence de l'algorithme COMPASS* que nous avons développé dans la chapitre 3, dans le cas d'un système de production distribution avec une politique (R, nQ) et prise en compte du coût de transport.

Table au 4.14. Solution localement optimale pour 5, 8 et 10 noeuds

	1		2		5		8		1		2		5		8	
	Sans coût de		Sans coût		Sans coût		Sans coût de		Avec coût		Avec coût		Avec coût		Avec coût de	
	transport		de transp.		de transp.		transp		de transp		de transp		de transp		transp	
	R	Q	R	Q	R	Q	R	Q	R	Q	R	Q	R	Q	R	Q
1	70	1	114	1	167	1	128	1	70	7	93	6	175	6	183	10
2			0	1	61	1	67	1			16	4	74	6	49	1
3					1	1	63	1					1	1	49	1
4					0	1	62	1					0	2	29	1
5					0	1	10	1					0	1	10	1
6							8	1							4	10
7							0	1							0	2
8							0	1							0	1
Cost	34.2		49.7		78.8		102.1		39.7		59.8		116.3		147.0	
FR	90.57		90.01		90.12		90.43		90.06		90.65		90.43		91.00	

4.4 Optimisation d'un système de production-distribution avec politique (s, S)

Dans les sections 4.2. et 4.3. nous avons utilisé la méthode COMPASS* pour optimiser le coût total de gestion du stock d'un système de production-distribution avec les politique base-stock et (R, nQ) . Par la suite, nous abordons le même problème pour un système de production-distribution avec politique (s, S).

4.4.1 Politique de gestion du stock (s, S)

La politique de gestion du stock (s, S) garantit que la position du stock IP_i de chaque entrepôt i ne diminue au-dessous du point de commande s_i . Lorsque le point de commande est atteint, un ordre d'approvisionnement d'une quantité $S_i - IP_i$ est lancé. Par conséquent, chaque fois qu'une commande client de taille X arrive, si la position du stock diminue au-dessous le point de commande s_i , un ordre d'approvisionnement pour une quantité $S_i - IP_i$ est lancé. En effet, le cas spécial de cette politique (S-1, S) est identique à la politique base-stock.

4.4.2. Optimisation sans prise en compte du coût de transport

Dans cette section on considère seulement le coût de possession du stock. Le problème à résoudre consiste à déterminer le point de commande de chaque entrepôt $s = [s_1, s_2, s_3, \dots]^T$ et la quantité à commander $S = [S_1, S_2, S_3, \dots]^T$ afin de minimiser le coût global de gestion de stock tout en respectant un taux de service client supérieur à α , $0 < \alpha < 1$.

$$\text{Minimize } C(s, S) = \sum_{i=1,2,\dots} C_{hi} E\left[\int_0^T I_i(t) dt\right]$$

$$\text{t.q.} \quad P(s, S) > \alpha$$

où C_{hi} est le coût unitaire de possession et $I_i(t)$ est le stock physiquement disponible à l'entrepôt i . $C(s, S)$ est le coût total moyen de gestion du stock pendant l'horizon de temps $[0, T]$. $P(R, S)$ dénote le taux de service client pour le même horizon de temps.

Des expérimentations de simulation sont conduites pour un système de production-distribution similaire à celui décrit dans la section 4.2. La taille des commande client est uniformément distribuée dans l'intervalle (3, 9). Le temps de service unitaire est exponentiellement distribué avec une moyenne de 0.1. Le temps de transport est de 3.0 unités de temps. Le coût unitaire de possession est égal à 1 et le taux de service client est de 90%. Le point de commande et la taille de commande sont choisies à l'intérieur des intervalles de temps suivantes : $s_i \in [0, 1000]$ et $S_i \in [0, 1000]$. L'horizon du temps considéré est de $[0, 1000]$. Nous utilisons la méthode COMPASS*, algorithme 3 développé dans le chapitre 3. La solution trouvée est illustrée dans le tableau 4.15.

Table au 4.15. Solution optimale locale pour 5, 8 et 10 noeuds, $\lambda = 1.3$

	5		8		10	
	s	S	s	S	s	S
1	194	201	283	337	564	566
2	56	61	180	182	424	426
3	0	1	175	186	302	309
4	0	2	113	121	216	218
5	0	3	51	64	183	186
6			32	117	3	144
7			7	10	3	8
8			1	3	2	312
9					2	10
10					1	21
Coût	92.60		248.07		253.37	
FR	91.16		90.92		90.85	

De la même manière que pour les politiques base-stock et (R,nQ) , dans le cas d'une politique (s, S) l'entrepôt de produits finis est plus sensible à la demande et a une influence plus importante sur le taux de service client.

4.4.3 Optimisation avec prise en compte du coût de transport

Dans ce cas, nous prenons en compte le coût unitaire de possession ainsi que le coût de transport afin de calculer le coût total de gestion du stock. Le problème est de déterminer les paramètres de la politique de gestion du stock (s, S) , où $s = [s_1, s_2, s_3, \dots]^T$ et $S = [S_1, S_2, S_3, \dots]^T$, qui minimisent le coût total de possession et de transport, tout en respectant un taux de service client minimal donné α , $0 < \alpha < 1$.

$$\text{Minimize } C(s, S) = \sum_{i=1,2,\dots} (C_{hi} E[\int_0^T I_i(t) dt] + C_{Ti} E[\int_0^T CT_i(t) dt])$$

$$\text{t.q.} \quad P(s, S) > \alpha$$

Les expérimentations numériques sont effectuées avec les mêmes paramètres que dans la section 4.4.2. De plus, le coût de transport est égal à 5 et la capacité d'un moyen de transport est de 7 unités. Le coût unitaire de possession du stock est de 1 et le taux de service client est de 90%. Les résultats obtenus sont donnés dans le tableau 4.16.

Tableau 4.16. Solution pour 5, 8 and 10 noeuds, $\lambda = 1.3$

	5		8		10	
	s	S	s	S	s	S
1	195	203	308	338	362	364
2	129	137	143	146	342	343
3	25	69	135	241	300	353
4	0	3	112	117	275	295
5	0	3	96	135	187	189
6			63	146	35	214
7			31	37	7	13
8			4	4	7	24
9					7	11
10					3	56
Coût	128.82		308.47		446.08	
FR	90.43		90.51		90.63	

On peut remarquer que le point de commande, ainsi que la taille de lot

d'approvisionnement diminuent au fur et à mesure qu'on remonte le système en s'éloignant de l'entrepôt de produit finis.

4.5 Conclusion

Dans ce chapitre nous avons traité le problème d'optimisation du coût total de gestion du stock dans des systèmes de production distribution caractérisé par :i) demande aléatoire, ii) capacité de production finie, iii) taille de lot aléatoire ou fixé, iv) délai de transport constant entre une unité de production et son entrepôt aval. Différentes politiques de gestion de stock ont été analysées. De même, nous avons considéré également le cas où le coût de transport est important et ne peut pas être négligé.

Les expérimentations numériques réalisées nous ont permis de tirer les conclusions suivantes:

Conclusion 1: Le niveau de base-stock diminue lorsqu'on remonte le système en s'éloignant de l'entrepôt de produits finis.

Cette remarque est valable si le coût de possession du stock est le même pour les différents entrepôts du système. L'entrepôt final est plus sensible au taux de service client imposé. Par conséquent, cet entrepôt est caractérisé par un niveau de stock le plus important. Au contraire, les entrepôts à l'autre bout du système peut être caractérisé par une capacité réduite. Nous rappelons qu'il s'agit d'un système de production-distribution qui gère un seul type de produit. Par conséquent, on n'est pas concerné par une différenciation des produits plus ou moins en amont dans la chaîne de production.

Conclusion 2. La politique de base-stock donne le meilleur résultat lorsque le coût du transport est négligé.

La politique de base-stock que nous considérons dans notre travail est caractérisée par le fait que l'ordre d'approvisionnement est généré par des commandes d'une taille qui

peut être supérieure à l'unité. Ainsi, l'entrepôt lance un ordre d'approvisionnement de la même quantité X que la commande du client. Ceci confère de la réactivité au réapprovisionnement de l'entrepôt et permet de réduire efficacement le niveau du stock.

Conclusion 3. La politique (R,nQ) est la plus économique lorsqu'on prend en compte le coût de transport.

Conclusion 4. Du point de vue économique la politique (s, S) est la moins performante.

Cette conclusion est illustrée à travers les résultats numériques présentés dans le tableau 4.17. Le taux de service client minimal demandé est de 90% et le taux moyen d'arrivée est $\lambda = 1.3$. Les autres paramètres ont les valeurs définies dans les sections 4.2. et 4.3.

Tableau 4.17. Solutions locale optimales pour différentes politiques de gestion du stock (5 noeuds)

	R	(R,nQ)		(s,S)		R	(R,nQ)		(s,S)	
	Sans coût	Sans coût		Sans coût		Avec coût	Avec coût		Avec coût	
	transp.	transp.		transp.		transp.	transp.		transp.	
	R	R	Q	s	S	R	R	Q	s	S
1	167	224	1	194	201	266	266	10	195	203
2	61	8	1	56	61	0	0	10	129	137
3	1	4	1	0	1	0	0	3	25	69
4	0	1	1	0	2	0	0	1	0	3
5	0	1	1	0	3	0	0	2	0	3
Coût	78.8	88.33		92.60		139.2	114.8		128.82	
FR	90.12	90.44		91.16		90.01	91.72		90.43	

Une raison du coût de gestion de stock élevé fournit par la politique (s, S) est le manque de flexibilité concernant le choix de la quantité à approvisionner..

Conclusion 5. Dans une politique (R, nQ) l'interdépendance entre le niveau de base-stock R et la taille lot Q est très faible.

Conclusion générale

Conclusions et perspectives

Les travaux de recherche présentés dans ce mémoire traitent le problème d'évaluation de performances et d'optimisation des systèmes de production-distribution multi-niveaux. Plus précisément, nous avons proposé une méthode analytique ainsi qu'une méthode basée sur la simulation pour évaluer des performances telles que le délai d'approvisionnement des entrepôts, le taux de service client et le coût total de gestion du stock. Notre objectif est de trouver le paramétrage du système de gestion de stock qui minimise le coût total de gestion de stock tout en respectant les contraintes imposées sur le fonctionnement du système.

Dans un premier temps nous développons une approche analytique pour l'évaluation des performances et l'optimisation des systèmes de production-distribution à deux niveaux. Le système considéré est composé d'un entrepôt alimenté par une unité de production. Il est caractérisé par : i) une demande aléatoire, ii) une capacité de production finie, iii) taille aléatoire des commandes client, et des ordres de re-approvisionnement, iv) délai de transport constant entre l'unité de production et l'entrepôt. L'approche analytique que nous proposons est basée sur des approximations techniques qui ont prouvé leur efficacité pour la plupart des systèmes de production : i) caractérisation des variables aléatoires par les deux premiers moments, ii) approximation markovienne des processus généraux, iii) approximation de processus dépendantes pas de processus indépendantes. Plus spécialement, cette approche utilise les deux premiers ordres des variables aléatoires pour évaluer le délai de re-approvisionnement de l'entrepôt, le niveau du stock et les coûts de possession et de rupture de stock. Le coût total de gestion du stock est minimisé à l'aide d'une méthode de gradient. Des expérimentations numériques nous ont permis de valider l'approche et de montrer son efficacité.

Cependant, une approche analytique est difficile à mettre en place pour résoudre le problème que nous avons formulé lorsque le système considéré est composé de plusieurs entrepôts et unités de production. Par conséquent, nous proposons une

approche d'évaluation de performances et d'optimisation des systèmes de production-distribution multi-niveaux basée sur la simulation. Dans ce cas, la performance qui nous intéresse ainsi que certaines contraintes sont évaluées par la simulation. Nous avons prouvé que les algorithmes que nous proposons convergent avec la probabilité 1 vers un ensemble de solutions localement optimales si certaines hypothèses sont respectées. Des expérimentations numériques montrent l'efficacité ainsi que la propriété de convergence rapide des nos algorithmes.

Finalement, nous avons étudié différentes politiques de gestion de stock dans le cadre du système de production-distribution considéré précédemment où on prend en compte également le coût du transport entre différentes entités du système. Les politiques de gestion de stock qui ont retenu notre attention sont : la politique base-stock, la politique (R, nQ) et la politique (s, S) . Notre objectif est de déterminer les paramètres des différentes politiques de gestion de stock de chaque entrepôt afin de minimiser le coût total de gestion du stock. Suite à un nombre important des expérimentations numériques, nous avons tiré les conclusions suivantes : i) le niveau de base-stock diminue au fur et à mesure qu'on remonte le système en s'éloignant de l'entrepôt de produits finis ; ii) lorsqu'on néglige le coût du transport, la politique base-stock est la plus économique ; iii) la politique (R, nQ) est la plus économique lorsqu'on prend en compte le coût du transport ; iv) la politique (s, S) est la moins économique parmi les trois politiques étudiées et v) le niveau de base-stock dépend peu de la taille de lot d'approvisionnement.

En ce qui concerne la poursuite de nos travaux de recherche, les pistes qui nous semblent les plus intéressantes sont les suivantes :

1. Continuer les travaux en vue de développement d'une approche analytique pour l'optimisation et l'évaluation de performances des systèmes de production-distribution multi-niveaux.
2. Développer une méthode d'allocation du budget de simulation intelligente pour la méthode COMPASS* présentée dans le chapitre 3. L'objectif est d'accélérer la convergence de l'algorithme. La méthode OCBA (optimal computational

budget allocation) proposée par C. H. Chen, L. Dai, H. C. Chen, et E. Yucesan(1998) fournit une bonne base pour ces travaux.

3. Prendre en compte des réseaux de production-distribution avec une structure plus complexe et éventuellement plusieurs produits.

Références

Références

- [Alrefaei et Andradottir, 1999] Alrefaei, M.H. and S. Andradottir. (1999). A simulated annealing algorithm with constant temperature for discrete stochastic optimization. *Management science*, 45(5),748-764
- [Andradottir, 1996] Andradottir, S. (1996). A global search method for discrete stochastic optimization. *SIAM Journal on Optimization* 6:513-530.
- [Andradottir, 1998a] Andradottir, S. (1998a). A review of simulation optimization techniques. *Proceedings of the 1998 Winter Simulation Conference*, 151-158.
- [Andradottir, 1998b] Andradottir, S. (1998b). Simulation optimization. In *Handbook on Simulation*, ed. J. Banks, 307-333. John Wiley & Sons, Inc., New York.
- [Andradottir, 1999] Andradottir, S. (1999). Accelerating the convergence of random search methods for discrete stochastic optimization. *ACM Transactions on Modeling and Computer Simulation*, 9:349–380.
- [Axsater, 1993] Axsater, S., (1993), Continuous review policies for multi-level inventory systems with stochastic demand. *Handbooks in Operations Research and Management Science*, vol 4, Logistics of Production and Inventory. Editors: Graves, S.C., Rinnooy Kan, A.H.G. and Zipkin, P.H. Elsevier North-Holland.
- [Axsater et Rosling 1993] Axsater, S. and Rosling, K., (1993), Notes: Installation vs. Echelon Stock Policies for Multilevel Inventory Control, *Management Science*, 39 (10), pp.1274-1280
- [Ballou,1992] Ballou, R.H. (1992), *Business logistics management*. third edition, Prentice-Hall.
- [Baretto et al., 1999] Baretto,M.R.P.,Chwif,L.,Eldabi,T. and Paul,R.J. (1999) simulation optimization with the linear move and exchange move optimization algorithm, in *proceedings of winter simulation conference*, IEEE,Piscataway,NJ,pp.806-814
- [Beamon, 1998] Beamon, B.M., (1998), Supply Chain Design and Analysis: Models and Methods, *International Journal of Production Economics*, 55, pp.281-294
- [Bhaskaran, 1998] Bhaskaran, S., (1998), Simulation Analysis of a Manufacturing

- Supply Chain, *Decision Sciences*, 29:(3), pp.633-657
- [**Bulgak et Sanders, 1988**] A. A. Bulgak and J. L. Sanders. (1988). Integrating a modified simulated annealing algorithm with the simulation of a manufacturing system to optimize buffer sizes in automatic assembly systems. *Proceedings of the 1988 Winter Simulation Conference*, pages 684-690, .
- [**Buzacott, et al., 1993**] Buzacott, J., Shanthikumar, J., (1993). "Stochastic Models of Manufacturing Systems" . Prentice Hall, February.
- [**Cassandras, 1993**] Cassandras, C. G. (1993) *Discrete Event System: Modeling and Performance Analysis*. Irwin,
- [**Chen, Lin, Yücesan, et Chick S. E, 2000**] Chen, C. H., J. Lin, E. Yücesan and S. E. Chick. (2000). Simulation budget allocation for further enhancing the efficiency of ordinal optimization. *Journal of Discrete Event Dynamic System: Theory and Applications*, 10:251-270.
- [**Chen et Shi, 2000**] Chen, C. H. and Shi L. (July 2000). A New Algorithm for Stochastic Discrete Resource Allocation Optimization,, *Journal of Discrete Event Dynamic Systems: Theory and Applications*
- [**Chen, Chen et Yücesan**] Chen C. H. Chen H. C. and Yücesan E. "Optimal Computing Budget Allocation in Selecting The Best System Design via Discrete Event Simulation," , To appear in *European Journal of Operations Research*.
- [**Chen, Dai, Chen et Yücesan, 1998**] C. H. Chen, L. Dai, H. C. Chen, and E. Yucesan, (1998). Efficient computation of optimal budget allocation for discrete event simulation experiment, *IIE Trans. Operat. Res.*.
- [**Clark et Scarf, 1960**] Clark, A. and Scarf, H., (1960), Optimal Policies for a Multi-echelon Inventory Problem, *Management Science*, 6, pp.475-490
- [**Cohen et Lee, 1988**] Cohen, M.A. and Lee, H.L., (1988), Strategic Analysis of Integrated Production-distribution Systems: Models and Methods, *Operations Research*, 36 (2), pp.216-228
- [**Cooper, et al., 1981**] Cooper, Robert B. (1981). "Introduction to Queuing Theory", Elsevier North Holland
- [**Dai, 1995**] Dai, L. (1995). Convergence properties of ordinal comparison in the simulation of discrete event dynamic systems. *Proceedings of the IEEE*

Conference on Decision and Control, 2604-2609.

[De Kok, 1991] Kok, A.G. de .(1991). *Basics of Inventory Management: part 1-5*. FEW working papers 520-525, Tilburg University, The Netherlands.

[Deuermeyer et Schwartz, 1981] Deuermeyer, B.L., Schwarz, L., (1981), A Model for the Analysis of System Service Level in Warehouse/Retailer Distribution Systems: the Identical Retailer Case. In *Multi-Level Production/Inventory Control Systems: Theory and Practice*. L.B. Schwarz (eds.), North Holland, Amsterdam.

[Diks et al., 1996] Diks, E.B., Kok, A.G. de and Lagodimos, A.G., (1996) Multi-echelon systems: a measure perspective. *European Journal of Operational Research*, vol 95,nr. 2, p 241-264.

[Ding et al. 2003a] Ding H., Benyoucef L., et Xie X. (2003) Simulation Model of a Mixed Make-to-Order and Make-to-Stock Distribution Network. *Proceedings of the 6th International Conference on Industrial Engineering and Production Management*. Porto, Portugal..

[Ding et al. 2003b] Ding H., Benyoucef L., et Xie X. (2003) A simulation-optimization approach using genetic search for supplier selection. *Proceedings of the 2003 Winter Simulation Conference*. New Orleans, U.S.A..

[Ding et al. 2004a] Ding H., Benyoucef L., et Xie X.(2004) A modeling and simulation framework for supply chain network design. *Supply Chain Optimization : product/process design, facility location and flow control*, Chapter 16, pp. 219-232. Springer. ISBN : 0-387-23566-3.

[Ding et al. 2004b] Ding H., Benyoucef L., et Xie X. (2004) A simulation optimization methodology for supplier selection problem. *International Journal of Computer Integrated Manufacturing*.. To appear.

[Ding et al. 2004c] Ding H., Benyoucef L., et Xie X.(2004) A simulation-based optimization method for production-distribution network design. *Proceedings of the 2004 IEEE International Conference on Systems, Man & Cybernetics*. The Hague, Netherlands.

[Ding et al. 2004d] Ding H., Benyoucef L., et Xie X.(2004) Multiobjective optimization on facility location and inventory deployment with customer service consideration. *Proceedings of the 2004 International Conference on Service*

- Systems and Service Management*. Beijing, China.
- [**Dudewicz et Dalal, 1975**] Dudewicz, E. J. and Dalal, S. R. (1975). Allocation of Observations in Ranking and Selection with Unequal Variances. *Sankhya*, B37:28-78.
- [**Eglese,1990**] Eglese,R.W.(1990) Simulated annealing: a tool for operational research. *European Journal of operational Research*, 46, 271-281
- [**Ettl, et al., 2000**] Ettl, M., Feigin, G.E., Lin, G.Y., and Yao, D.D., (2000), A Supply Network Model with Base-Stock Control and Service Requirements, *Operations Research*, 48, pp.216-232
- [**Federgruen, 1993**] Federgruen, A. (1993) Centralized planning models for multi-echelon inventory systems under uncertainty *Handbooks in Operations Research and Management* .Science, vol 4, Logistics of Production and Inventory. Editors: Graves, S.C., Rinnooy Kan, A.H.G. and Zipkin, P.H. Elsevier North-Holland.
- [**Federgruen et Zipkin, 1984**] Federgruen, A. and Zipkin, P., (1984), Computational Issues in an Infinite-Horizon Multi-echelon Inventory Model, *Operations Research*, 32 (4), pp.818-836
- [**Flescher,1995**] Fleischer, M.A. (1995). Simulated annealing: past, present, and future. *Proceedings of the 1995 Winter Simulation Conference*, 155-161.
- [**Fu, 2002**] Fu, M. C. (2002). Optimization for simulation: Theory vs. practice. *INFORMS Journal on Computing*, 14:192–215.
- [**Garg, 1999**] Garg, Amit., (1999), An Application of Designing Products and Processes for Supply Chain Management, *IIE Transactions*, 31, pp.417-429
- [**Glover et Laguna,1997**] Glover, F. and Laguna, M.(1997) Tabu Search, Kluwer, Norwell, MA.
- [**Goldsman et Nelson,1998**] Goldsman,D and Neson,B.L (1998) Statistical screening ,selection and multiple comparison procedures in computer simulation, in proceeding of the 1998 Winter Simulation conference,IEEE,Piscataway,NJ,pp.159-166.
- [**Gong et al. 1999**] Gong, W. B., Y. C. Ho and W. Zhai. (1999). Stochastic comparison algorithm for discrete optimization with estimation. *SIAM Journal on*

- Optimization*, 10:384–404.
- [Graves, 1985] Graves, S., (1985), A Multi-echelon Inventory Model for a Repairable Item with One-for-one Replenishment, *Management Science*, 31, pp.1247-1256
- [Graves, 1988] Graves, S.C., (1988), Safety Stocks in Manufacturing Systems, *Journal of Manufacturing and Operations Management*, 1 (1), pp.67-101
- [Haddock et Mittenthal,1992] J. Haddock and J. Mittenthal. (1992). Simulation optimization using simulated annealing. *Comput. and Indust. Engrg.*, 22:387-395 .
- [Hadley et Whitin 1963] Hadley, G. and Whitin, T.M., (1963) *Analysis of inventory systems*. Prentice-Hall, Englewood Cliffs, NJ.
- [Harris, 1913] Harris, F.W. (1913), how many parts to make at once. *Factory, magazine of management*, vol 10, nr. 1, p 135-136.
- [Ho, Sreeniva, et Vakili,1992] Ho,Y.C.,Sreeniva,R. and Vakili,P.(1992) Ordinal optimization of discrete event dynamic systems. *Discrete Event Dynamical Systems*,2(2),61-88
- [Ho,1994] Ho,Y.C.,(1994), Overview of ordinal optimization, in *Proceeding of the IEEE conference on Decision and Control*.IEEE,Piscataway,NJ,pp.2199-2204
- [Ho,2002] Ho,T.C., D. Li and L.H. Lee (2002). "Constraint Ordinal Optimization", *Information Sciences*, Vol. 148, 201-220
- [Ho, et al.2000] Ho, Y. C., C. G. Cassandras, C. H. Chen and L. Y. Dai. (2000). Ordinal optimization and simulation. *Journal of the Operational Research Society*, 51:490–500.
- [Hong et Nelson 2006] Hong L.J. and Nelson Barry L.(2004) discrete optimization via simulation using COMPASS. *Operations Research*, 54.1 :115-129.
- [Houlihan, 1985] Houlihan, J.B., (1985), International Supply Chain Management. *International Journal of Physical Distribution & Materials Management* 15: pp.22-38
- [Houtum et al., 1996] Houtum, van G.J., Inderfurth, K., and Zijm, W.H.M, (1996) Material coordination in stochastic multi-echelon systems. *European Journal of*

- Operational Research, vol 95, nr. 1, p 1-23.
- [James, Paul etc.2004] James R. Swisher, Paul D. Hyden, Sheldon H. Jacobson and Lee W. Schruben (2004), A survey of recent advances in discrete input parameter discrete-event simulation optimization, *IIE Transactions* 36, 591-600
- [Jia et al., 2005] Jia, Q.-S., Ho, Y.-C., and Zhao, Q.-C., (2005) "Comparison of selection rules for ordinal optimization," *Mathematical and Computer Modelling*, to appear.
- [Kleinrock, 1975] Kleinrock, L. (1975), *Queueing Systems, Vol. I: Theory*. John Wiley & Sons, New York,
- [Lamming, 1996] Lamming, R., (1996) Squaring Lean Supply with Supply Chain Management. *International Journal of Operations & Production Management* 16(2): pp.183-196
- [Lee et Billington, 1993] Lee, H.L. and Billington, C., 1993, Material Management in Decentralized Supply Chains. *Operations Research* 41: pp.835-847
- [Lee et Moinzadeh, 1987a,] Lee, H.L. and Moinzadeh, (1987a), Operating Characteristics of a Two-echelon Inventory System for Repairable and Consumable Items Under Batch Ordering Policy and Shipment Policy, *Naval Research Logistics*, 34, pp.365-380
- [Lee et Moinzadeh, 1987b] Lee, H.L. and Moinzadeh, (1987b), Two-parameter Approximations for Multi-echelon Repairable Inventory Models with Batch Ordering Policy, *IIE Transaction*, 19, pp.140-149
- [Lee et al., 1999] Lee, L.H., T.W.E. Lau., and Y.C. Ho. (1999). Explanation of goal softening in ordinal optimization. *IEEE Transactions on Automatic Control* 44:94-98.
- [Lee et Nahmias, 1993] Lee, H.L., and Nahmias, S. (1993) *Single product, single location models*. Handbooks in Operations Research and Management Science, vol 4, Logistics of Production and Inventory. Editors: Graves, S.C., Rinnooy Kan, A.H.G. and Zipkin, P.H. Elsevier North-Holland.
- [LI, SAVA et XIE, 2004] Jie Li, alexandru SAVA, Xiaolan XIE. , (2004) "Evaluation des performances des systèmes de production-distribution à deux niveaux – une approche analytique," *Proc. MOSIM*, Nantes, France.

- [LI,SAVA et XIE, 2005a]Jie LI, Alexandru SAVA, Xiaolan XIE. (2005) "Performance Evaluation and Optimization of a Two-stage Production-distribution System with Batch Orders and finite transportation time," *Proc. 16th IFAC-World Congress*, Prague, Czech Republic.
- [LI,SAVA et XIE, 2005b] Jie LI, Alexandru SAVA, Xiaolan XIE.(2005) "A performance evaluation approach of an inventory system with batch orders," *Proc. 1st I*PROMS VIRTUAL INTERNATIONAL CONFERENCE ON INTELLIGENT MACHINES AND SYSTEMS*, CARDIFF (UK), 4 AU 15 JUILLET 2005, ISBN-10:0080447309
- [LI,SAVA et XIE, 2006a]Jie LI, Alexandru SAVA, Xiaolan XIE.(2006) "An Analytical Approach for Performance Evaluation and Optimization of a Two-stage Production-distribution System," submitted to *International Journal of Production Research*.
- [LI,SAVA et XIE, 2006b]Jie LI, Alexandru SAVA, Xiaolan XIE.(2006) "A Simulation Based Approach For Optimization Of Multi-Stage Production-Distribution Systems With Batch Order," *Proc. INCOM'06*, Saint Etienne, France, May 17-19.
- [LI,SAVA et XIE, 2006c]Jie LI, Alexandru SAVA, Xiaolan XIE.(2006) "A Simulation Based Optimization Method for Multi-Stage Production-Distribution Systems with Service Level Constraints," INRIA report
- [LI,SAVA et XIE, 2006d]Jie LI, Alexandru SAVA, Xiaolan XIE.(2006) "Discrete Optimization via Simulation based on Most Promising Area Stochastic Search," under preparation.
- [LI,SAVA et XIE, 2006e] Jie LI, Alexandru SAVA, Xiaolan XIE.(2006) "Discrete Optimization via Simulation method with Simulated Constraint," under preparation.
- [Liepins et Hilliard,1989] Liepins,G.E. and Hilliard,M.R. (1989) Genetic algorithms: foundations and applications. *Annals of Operations Research*,21,31-58
- [Liu, et al.,1990] Liu, L., Kashyap, B.R.K., and Templeton, J.G.C,(1990) "On the $GI^x/G/\infty$ Queueing System", *Journal of Applied Probability*, 27, 671 – 683.
- [Liu, et al., 2004] L. Liu, X. Liu, D. D. Yao.(2004) "Analysis and Optimization of Multi-Stage Inventory-Queues", *Management Science* 50, 365-380

- [**Mabert et Venkataramanan ,1998**] Mabert, V.A. and Venkataramanan, M.A., (1998), Special Research Focus on Supply Chain Linkages: Challenges for Design and Management in the 21st Century. *Decision Sciences*, 29: (3) pp.537-552
- [**Masuyama, and Takine,2002**] Hiroyuki Masuyama and Tetsuya Takine.(2002) “Analysis of an Infinite-Server Queue with Batch Markovian Arrival Streams”, *Queueing Systems* 42 (3) p.269-296
- [**Muckstadt 1973**] Muckstadt, J., (1973), A Model for a Multi-item, Multi-echelon, Multi-indenture Inventory System, *Management Science*, 20, pp.472-481
- [**Muhlenbein, 1997**] Muhlenbein,H.(1997) Genetic algorithms, in Local Search in Combinatorial Optimization,Aarts,E. and Lenstra,J.K.(eds.),pp.137-172.
- [**Olafsson et Kim, 2002**] Olafsson, S. and J. Kim. (2002). Simulation optimization. *Proceedings of the 2002 Winter Simulation Conference*, 79–84.
- [**Olafsson et Shi, 2002**] Olafsson,S. and Shi,L. (2002) Ordinal comparison via the nested Partitions method. Discrete event dynamic systems: Theory and Applications, 12, pp.211-239
- [**Pichitlamken et Nelson, 2003**] Pichitlamken, J. and B. L. Nelson. (2003). A combined procedure for optimization via simulation. *ACM Transactions on Modeling and Computer Simulation*, 13:155-179.
- [**Rosling, 1989**] Rosling, K., (1989), Optimal Inventory Policies for Assembly Systems Under Random Demands, *Operations Research*, 37 (4), pp.565-579
- [**Saporta, 1990**] Saporta, G.,(1990), *Probabilités Analyse des Données et Statistique*, Technip Press,W.H.,
- [**Scarf et Karlin 1958**] Scarf, H.E. and Karlin, H. (1958) A min-max solution of an inventory problem. *Studies in the mathematical theory of inventory and production*. Editors: Arrow, k., Karllin, S. and Scarf, H.
- [**Scarf, 2002**] Scarf, H.E. (2002) Inventory theory. *Operations Research*, vol 50, nr. 1, p 186-191.
- [**Sherbrooke 1968**] Sherbrooke, C., (1968), METRIC: a Multi-echelon Technique for Recoverable Item Control, *Operations Research*, 16, pp.122-141
- [**Sherbrooke, 1986**] Sherbrooke, C., (1986), VARI-METRIC: Improved

- Approximations for Multi-indenture, Multiechelon Availability Models, *Operations Research*, 34, pp.311-319
- [**Shi et Olafsson 2000a**] Shi, L. and Olafsson, S. (2000a) Nested partitions method for global optimization. *Operations Research*, Vol.48.No.3, May-June pp.390-407
- [**Shi et Olafsson 2000b**] Shi, L. and S. ´ Olafsson. (2000b). Nested partitions method for stochastic optimization. *Methodology and Computing in Applied Probability*, 2:271–291.
- [**Silver et al., 1998**] Silver, E.A., Pyke, D.F., and Peterson, R. (1998) *Inventory management and production Planning and Scheduling*. third edition, Wiley, New York.
- [**Stevens, 1989**] Stevens, G.C., (1989), Integrating the Supply Chain. *International Journal of Physical Distribution & Materials Management* 19: pp.3-8
- [**Svoronos et Zipkin, 1988**] Svoronos, A. and Zipkin, P., (1988), Estimating the Performance of Multi-level Inventory Systems, *Operations Research*, 36, pp.57-72
- [**Svoronos et Zipkin, 1991**] Svoronos, A., Zipkin, P., (1991), Evaluation of One-for-one Replenishment Policies for Multiechelon Inventory Systems, *Management Science*, 37, pp.68-83
- [**Swaminathan, Smith et Sadeh, 1998**] Swaminathan, J.M., Smith, S.F. and Sadeh, N.M., (1998), Modeling Supply Chain Dynamics: A Multi agent Approach. *Decision Sciences*, 29: (3) pp. 607-632
- [**Tijms et Groeneveld, 1984**] Tijms, H.C. and Groeneveld H. (1984) Simple approximations for the reorder point in periodic and continuous review (s,S) inventory systems with service level constraints. *European Journal of Operational Research*, vol 17, nr. 2, p 175-190.
- [**Towill, 1991**] Towill, D.R., (1991), Supply Chain Dynamics, *International Journal of Computer Integrated Manufacturing* 4(4), pp.197-208
- [**Towill et al., 1992**] Towill, D.R., Naim, N.M. and Wikner, J., (1992), Industrial Dynamics Simulation Models in the Design of Supply Chains, *International Journal of Physical Distribution and Logistics Management* 22, pp.3-13
- [**Wagner et Whitin 1958**] Wagner, H., and Whitin, T.M. (1958) Dynamic version of

the economic lotsize model. *Managemnent Science*, vol 5, nr. 1, p 89-96.

[Xie, 1997] Xie, X. (1997). Dynamics and convergence rate of ordinal comparison of stochastic discrete-event systems. *IEEE Transactions on Automatic Control* 42:586-590.

[Yucesan et Jacobson, 1996] Yucesan, E., Jacobson, SH, (1996), "Computational Issues for ACCESSIBILITY in Discrete Event Simulation," *ACM Transactions on Modeling and Computer Simulation*, 6(1), 53-75

[Zipkin,2000] Zipkin, P. H.(2000) "*Foundations of inventory management*". McGraw Hill.