

Table des Matières

Table des Matières	i
Table des Figures.....	v
INTRODUCTION GÉNÉRALE.....	1
CHAPITRE I	5
1. Pilotage de flux dans les chaînes logistiques.....	5
1.1. INTRODUCTION	5
1.2. GESTION DE CHAÎNES LOGISTIQUES	5
1.2.1. Externalisation et les chaînes logistiques	7
1.2.2. Définition de la chaîne logistique.....	8
1.2.3. Définition de la gestion de chaînes logistiques	9
1.3. PILOTAGE DE FLUX DANS LES CHAÎNES LOGISTIQUES	9
1.3.1. Le système élémentaire de flux	11
1.3.1.1. Le système d'approvisionnement	11
1.3.1.2. Le système de demande	12
1.3.1.3. Le stock	13
1.4. POLITIQUES DE PILOTAGE DE FLUX.....	15
1.4.1. Politiques d'approvisionnement et de stockage	15
1.4.1.1. Politiques d'approvisionnement et de stockage dans les systèmes multi-étages.....	19
1.4.2. Politiques de stock nominal.....	22
1.4.2.1. Systèmes de stock nominal multi-étages	23
1.4.3. Politiques Kanban	26
1.4.3.1. Extensions de la politique Kanban	29
1.4.4. Les politiques du type MRP	31
1.4.5. Classification des politiques de pilotage de flux	35
1.5. MODÉLISATION DYNAMIQUE DES FLUX.....	37
1.5.1. Modélisation des systèmes stock nominal.....	38
1.6. CONCLUSIONS.....	42
CHAPITRE II.....	43
2. Optimisation des décisions dans les chaînes logistiques	43
2.1. INTRODUCTION	43

2.2.	CONCEPTS DE BASE DE LA THÉORIE DES JEUX.....	44
2.3.	JEUX STATIQUES	45
2.3.1.	Existence de l'équilibre	46
2.3.2.	Unicité de l'équilibre.....	47
2.3.3.	Pareto optimalité de l'équilibre	48
2.3.4.	Jeux statiques et compétition dans les chaînes logistiques.....	52
2.4.	JEUX DYNAMIQUES	54
2.4.1.	Jeux de Stackelberg.....	56
2.4.2.	Jeux répétés	61
2.4.3.	Jeux stochastiques	62
2.5.	JEUX AVEC INFORMATION ASYMÉTRIQUE.....	63
2.6.	JEUX COOPÉRATIFS	66
2.7.	CONCLUSIONS.....	69
CHAPITRE III		71
3.	Politique d'approvisionnement dans un système à plusieurs fournisseurs	71
3.1.	INTRODUCTION	71
3.2.	MODÉLISATION DU PROBLÈME DANS LE CAS DE PRODUCTION POUR STOCK ..	75
3.3.	POLITIQUE D'APPROVISIONNEMENT OPTIMAL DANS LE CAS DE PRODUCTION À LA COMMANDE.....	79
3.3.1.	Résolution du problème relaxé.....	80
3.3.2.	Problème de choix restrictif.....	82
3.4.	MÉTHODE APPROXIMATIVE DANS LE CAS DE PRODUCTION POUR STOCK	84
3.4.1.	Calcul des paramètres de Bernoulli.....	85
3.4.2.	Calcul du niveau de stock nominal.....	85
3.5.	ANALYSES NUMÉRIQUES.....	86
3.6.	CONCLUSIONS.....	88
CHAPITRE IV		91
4.	Optimisation des décisions dans une chaîne logistique décentralisée à deux étages de production/stockage	91
4.1.	INTRODUCTION	91
4.2.	MODÈLE DE PILOTAGE DE FLUX	94
4.3.	CALCULS DES MESURES DE PERFORMANCES.....	98
4.3.1.	Le système de fabrication de produits intermédiaires	98
4.3.2.	Le système de fabrication de produits finis	99
4.3.2.1.	La méthode approximative LZ	101
4.4.	SYSTÈME DÉCENTRALISÉ ET ÉQUILIBRE DE STACKELBERG	103
4.4.1.	Problème d'optimisation du fournisseur	105

4.4.2.	Problème d'optimisation du producteur	107
4.4.2.1.	Les valeurs optimales des paramètres du contrat	108
4.4.2.2.	La valeur optimale du niveau de stock nominal	111
4.5.	PROBLÈME CENTRALISÉ ET PERFORMANCE DU CONTRAT	114
4.5.1.	Prise en compte le pouvoir de négociation du fournisseur	116
4.6.	CONCLUSIONS	117
	CONCLUSIONS GÉNÉRALES ET PERSPECTIVES.....	119
	ANNEXE A	125
A.	Processus Stochastiques	125
A.1.	LOIS DE PROBABILITÉ ET PROCESSUS STOCHASTIQUES	125
A.1.1.	La loi géométrique	125
A.1.2.	La loi de Poisson	125
A.1.3.	La loi exponentielle	125
A.1.4.	Processus stochastique	126
A.1.5.	Chaîne de Markov à temps discret	126
A.1.6.	Chaîne de Markov à temps continu	127
A.1.7.	Processus de comptage	127
A.1.8.	Processus de Poisson	128
A.1.9.	Processus de Poisson composé	129
A.2.	MODÈLES DE FILES D'ATTENTE	129
A.2.1.	Caractérisation des modèles de file d'attente	130
A.2.1.1.	Processus d'arrivée	130
A.2.1.2.	Distribution du temps de service	131
A.2.1.3.	Nombre de serveurs	132
A.2.1.4.	Capacité du système	132
A.2.1.5.	Discipline de service	132
A.2.2.	Notation des modèles de file d'attente	132
A.2.3.	Évaluation de performances	133
A.2.4.	La file M/M/1	135
A.2.5.	Réseaux des files d'attente	136
A.2.5.1.	Réseaux à forme produit	138
A.2.5.2.	Réseaux des files d'attente avec blocage	139
A.2.5.3.	Réseaux des files d'attente avec stations de synchronisation	140
A.3.	FORMULATION DE LA FONCTION $C(s)$	141
	ANNEXE B	143
B.	Démonstrations pour le modèle à plusieurs fournisseurs	143
B.1.	FORMULATION DE LA FONCTION P_v	143

B.1.1.	Fonctions de génération de probabilité.....	143
B.1.2.	Obtention de la fonction P_v	144
B.2.	DÉMONSTRATION DU LEMME 3.1	145
B.3.	DÉMONSTRATION DU LEMME 3.2	146
B.3.1.	Existence de l'indice m^*	146
B.3.2.	Unicité de l'indice m^* et évolution du paramètre τ_m	147
B.4.	DÉMONSTRATION DE LA PROPRIÉTÉ 3.2	147
B.4.1.	Admissibilité de la politique $\alpha^*(m^*)$	147
B.4.2.	Optimalité de la politique $\alpha^*(m^*)$	148
ANNEXE C		151
C. Lois de type phase et la méthode LZ		151
C.1.	LOIS DE TYPE PHASE.....	151
C.1.1.	Lois de type phase à temps continu.....	151
C.1.2.	Lois de type phase à temps discret	153
C.2.	MÉTHODE APPROXIMATIVE DE LEE ET ZIPKIN	154
RÉFÉRENCES		157

Table des Figures

Figure 1.1.	Pilotage de flux dans une chaîne logistique	10
Figure 1.2.	Système élémentaire de flux	11
Figure 1.3.	Évolution du stock avec la politique (T, S)	17
Figure 1.4.	Évolution du stock avec la politique (T, s, S)	17
Figure 1.5.	Évolution du stock avec la politique (R, Q)	18
Figure 1.6.	Évolution du stock avec la politique (s, S)	18
Figure 1.7.	Systèmes multi-étages	19
Figure 1.8.	Évolution du stock avec la politique $(S - 1, S)$	23
Figure 1.9.	Politique de stock nominal $(S - 1, S)$ dans un système à deux étages	24
Figure 1.10.	La politique Kanban dans un système à deux étages	27
Figure 1.11.	La politique Kanban généralisée dans un système à deux étages	30
Figure 1.12.	La politique Kanban étendue dans un système à deux étages	31
Figure 1.13.	Exemple d'un plan MRP	33
Figure 1.14.	Système mono-étage de production/stockage	38
Figure 1.15.	Représentation du système par les files d'attente	40
Figure 2.1.	Bataille des sexes	48
Figure 2.2.	Dilemme du prisonnier	49
Figure 2.3.	Représentation d'un jeu dynamique par un arbre de décision	55
Figure 2.4.	Représentation d'un jeu statique à deux joueurs avec un arbre de décision	56
Figure 2.5.	Exemples des jeux à deux joueurs	58
Figure 3.1.	La chaîne logistique à deux étages	72
Figure 3.2.	Le réseau ouvert de files d'attente avec n files en parallèle	76
Figure 3.3.	Représentation du système par les files d'attente	78
Figure 4.1.	La chaîne logistique à deux étages de production/stockage	92
Figure 4.2.	Représentation du système par les files d'attente	95
Figure 4.3.	Représentation du système sans la répartition physique des produits intermédiaires	96
Figure 4.4.	Représentation de D_1 par une phase	99
Figure 4.5.	Représentation de L_2 par une loi de type phase	101
Figure 4.6.	Représentation de D_2 par une loi de type phase	103

Figure 4.7.	Les niveaux de stocks nominaux S_2^* et $S_1^*(b_1^*(S_2^*))$ pour $h_2 = 0.8$, $b_2 = 10$, $\rho_2 = 0.5$	113
Figure A.1.	Représentation graphique d'un système de file d'attente simple	130
Figure A.2.	Diagramme de transition d'un système de file d'attente M/M/1 (λ, μ)	135
Figure A.3.	Réseau ouvert (acyclique) des files d'attente en tandem.....	136
Figure A.4.	Réseau ouvert des files d'attente en tandem avec feed-back.....	137
Figure A.5.	Réseau fermé des files d'attente en tandem	137
Figure A.6.	Réseau mixte des files d'attente en tandem.....	137
Figure A.7.	Acheminement probabiliste dans un réseau ouvert des files d'attente	137
Figure A.8.	Système de file d'attente avec synchronisation d'arrivées.....	140
Figure A.9.	Réseau Fork/Join avec synchronisation.....	140
Figure C.1.	Représentation de la distribution hypo-exponentielle	152
Figure C.2.	Représentation de la distribution hyper-exponentielle	152
Figure C.3.	Représentation de la distribution de Cox.....	153

INTRODUCTION GÉNÉRALE

Face à des marchés fluctuants et à une intensification de la compétition, la recherche de flexibilité et de maîtrise des coûts a conduit de nombreuses entreprises à externaliser certaines activités jugées périphériques ou nécessitant des ressources spécifiques que d'autres entreprises détiennent, et à se centrer sur ce qu'elles définissent comme le cœur de leur métier. Suivant cette démarche, les entreprises d'aujourd'hui construisent leur offre de produit ou de service en s'appuyant (plus ou moins durablement) sur d'autres entreprises dont elles mobilisent les ressources et les compétences. Ainsi, les entreprises d'aujourd'hui appartiennent généralement à des *réseaux d'entreprises*, autrement appelées des *chaînes logistiques inter-organisationnelles*, reposant sur des nombreux acteurs spécialisés : les entreprises fabriquant les composants, les produits intermédiaires ou les produits finis, les sous-traitants, les distributeurs, les détaillants, les prestataires de services logistiques, etc.

Savoir gérer le flux au sein de ces chaînes logistiques est essentiel au niveau opérationnel et est porteur de nombreux enjeux. Si la gestion de flux n'est pas toujours considérée comme un vecteur stratégique à part entière, elle est néanmoins de plus en plus sollicitée en tant que support aux stratégies organisationnelles pour satisfaire effectivement une demande et créer de la valeur à la fois pour le client et l'actionnaire, mais aussi pour les parties prenantes du réseau. Puisque les entreprises d'aujourd'hui doivent gérer de multiples interfaces avec d'autres entreprises et que leur réussite individuelle est largement liée aux réactions, aux compétences et à la réussite de chaque participant, l'importance de la dimension inter-organisationnelle du pilotage de flux n'est plus contestée. Elle est, d'ailleurs, essentielle pour faire face aux besoins fluctuants et personnalisés des clients et à la chrono-compétition qui se développe depuis la fin des années 1990 visant à raccourcir les cycles de vie des produits/services et des projets de conception/production/diffusion/retrait, mais aussi à accélérer les flux physiques et d'information. En analogie avec la gestion de projet, les chaînes logistiques deviennent des *chemins critiques* où le moindre incident peut, par propagation, se répercuter jusqu'au client final. Pour faire face à ces changements, l'ambition affichée est de répondre au double objectif d'amélioration des niveaux de service et de réduction des coûts, en synchronisant des flux tout au long de la chaîne logistique.

Les travaux que nous avons développés s'intègrent dans le cadre de pilotage de flux inter-organisationnelle dans les chaînes logistiques multi-acteurs. Nous nous intéressons aux politiques de pilotage en flux tirés par les demandes finales, en particulier à la *politique de stock nominal* dont

l'application est appropriée dans les cas où les coûts de commande sont relativement faibles. Le pilotage des flux physiques au sein des réseaux d'entreprises reste fragile à cause du caractère aléatoire des variations dues au marché et aux partenaires commerciaux. Ainsi, l'analyse de ces caractéristiques conduit à étudier des modèles qui intègrent des représentations probabilistes des phénomènes aléatoires liés à l'offre et à la demande. En outre, les modèles proposés doivent de préférence être capables d'évoquer les problèmes liés aux ressources à capacité limitée des entreprises industrielles. Dans ce cadre, nous analysons les modèles simplifiés des chaînes logistiques à deux niveaux ayant des demandes finales aléatoires et des systèmes de fabrication à capacité limitée avec des aléas liés aux temps de fabrication. Afin de modéliser les systèmes étudiés, nous nous appuyons sur la théorie des files d'attente qui permet d'analyser analytiquement des effets de congestion provoqués par des aléas et des limites de capacité dans les réseaux de production de biens et de services.

Pour les chaînes logistiques étudiées, nos travaux visent à réduire globalement les stocks tout en gardant un niveau de service satisfaisant en termes de délai de livraison aux clients finaux. Dans ce but, nous adoptons deux approches.

Dans la première approche, nous analysons les effets des stratégies multi-fournisseurs sur les performances des chaînes logistiques. Malgré leurs avantages potentiels, les études théoriques sur l'implémentation des stratégies multi-fournisseurs sont peu fréquentes dans la littérature. Dans un environnement aléatoire où les variations des délais de livraison des fournisseurs ne sont pas prévisibles, le délai moyen d'approvisionnement et les coûts moyens de stockage et de rupture peuvent être réduits en adoptant une stratégie multi-fournisseurs.

Dans la deuxième approche, nous analysons des relations contractuelles entre les entreprises adjacentes des chaînes logistiques permettant aux entreprises d'améliorer le niveau de service tout en minimisant les coûts encourus par chacun des partenaires. Nous nous focalisons sur le caractère multi-acteurs des chaînes logistiques où chaque entreprise particulière a le but principal d'optimiser sa politique d'approvisionnement par rapport à ses critères locaux. Ceci se traduit par une optimisation individuelle souvent effectuée d'une façon concurrentielle. Nous parlons alors de la décentralisation des décisions conduisant souvent à une perte d'efficacité pour l'ensemble de la chaîne. Pour comprendre et maîtriser l'organisation des transactions entre partenaires d'une chaîne logistique, il est donc essentiel de représenter les antagonismes entre leurs objectifs économiques, ainsi que les éventuelles relations de dominance entre les entreprises concernées. L'outil mathématique privilégié pour cette analyse est la théorie des jeux. Dans ce contexte, nous analysons la gestion compétitive des stocks dans les chaînes logistiques à deux niveaux. L'objectif de ce travail est d'exploiter la capacité de prévision et d'anticipation de la théorie des jeux pour évaluer les performances globales de la chaîne logistique étudiée et leur dégradation due à la décentralisation des décisions. Nous utilisons la notion de

coordination qui consiste à élaborer des contrats entre les acteurs afin d'améliorer les performances du fonctionnement global.

L'exposé de nos travaux est organisé en quatre chapitres.

Le premier chapitre présente les concepts généraux de la gestion de chaînes logistiques et un état de l'art sur les principales politiques de pilotage de flux, plus particulièrement les politiques de gestion des stocks à point de commande, les politiques de stock nominal, les politiques Kanban et les politiques MRP. Nous décrivons le principe de fonctionnement de ces différentes politiques de pilotage de flux et mettons en évidence leurs similarités et leurs différences ainsi que leurs avantages et inconvénients dans un contexte multi-échelons. Ensuite, nous nous consacrons à la modélisation des systèmes de stock nominal avec capacité finie de production par le formalisme des files d'attente. Nous développons le modèle de base, nommé la file d'attente de production pour stock, pour un système de production et de stockage mono-étage et mono-produit. Nous reprenons le modèle de base présenté et nous l'étendons dans le troisième chapitre en considérant un système de production qui fonctionne comme un réseau ouvert des files d'attente en parallèle et dans le quatrième chapitre pour un système à deux étages de production et de stockage.

Dans le deuxième chapitre, nous présentons les concepts de base de la théorie des jeux et nous effectuons un état de l'art sur les applications de cette théorie dans le domaine de la gestion de chaînes logistiques. Nous exposons les problèmes qu'on peut rencontrer en raison du caractère distribué des décisions dans le contexte multi-acteurs des chaînes logistiques et les apports de la théorie des jeux dans différentes situations d'information et d'interaction entre les acteurs. Nous nous concentrons sur les applications des jeux non-coopératifs, plus particulièrement des jeux statiques, des jeux dynamiques, et des jeux avec information asymétriques.

Nous consacrons le troisième chapitre à nos contributions sur les stratégies multi-fournisseurs dans les chaînes logistiques. Nous analysons le problème d'approvisionnement d'une entreprise ayant une demande aléatoire d'un produit. Les fournisseurs disponibles pour l'approvisionnement de produit sont homogènes en termes de prix et de qualité mais hétérogènes en termes de taux moyen de production des commandes. L'entreprise a l'option d'envoyer chaque commande d'approvisionnement à un fournisseur différent. L'application d'une politique d'acheminement de commande probabiliste avec laquelle chaque commande est affectée à un des fournisseurs selon des probabilités fixées en avance est analysée. Nous fournissons l'expression analytique de l'espérance de la somme des coûts de stockage et de rupture en fonction des variables de décision, notamment le niveau de stock nominal et les probabilités d'affectation. Pour le cas de la production à la commande où le niveau de stock nominal de l'entreprise est nul, nous déterminons les valeurs optimales des probabilités d'affectation des commandes aux fournisseurs. Pour le

cas de la production pour stock, nous proposons une méthode approximative pour résoudre le problème d'optimisation.

Le quatrième chapitre expose nos contributions sur la coordination des chaînes logistiques décentralisées. Nous analysons un maillon élémentaire d'une chaîne logistique qui est constitué de deux étages de production et de stockage gérés par deux acteurs différents : un producteur ayant une demande aléatoire d'un produit fini et son fournisseur de produit intermédiaire. Chaque étage est confronté aux effets de congestion à cause du système de fabrication à capacité limité et des demandes et des temps de fabrication aléatoires. Nous utilisons les niveaux moyens de stock possédé et les niveaux moyens de rupture de stock des entreprises comme mesures de performances du système analysé. L'analyse exacte des mesures de performances est possible seulement dans des cas particuliers. Nous adoptons une méthode approximative issue de la littérature afin de développer des résultats analytiques. Dans cette chaîne logistique décentralisée, chaque entreprise est une entité individuelle qui décide de son niveau de stock nominal dans le but de maximiser son profit moyen. Nous analysons le résultat d'un jeu de Stackelberg entre les acteurs en supposant que le producteur est l'acteur dominant dans le jeu. Étant l'acteur dominant, le producteur propose un contrat à son fournisseur. Nous étudions l'application d'un contrat de coordination fixant le prix d'achat de produits intermédiaires et imposant une pénalité pour les livraisons retardées du fournisseur.

Finalement, nous présentons les conclusions obtenues pendant les études effectuées et les portes que ces travaux ont ouvertes pour la continuité des recherches sur le sujet.

CHAPITRE I

1. Pilotage de flux dans les chaînes logistiques

1.1. INTRODUCTION

Selon la terminologie actuelle, les chaînes logistiques coordonnent les séquences d'activités nécessaires au fonctionnement de toute une filière industrielle : l'approvisionnement en composants et matières premières, la production de produits semi-finis, l'assemblage de produits finis, la livraison aux clients. Les objectifs principaux dans les chaînes logistiques sont la compétitivité et la réactivité croissante pour faire face aux exigences des clients. Pour garantir un niveau de service satisfaisant vis-à-vis des clients, tous les flux physiques traversant les différents niveaux de la chaîne logistique doivent être pilotés efficacement. Dans ce contexte, différentes politiques de pilotage de flux ont vu le jour depuis plusieurs décennies.

Dans ce chapitre introductif, nous présentons un état de l'art des principales politiques de pilotage de flux étudiées dans la littérature et nous définissons les concepts utilisés lors de cette thèse. La deuxième section de ce chapitre est consacrée à la notion de gestion de chaînes logistiques. Nous présentons ensuite les principaux concepts liés au pilotage de flux dans les chaînes logistiques. Nous effectuons dans la quatrième section une description des principales politiques de pilotage de flux issues de la littérature. La cinquième section est consacrée à la modélisation des systèmes de pilotage de flux étudiés dans cette thèse par le formalisme des files d'attente.

1.2. GESTION DE CHAÎNES LOGISTIQUES

Jusqu'au milieu des années 1970, le produit a été le centre d'intérêt des entreprises industrielles. La tendance générale dans l'industrie a été de fournir des produits répondant aux spécifications des concepteurs, lesquelles étaient établies pour réaliser des fonctionnalités bien précises, et de pousser la production dans l'objectif d'inonder le marché. L'analyse des systèmes de production de bien et de service était conduite sur la base de regroupement d'activités fédérées par les départements (réapprovisionnement, gestion des stocks, comptabilité, fabrication, contrôle de la qualité, expédition, maintenance, etc.). Pour des raisons organisationnelles (périmètre de responsabilité lié aux départements) et intellectuelles (réduction de la complexité), les responsables d'activités analysaient et résolvaient les

problèmes concernant leurs activités de manière indépendante, sans se préoccuper des répercussions de ces décisions sur l'ensemble des activités de l'entreprise. Cette vision *verticale* a été considérée suffisante jusqu'aux années 1980.

La compétition mondiale s'est considérablement renforcée depuis le début des années 1980, dû aux progrès techniques et économiques comme la disparition de nombreuses frontières douanières, l'amélioration considérable des moyens de transport et de diffusion de l'information, et la dissémination des technologies et des connaissances. Cette concurrence intense et la saturation des marchés ont créé une économie de l'offre dont le but ultime est la satisfaction des clients. Afin de survivre, les entreprises se trouvaient dans l'obligation de fournir des produits plus variés et d'accentuer la notion de service (service après-vente, échange et remboursement, prise en compte des risques de vol ou de détérioration, livraison à domicile, formation de l'utilisateur, etc.) et de qualité tout en maintenant des prix compétitifs.

L'industrie a d'abord réagi par l'automatisation, gage de productivité et de régularité de la qualité. Même si l'automatisation s'est révélée très efficace pour la fabrication de masse, le niveau des investissements à consentir et la rigidité des systèmes de fabrication automatisés ont rapidement montré leurs limites face à la variabilité croissante de la demande des consommateurs et à l'évolution rapide des technologies. Désormais, les variations de la demande concernent la nature même des produits (fonctionnalités, technologies, esthétique etc.), réduisant ainsi considérablement la durée de vie des produits et exigeant des systèmes de production capables de s'adapter à ces variations en profondeur de la demande. En outre, la date de disponibilité de l'objet est devenue un nouvel attribut de la concurrence qui joue à la fois sur la rapidité de mise sur le marché de produits nouveaux et sur celle de livraison des commandes en produits existants. Ainsi, nous observons un changement dans les marchés vers des produits personnalisés exigeant des systèmes de fabrication *agiles* pour répondre à l'évolution des demandes et de la technologie. Ceci nécessite flexibilité des processus et coordination entre les sites. L'organisation classique qui accentue la fragmentation des processus et la spécialisation des acteurs induit des besoins croissants de coordination pour faire face à ce durcissement de la concurrence et aux exigences de la clientèle.

Entre les années 1975 et 1990, la plupart des entreprises ont commencé à cartographier les processus dans le but d'évaluer leur efficacité, sans changer l'organisation classique, centrée autour d'activités. L'industrie a ensuite réalisé les avantages de l'intégration des activités, aussi bien en conception de produits qu'en fabrication. À partir des années 1980, un mouvement s'appuyant sur une vision *horizontale* centrée sur le *processus* a fait son apparition. L'industrie a adopté les techniques ayant une vision processus comme les normes *ISO*, la *Qualité Totale*, et le *Juste-à-Temps*. Les études sur la coordination des unités organisationnelles ont débuté par les contributions sur l'effet de coup de fouet (*bullwhip effect*), la planification de production hiérarchisée, la gestion des stocks dans les réseaux de production/distribution, et la différenciation retardée.

Au début des années 1990, à l'organisation classique par départements autour de métiers s'est substitué un mode de fonctionnement par *réseau d'unités organisationnelles*, dans le but d'avoir une structure globale cohérente, capable de s'ajuster rapidement à la demande du client final. Cette démarche est fortement liée à la prise de conscience que les objectifs individuels des différentes unités organisationnelles peuvent conduire à une perte d'efficacité et nécessitent des mécanismes de coordination permettant d'améliorer les performances globales. Ce concept a donné naissance à la notion de *gestion de chaînes logistiques (supply chain management)* dont le but ultime est la satisfaction du consommateur résultant de la performance d'un enchaînement de processus à considérer dans leur ensemble et non de façon individuelle.

Cette modification de l'organisation n'a été rendue possible que grâce aux progrès de l'informatique et de la communication. Depuis le début des années 1990, les entreprises s'intéressent au dialogue entre les activités au travers de progiciels intégrés tels que les ERP (*Enterprise Resource Planning*). Les relations instantanées avec des fournisseurs offrant le meilleur prix sont alors remplacées par une vision du coût total depuis les sources d'un produit jusqu'à sa consommation. Les entreprises dépendent de plus en plus des processus en amont et en aval et accroissent les échanges d'information avec leurs fournisseurs et leurs clients. Les améliorations des moyens de communication informatisés (internet, intranet, réseaux locaux (LAN), réseaux métropolitains (MAN), réseaux grand distance (WAN), etc.) et des techniques d'échange électronique d'information (EDI : *Electronic Data Interchange*, XML : *Extensible Markup Language*, etc.) permettent désormais à un système d'information de communiquer avec un autre système d'information avec un minimum d'interventions humaines. Afin d'automatiser le partage d'information, les partenaires utilisent de plus en plus les plateformes du commerce électronique.

1.2.1. Externalisation et les chaînes logistiques

À sa naissance, la *gestion de chaînes logistiques* était un concept de gestion de l'entreprise. Les cadres dirigeants étaient considérés comme les seules capables de concilier les objectifs antagonistes des unités organisationnelles. Cependant, bien que la coordination des flux physiques, des flux d'information et des flux financiers au sein d'une entreprise de grande taille soit déjà une tâche ambitieuse, la coordination d'une chaîne logistique constituée de différentes entreprises est évidemment encore plus difficile. De nos jours, les gains potentiels de la coordination des unités organisationnelles, et de l'intégration des flux d'information et des efforts de planification tout au long des chaînes logistiques sont impressionnants. Pourtant, ces gains ne peuvent plus être accomplis au sein d'une seule entreprise car les entreprises se focalisent plus en plus sur le cœur de leurs compétences en externalisant la plupart des activités, y compris les activités de support (gestion du personnel, gestion du système d'information, transport, etc.).

L'*externalisation (outsourcing)* est la délégation sur une période pluriannuelle de la gestion d'une ou de plusieurs fonctions de l'entreprise à un prestataire extérieur. L'externalisation diffère des pratiques de

sous-traitance (*subcontracting*). La sous-traitance est utilisée lorsqu'une entreprise (donneur d'ordre) confie une ou plusieurs fonctions, qu'elle est pourtant capable de réaliser en interne, à une autre entreprise (sous-traitant ou preneur d'ordres) pour des raisons diverses comme, par exemple, une surcharge ponctuelle de travail, une indisponibilité de machine ou de personnel qualifié. En revanche, l'externalisation est un modèle de gestion stratégique où l'entreprise ne dispose pas de compétences nécessaires pour réaliser les fonctions déléguées aux prestataires. Les décisions concernant l'externalisation sont parmi les plus stratégiques car elles déterminent la structure organisationnelle de l'entreprise en définissant les fonctions pour lesquelles les compétences nécessaires sont à développer et les fonctions qui sont à acheter.

L'externalisation permet aux entreprises de se focaliser sur le cœur de leur métier, de réduire les coûts fixes et d'améliorer la qualité (Chopra et Meindl, 2007). En effet, les impératifs croissants de réactivité et de compétitivité par les coûts et la qualité ne permettent plus aux entreprises la détention en leur sein de toutes les compétences requises. En outre, la complexité de certaines productions est telle qu'aucune entreprise n'est en mesure d'assurer seule la maîtrise d'œuvre de l'opération pour aboutir, en temps utile et à un coût acceptable, à un résultat technique satisfaisant, compte tenu de l'état de la concurrence internationale. Pour ces raisons, on assiste à une montée en puissance d'alliances plus ou moins stable conduisant à la création de *réseaux d'entreprises*, définis comme une structure flexible et adaptative mobilisant un ensemble coordonné et stabilisé de compétences (Stadtler, 2005).

Par conséquent, les caractéristiques et la qualité d'un produit ou d'un service vendu aux clients dépendent largement des différentes entreprises qui contribuent à sa création. Ceci pose de nouveaux défis managériaux, en intégration d'entreprises juridiquement séparées et en coordination de divers flux au delà des frontières juridiques des entreprises.

1.2.2. Définition de la chaîne logistique

L'objet de la gestion de chaînes logistiques est évidemment la *chaîne logistique* (*supply chain*). Une chaîne logistique est un réseau d'organisations qui contribuent aux différents processus et activités, à travers les interactions en amont et en aval, apportant une valeur ajoutée sous la forme de produits et de services pour les clients finaux. D'un point de vue conceptuel, une chaîne logistique peut être considérée comme une succession de processus d'approvisionnement, de fabrication, de distribution et de vente d'un produit, depuis le premier des fournisseurs jusqu'au client final. Une chaîne logistique est donc constituée de fournisseurs, de centres de production, d'entrepôts de stockage, de centres de distribution et de points de vente, le tout traversé par un flux physique qui transforme progressivement les matières premières et composantes en produits finis.

Au sens large, une chaîne logistique est constituée de deux ou plusieurs organisations juridiquement séparées, par ailleurs liées par des flux physiques, des flux d'information correspondants aux échanges d'information entre les organisations et des flux financiers ou monétaires associés aux flux physiques. Ces organisations peuvent être les fournisseurs de matières premières, les entreprises fabriquant les composants, les produits intermédiaires ou les produits finis, les prestataires de services logistiques, et même le client final. Au sens strict, le terme chaîne logistique est utilisé pour les entreprises de grande taille ayant souvent des sites géographiquement séparés. Une chaîne logistique au sens large est aussi appelée chaîne logistique *inter-organisationnelle*, tandis que le terme *intra-organisationnelle* se réfère aux chaînes logistiques au sens stricte.

1.2.3. Définition de la gestion de chaînes logistiques

La gestion de chaînes logistiques reprend l'idée que la satisfaction d'un client est le résultat de la mise en œuvre d'une succession de processus, qu'il s'agit de prendre en considération dans une approche systématique, sans trop se préoccuper du périmètre juridique de l'entreprise, en remontant, si nécessaire, jusqu'à l'approvisionnement des matières premières (Giard, 2003). Le but de la gestion de tous les efforts au sein d'une chaîne logistique est la compétitivité croissante. Il faut maintenant être capable de mettre sur le marché, avant les concurrents, des produits de bonne qualité, de faible coût, et qui répondent aux désirs exprimés ou latents des clients. Il faut aussi être capable d'accompagner le client tout au long de la vie de produits, et en recherche permanente d'innovation dans la production et les services.

Selon la définition de Smichi-Levi *et al.* (2003), la *gestion de chaînes logistiques* consiste à coordonner efficacement les fournisseurs, les producteurs, les entrepôts et les détaillants afin de produire et distribuer les produits en bonne quantité, au bon endroit et au bon moment, et de minimiser le coût global, tout en obtenant un niveau de service suffisant. Les outils nécessaires à cette coordination relèvent de la recherche opérationnelle et empruntent aux techniques des systèmes d'information et de communication.

1.3. PILOTAGE DE FLUX DANS LES CHAÎNES LOGISTIQUES

De manière traditionnelle, les décisions de gestion de chaînes logistiques (et les problématiques scientifiques correspondantes) sont classées par niveau, chaque niveau concerne son horizon de temps et il a son degré de précision des données. Une décomposition de ce type a les trois niveaux suivants : le niveau stratégique, le niveau tactique et le niveau opérationnel. Les décisions opérationnelles assurent la flexibilité pour faire face aux fluctuations de la demande et des disponibilités de ressources à court et à très court terme. Ainsi, le niveau opérationnel regroupe l'ensemble de décisions de pilotage de flux à court terme telles que les décisions de lancement des ordres d'approvisionnement, de fabrication ou d'assemblage et l'ensemble de décisions d'ordonnancement à très court terme consistant à organiser la production au sein des ateliers et à gérer l'affectation des opérations sur les postes de travail. Par la suite,

nous nous intéressons à ce troisième niveau décisionnel, plus particulièrement au pilotage de flux physiques dans les chaînes logistiques.

La notion de pilotage de flux a suivi l'évolution de la notion de chaîne logistique à travers le temps. Le pilotage de flux se limitait au début, à l'ensemble des règles de gestion des stocks. Par la suite, il a évolué pour intégrer plusieurs caractéristiques endogènes des systèmes de production comme les besoins de coordination des différents flux physiques et les capacités de production limitées. Actuellement, cette notion s'étend de plus en plus pour englober toute la chaîne logistique depuis l'approvisionnement en matières premières jusqu'à la distribution aux clients finaux. Le pilotage de flux consiste aujourd'hui à coordonner tous les flux physiques traversant les différents niveaux de la chaîne logistique dans l'objectif de garantir un certain niveau de service vis-à-vis du client tout en minimisant les coûts.

D'un point de vue pratique, piloter les flux dans la chaîne logistique consiste à prendre des décisions, à chaque étape de la chaîne (depuis les fournisseurs jusqu'au client final) et pour chaque entité (matière première, composant, produit intermédiaire ou produit fini), permettant de déterminer quand et en quelle quantité lancer une activité telle que l'activité d'approvisionnement, de transport, de fabrication ou d'assemblage. La Figure 1.1 visualise des flux physiques résultant de décisions de pilotage de flux dans une chaîne logistique. Généralement, les décisions de pilotage de flux sont concrétisées par des ordres d'approvisionnement, de fabrication, ou d'assemblage. Les entreprises organisées en chaîne logistique ou les ateliers de la même entreprise communiquent en permanence à travers des décisions de passation d'ordres.

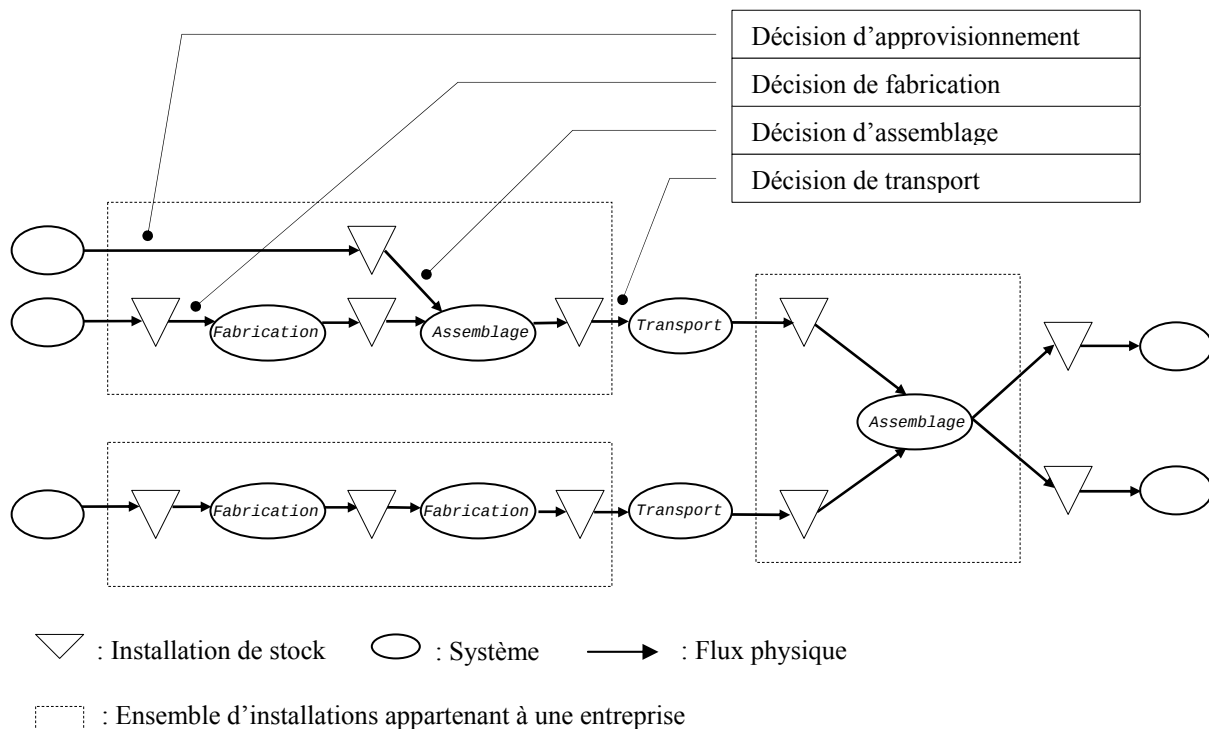


Figure 1.1. Pilotage de flux dans une chaîne logistique

1.3.1. Le système élémentaire de flux

Le mécanisme de décision de passation d'ordres nécessite l'information sur le système d'approvisionnement, sur la demande et l'état du stock de chaque produit considéré. Nous définissons un système élémentaire de flux comme un système constitué du système d'approvisionnement, de l'installation de stock et du système de demande d'un produit (Figure 1.2). Par la suite, nous présentons les différentes caractéristiques des systèmes élémentaires de flux dont les analyses sont essentielles pour la prise de décision de pilotage de flux.

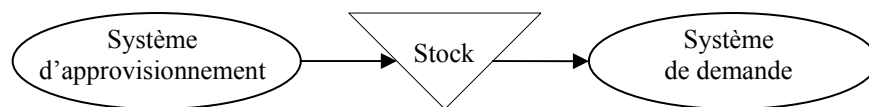


Figure 1.2. Système élémentaire de flux

1.3.1.1. Le système d'approvisionnement

Pour un produit donné (matière première, composant, produit intermédiaire, produit fini, etc.), l'origine de l'approvisionnement peut être interne à l'entreprise (système d'assemblage ou de fabrication) ou externe (fournisseur externe). Un paramètre important de l'approvisionnement est le délai d'approvisionnement (délai d'obtention), c'est-à-dire le temps qui s'écoule entre le moment où une commande (une commande d'approvisionnement dans le cas d'approvisionnement externe et un ordre de fabrication ou d'assemblage dans le cas d'approvisionnement interne), et celui où les unités demandées sont effectivement disponibles. Ce délai d'approvisionnement comporte des temps de lancement de production et de production (à l'intérieur de l'entreprise ou non), de transport, de réception, etc.

Le système d'approvisionnement peut être considéré comme un processus exogène ou endogène. Dans le cas d'un système d'approvisionnement exogène, le système d'approvisionnement est à capacité illimitée et la charge du système n'affecte pas le délai d'approvisionnement. Dans le cas d'un système d'approvisionnement endogène, le système d'approvisionnement est à capacité limitée et le délai d'approvisionnement comporte aussi les délais (les temps d'attente) liés à la charge et la capacité du système d'approvisionnement. Dans la littérature, un système avec un processus d'approvisionnement exogène est souvent appelé *système de stockage*. En générale, les systèmes de transport sont traités comme des systèmes d'approvisionnement exogène. Quand le système d'approvisionnement est supposé endogène, le système étudié est souvent appelé *système de production et de stockage* (référéncé comme un *système de production/stockage*).

Les systèmes d'approvisionnement sont souvent confrontés aux aléas qui peuvent être dûs par exemple, aux pannes des machines, indisponibilités des opérateurs, problèmes de transport, etc. Les différents aléas dans les systèmes d'approvisionnement sont souvent traduits sous forme d'incertitude sur le délai

d'approvisionnement et/ou la quantité approvisionnée. L'incertitude sur la quantité approvisionnée est souvent liée aux produits défectueux qui résultent par exemple des problèmes dans le processus de production.

Les coûts associés à l'approvisionnement d'un produit sont constitués des prix d'achat (dans le cas d'approvisionnement externe) et des coûts de commande. Dans le cas d'un approvisionnement externe, le coût de commande comporte des frais d'administration, de transport, de réception, etc. Dans le cas d'un approvisionnement interne, il s'agit le coût de lancement de la production (les coûts de réglage, les rebuts de ces réglages, les coûts de gestion de l'ordre de fabrication, etc.). Le coût de commande est considéré indépendant de la quantité commandée. Les économies d'échelle jouent un rôle important dans la prise de décision de pilotage de flux. L'évaluation des coûts de commande peut conduire à la mise en production ou à l'achat de certains produits par lots de plusieurs unités plutôt que par unité. Dans le cas d'approvisionnement externe, l'achat par lots de grande quantité peut aussi être incité par les rabais sur quantité.

1.3.1.2. Le système de demande

Les demandes d'un produit peuvent provenir de l'extérieur de l'entreprise (client externe) ou de l'intérieur (système d'assemblage ou de fabrication en aval). Dans le cas de demande interne, la demande est souvent considérée dépendante, c'est-à-dire fonction de la demande pour d'autres produits (typiquement les produits finis vendus sur le marché). En pratique, la demande est souvent considérée comme un processus exogène, ce qui veut dire que l'état du système étudié n'influence pas les demandes.

L'information sur les demandes futures peut être obtenue directement sous la forme de commandes fermes. Les commandes fermes représentent une information fiable sur la demande, tant sur les quantités que sur les dates. Les commandes fermes sont définitives et ne comportent pas d'incertitude, c'est-à-dire qu'elles représentent un engagement de la part des clients. En l'absence des commandes fermes, l'incertitude sur la demande peut porter sur les quantités de demande et/ou sur les dates de demande.

L'information sur les demandes futures peut aussi être obtenue sous formes de prévisions. Les prévisions représentent une information incertaine sur la demande. Elles peuvent provenir des commandes prévisionnelles et/ou des prévisions de la demande. Les commandes prévisionnelles émanent des clients finaux mais, contrairement aux commandes fermes, elles contiennent une incertitude sur la quantité et/ou sur la date de la demande. Les prévisions de la demande sont des données provenant d'études prospectives de marketing ou obtenues à l'aide d'autres méthodes qualitatives ou quantitatives de prévision basées sur des historiques de ventes.

L'information sur les demandes futures (sous forme de commandes fermes ou de prévisions) est moins fiable quand on s'éloigne dans le temps. S'il n'y a pas d'information disponible sur les demandes futures,

la demande est souvent modélisée en utilisant des lois statistiques construites soit à partir d'informations sur le passé soit à partir de probabilités à priori.

Le mode de production du produit influence la nature de l'information sur la demande. Nous pouvons classer les modes de production en deux types : production à la commande (*make-to-order*) et production pour stock (*make-to-stock*). Dans le cas de production à la commande, on ne constitue pas de stock de sortie à l'avance. Tout ou partie de la fabrication (et/ou de l'assemblage) est déclenché afin de satisfaire une commande ferme ou une commande actuelle. Les systèmes de production à la commande sont souvent divisés en deux sous-types : assemblage à la commande (*assemble-to-order*) et fabrication à la commande (*engineer-to-order*). On parle d'assemblage à la commande lorsque les composants existants (fabriqués pour stock) sont assemblés pour exécuter un produit en réponse à une commande ferme. On parle de fabrication à la commande quand, en réponse à une commande ferme, il faut exécuter un travail de conception pouvant ou non nécessiter la création de nouveaux composants. Dans le cas de la production pour stock, on constitue un stock à partir duquel les clients vont être servis. Dans ce cas, on n'a pas d'information fiable sur les demandes futures. La production pour stock est déclenchée soit pour renouveler la consommation du stock soit pour satisfaire les demandes anticipées. La production pour stock est un mode de production souvent utilisé pour les produits standardisés.

En fonction de l'information disponible, la demande est considérée comme déterministe ou aléatoire. Les demandes au fil de temps peuvent être stationnaires ou non-stationnaires. Dans le cas de la demande stationnaire, les caractéristiques de la demande sont les mêmes dans le temps. Si de plus la demande est déterministe, son niveau est constant. Si la demande est aléatoire, la loi suivie est la même et conserve les mêmes valeurs pour ses paramètres caractéristiques. Dans le cas d'une demande non-stationnaire, les caractéristiques de la demande évoluent au cours du temps pour une raison quelconque (saisonnalité de la demande, etc.).

1.3.1.3. Le stock

Dans un système élémentaire de flux, les stocks correspondent aux produits immobilisés, en attente de transfert ou de traitement. Les raisons pour constituer des stocks sont nombreuses : absorber le délai d'approvisionnement, lisser les charges d'un système d'approvisionnement à capacité limitée face aux saisonnalités de la demande (stocks saisonniers), faire face aux incertitudes liées au processus d'approvisionnement et/ou au processus de demande (stocks de sécurité), utiliser les économies d'échelle liées au système d'approvisionnement (stock de cycle), etc.

Le but des politiques de pilotage de flux est de maîtriser le stock dans l'espace et dans le temps, de façon à obtenir le meilleur compromis entre les coûts liés à sa présence et ceux résultants de son insuffisance. La présence de stocks engendre des coûts de stockage liés à l'immobilisation d'un capital, à l'occupation

d'un espace de stockage, aux équipements et aux frais permettant d'assurer le stockage dans de bonnes conditions. L'insuffisance du stock provoque une rupture de stock qui qualifie le fait qu'un stock n'est pas en mesure de satisfaire immédiatement la totalité de la demande qui s'adresse à lui. Les conséquences de cette rupture sont différentes selon que la demande est interne ou externe. En cas de demande externe, la demande non satisfaite peut être perdue (on parle de *ventes manquées*) ou reportée (on parle de *demandes retardées*). Dans le cas de ventes manquées, le coût associé est le manque à gagner de non fourniture d'une unité, généralement la marge bénéficiaire. Dans le cas de demandes retardées, le coût de rupture n'inclut pas la marge car la vente sera réalisée plus tard. Ce coût de rupture est le coût administratif d'ouverture d'un dossier et le coût commercial correspondant à une ristourne accordée pour livraisons tardives. Dans le cas de demande interne, la rupture entraîne un chômage technique des postes en aval et le coût de rupture correspond au coût financier de ce chômage technique. Notons qu'il est souvent difficile d'estimer le coût de rupture de manière fiable. Dans ces cas, il est courant de modéliser l'objectif en termes de qualité de service sous forme d'une contrainte sur le niveau de service.

Dans un système de stockage, la quantité en stock peut être contrôlée à tout instant par la technique de l'inventaire permanent (*continuous review*) ou de façon périodique par la technique de l'inventaire périodique (*periodic review*). Le système d'information à inventaire permanent est, en général, plus coûteux que celui à inventaire périodique mais permet de réagir instantanément à des situations inattendues. L'inventaire périodique permet de détecter des détériorations, des erreurs, ou des vols en particulier dans le cas de produits à rotation très lente. Le système d'information approprié est installé selon les caractéristiques du système de stockage et de la politique de pilotage de flux appliquée.

La mesure de l'état du stock utilisée varie aussi selon les caractéristiques du système de stockage et de la politique de pilotage de flux appliquée. Dans le cas de demandes retardées, l'état du stock est caractérisé par le niveau de stock qui représente la quantité nette en stock à un moment donné :

Niveau de stock à l'instant t = niveau de stock possédé à l'instant t

– niveau de rupture de stock à l'instant t

Le niveau de stock possédé correspond à la quantité en stock, c'est-à-dire le nombre d'unités physiquement présentes en stock. Le niveau de rupture de stock correspond aux demandes retardées, c'est-à-dire les unités demandées mais qui ne sont pas encore livrées. Le niveau de stock représente donc le stock utilisable pour satisfaire les demandes ultérieures à l'instant t .

Dans certaines situations, les décisions d'approvisionnement ne peuvent pas être basées uniquement sur le niveau de stock. Il peut être nécessaire de prendre en compte les approvisionnements attendus correspondant aux commandes d'approvisionnement ou aux ordres de fabrication déjà passés dont la

livraison est encore attendue à l'instant t . Dans ces cas, l'état du stock est caractérisé par la position de stock :

Position de stock à l'instant t = niveau de stock à l'instant t + commandes attendues à l'instant t

1.4. POLITIQUES DE PILOTAGE DE FLUX

Dans cette section, nous présentons les politiques de pilotage de flux en les classifiant en deux familles : les politiques réactives et les politiques proactives. Les politiques réactives sont des politiques de pilotage de flux qui déclenchent un ordre d'approvisionnement, de fabrication ou d'assemblage à l'arrivée d'une demande. Dans le cas où le système comprend un stock de sortie, un ordre est déclenché pour renouveler la consommation de ce stock. Dans le cas où on ne dispose pas d'un stock de sortie, un ordre est déclenché afin de satisfaire la demande réalisée. Les politiques réactives réagissent donc aux réalisations des demandes. Les politiques réactives sont notamment les politiques d'approvisionnement et de stockage, les politiques de stock nominal et les politiques Kanban. Les politiques proactives, notamment les politiques du type MRP, déclenchent les ordres d'approvisionnement, de fabrication ou d'assemblage afin de satisfaire les demandes futures. Les politiques proactives doivent être utilisées, généralement, dans le cas où on dispose d'information sur les demandes futures. En revanche, si on ne dispose pas d'information sur les demandes futures, on optera pour une politique réactive.

1.4.1. Politiques d'approvisionnement et de stockage

Nous appelons politiques d'approvisionnement et de stockage les premières politiques de gestion des stocks développées depuis les années 30. Le type de régulation adopté pour déterminer quand et en quelle quantité il faut commander différencie fondamentalement les politiques utilisées. Trois principaux types de régulation peuvent être adoptés pour déterminer quand déclencher le réapprovisionnement du stock :

1. Gestion calendaire : Le réapprovisionnement du stock se fait à intervalles réguliers de périodes T . En pratique, cette période est souvent un nombre fixé de jours, de semaines, voir de mois. Ce type de gestion se rencontre dans les systèmes à inventaire périodique. L'intervalle de commande coïncide typiquement avec celui de l'inventaire périodique.

2. Gestion à point de commande : Le réapprovisionnement du stock est déclenché lorsque la position de stock devient inférieure ou égale à un niveau s appelé le point de commande (le point de commande s peut aussi être représenté par la notation R). Ce type de gestion se rencontre dans les systèmes à inventaire permanent.

3. Gestion calendaire conditionnelle : Ce dernier cas de figure mélange les deux techniques précédentes. Le réapprovisionnement du stock est déclenché si au terme d'un temps T la position de stock

devient inférieure ou égale au point de commande s . Ce type de gestion se rencontre dans les systèmes à inventaire périodique où le coût de commande est relativement important par rapport aux autres coûts.

La gestion calendaire permet de regrouper les commandes par fournisseur, ce qui peut réduire les coûts de commande et de transport. Par contre, la gestion calendaire est aveugle à l'intérieur d'une période. Elle rend le système insensible à des situations inattendues se produisant entre deux instants de commande, et donc nécessite des stocks de sécurité élevés. La gestion à point de commande permet de mieux contrôler le niveau de stock, et donc de diminuer le niveau de rupture tout en maintenant des stocks de sécurité plus faibles.

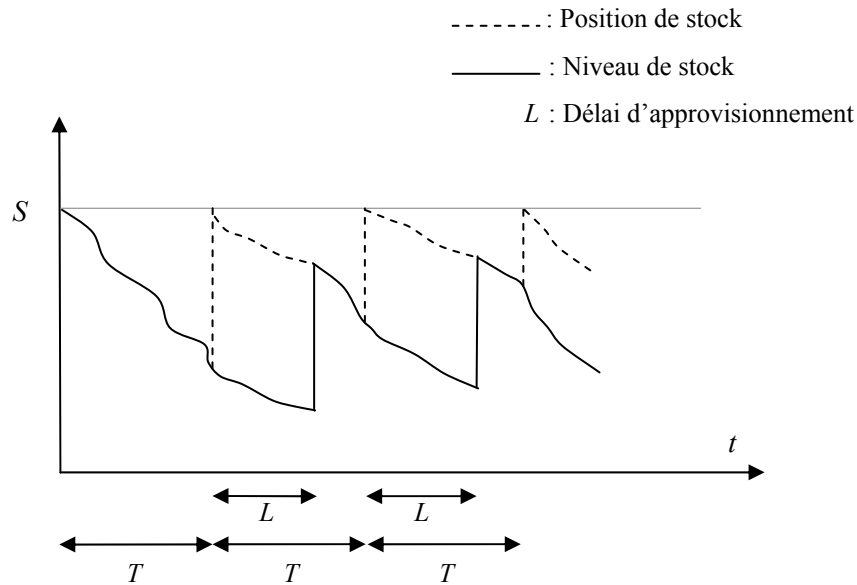
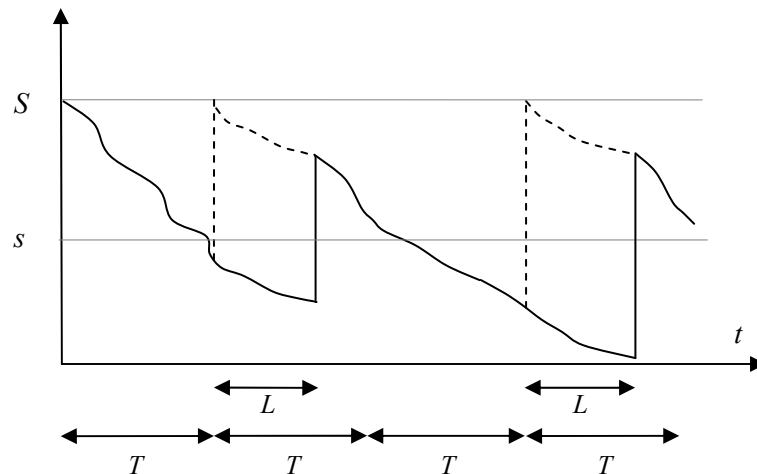
Les types de régulation utilisés pour déterminer en quelle quantité il faut commander sont :

1. Quantité fixe de commande : La commande porte sur une quantité Q fixée à l'avance par calcul ou par règle empirique.

2. Niveau de reapprovisionnement : La quantité commandée est égale à la différence entre le niveau de stock S appelé le niveau de reapprovisionnement et la position de stock.

N'importe quel type de régulation répondant à la question « *quand* » peut être combinée avec n'importe quel type de régulation répondant à la question « *combien* ». En pratique, certaines combinaisons présentent des avantages sur les autres et s'imposent donc plus fréquemment.

La politique de gestion calendaire à niveau de reapprovisionnement, notée (T, S) , et la politique de gestion calendaire conditionnelle à niveau de reapprovisionnement, notée (T, s, S) , sont les politiques classiques utilisées dans les systèmes à inventaire périodique. La politique (T, S) (Figure 1.3) peut déclencher des commandes en petites quantités car elle déclenche une commande même si, au moment de déclenchement de commande, la différence entre le niveau de reapprovisionnement et la position de stock est très petite. La politique (T, s, S) (Figure 1.4) permet d'éviter le déclenchement d'une commande de trop petite taille si la demande pendant la période a été très faible.

Figure 1.3. Évolution du stock avec la politique (T, S) Figure 1.4. Évolution du stock avec la politique (T, s, S)

Pour les systèmes à inventaire permanent, les deux politiques les plus souvent utilisées sont : la politique à point de commande et quantité fixe de commande, notée (R, Q) , et la politique à point de commande et niveau de reapprovisionnement, notée (s, S) . En appliquant la politique (R, Q) , si les demandes sont d'une unité à chaque fois, la position de stock devient égale exactement au point de commande R à chaque déclenchement de commande (Figure 1.5). Sinon, la position de stock est souvent inférieure au point de commande R au moment du déclenchement de commande. Dans ce cas, le position de stock n'atteint plus le niveau $R + Q$. Si la position de stock est suffisamment basse au moment du déclenchement d'une commande, il peut être nécessaire de commander plus qu'un lot afin de ramener la position de stock au dessus du point de commande R et de telle sorte que la position de stock résultant ne dépasse pas le niveau $R + Q$. Pour cette raison, cette politique est parfois notée comme (R, nQ) (avec $n = 1, 2, \dots$).

Si les demandes sont d'une unité à chaque fois, la politique (s, S) est équivalente à la politique (R, Q) avec $s = R$ et $S = R + Q$. Si les demandes ne sont pas unitaires, la quantité commandée pour reconstituer le niveau de stock à S est variable (Figure 1.6). En appliquant la politique (s, S) , la position de stock atteint le niveau S à chaque déclenchement de réapprovisionnement du stock. Pour la plupart des systèmes mono-étage ayant des demandes aléatoires et stationnaires, la politique optimale par rapport à un critère de minimisation des coûts moyens est en effet du type (s, S) . L'inconvénient de la politique (s, S) est la complexité de la procédure de détermination des valeurs optimales de ces paramètres s et S . En pratique, il est souvent plus facile d'utiliser une politique (R, Q) avec une taille de lot fixe.

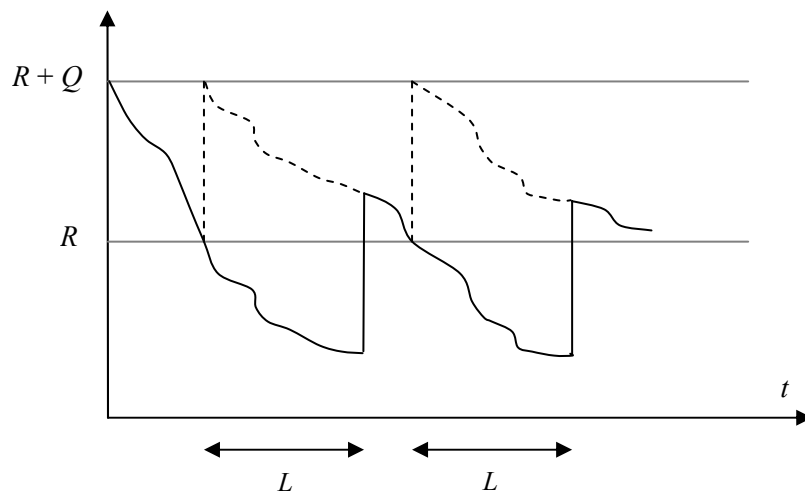


Figure 1.5. Évolution du stock avec la politique (R, Q)

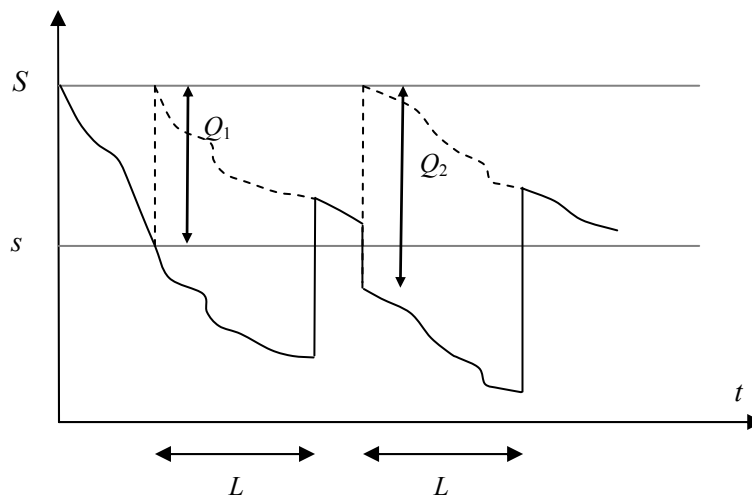


Figure 1.6. Évolution du stock avec la politique (s, S)

Ayant choisi une politique, le problème est de fixer les valeurs de ses paramètres de contrôle (c'est-à-dire les valeurs de s, S, R, Q, T) de façon à minimiser la fonction de coût considérée. La littérature sur cette problématique se focalise sur les performances des systèmes avec demande aléatoire stationnaire et délai

d'approvisionnement fixe. Quand le délai d'approvisionnement est considéré fixe, on sous-entend que le système d'approvisionnement est à capacité illimité car la charge du système n'affecte pas le délai d'approvisionnement. Quelques résultats de base sont fournis par Rao (2003) pour la politique (T, S) , Federgruen et Zheng (1992) pour la politique (R, Q) , et Zheng et Federgruen (1991) pour la politique (s, S) .

Dans la littérature, il existe des travaux analysant le cas de délai d'approvisionnement aléatoire quand le système d'approvisionnement est à capacité illimitée. Dans ces travaux, le processus d'approvisionnement est supposé soit exogène et séquentiel (*exogenous sequential supply*) soit exogène et parallèle (*exogenous parallel supply*). Comme décrit par Sovorons et Zipkin (1991) et Song (1994), dans le cas d'un processus *séquentiel*, les ordres sont exécutés dans la séquence même où ils ont été générés, ce qui crée une dépendance entre les délais d'approvisionnement successifs. Dans le cas d'un processus *parallèle*, le système d'approvisionnement est souvent modélisé comme un système de file d'attente avec un nombre infini de serveurs. Dans ce cas, les délais d'approvisionnement successifs sont indépendants. Un système d'approvisionnement endogène confronté aux aléas est souvent modélisé comme un système de file d'attente avec un nombre fini de serveurs. Une bonne discussion sur les modèles utilisés dans le cas de délai d'approvisionnement aléatoire est présentée par Zipkin (2000).

1.4.1.1. Politiques d'approvisionnement et de stockage dans les systèmes multi-étages

Les politiques d'approvisionnement et de stockage citées sont utilisées pour gérer des stocks mono-étages ainsi que multi-étages. Dans la littérature, les systèmes multi-étages sont souvent classifiés en trois groupes élémentaires : systèmes à structure linéaire, systèmes de distribution et systèmes d'assemblage (Figure 1.7).

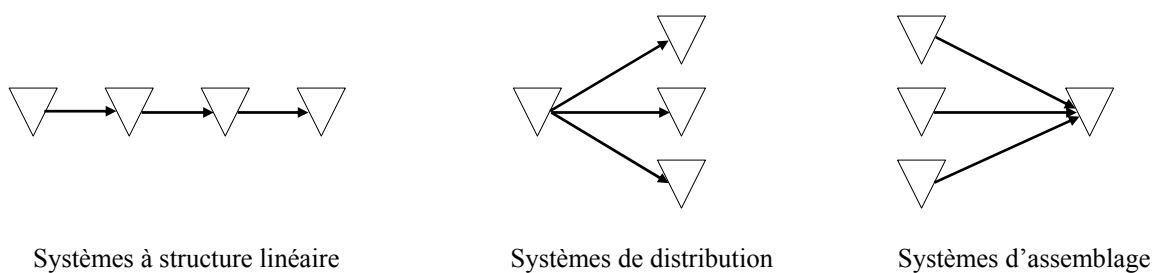


Figure 1.7. Systèmes multi-étages

Dans les systèmes à structure linéaire, chaque étage correspond à l'installation de stock d'un produit ayant subi une transformation (fabrication et/ou transport). Les étages peuvent représenter les ateliers de la même entreprise ou les différentes entreprises au long de la chaîne logistique. Chaque étage fonctionne comme un système de demande pour l'étage en amont et un système d'approvisionnement pour l'étage en aval. Le délai d'approvisionnement d'un étage comporte souvent des délais liés au système de transformation en amont. Des délais supplémentaires peuvent être provoqués si l'installation de stock en

amont n'est pas en mesure de satisfaire immédiatement la totalité d'une demande qui s'adresse à lui. Le système de demande du dernier étage (système de demande finale) est souvent considéré exogène au réseau logistique analysé. Les systèmes de distribution généralisent cette structure linéaire à une structure divergente. En termes de production, le premier étage correspond par exemple à l'installation de stock d'une matière première qui est ensuite transformée en plusieurs produits. En termes de transport, le premier étage correspond à un entrepôt et les étages suivants aux détaillants. Les systèmes d'assemblage généralisent la structure linéaire à une structure convergente. Dans ces types de systèmes, plusieurs composants sont assemblés pour fabriquer un produit. Notons que les réseaux logistiques complexes peuvent être représentés en combinant les différentes caractéristiques des systèmes élémentaires présentés.

Les politiques du type (R, Q) sont souvent utilisées pour gérer les systèmes multi-étages où les coûts de commande sont relativement importants. Sovorons et Zipkin (1988) font partie des premiers à avoir étudié un système de distribution à deux étages, constitué d'un entrepôt et de n détaillants identiques, dans lequel chaque installation de stock est gérée par une politique du type (R, Q) . La demande arrive chez chaque détaillant selon un processus de Poisson. Les auteurs proposent des approximations pour les mesures de performances du système, plus précisément pour les espérances mathématiques du niveau de stock possédé et du niveau de rupture de stock de chaque installation, en supposant des délais de transport fixes. Axsater (1993) analyse le même système de distribution ainsi qu'un système à structure linéaire ayant un seul détaillant et présente des évaluations exactes des mesures de performances. Axsater (2000b) généralise ces résultats pour un système ayant n détaillants non-identiques et dans lequel la demande arrive chez chaque détaillant selon un processus de Poisson composé. Cachon (2001a) expose des évaluations exactes des mesures de performances d'un système de distribution où chaque installation est gérée par une politique du type (R, Q) à inventaire périodique, autrement dit, par une politique du type $(T=1, R, Q)$.

Dans la littérature, il existe d'autres politiques utilisées pour gérer des stocks multi-étages. Axsater et Rosling (1993) montrent que les politiques utilisées pour gérer les stocks multi-étages peuvent être classées en deux types : les politiques du type « installation » et les politiques du type « échelon ». Cette différenciation dépend de la définition de la position de stock utilisée dans chaque installation de stock. Dans les politiques du type installation, la position de stock est définie pour chaque installation comme étant la position de stock définie classiquement pour l'installation en question. Dans les politiques du type échelon, la position de stock d'une installation est définie comme étant la somme des approvisionnements attendus de l'installation en question, du nombre de produits présents dans l'installation en question et dans tous les installations en aval (y compris les produits en transit dans les systèmes de transformation) moins le nombre de demandes finales retardées. Cette notion du stock échelon a été introduite par Clark et Scarf (1960). L'idée de base des politiques du type échelon est que si les installations en aval ont des

niveaux de stock faibles, alors elles vont passer des commandes d'approvisionnement dans le futur et donc l'installation en question peut avoir besoin davantage de produits en stock. D'autre part, si les installations en aval ont des niveaux de stock élevés, alors l'installation en question n'a pas besoin d'un réapprovisionnement immédiat.

Considérons un système à structure linéaire ayant des demandes finales aléatoires et dans lequel chaque étage est géré par une politique (R, Q) du type installation. Dans un tel système, un étage déclenche le réapprovisionnement du stock seulement si l'étage en aval vient de passer une commande d'approvisionnement. Ce n'est pas toujours le cas dans un système de stock échelon, car la position de stock échelon d'un étage n'est pas diminuée par les commandes de réapprovisionnement de l'étage en aval mais par les demandes finales. Pour ces raisons, une politique (R, Q) du type installation donnée peut toujours être remplacée par une politique (R, Q) du type échelon qui déclenche les commandes de réapprovisionnement en même temps et conserve la même évolution du stock pour chaque étage. Par contre, une politique (R, Q) du type échelon, qui peut déclencher le réapprovisionnement du stock d'un étage sans qu'il y a une commande venant de l'étage en aval, ne peut pas être remplacée par une politique (R, Q) du type installation équivalente (Axsater, 2000a ; Axsater et Rosling, 1993).

Ces résultats montrent que les politiques (R, Q) du type installation peuvent être considérés comme un sous-ensemble des politiques (R, Q) du type échelon. Pour un critère de performance donné, une optimisation au niveau des politiques du type échelon donne un résultat au moins autant performant que celui d'une optimisation au niveau des politiques du type installation (Axsater, 2003). Pour les cas où le critère est la minimisation des coûts de stockage et de rupture, la meilleure politique (R, Q) du type échelon est en général plus performante que la meilleure politique (R, Q) du type installation. Chen (2000) fournit plus de résultats sur l'optimalité des politiques (R, Q) du type échelon. Dans le but de déterminer la meilleure politique (R, Q) du type échelon, Chen et Zheng (1994a) présentent des évaluations exactes des mesures de performances des systèmes à structure linéaire ayant des demandes finales aléatoires.

Ces observations peuvent être généralisés pour les systèmes d'assemblage (Axsater et Rosling, 1993 ; Chen, 2000) mais pas pour les systèmes de distribution. Dans un système de distribution, la meilleure politique (R, Q) du type échelon peut être plus performante que la meilleure politique (R, Q) du type installation, ou contrairement la meilleure politique (R, Q) du type installation peut être plus performante que la meilleure politique (R, Q) du type échelon (Axsater et Juntti, 1996). Il existe des travaux cherchant à déterminer la meilleure politique (R, Q) du type échelon pour les systèmes de distribution ayant des demandes aléatoires. Chen et Zheng (1997) fournissent des évaluations exactes des espérances mathématiques des niveaux de stock possédés et des niveaux de rupture de stock d'un système de distribution ayant n détaillants identiques et dans lequel la demande arrive chez chaque détaillant selon un processus de Poisson.

Notons que la mise en application d'une politique du type installation à un étage nécessite seulement l'information locale sur l'état du stock et les demandes de l'étage en question. Par contre, la mise en application d'une politique du type échelon nécessite l'information sur l'état du stock de toutes les installations en aval. En principe, l'évolution de la position de stock échelon d'un étage peut être déterminée en utilisant, si disponible, l'information sur l'état initial du système et les demandes finales arrivées. Mais en pratique, cette approche peut surestimer l'état réel du système, car elle ne prend pas en compte des changements variés d'état dûs par exemple aux détériorations des produits en stock.

Les politiques d'approvisionnement et de stockage citées ainsi que d'autres politiques sont décrites plus en détail par Axsater (2000a), Axsater (2003) et Zipkin (2000).

1.4.2. Politiques de stock nominal

La politique de stock nominal (*base stock policy*) est parmi les politiques d'approvisionnement et de stockage les plus souvent étudiées dans la littérature. En appliquant une politique de stock nominal, le réapprovisionnement du stock est déclenché pour ramener la position de stock à un niveau S en permanence. La politique de stock nominal est aussi appelée politique avec niveau de rechargement S (*order-up-to-S policy*) ou politique S (*S policy*).

La politique de stock nominal à inventaire permanent déclenche le réapprovisionnement du stock lorsque la position de stock devient inférieure au niveau S , c'est-à-dire chaque fois qu'une demande arrive. La quantité commandée dans chaque déclenchement est égale à la différence entre la position de stock et le niveau S . Par conséquent, la quantité commandée est exactement la quantité de demande. Dans le cas de demande unitaire, la politique de stock nominal devient un cas particulier de la politique (s, S) avec $s = S - 1$ et de la politique (R, Q) avec $R = S - 1$ et $Q = 1$. Notons qu'ici une demande unitaire peut aussi être définie comme une quantité entière et fixe de produits. Appliquée à ces types de systèmes, la politique de stock nominal est souvent notée $(S - 1, S)$ et parfois appelée politique d'approvisionnement un-à-un (*one-to-one replenishment policy*).

L'application de la politique de stock nominal se rencontre dans les systèmes où le coût de commande est négligeable par rapport aux autres coûts. Par exemple, quand chaque unité dans le stock est de valeur, les coûts de stockage et de rupture dominent forcément les coûts fixes de commande. Parallèlement, pour les produits ayant un taux de demande faible, les économies liées au système rendent l'utilisation de lots de grande taille peu intéressante. En outre, dans certains cas, les demandes et les livraisons provenant des fournisseurs sont d'une quantité fixe déterminée par exemple par les contraintes du système de transport. Pour ces types de systèmes, les demandes unitaires et les commandes unitaires avec $Q = 1$ ont un sens en termes de cette quantité fixe. Pour la plupart des systèmes mono-étage ayant des coûts de commande

négligeables et des demandes aléatoires et stationnaires, une politique de stock nominal est optimale par rapport aux coûts moyens de stockage et de rupture.

Selon la politique de stock nominal, la position de stock reste constante au niveau S appelé le niveau de stock nominal (Figure 1.8) : à l'état initial, le stock contient un nombre de produits égal au niveau de stock nominal S (la position de stock est aussi égale au niveau de stock nominal S), ensuite, l'arrivée de chaque demande déclenche instantanément une commande dont la quantité est exactement la quantité de demande pour ramener la position de stock au niveau de stock nominal S . Le niveau de stock nominal S détermine ainsi le niveau maximal du stock.

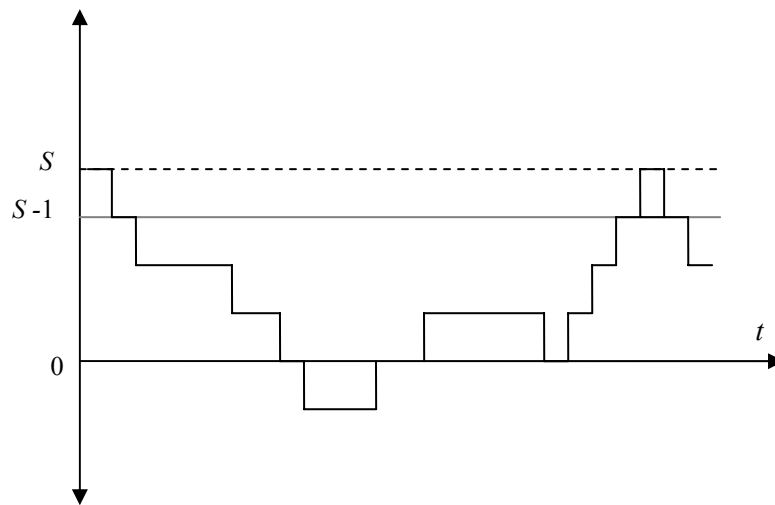


Figure 1.8. Évolution du stock avec la politique $(S-1, S)$

1.4.2.1. Systèmes de stock nominal multi-étages

La politique de stock nominal est utilisée pour gérer des stocks mono-étage ainsi que multi-étages. En appliquant une politique de stock nominal du type installation pour chaque étage d'un système à structure linéaire, l'arrivée d'une demande finale déclenche simultanément une demande pour chaque étage en amont. La Figure 1.9 illustre le fonctionnement de la politique de stock nominal $(S-1, S)$ dans un système à deux étages de production/stockage¹. Donc, la politique de stock nominal permet de réagir rapidement à la demande finale provenant des clients. Par contre, si un étage est perturbé pour une raison quelconque, la politique de stock nominal augmente inutilement le nombre d'en-cours dans le système.

¹ Dans les systèmes de stockage multi-échelons, l'étage « 1 » représente souvent le dernier étage donc l'étage des produits finis. Similairement, dans les systèmes MRP et pour les nomenclatures, les produits finis sont au niveau « 0 ». Par la suite, nous utilisons la convention inverse : pour un système à n étage, l'étage « 0 » représente les matières premières et l'étage « n » représente les produits finis. Pour un étage i , les étages en amont sont les étages $i-1, i-2, \dots, 0$ et les étages en aval sont les étages $i+1, i+2, \dots, n$.

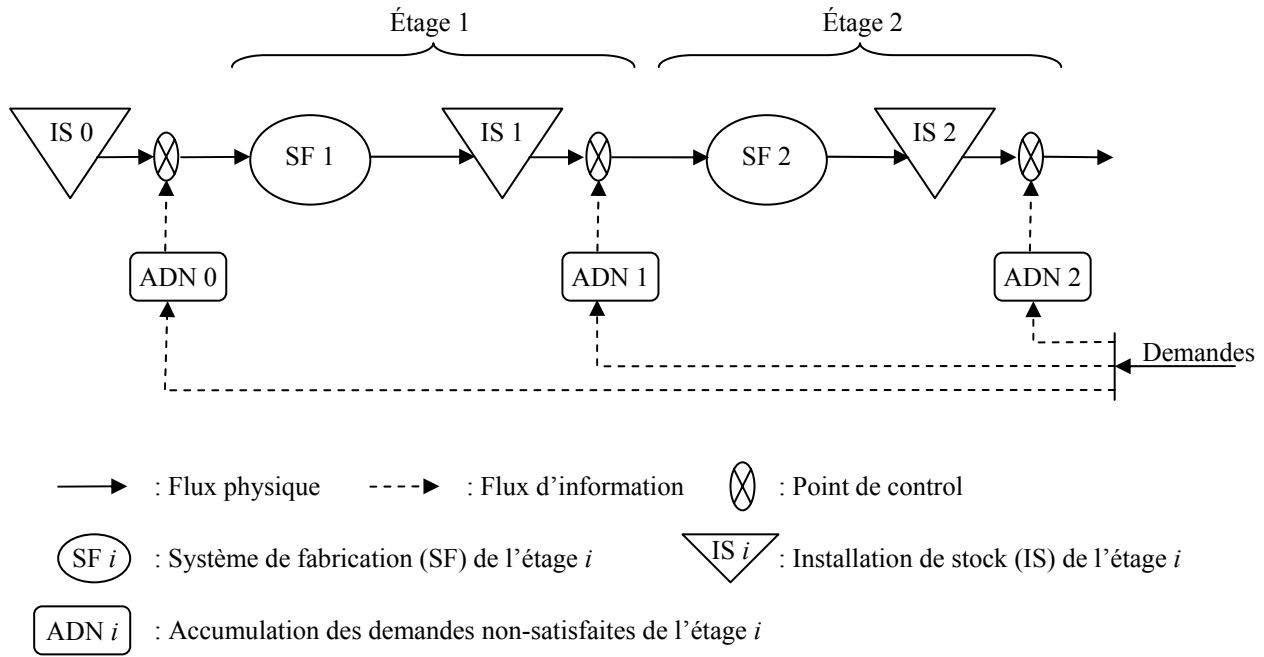


Figure 1.9. Politique de stock nominal $(S-1, S)$ dans un système à deux étages

Dans un système à structure linéaire, une politique de stock nominal du type installation donnée peut toujours être remplacée par une politique de stock nominal du type échelon équivalente en termes de dates de déclenchement des commandes de réapprovisionnement. En outre, une politique de stock nominal du type échelon déclenche aussi simultanément une demande pour chaque étage à l'arrivée d'une demande finale. Par conséquent, contrairement à une politique (R, Q) du type échelon, une politique de stock nominal du type échelon donnée peut toujours être remplacée par une politique de stock nominal du type installation équivalente (Axsater et Rosling, 1993). Pour un critère de performance donné, la meilleure politique du type échelon et la meilleure politique du type installation sont alors équivalentes. Considérons un système à n étages. Soit $(S_2)_{i=1}^n$ une politique de stock nominal du type échelon avec S_i non-décroissant pour $i = n, \dots, 1$. La politique de stock nominal du type installation équivalente à $(S_2)_{i=1}^n$ est caractérisé par $(\bar{S}_2)_{i=1}^n$ où $\bar{S}_n = S_n$ et $\bar{S}_i = S_i - S_{i+1}$ pour $i = n-1, \dots, 1$.

Clark et Scarf (1960) analysent un système à structure linéaire ayant des délais de transport fixes. À chaque étage, la gestion des stocks est accomplie suivant une politique de stock nominal du type échelon à inventaire périodique. Dans un système à inventaire périodique, la politique de stock nominal fonctionne comme une politique (T, S) avec $T = 1$. La demande finale qui s'étale sur plusieurs périodes est aléatoire. La distribution de probabilité de la demande peut différer d'une période à l'autre. Les demandes finales qui ne sont pas satisfaites sont retardées. Chaque unité de demande finale retardée induit un coût unitaire de rupture par période. Le critère est de minimiser des coûts moyens de stockage et de rupture du système multi-étages sur un nombre fini de périodes. Clark et Scarf (1960) montrent que les valeurs optimales des niveaux de stock nominaux peuvent être obtenues par programmation dynamique.

Le niveau de stock nominal de l'étage n est calculé en supposant que l'étage $n - 1$ est toujours en mesure de satisfaire immédiatement la totalité d'une demande qui s'adresse à lui, c'est-à-dire en supposant que le niveau de stock nominal de l'étage n est indépendant du niveau de stock nominal de l'étage $n - 1$. Les occurrences des ruptures de stock à l'étage $n - 1$ ne sont pas exclues. Les coûts associés à ces ruptures sont pris en compte en déterminant le niveau de stock nominal de l'étage $n - 1$. En supposant que le niveau de stock nominal de l'étage $n - 1$ est indépendant du niveau de stock nominal de l'étage $n - 2$, le niveau de stock nominal de l'étage $n - 1$ est calculé en minimisant la somme des coûts moyens de stockage du système et des augmentations des coûts moyens de stockage et de rupture quand l'étage $n - 1$ n'est pas en mesure de satisfaire immédiatement la totalité d'une demande. En continuant la même procédure, les niveaux de stock nominaux des étages $i = n - 2, n - 3, \dots, 1$ peuvent être obtenus. Notons que l'étage 0 est considéré comme un stock infini. Clark et Scarf (1960) montrent que la politique de stock nominal du type échelon obtenue avec cette procédure qui se base sur la minimisation de n fonctions convexes est en effet la meilleure politique qui minimise les coûts moyens du système. La politique de stock nominal du type installation équivalant à cette politique du type échelon est alors la meilleure politique du type installation.

Federgruen et Zipkin (1984) généralisent les résultats de Clark et Scarf (1960) pour un problème comportant un nombre infini de périodes et des demandes stationnaires. Ils montrent qu'une politique de stock nominal du type échelon stationnaire est optimale. Chen et Zheng (1994b) simplifient les démonstrations d'optimalité et généralisent les résultats pour les systèmes à inventaire permanent. Notons que l'algorithme de Clark et Scarf est applicable dans le cas des systèmes d'approvisionnement exogènes et séquentiels (Zipkin, 2000 ; Gallego et Zipkin, 1999). L'algorithme de Clark et Scarf est utile pour démontrer l'optimalité des politiques de stock nominal et des politiques similaires, mais limité sur le plan informatique pour le calcul des niveaux de stock nominaux. Federgruen et Zipkin (1984) montrent que dans un système à deux étages avec demande finale aléatoire suivant une loi normale, les valeurs optimales des niveaux de stock nominaux peuvent être calculées. Pour les systèmes à structure linéaire plus généraux, van Houtum *et al.* (1996) et Shang et Song (2003) proposent des techniques approximatives.

Puisqu'un système d'assemblage peut être remplacé par un système à structure linéaire équivalent, les résultats des systèmes à structure linéaire peuvent être généralisés pour les systèmes d'assemblage (Chen et Zheng, 1994b). Les politiques de stock nominal ne sont pas nécessairement optimales pour les systèmes de distribution similaires. Il existe des travaux analysant les systèmes de distribution sur base de la technique proposée par Clark et Scarf (van Houtum, 2006). Une autre problématique liée est de déterminer les valeurs optimales des niveaux de stock nominaux pour un système de distribution dans lequel chaque installation de stock est gérée par une politique de stock nominal du type installation. Dans

ce but, Sherbrooke (1968) propose une méthode approximative (méthode *METRIC*). Pour plus de détails, voir Axsater (2000a), Axsater (2003) et Zipkin (2000).

1.4.3. Politiques Kanban

La philosophie *Juste-à-Temps* (JAT), dont l'intention principale est de réduire des stocks, a été développée dans les années 70. La politique Kanban a été développée comme un système d'information pour implémenter la philosophie JAT chez Toyota Motors. Le mot « Kanban » signifie « étiquette » (ou « carte ») en japonais. Dans les systèmes Kanban, les étiquettes sont utilisées pour transmettre les commandes d'approvisionnement entre les étages. La politique Kanban est une mode de pilotage efficace pour traiter les systèmes de fabrication à capacité limitée.

Dans un système multi-étages piloté par un système Kanban, chaque étage a un nombre fini d'étiquettes donnant la description du produit fabriqué et stocké à cet étage. Chaque unité dans le stock doit avoir une étiquette attachée sur elle. Les étiquettes qui ne sont pas attachées à une pièce dans le système de fabrication ou dans le stock d'un étage se trouvent dans le panier de Kanban de l'étage en amont. Considérons un système constitué de n étages en série. Supposons que la demande arrive à l'étage n à chaque fois en quantité unitaire et le système de fabrication de l'étage n nécessite une unité de produit $n - 1$ pour fabriquer une unité de produit n . Soit K_n le nombre de Kanban (le nombre d'étiquettes) de l'étage n . Lorsqu'une demande se présente à l'étage n , si une unité de produit n avec une étiquette attachée est présente dans le stock, la demande est satisfaite et l'étiquette est détachée pour être transmise à l'étage $n - 1$. Une fois l'étiquette est transmise à l'étage $n - 1$, s'il n'y a pas une unité de produit $n - 1$ disponible dans le stock de l'étage $n - 1$, l'étiquette reste en attente dans le panier de Kanban de l'étage $n - 1$. Si une unité de produit $n - 1$ avec une étiquette de l'étage $n - 1$ attachée est présente dans le stock, l'étiquette de l'étage $n - 1$ est détachée pour être transmise à l'étage $n - 2$ et remplacée par l'étiquette de l'étage n qui a transmis la demande. La pièce est ensuite poussée au système de fabrication de l'étage n . Quand le processus de fabrication est terminé, la pièce prend sa place dans le stock de l'étage n avec l'étiquette attachée sur elle.

Les étages en amont fonctionnent de la même façon : une étiquette commence son cycle attachée à une pièce dans le stock. Une fois la pièce consommée, l'étiquette est transmise à l'étage en amont. L'étiquette autorise l'étage en amont pour une livraison unitaire. Quand une livraison se produit, l'étiquette est attachée à la pièce livrée et descend le flux de pièce. Cependant, s'il arrive qu'une pièce ne soit pas disponible dans le stock d'un étage à l'arrivée d'une demande, aucune étiquette n'est transmise à l'étage en amont et l'information sur la demande qui déclenche la production est bloquée à ce niveau de la chaîne. La Figure 1.10 illustre le fonctionnement de la politique Kanban dans un système à deux étages. Un cas particulier de la politique Kanban est la politique CONWIP (*CONstant Work In Process*). En appliquant la politique CONWIP dans un système multi-étages de production/stockage, une commande de matière

première est déclenchée lorsqu'une demande finale est satisfaite au dernier étage. Donc, la politique CONWIP est équivalente à une politique Kanban avec un seul étage qui représente toute la chaîne de production (Spearman et Zazanis, 1992).

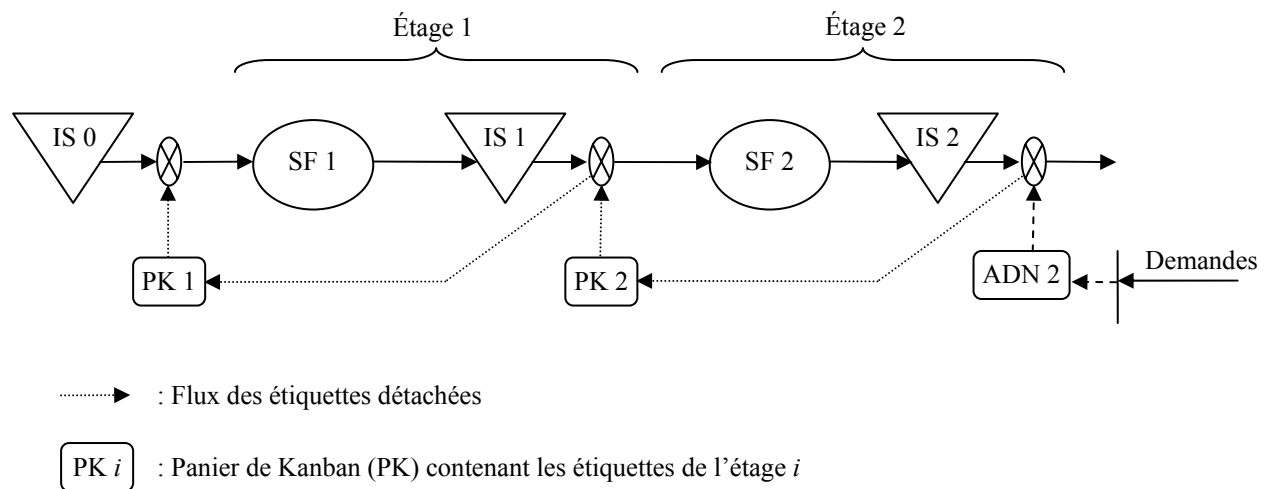


Figure 1.10. La politique Kanban dans un système à deux étages

La politique Kanban est similaire à la politique de stock nominal en plusieurs façons. Toutes les deux politiques sont des politiques réactives : l'arrivée d'une demande finale déclenche directement ou indirectement tous les autres événements. Les deux politiques nécessitent seulement l'information locale. Le flux d'information est dans le sens opposé du flux physique : l'information sur la demande finale se propage en arrière dans le réseau étage par étage à travers les commandes d'approvisionnement. Les deux politiques ont un seul paramètre de contrôle pour chaque étage i , le niveau de stock nominal, S_i , ou le nombre de Kanban, K_i , qui détermine le niveau maximal de stock de l'étage.

Les différences entre les deux politiques sont aussi importantes. En appliquant une politique de stock nominal, l'arrivée d'une demande à un étage i déclenche toujours une demande pour l'étage $i - 1$. La politique Kanban, au contraire, déclenche une demande à l'étage $i - 1$ quand l'étage i satisfait une de ses demandes. Par conséquent, les mécanismes de propagation de demande ne sont identiques que si l'étage a un niveau de stock positif. En appliquant la politique Kanban, les demandes qui sont retardées ne sont pas propagées vers l'amont. Le nombre de Kanban a alors une autre fonction importante pour le système : le nombre de produits dans le système de fabrication (en attente de fabrication et en fabrication) d'un étage i , c'est-à-dire le nombre d'en-cours d'un étage est limité par le nombre de Kanban. En appliquant la politique de stock nominal, le nombre d'en-cours est, au contraire, illimité.

Le fait qu'elle limite le nombre d'en-cours est un des points forts de la politique Kanban. Considérons par exemple un système à structure linéaire à deux étages. Le système de fabrication de l'étage 2 est plus lent que celui de l'étage 1. En appliquant la politique de stock nominal, l'étage 1 continue à fabriquer des produits et à les transmettre au système de fabrication de l'étage 2 tant qu'il y a des demandes finales. Par

conséquent, le nombre d'en-cours de l'étage 2 augmente rapidement. L'idée de base de la politique Kanban est qu'il n'est pas nécessaire de continuer à charger un système de fabrication s'il est déjà congestionné. La politique Kanban bloque la transmission des demandes à l'étage 1 et donc le flux de matière vers l'étage 2 dès que le système de fabrication de l'étage 2 devient trop congestionné. Puisqu'il n'y a plus de demandes, l'étage 1 ne fabrique plus des produits. De cette façon, la politique Kanban limite le nombre total d'en-cours du système qui comprend les produits en attente de fabrication et en fabrication dans l'étage 2 et aussi les produits dans le stock de l'étage 1 par le nombre total de Kanban dans le système, $K_1 + K_2$.

Cependant, la politique Kanban peut générer des délais de livraison plus élevés pour les demandes finales. Considérons encore un système à structure linéaire à deux étages, mais cette fois le système de fabrication de l'étage 1 est plus lent que celui de l'étage 2. L'étage 2 ne nécessite pas un stock élevé. Donc, il est logique de lui attribuer un nombre faible de Kanban. Par contre, un nombre faible de Kanban limite le nombre d'en-cours de l'étage 2 et l'étage 1 peut arrêter de fabriquer des produits même s'il y a des demandes finales retardées. Puisque le système de fabrication de l'étage 1 est lent, le système de fabrication de l'étage 2 peut être bloqué à cause d'une rupture à l'étage 1. Afin d'éviter une telle situation, le nombre de Kanban de l'étage 2 peut être augmenté mais ça augmente aussi le stock de l'étage 2. La politique de stock nominal, au contraire, permet d'ajuster le niveau maximal de stock de chaque étage en transmettant librement des demandes finales aux étages goulots.

Dans la section précédente, nous avons indiqué que la politique de stock nominal est un cas particulier de la politique (R, Q) avec $R = S - 1$ et $Q = 1$. La politique Kanban est aussi similaire à la politique (R, Q) . La politique Kanban est souvent appliquée en utilisant des conteneurs standards de taille de lots fixe, notée Q , pour la circulation entre les étages. Les étiquettes restent attachées aux conteneurs. Une étiquette est transmise à l'étage en amont chaque fois qu'un container est vidé. La politique Kanban avec une taille de lots fixe peut être interprété comme une politique du type (R, Q) avec $R = (K - 1)Q$ et une contrainte additionnelle qui bloque les commandes pour l'étage en amont dans le cas d'une rupture de stock. Cette contrainte peut être mise en application en utilisant une définition différente de la position de stock dans laquelle la position de stock n'est pas diminuée par le niveau de rupture de stock (Axsater et Rosling, 1993).

Beaucoup de travaux traitant de la politique Kanban ont été développés depuis son apparition. Une revue détaillée de la littérature sur la politique Kanban est donné par Berkley (1992). La plupart des travaux de recherche se basent sur la théorie des files d'attente pour évaluer les performances des politiques Kanban dans des systèmes à capacité limitée. Di Mascolo *et al.* (1996) utilisent le formalisme des files d'attente pour modéliser un système multi-étages géré par la politique Kanban. Le système de fabrication de chaque étage est constitué de plusieurs machines exécutants des opérations différentes. Les auteurs proposent une méthode approximative pour la détermination des mesures de performances du système.

Pour un système similaire, Baynat *et al.* (2001) proposent une méthode approximative donnant les mêmes résultats mais qui résulte en un algorithme simplifié. Matta *et al.* (2005) analysent les systèmes d'assemblage gérés par la politique Kanban en utilisant la théorie des files d'attente. Baynat *et al.* (2002) analysent les systèmes Kanban multi-étages où chaque étage fabrique plusieurs produits (système multi-produits).

1.4.3.1. Extensions de la politique Kanban

La politique Kanban a un seul paramètre de contrôle, le nombre de Kanban. Ceci implique une certaine simplicité d'application mais aussi des restrictions. Les extensions de la politique Kanban telles que la politique Kanban généralisée et la politique Kanban étendue ayant deux paramètres de contrôle ont été proposées dans la littérature.

La politique Kanban généralisée (Frein *et al.*, 1995) est une politique hybride de la politique de stock nominal et de la politique Kanban. Elle utilise deux paramètres pour chaque étage i : le niveau de stock nominal S_i et le nombre de Kanban K_i . Chaque étage i a deux paniers de Kanban, un pour les étiquettes de l'étage i et un pour les étiquettes de l'étage $i + 1$ représentant les commandes de l'étage $i + 1$. À l'état initial, le stock de l'étage i contient un nombre de produits égal au niveau de stock nominal S_i et toutes les étiquettes de l'étage i se trouvent dans le panier de Kanban de l'étage i . Lorsqu'une demande se présente, si une unité de produit i est présente dans le stock, la demande est satisfaite. Sinon, la demande est retardée. La demande est transmise à l'étage $i - 1$, si une étiquette de l'étage i est présente dans le panier de Kanban. S'il y a une unité de produit $i - 1$ disponible dans le stock, le produit est envoyé au système de fabrication de l'étage i avec l'étiquette de l'étage i qui a transmis la demande attachée sur lui. Quand le processus de fabrication est terminé, l'étiquette est détachée et remise dans le panier de Kanban et le produit prend sa place dans le stock de l'étage i . S'il arrive qu'une étiquette ne soit pas disponible dans le panier de Kanban d'un étage à l'arrivée d'une demande, aucune étiquette n'est transmise à l'étage en amont et la remontée de la demande est bloquée. La Figure 1.11 illustre le fonctionnement de la politique Kanban généralisée dans un système à deux étages.

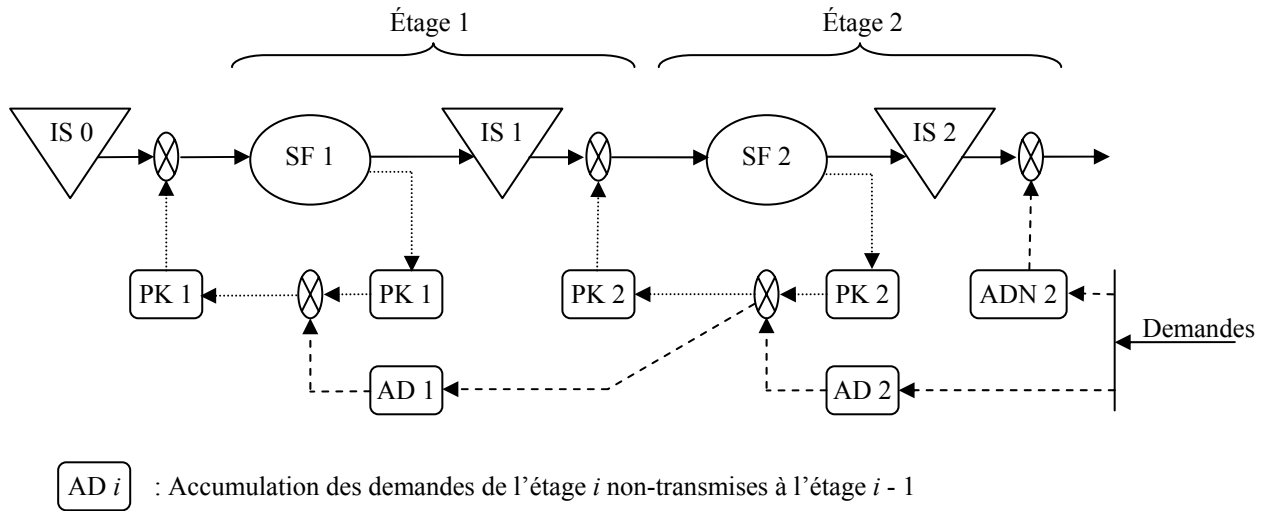


Figure 1.11. La politique Kanban généralisée dans un système à deux étages

La politique Kanban bloque la transmission des demandes vers l'amont quand les produits dans le stock d'un étage sont épuisés. La politique Kanban généralisée relâche cette contrainte et bloque la transmission des demandes vers l'amont quand les étiquettes dans le panier de Kanban sont épuisées. Dans le cas où $S_i = K_i$ pour chaque étage i , la politique Kanban généralisée fonctionne comme la politique Kanban. La politique de stock nominal est équivalente à une politique Kanban généralisée ayant $K_i = \infty$ pour chaque étage i . Pour $S_i = 0$, un système Kanban généralisé fonctionne comme un réseau des systèmes de fabrication en tandem dans lequel le nombre de produit en attente de fabrication et en fabrication à chaque étage i est limité par le nombre de Kanban K_i .

La politique de Kanban étendue est, comme la politique de Kanban généralisée, une politique hybride de la politique de stock nominal et de la politique Kanban. Elle a été développée par Dallery et Liberopoulos (2000). La politique de Kanban étendue utilise deux paramètres de contrôle pour chaque étage i : le niveau de stock nominal S_i et le nombre de Kanban K_i où $K_i \geq S_i$. À l'état initial, le stock de l'étage i contient un nombre de produits égal au niveau de stock nominal S_i . Chaque unité dans le stock doit être attachée à une étiquette. Les étiquettes restantes ($K_i - S_i$) se trouvent dans le panier de Kanban de l'étage $i - 1$. Lorsqu'une demande se présente à l'étage i , la demande est automatiquement transmise à tous les étages en amont. Pour qu'une pièce soit transmise au système de fabrication de l'étage i , outre la présence d'une demande, une étiquette de l'étage i doit être disponible dans le panier de Kanban de l'étage $i - 1$. Dans ce cas, l'étiquette de l'étage $i - 1$ est détachée pour être transmise à l'étage $i - 2$ est remplacée par une étiquette de l'étage i . Ensuite, la pièce entre dans le système de fabrication de l'étage i . La Figure 1.12 illustre le fonctionnement de la politique Kanban étendue dans un système à deux étages.

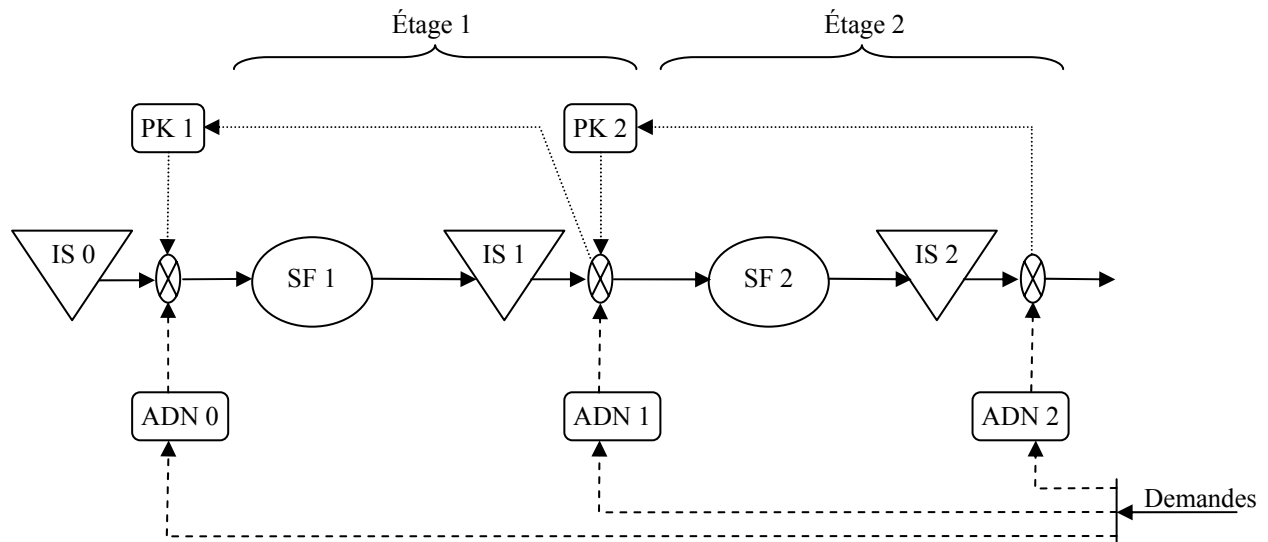


Figure 1.12. La politique Kanban étendue dans un système à deux étages

Du fait de l'existence d'un nombre d'étiquettes de l'étage i en avance dans le panier de Kanban de l'étage $i - 1$, la politique Kanban étendue peut délivrer une pièce à l'étage i , dès qu'une demande arrive, même si aucune étiquette n'a été libérée à l'étage i . La politique Kanban étendue fonctionne comme la politique Kanban quand $S_i = K_i$ et comme la politique de stock nominal quand $K_i = \infty$ pour chaque étage i .

Dans la littérature, il existe des extensions considérant les politiques hybrides de la politique Kanban et de la politique (R, Q) . Axsater et Rosling (1999) proposent des politiques (R, Q) du type échelon avec les contraintes de blocage. Liberopoulos et Dallery (2003) proposent des politiques Kanban du type échelon et des politiques hybrides de pilotage de flux qui combinent les politiques du type Kanban et les politiques du type (R, Q) .

1.4.4. Les politiques du type MRP

Les principes de base de la politique MRP (*Material Requirements Planning*, Planification des Besoins en Composants) ont été développés dans les années 60, avec pour objectif de mieux gérer l'approvisionnement des matières et composants nécessaires à la fabrication des produits finis. La politique MRP (détaillée par Vollman *et al.* (1997)) consiste à déterminer, pour chaque produit (matière première, composant, produit intermédiaire, produit fini, etc.), les dates et les quantités des lancements de production ou des commandes d'approvisionnement sur un horizon de planification donné dans le but de satisfaire les besoins exprimés dans le Plan Directeur de Production (PDP) pour chaque période. Les besoins exprimés dans le PDP couvrent généralement les commandes fermes et peuvent s'étendre aux demandes prévisionnelles selon le mode de production.

Dans la logique MRP, les calculs des besoins en produits sont réalisés en cascade au départ des besoins en produits finis. Les calculs des besoins en cascade nécessitent l'information sur la nomenclature de chaque

produit fini qui spécifie les différents produits intervenant dans sa composition aux différentes étapes de fabrication. Dans une nomenclature, les produits sont ordonnés par niveau sous forme d'arborescence. Cette structure permet d'indiquer les produits d'un niveau qui entrent dans la fabrication d'un produit de niveau directement inférieur. En planifiant d'abord les lancements de production des produits finis, la technique MRP enchaîne successivement les calculs pour chaque produit :

1. Explosion de nomenclature : Les besoins bruts exprimant la demande totale au début de chaque période sont dérivés des demandes prévisionnelles du produit et des lancements de production de chacun des produits de niveau directement inférieur. Les besoins bruts traduisent alors des demandes indépendantes et des demandes dépendantes du produit.

2. Calcul des besoins nets : Les besoins nets exprimant la quantité effectivement requise au début de chaque période sont déterminés sur la base des besoins bruts, du niveau de stock au début de l'horizon de planification, et des livraisons attendues correspondants aux ordres passés avant le début de l'horizon de planification.

3. Calcul des livraisons planifiées : Les besoins nets d'une période peuvent être couverts soit par une livraison au début de la même période soit par une livraison antérieure planifiée pour couvrir les besoins nets de plusieurs périodes consécutives. La détermination de la quantité à livrer pour satisfaire les besoins d'une ou plusieurs périodes repose normalement sur un arbitrage entre les coûts de stockage et les coûts de commande. Le problème de regroupement des ordres de production ou d'approvisionnement est connu dans la littérature sous le nom de problème de calcul des tailles de lots (*lot-sizing problem*). Les règles de calcul des tailles de lots utilisés dans les systèmes MRP s'étendent des règles de décision simples et des procédures heuristiques aux procédures d'optimisation. Parmi les nombreuses règles de calcul des tailles de lots utilisées, nous citons (détaillé par Vollman *et al.* (1997)) : la politique lot-pour-lot ; la politique à quantité fixe de commande ; la politique à nombre de période fixe ; l'heuristique *Part-Période-Balancing* ; l'heuristique de *Silver-Meal* ; l'algorithme de *Wagner-Whitin*. L'algorithme de *Wagner-Whitin* se base sur la programmation dynamique et donne une solution optimale pour le problème de calcul des tailles de lots. Il présente toutefois l'inconvénient d'une implémentation pratique complexe et difficile à mettre en œuvre.

4. Absorption des délais : Les dates de lancement de production (ou de commande d'approvisionnement) sont obtenues en retranchant les délais d'approvisionnement des dates de livraisons planifiées.

En d'autres termes, la technique de base de MRP considère un système multi-niveaux (multi-étages) à inventaire périodique avec des délais d'approvisionnement fixes et des demandes déterministes et non-stationnaires. Puisqu'elle détermine les tailles de lots des produits appartenant à un niveau en ignorant les relations avec les niveaux supérieurs, elle peut être classée comme une approche heuristique pour

résoudre le problème de calcul des tailles de lots multi-niveaux (*multi-level lot sizing problem*) connu comme un problème difficile à résoudre (Zipkin, 2000).

ARTICLE X	SEMAINE	1	2	3	4	5	6	7	8
Délai d'approvisionnement (L) = 1 Quantité fixe de commande (Q) = 25	Besoins bruts	10		25	10	20	5		10
	Approvisionnements attendus		25						
	Niveaux de stock planifiés	22	12	37	12	2	7	2	2
	Besoins nets					18	5		10
	Approvisionnements planifiés					25			25
	Lancements planifiés				25			25	

Figure 1.13. Exemple d'un plan MRP

La technique MRP ne tient pas en compte les contraintes de capacité du système d'approvisionnement lorsqu'elle calcule les ordres à planifier. Ceci est dû aux calculs par produit qui ne permette pas de prendre en compte les niveaux d'utilisation des ressources dans leur totalité. Ainsi, les délais d'approvisionnement ne reflètent pas les temps d'attente liés à la charge et la capacité du système d'approvisionnement. Dans le but de relier le plan MRP à la capacité disponible, la plupart des logiciels utilisent un module CRP (*Capacity Requirements Planning*, Planification de Besoins en Capacité) qui permet d'établir dans le temps des projections des besoins de capacité des ressources. Si la capacité requise pour certaines périodes dépasse la capacité disponible, il faut procéder à des ajustements en modifiant le plan MRP, en augmentant la capacité disponible ou en modifiant le PDP. Dans la littérature, une attention importante est portée sur les problèmes de calcul de tailles de lots multi-niveaux à capacités limitées (Billington *et al.*, 1983 ; Ozdamar et Barbarosoglu, 2001 ; Belvaux et Wolsey, 2001).

En pratique, les prévisions de demande sont mises à jour afin de tenir compte des changements affectant la demande et le PDP est actualisé avec une périodicité de révision. Lorsqu'il y a des changements dans le PDP, il devient nécessaire de recalculer les plans MRP. Des changements mineurs dans le PDP peuvent provoquer des changements relativement importants dans les plans MRP. Si ceci se produit trop fréquemment, les plans deviennent instables. Cette instabilité est souvent appelée la nervosité du système. Les changements dans les plans MRP peuvent aussi être provoqués par les arrivées des demandes non-anticipées, les arrivées tardives des approvisionnements planifiés, ou les arrivées d'approvisionnements avec des produits défectueux. La règle de calcul des tailles de lots utilisée a un impact sur la sensibilité des plans MRP aux changements des données. En comparaison avec les heuristiques utilisées dans les systèmes MRP, l'algorithme de Wagner-Whitin fournit, en général, des plans MRP plus sensibles aux changements.

En présence des incertitudes (liées aux quantités de demande prévisionnelle, dates de demande prévisionnelle, délais d'approvisionnement, quantités approvisionnés), la politique MRP nécessite la mise en place de paramètres de sécurité comme les stocks de sécurité et les délais de sécurité. En utilisant le délai de sécurité, le délai d'approvisionnement est déterminé en amplifiant le délai d'approvisionnement actuel avec le délai de sécurité. Pour des cas où les délais actuels sont aléatoires, Dolgui et Ould-Louly

(2002) et Ould-Louly et Dolgui (2004) cherchent à déterminer les délais d'approvisionnement planifiés qui minimisent les coûts de fonctionnement. La majorité de travaux de recherche sur les politiques MRP s'intéresse à l'analyse des performances des différents paramètres de sécurité. Koh *et al.* (2002) et Guide *et al.* (2000) donnent des revues de littérature sur cette problématique.

Buzacott et Shantikumar (1994) sont parmi les premiers à utiliser des modèles stochastiques pour analyser les performances de la politique MRP. Ils étudient un système mono-étage piloté par une politique MRP à inventaire permanent. Le système de fabrication en amont est à capacité limitée avec des temps de fabrication incertains. Les demandes arrivent selon un processus stochastique stationnaire. Les demandes sont connues à l'avance sur un horizon de temps donné. Ils étudient l'impact du choix entre l'utilisation des stocks de sécurité et l'exploitation des délais de sécurité. Ils montrent que quand les prévisions de la demande ne sont pas fiables, l'utilisation des stocks de sécurité est préférable. En outre, dans ces types de situations, les politiques de pilotage de flux qui ne prennent pas en compte les prévisions sont préférables aux politiques du type MRP.

Hennet (2003) analyse l'application d'une politique similaire dans un système multi-étages dont les demandes finales sont déterministes pour les premiers T_S périodes et aléatoires et stationnaires pour les périodes suivantes, $T_S + 1$ à T_L . Les temps de fabrication sont supposés certains. Pour chaque période, les ressources de production ainsi que de stockage sont à capacité limitée et donc soumises aux contraintes de capacité agrégée. L'auteur montre qu'un plan de production à long terme, pour les périodes $k = T_S + 1, \dots, T_L$, qui minimise le coût moyen de production, de stockage et de rupture peut être obtenu. Un plan à court terme, pour les périodes $k = 0, \dots, T_S$, qui aboutit à un état imposé par le plan à long terme et qui minimise le coût moyen à court terme peut ensuite être déterminé.

En comparaison avec les politiques d'approvisionnement et de stockage, la technique MRP est plus adaptée pour les systèmes ayant des demandes fortement non-stationnaires, spécialement pour les systèmes de production multi-étages avec lotissement. Notons que dans les systèmes de production multi-étages, le caractère non-stationnaire de la demande d'un produit est amplifié par les décisions de lotissement des produits de niveau inférieur. Comparons par exemple une politique MRP avec une politique du type (R, Q) à inventaire périodique. En utilisant une politique à quantité fixe de commande, la technique MRP planifie un ordre de fabrication ou une commande d'approvisionnement au début d'une période t si la position de stock (après la réalisation de la demande de la période t) est inférieure à la demande pendant l'intervalle $(t, t + L]$. Une politique du type (R, Q) peut générer les mêmes commandes seulement si le point de commande R est actualisé pour chaque période et fixé égal à la demande pendant l'intervalle $(t, t + L]$ moins un (Axsater, 2000a). Les politiques d'approvisionnement et de stockage dont les paramètres évoluent dans le temps sont souvent appelées les politiques dynamiques ou non-stationnaires. Pour le cas analysé, une politique stationnaire du type (R, Q) peut amplifier les stocks par exemple quand la demande se produit rarement mais est de quantité élevée (Axsater, 2000a).

En appliquant la politique MRP conventionnelle dans un système multi-étages, toutes les données nécessaires sont transmises à l'unité de contrôle où toutes les décisions sont prises. Les décisions prises sont ensuite communiquées à chaque étage. Buzacott (1997) analyse une implémentation dans laquelle les étages communiquent à travers les commandes d'approvisionnement. Chaque étage a seulement accès à l'information locale. Il montre que les changements possibles de données et donc la nervosité des systèmes MRP rendent cette implémentation difficile. Nous pouvons aussi noter qu'une telle implémentation peut amplifier l'effet de coup de fouet dans le système (voir la section suivante).

1.4.5. Classification des politiques de pilotage de flux

Nous avons classifié les politiques de pilotage de flux en deux familles : les politiques réactives et les politiques proactives. Dans la littérature, les politiques de pilotage de flux sont classifiées à l'aide de plusieurs critères. Dans cette section, nous présentons les classifications les plus connues en montrant qu'elles comportent certains inconvénients.

Dans la littérature, les politiques de pilotages de flux sont souvent classifiées en politiques à flux tiré et politiques à flux poussé. La distinction classique entre politiques à flux tiré et politiques à flux poussé est basée sur l'instant de déclenchement de la production en réponse à la demande. Suivant cette définition, les politiques à flux tiré sont celles où la demande réalisée tire la production. Quant aux politiques à flux poussé, le déclenchement de la production est fait avant que la demande n'arrive. Selon cette distinction, la classe des politiques réactives correspond à la classe des politiques à flux tiré et la classe des politiques proactives correspond à la classe des politiques à flux poussé. Par contre, cette distinction n'est pas la seule utilisée dans la littérature.

Une deuxième distinction entre politiques à flux tiré et politiques à flux poussé classe les systèmes de production en deux groupes : les systèmes qui tiennent compte de la demande et les systèmes qui ne tiennent pas compte de la demande. Selon cette définition, la notion de flux tiré et flux poussé est liée respectivement à la notion de boucle fermée et boucle ouverte reliant la demande au système de production. Dans le système en boucle fermée, la production est déclenchée en réponse à une demande présente ou future alors que dans le système à boucle ouverte, la production est déclenchée en ignorant la demande. Suivant cette classification, même la politique MRP devient alors une politique à flux tiré dans laquelle le PDP tire la production pour avoir les produits juste-à-temps (Liberopoulos et Dallery, 2003).

Les politiques à flux tiré sont parfois définies comme les politiques qui limitent le nombre d'en-cours dans le système. Selon cette définition, seules les politiques Kanban entrent dans la classe des politiques en flux tiré par l'aval. La politique de stock nominal devient une politique à flux poussé car elle déclenche la production dans tous les étages du système en réponse à une réalisation de demande finale.

Dans la littérature, il existe des travaux analysant les politiques réactives dans le cas où l'information sur les demandes futures est disponible sous forme de prévisions ou de commandes fermes. Babai (2005) analyse les politiques non-stationnaires du type (R_n, Q) et (T, S_t) à inventaire périodique quand l'information sur les demandes futures est disponible sous forme de prévisions et d'incertitudes prévisionnelles obtenues à l'avance sur un horizon donné. Il montre l'équivalence entre les politiques analysées et les politiques du type MRP. Les politiques non-stationnaires d'approvisionnement et de stockage sur prévisions sont souvent utilisées pour analyser l'effet de coup de fouet (*Bullwhip effect*). L'effet de coup de fouet est un phénomène d'amplification de la variabilité de la demande lorsque l'on remonte vers l'amont de la chaîne logistique. Lee *et al.* (1997) et Chen *et al.* (2000) quantifient l'effet de coup de fouet quand les commandes sont générées par une politique de stock nominal non-stationnaire. Le paramètre de contrôle de la politique utilisée est actualisé à chaque période en fonction des prévisions disponibles. La quantité de commande à chaque période est constituée de deux termes. Le premier terme correspond à la dernière demande réalisée. Le deuxième terme sert à ajuster le niveau de stock nominal pour s'adapter au changement de la demande prévisionnelle durant le délai d'approvisionnement. Lee *et al.* (2000) montrent que l'effet de coup de fouet peut être réduit si l'information sur la demande finale est disponible à chaque étage.

Karaesmen *et al.* (2002) étudient un modèle similaire à celui de Buzacott et Shantikumar (1994) en temps discret. Les auteurs analysent un système mono-étage de production/stockage dans lequel les demandes unitaires arrivent selon un processus stochastique stationnaire. L'arrivée de la $n^{\text{ième}}$ demande est signalée l_n unités de temps avant sa date de réalisation. La politique de pilotage de flux utilisée est une politique de stock nominal à deux paramètres de contrôle : le niveau de stock nominal, nommé S , et le délai d'approvisionnement planifié, nommé L . L'intervalle de temps entre la signalisation de la $n^{\text{ième}}$ demande et le déclenchement du $n^{\text{ième}}$ ordre de fabrication est $\max\{0, l_n - L\}$. Le paramètre L peut donc être interprété comme le délai d'approvisionnement planifié d'un système MRP. Si l'information avancée sur la demande n'est pas disponible ($l_n = 0$), la politique étudiée devient une politique classique de stock nominal. Les analyses numériques montrent que les performances de la politique étudiée sont très proches de celles de la politique optimale. En outre, les diminutions des coûts de stockage et de rupture sont significatives quand l'information avancée sur la demande est disponible. Liberopoulos et Koukoumialos (2005) analysent numériquement un système similaire et un système de Kanban généralisé avec information avancée sur la demande.

Les politiques de pilotage de flux citées sont des extensions des politiques réactives. Elles sont modifiées afin de tenir compte l'information sur les demandes futures. Nous appelons ces politiques les politiques réactives modifiées. Selon la distinction présentée entre les politiques réactives et les politiques proactives, les politiques réactives modifiées entrent dans la classe des politiques proactives.

1.5. MODÉLISATION DYNAMIQUE DES FLUX

Selon Cassandras (1993), un système à événements discrets est un système dans lequel les variables qui représentent l'état évoluent de manière brusque d'une valeur à l'autre, suite à des événements, et gardent une valeur constante dans les autres cas. Le mot discret ne signifie ni « temps discret », ni « état discret » mais se réfère au fait que la dynamique est introduite par l'existence d'événements dont les dates d'occurrence n'ont pas une importance primordiale. Ce sont surtout les conditions et l'ordre d'occurrence de ces événements qui importent.

Dans le cadre cette thèse, nous analysons les systèmes de production et de stockage qui peuvent être considérés comme des systèmes à événements discrets. Dans les systèmes de production et de stockage étudié, les pièces restent en attente dans un stock afin de satisfaire les demandes du produit stocké. Les demandes qui ne peuvent pas être immédiatement satisfaites dès leurs arrivées sont retardées. Les commandes d'approvisionnements et/ou les ordres de fabrication sont déclenchés pour renouveler la consommation du stock.

Le système de fabrication de chaque produit est à capacité limitée. Le délai d'approvisionnement pour un ordre de fabrication donné dépend alors, outre de la disponibilité des produits nécessaires pour exécuter le processus de transformation, de la charge du système de fabrication à son arrivée et donc des ordres précédents. Quand le système de fabrication est congestionné, les pièces disponibles pour l'exécution du processus de transformation et les ordres de fabrication restent en attente d'exécution. Les ordres de fabrication peuvent aussi rester en attente dans le cas où les pièces qui seront transformées ne sont pas encore disponibles.

L'évolution des systèmes étudiés peut être définie en utilisant des variables d'état discrètes comme le nombre de pièces en attente dans les différents stocks, le nombre de demandes retardées des différents stocks, le nombre de pièces en attente de fabrication dans les différents systèmes de fabrication, le nombre d'ordres de fabrication en attente d'exécution dans les différents systèmes de fabrication, et l'état de différents systèmes de fabrication (libre ou occupée). Les dynamiques des systèmes étudiés déterminent comment ces variables d'état évoluent dans le temps. Les événements définissant la dynamique des systèmes étudiés correspondent alors à l'arrivée d'une demande, la satisfaction d'une demande, le début de traitement d'un ordre de fabrication dans un système de fabrication, la fin de traitement d'un ordre de fabrication dans un système de fabrication, etc.

Nous analysons les systèmes de production et de stockage confronté aux aléas. Nous supposons que les demandes finales et les temps de fabrication des produits dans les différents systèmes de fabrication sont incertains. Les outils les plus souvent utilisés pour la modélisation des systèmes de production et de stockage dans un cadre stochastiques sont les variantes des réseaux de Petri (*Petri Nets*, PNs) et les files

d'attente. Le travail de Di Mascolo *et al.* (1991) est l'un des premiers qui utilisent les PN pour la modélisation des systèmes Kanban. Chaouiya et Dallery (1997) utilisent les PN stochastiques (*Stochastic Petri Nets*, SPNs) afin de modéliser l'application de la politique de stock nominal, la politique Kanban, la politique Kanban généralisée et la politique Kanban extensive dans des systèmes d'assemblage. Moore et Gupta (1999) utilisent les SPNs colorés pour la modélisation des systèmes Kanban. En utilisant les SPNs comme outil de modélisation, l'analyse exacte des performances est possible seulement sous des hypothèses restrictives. Les PN sont en général bien adaptés pour la modélisation des systèmes complexes dont l'analyse des performances se base sur des méthodes de simulation.

Dans le cadre cette thèse, nous utilisons les files d'attente afin de modéliser les systèmes de pilotage de flux étudiés et d'analyser leurs performances analytiquement. Les définitions des lois de probabilité et des processus stochastiques principaux ainsi que les fondements de la théorie des files d'attente sont donnés dans l'annexe A. Pour plus de détails sur les processus stochastique et la théorie des files d'attente, nous citons Ross (2000) et Kleinrock (1975). Par la suite, nous présentons les généralités sur la modélisation des systèmes de production et de stockage étudiés par les files d'attente.

1.5.1. Modélisation des systèmes stock nominal

Dans le cadre cette thèse, nous utilisons le formalisme des files d'attente pour modéliser les systèmes de stock nominal. Considérons un système mono-étage de production/stockage d'un produit (Figure 1.14) géré par une politique de stock nominal $(S - 1, S)$ à inventaire permanent.

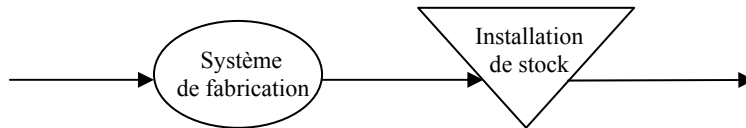


Figure 1.14. Système mono-étage de production/stockage

Selon la politique de stock nominal $(S - 1, S)$, le stock contient le niveau de stock nominal $S \geq 0$ à l'état initial. Nous supposons que les demandes externes arrivent chaque fois en quantité unitaire et selon un processus de Poisson ayant le taux λ . Quand la demande arrive, elle est satisfaite s'il y a des produits dans le stock, sinon elle est retardée. En appliquant la politique de stock nominal $(S - 1, S)$, un ordre de fabrication unitaire est déclenché lorsque la position de stock devient $S - 1$, c'est-à-dire à chaque fois qu'une demande arrive.

L'évolution du nombre de produits dans le stock est décrite en utilisant les variables aléatoires suivantes :

$I(t)$: le niveau de stock possédé à l'instant t

$B(t)$: le niveau de rupture de stock à l'instant t

$IN(t)$: le niveau de stock à l'instant t , $IN(t) = I(t) - B(t)$

$IO(t)$: le nombre de commandes attendues à l'instant t

$IP(t)$: la position de stock à l'instant t , $IP(t) = IO(t) + I(t) - B(t)$

Pour chaque variable aléatoire $X = I, B, IN, IO, IP$, $\{X(t), t \geq 0\}$ définit un processus stochastique à temps continu. Supposons maintenant que le système de fabrication traite un seul ordre de fabrication à la fois, selon la discipline de service FIFO. Les temps de fabrication successifs des produits sont indépendants et suivent tous une loi exponentielle de taux μ où $\rho = \lambda / \mu$ satisfait la condition $\rho < 1$. En outre, nous supposons que les matières nécessaires pour exécuter la fabrication des produits sont toujours disponibles. Selon ces hypothèses et puisqu'un ordre de fabrication est déclenché chaque fois qu'une demande arrive, le système de fabrication est un modèle de file d'attente M/M/1 (μ, λ).

Soit $K(t)$ le nombre de commandes en attente de fabrication et en fabrication dans le système de fabrication à l'instant t . Le nombre de commandes attendues de l'installation de stock est exactement le nombre de commandes en attente de fabrication et en fabrication dans le système de fabrication : $IO(t) = K(t)$ pour $t \geq 0$. Le processus stochastique $\{K(t), t \geq 0\}$ est indépendant du niveau de stock nominal S . Si $S = 0$, nous avons alors un système de production à la commande. Si $S > 0$, le système analysé est un système de production pour stock.

Le processus stochastique $\{K(t), t \geq 0\}$ atteint à long terme un régime permanent indépendant de son état initial. En outre, selon la politique de stock nominal, la position de stock reste constante au niveau du stock nominal S : $IP(t) = S$ pour $t \geq 0$. Par conséquent, le processus stochastique $\{IN(t) = S - K(t), t \geq 0\}$ atteint aussi un régime permanent. Nous décrivons l'évolution du nombre de produits dans le stock en utilisant les variables d'état en régime permanent suivantes :

P : le nombre de commandes en attente de fabrication

K : le nombre de commandes en attente de fabrication et en fabrication

I : le niveau de stock possédé

B : le niveau de rupture de stock

Le fonctionnement du système est représenté par un réseau des files d'attente avec stations de synchronisation (Figure 1.15). Les stations de synchronisation permettent de relier le flux d'information et le flux physique. Dans la Figure 1.15, les flux d'information correspondent aux arrivées de demandes finales et aux ordres de fabrication. L'arrivée d'une demande finale déclenche instantanément un ordre de fabrication. La modélisation des systèmes de stock nominal et des systèmes Kanban par des réseaux des files d'attente avec stations de synchronisation est décrite par plusieurs auteurs (voir Di Mascolo *et al.* (1996), Baynat *et al.* (2001), Matta *et al.* (2005), Frein *et al.* (1995), Liberopoulos et Dallery (2003), Baynat *et al.* (2002) pour la modélisation des systèmes Kanban multi-étages, Karaesmen et Dallery

(2000), Duri *et al.* (2000), Karaesmen *et al.* (2002) pour la modélisation des systèmes à stock nominal multi-étages).

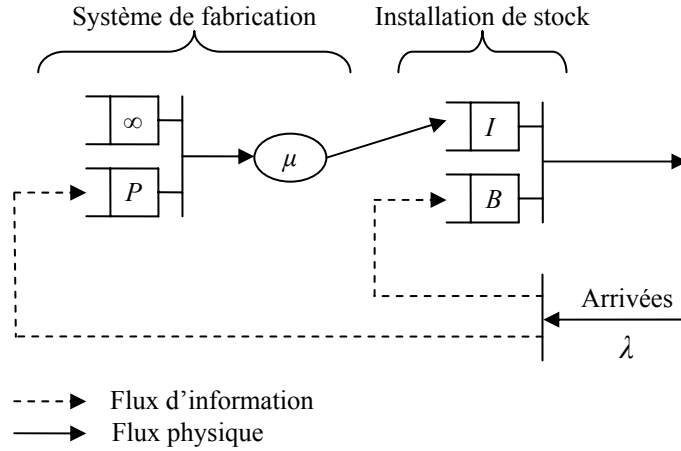


Figure 1.15. Représentation du système par les files d'attente

Supposons que chaque unité dans le stock induit un coût de stockage « h » et chaque demande retardée induit un coût de rupture de stock « b » par unité de temps. La fonction de coût considéré est la somme des coûts moyens de stockage et de rupture :

$$C(S) = E[hI + bB] = hE[I] + bE[B]$$

En régime permanent, $S = K + I - B$. Par conséquent, $I = [S - K]^+$ et $B = [K - S]^+$ où $[x]^+$ dénote $\max\{x, 0\}$. Soit $P_k = \Pr\{K = k\}$ la probabilité d'avoir k commandes dans la file d'attente M/M/1 (μ, λ) en régime permanent. Les mesures de performances qui sont le niveau moyen de stock possédé $E[I]$ et le niveau moyen de rupture de stock $E[B]$ sont calculées par les équations

$$E[I] = \sum_{k=0}^S (S - k)P_k,$$

$$E[B] = \sum_{k=S}^{\infty} (k - S)P_k.$$

La fonction de coût s'écrit alors

$$C(S) = h \sum_{k=0}^S (S - k)P_k + b \sum_{k=S}^{\infty} (k - S)P_k. \quad (1.1)$$

La fonction (1.1) a la même forme que celle de la fonction de coût d'un modèle de marchand de journaux. Le modèle de marchand de journaux est un modèle classique de gestion des stocks sur une seule période. Le stock est approvisionné une seule fois afin de satisfaire la demande de la période, qui est une variable

aléatoire de distribution connue (Khouja, 1999). La fonction (1.1) a la forme d'un modèle de marchand de journaux avec S étant la quantité de commande et K étant la demande par période (Shang et Song, 2003). La fonction (1.1) est développée dans l'annexe A.

Dans la littérature, le modèle présenté est nommé la file d'attente M/M/1 de production pour stock. En appliquant la politique de stock nominal ($S - 1, S$), le système de fabrication est en marche quand le niveau de stock est inférieur au niveau de stock nominal S . Sinon, quand le niveau de stock est égal au niveau de stock nominal S , le système de fabrication reste en état d'arrêt. Veatch et Wein (1996) ont prouvé l'optimalité de la politique de stock nominal pour le système étudié si le coût de lancement de la production est négligeable.

Il existe des travaux de recherche analysant les modèles de file d'attente de production pour stock avec produits multiples. Dans ce cas, le problème est de déterminer quand le système de fabrication est en marche et quel produit il doit fabriquer quand il est en marche (Zheng et Zipkin, 1990 ; Veatch et Wein, 1996 ; Ha, 1997a ; de Vericourt *et al.*, 2000). Pour le cas de deux produits ayant des taux moyens de fabrication équivalents, Ha (1997a) caractérise la politique optimale par un niveau de stock nominal pour chaque produit et une courbe de changement (en fonction des niveaux de stock des produits) qui détermine le produit à fabriquer. Pour le cas de deux produits ayant des taux moyens de fabrication différents, de Vericourt *et al.* (2000) déterminent, dans une région restreinte de l'espace d'état, la politique optimale.

En présence de plusieurs classes de clients, le problème est souvent nommé le problème de rationnement de stock. Ha (1997b) est parmi les premiers qui analysent un problème de rationnement de stock en supposant que les commandes qui ne sont pas servies sont perdues. Les classes de clients diffèrent par leur taux moyen d'arrivée et la pénalité générée par une vente perdue. L'auteur montre que la politique optimale est du type stock nominal, couplée par un niveau de rationnement pour chaque classe de clients ; un client est servi seulement si le niveau de stock possédé est supérieur au niveau de rationnement de la classe correspondante. Ha (2000) généralise ces résultats pour le cas où les temps de fabrication suivent une distribution d'Erlang. De Vericourt *et al.* (2002) analysent un système où les demandes qui ne sont pas satisfaites immédiatement sont retardées. La politique optimale est encore une politique de stock nominal, couplée par un niveau de rationnement pour chaque classe de client. Selon cette politique, une demande est satisfaite immédiatement si le niveau de stock possédé est supérieur au niveau de rationnement de la classe correspondante. De la même façon, un produit qui vient d'être fabriqué est utilisé pour satisfaire une demande retardée d'une classe si le niveau de stock possédé est supérieur ou égal au niveau de rationnement de la classe correspondante. S'il existe plusieurs classes qui satisfont cette condition, le produit est attribué à la classe ayant le coût de rupture le plus élevé. Si aucune classe avec un niveau de rupture de stock positif ne satisfait cette condition, le produit prend sa place dans le stock de sortie. Notons que le niveau de rationnement pour la classe ayant le coût de rupture le plus élevé est nul.

1.6. CONCLUSIONS

Dans ce chapitre, nous avons présenté les concepts généraux relatifs aux chaînes logistiques et au pilotage de flux dans les chaînes logistiques. Dans la littérature, nous trouvons différentes politiques de pilotage de flux, que nous avons classifiées en deux familles : les politiques réactives et les politiques proactives. Dans la première classe, nous trouvons les politiques classiques de gestion des stocks, les politiques de stock nominal et les politiques Kanban. Les politiques MRP et les politiques réactives modifiées rentrent dans la classe proactive.

Les principaux concepts de chaque politique ont été présentés dans le but de mettre en évidence leurs similarités et leurs différences ainsi que les avantages et inconvénients de les utiliser dans le vaste contexte des chaînes logistiques.

Parmi les politiques de pilotage étudiées, les travaux développés au cours de cette thèse s'intéressent en particulier à la politique de stock nominal appliquée aux systèmes de production et de stockage. Les hypothèses adoptées considèrent que les coûts de lancement de la production sont négligeables et que l'information sur les demandes finales futures n'est pas disponible.

Nous examinons les systèmes de production et de stockage à capacité limitée. L'analyse du comportement des systèmes étudiés s'appuie sur la théorie des files d'attente. Dans le cadre cette thèse, nous allons généraliser le modèle de file d'attente M/M/1 de production pour stock présenté dans la section précédente.

CHAPITRE II

2. Optimisation des décisions dans les chaînes logistiques

2.1. INTRODUCTION

Les chaînes logistiques décentralisées sont constituées de différentes entreprises qui interviennent dans le processus de fabrication d'une famille de produits. Pour chaque entreprise particulière, il s'agit principalement d'optimiser sa politique de production et d'approvisionnement par rapport à ses propres critères économiques. Dans le contexte de la chaîne logistique, ceci se traduit par une optimisation individuelle souvent effectuée d'une façon concurrentielle. Puisque les objectifs des acteurs sont souvent antagonistes, le caractère distribué de la structure décisionnelle peut conduire à une perte d'efficacité pour l'ensemble de la chaîne et nécessite des mécanismes de coordination permettant d'améliorer les performances globales, tout en limitant les risques encourus par chacun des partenaires.

Pour comprendre et maîtriser l'organisation des transactions entre partenaires d'une chaîne logistique, il est donc essentiel de représenter les antagonismes entre leurs objectifs économiques, ainsi que les éventuelles relations de dominance entre les entreprises concernées. L'outil mathématique privilégié pour cette analyse est la théorie des jeux. Cette théorie permet en effet de prévoir les comportements d'acteurs rationnels dans différents contextes d'interaction et d'information : situations de conflit, de dominance ou de coopération, information partielle ou complète. Par la suite, nous nous intéressons davantage à l'interaction entre la théorie des jeux et la gestion de chaînes logistiques qu'aux fondements de cette théorie. Cachon et Netessine (2004) et Leng et Parlar (2005) fournissent des revues de littérature sur les applications des concepts de la théorie des jeux dans le domaine de la gestion de chaînes logistiques. Pour plus de détails sur les concepts de la théorie des jeux, nous citons Osborne et Rubinstein (1994), Myerson (1991) et Fudenberg et Tirole (1991).

La théorie des jeux possède deux grands axes : la théorie des jeux coopératifs et la théorie des jeux non-coopératifs. La théorie des jeux coopératifs implique un changement majeur sur les paradigmes analysés en comparaison avec la théorie des jeux non-coopératifs. La théorie des jeux coopératifs se focalise sur la valeur de la coopération, c'est-à-dire la valeur qu'un ensemble de joueurs peuvent créer en coopérant, sans préciser les actions spécifiques que les joueurs doivent entreprendre afin de créer cette valeur. En

revanche, la théorie des jeux non-coopératifs se focalise sur les actions spécifiques des joueurs individuels. Par nature, les jeux coopératifs sont bien adaptés pour analyser les interactions entre les entreprises pendant la phase de conception de la chaîne logistique et donc l'apport de la coopération sur les performances de la chaîne au niveau stratégique. La théorie des jeux non-coopératifs permet, quant à elle, de prévoir les comportements des acteurs individuels au sein d'une chaîne logistique décentralisée, en particulier lorsqu'ils ont des intérêts conflictuels. Les jeux non-coopératifs sont donc beaucoup mieux adaptés pour analyser les effets des décisions décentralisées sur les performances de la chaîne aux niveaux tactique et opérationnel. Dans ce chapitre, nous nous concentrons principalement sur les jeux non-coopératifs. La deuxième section est consacrée aux concepts de base de la théorie des jeux. La troisième section décrit les jeux statiques et les études utilisant ce concept dans le but d'analyser les comportements des acteurs au sein des chaînes logistiques décentralisées. La quatrième section traite les études employant les jeux dynamiques. Dans la cinquième section, nous exposons les applications des jeux à information asymétrique. Dans la sixième section de ce chapitre, nous exposons brièvement les interactions entre la théorie des jeux coopératifs et la gestion de chaînes logistiques.

2.2. CONCEPTS DE BASE DE LA THÉORIE DES JEUX

La théorie des jeux est un domaine scientifique fondé sur un ensemble d'outils analytiques permettant de comprendre certains phénomènes observés quand plusieurs centres de décision interagissent, en particulier lorsqu'ils ont des intérêts conflictuels. Ces outils sont efficaces pour analyser les situations dans lesquelles la décision d'un acteur a une influence sur le gain des autres acteurs du jeu. La théorie des jeux a pour hypothèses : (a) que les acteurs (les joueurs) sont *rationnels*, c'est-à-dire qu'ils prennent leurs décisions dans le but de maximiser leur satisfaction, et (b) que chaque acteur prend en compte l'information dont il dispose sur les autres acteurs.

Les acteurs participant à un jeu forment un ensemble N avec $\text{card}(N) = n$. Chaque acteur $i \in N$ possède un ensemble de stratégies disponibles, noté X_i . Un acteur i peut choisir une stratégie particulière $x_i \in X_i$ ou il peut choisir une stratégie au hasard parmi certaines des stratégies disponibles. Dans le premier cas, nous parlons de *stratégie pure*. Dans un jeu à stratégies pures, l'espace de stratégies (pures) s'écrit $\prod_{i \in N} X_i = X_1 \times X_2 \times \dots \times X_n$. Un profil de stratégie possible du jeu est alors $x = (x_1, x_2, \dots, x_n)$ où $\forall x \in \prod_{i \in N} X_i$. Dans le deuxième cas, nous parlons de *stratégie mixte*. Une stratégie mixte θ_i pour un acteur i est une distribution de probabilité sur l'ensemble X_i , c'est-à-dire $\theta_i : X_i \rightarrow [0,1]$ tel que $\sum_{x_i \in X_i} \theta_i(x_i) = 1$. On écrit $\theta_i \in \Delta(X_i)$ où $\Delta(X_i)$ est l'ensemble de toutes les distributions de probabilités définies sur X_i . L'ensemble de stratégies mixtes d'un acteur i est alors $\Delta(X_i)$. Pour ce cas, l'espace de stratégies (mixtes) s'écrit $\prod_{i \in N} \Delta(X_i) = \Delta(X_1) \times \Delta(X_2) \times \dots \times \Delta(X_n)$ où

$\forall \theta = (\theta_1, \theta_2, \dots, \theta_n) \in \prod_{i \in N} \Delta(X_i)$. Une stratégie pure x_i peut être représentée en termes d'une stratégie mixte $\theta_i \in \Delta(X_i)$ avec $\theta_i(x_i) = 1$. Les stratégies mixtes sont souvent utilisées quand le jeu ne possède pas d'équilibre en stratégies pures ou quand l'équilibre en stratégies pures n'est pas Pareto efficace. La littérature sur la gestion de chaînes logistiques décentralisées se focalise seulement sur les équilibres en stratégies pures.

Un acteur i détermine sa stratégie afin de maximiser sa *fonction de paiement* (*payoff function*), appelée aussi *fonction d'utilité*, $\pi_i(x_1, x_2, \dots, x_n)$ où $\pi_i : \prod_{i \in N} X_i \rightarrow R$. La fonction de paiement de chaque acteur est son gain à la fin du jeu qui est déterminé selon la stratégie qu'il a adoptée ainsi que les stratégies adoptées par les autres joueurs. Plusieurs grandeurs peuvent être représentées par la fonction de paiement telles que les profits ou les coûts dans un contexte de minimisation. Dans les *jeux à somme nulle*, le gain global des acteurs est constant (considéré comme nul par normalisation) ainsi, le gain de l'un des acteurs se traduit par une perte pour un ou plusieurs des acteurs. Dans les *jeux à somme non-nulle*, les stratégies des acteurs agissent sur le gain global du jeu. Les problèmes traités dans le cadre de la gestion de chaînes logistiques entrent dans la classe de jeux à somme non-nulle.

Il existe plusieurs formes de jeu qui indiquent les règles fixées par les joueurs en termes de succession d'étapes du jeu. La *forme normale* ou *stratégique*, appelée aussi *jeu statique*, se déroule en une seule étape dans laquelle les joueurs choisissent leurs stratégies simultanément. Les *formes extensives* ou *jeux dynamiques* se déroulent en deux ou plusieurs étapes dans lesquelles les joueurs choisissent leurs stratégies simultanément ou successivement dans un ordre prédéterminé.

Un jeu est dit à *information complète* (ou *information symétrique*) si les ensembles de stratégies disponibles des différents acteurs ainsi que leurs fonctions de paiement sont connues de tous les joueurs. Dans un jeu à *information asymétrique*, les connaissances des joueurs diffèrent de l'un à l'autre.

2.3. JEUX STATIQUES

En utilisant les notations de Cachon et Netessine (2004), un jeu statique est un tuple $\langle N, (X_i)_{i \in N}, (\pi_i)_{i \in N} \rangle$ où N est l'ensemble de joueurs avec $\text{card}(N) = n$ et pour chaque joueur $i \in N$, X_i est l'ensemble de stratégies disponibles, $\pi_i(x_1, x_2, \dots, x_n)$ est la fonction de paiement. Les joueurs choisissent leur stratégie simultanément en connaissant l'ensemble de stratégies et la fonction de paiement de tous les joueurs dans le jeu.

Dans un jeu de n joueurs, le joueur i cherche une stratégie $x_i^* \in X_i$ qui maximise sa fonction de paiement $\pi_i(x_i, x_{-i})$ compte tenu des stratégies des autres joueurs, notées x_{-i} . Donc, la *fonction de meilleure réponse* (appelée aussi la *courbe de réaction*) du joueur i est définie par :

$$x_i^*(x_{-i}) = \arg \max_{x_i \in X_i} \pi_i(x_i, x_{-i}). \quad (2.1)$$

Étant donné x_{-i} , le joueur i peut être indifférent entre certaines stratégies dans l'ensemble X_i . Par conséquent, la courbe de réaction $x_i^*(x_{-i})$ peut être une *correspondance* (plutôt qu'une fonction) qui spécifie, pour chaque x_{-i} , un ensemble $x_i^*(x_{-i}) \subset X_i$. Supposons que $\prod_{i \in N} X_i \subset R^n$ soit un ensemble convexe. Si $\pi_i(x_i, x_{-i})$ est continue en x , différentiable et quasi-concave en x_i alors la courbe de réaction $x_i^*(x_{-i})$ est déterminée par les conditions du premier ordre :

$$\frac{\partial \pi_i(x_i, x_{-i})}{\partial x_i} = 0 \quad i = 1, \dots, n$$

L'équilibre de Nash caractérise l'état permanent du jeu statique. Un profil de stratégie $(x_1^*, x_2^*, \dots, x_n^*)$ est un *équilibre de Nash* si

$$\pi_i(x_i^*, x_{-i}^*) \geq \pi_i(x_i, x_{-i}^*) \quad \forall x_i \in X_i, i = 1, \dots, n$$

ou également si x_i^* est une meilleure réponse à x_{-i}^* pour tous les joueurs $i \in N$:

$$x_i^* \in x_i^*(x_{-i}^*) = \arg \max_{x_i \in X_i} \pi_i(x_i, x_{-i}^*) \quad i = 1, \dots, n \quad (2.2)$$

Pour arriver à cet équilibre, chaque joueur i prédit les stratégies des autres joueurs, x_{-i}^* , et choisit x_i^* en conséquence. S'il existe un équilibre qui satisfait les conditions (2.2), aucun joueur n'aura intérêt à dévier unilatéralement de cet équilibre. Selon cette définition, deux problèmes peuvent se poser : la non-existence d'un équilibre et l'existence d'équilibres multiples.

2.3.1. Existence de l'équilibre

Les équilibres de Nash se trouvent à l'intersection des courbes de réactions des différents joueurs, ce qui correspond à la notion du point fixe. Ainsi, l'existence d'un équilibre de Nash découle des théorèmes de point fixe. Pour un jeu statique $\langle N, (X_i)_{i \in N}, (\pi_i)_{i \in N} \rangle$, soit $R : \prod_{i \in N} X_i \rightarrow \prod_{i \in N} X_i$ la correspondance telle que $R(x) = \prod_{i \in N} x_i^*(x_{-i})$ pour $\forall x = (x_1, x_2, \dots, x_n) \in \prod_{i \in N} X_i$. Le point fixe d'une correspondance $F : S \rightarrow S$ est défini comme un point $x \in S$ qui satisfait $x \in F(x)$. De la même façon, le point fixe d'une fonction $f : S \rightarrow S$ est un point $x \in S$ qui satisfait $f(x) = x$. Suivant cette définition, un équilibre de Nash est un point fixe de la correspondance $R(x)$. Si on peut démontrer que la correspondance $R(x)$ a un point fixe, alors le jeu $\langle N, (X_i)_{i \in N}, (\pi_i)_{i \in N} \rangle$ a au moins un équilibre de

Nash. Si pour chaque joueur i , la meilleure réponse à $x_{-i} \in \prod_{j \in N - \{i\}} X_j$ est unique, l'équilibre de Nash est le point fixe de la fonction $r : \prod_{i \in N} X_i \rightarrow \prod_{i \in N} X_i$ telle que $r(x) = (x_1^*(x_{-1}), x_2^*(x_{-2}), \dots, x_n^*(x_{-n}))$ pour $\forall x = (x_1, x_2, \dots, x_n) \in \prod_{i \in N} X_i$.

Les trois théorèmes majeurs de point fixe utilisés dans le but de montrer l'existence d'un équilibre de Nash sont ceux de Kakutani, Brouwer et Tarski (Cachon et Netessine, 2004). John F. Nash a montré en 1950 que tout jeu fini (avec N et X_i fini pour tous $i \in N$) a un équilibre de Nash en stratégies mixtes en utilisant le théorème de point fixe de Kakutani (Myerson, 1991). Néanmoins, l'application directe de ces théorèmes se montre difficile dans plupart des cas. Des méthodes alternatives, dérivées des théorèmes de point fixe, ont été proposées.

Le théorème de Debreu, Glicksberg et Fan (dérivé du théorème de point fixe de Kakutani) (Cachon et Netessine, 2004) indique qu'au moins un équilibre de Nash en stratégies pures existe si, pour chaque joueur $i \in N$, l'ensemble de stratégies X_i est compact (borné et fermé) et convexe et la fonction de paiement $\pi_i(x_i, x_{-i})$ est continue en x et quasi-concave en x_i . Notons que le théorème de Debreu, Glicksberg et Fan n'implique pas l'unicité de l'équilibre de Nash. Des équilibres multiples existent si les courbes de réaction se croisent plus qu'une fois. Parmi les travaux utilisant le théorème de Debreu, Glicksberg et Fan, nous citons Cachon et Zipkin (1999), Lippman et McCardle (1997), Mahajan et van Ryzin (2001) et Netessine *et al.* (2006).

Si la quasi-concavité des fonctions de paiement ne peut pas être vérifiée, la démonstration de l'existence d'un équilibre de Nash en stratégies pures peut se baser sur la notion de jeu super-modulaire. Un jeu est super-modulaire si la courbe de réaction de chaque joueur est non-décroissante par rapport aux stratégies des autres joueurs. Une fonction de paiement $\pi_i(x_i, x_{-i})$ deux fois continûment dérivable est super-modulaire si et seulement si $\partial^2 \pi_i / \partial x_i \partial x_j \geq 0$ pour tous $x \in \prod_{i \in N} X_i$ et tous $j \neq i$. Un jeu est super-modulaire si les fonctions de paiement de tous les joueurs sont super-modulaires. Selon le théorème de point fixe de Tarski, c'est une condition suffisante pour l'existence d'un équilibre de Nash en stratégies pures et donc chaque jeu super-modulaire possède au moins un équilibre de Nash. Notons que la notion de super-modularité ne nécessite pas la continuité des fonctions de paiement et il existe plusieurs façons de montrer qu'un jeu est super-modulaire (voir Fudenberg et Tirole (1991) pour plus de détails). Cachon (2001b), Cachon et Lariviere (1999), Lippman et McCardle (1997) ont utilisé cette propriété pour démontrer l'existence d'un équilibre de Nash.

2.3.2. Unicité de l'équilibre

L'existence d'un équilibre de Nash unique qui signifie un choix fixe pour chaque joueur est toujours recherchée. L'existence d'équilibres multiples peut générer des problèmes de cohérence, car les différents

acteurs peuvent opter pour des équilibres différents, ce qui peut aboutir à une solution qui n'est pas un équilibre de Nash.

Analysons un exemple où l'existence d'équilibres multiples génère une solution qui n'est pas un équilibre de Nash. La Figure 2.1 montre la *représentation normale* (ou *matricielle*) d'un jeu statistique à deux joueurs connu sous le nom de la bataille des sexes. Un garçon J_1 et une fille J_2 désirent sortir ensemble le samedi après-midi, néanmoins, le garçon préfère d'aller à un match de foot alors que la fille souhaite faire du shopping. Aucun d'entre eux ne prendra plaisir à sortir seul. Chaque joueur possède deux stratégies : aller à un match de foot ou faire du shopping.

$J_1 \backslash J_2$	<i>foot</i>	<i>shopping</i>
<i>foot</i>	$(\pi_1(f, f) = 3, \pi_2(f, f) = 1)$	$(\pi_1(f, s) = 0, \pi_2(f, s) = 0)$
<i>shopping</i>	$(\pi_1(s, f) = 0, \pi_2(s, f) = 0)$	$(\pi_1(s, s) = 1, \pi_2(s, s) = 3)$

Figure 2.1. Bataille des sexes

Dans le jeu de bataille des sexes, il existe deux équilibres en stratégies pures : $(foot, foot)$ et $(shopping, shopping)$. L'équilibre $(foot, foot)$ favorise le joueur J_1 et l'équilibre $(shopping, shopping)$ favorise le joueur J_2 . Donc, les joueurs préfèrent des équilibres de Nash différents, ce qui peut ramène le système au point $(foot, shopping)$ qui n'est pas un équilibre de Nash et qui est moins désiré par les joueurs.

Démontrer l'unicité de l'équilibre est généralement plus difficile que de démontrer l'existence. Cachon et Netessine (2004) présentent plusieurs méthodes utilisées dans cet objectif. Les méthodes présentées assument l'existence d'un équilibre. Les auteurs précisent qu'il n'existe pas des méthodes générales pour démontrer l'unicité de l'équilibre dans le cas de jeux super-modulaires.

2.3.3. Pareto optimalité de l'équilibre

Un point $a \in R^n$ est dit *Pareto inférieur* à ou *Pareto dominé* par un point $b \in R^n$, si au moins une composante de a (par exemple la $j^{\text{ème}}$ composante, noté a_j) est inférieure à la composante correspondante de b ($a_j < b_j$) tandis que toutes les autres composantes de a ne sont pas supérieures aux composantes correspondantes de b ($a_i \leq b_i$ pour tous $i \neq j$). En d'autres termes, on dit que b est *Pareto supérieur* à a ou que b domine a (au sens de Pareto). Un point $b \in S$ où $S \subset R^n$ est un point *Pareto optimal* ou *Pareto efficace* de S s'il n'est pas Pareto dominé par un autre point dans S . La *Pareto frontière* de S se réfère à l'ensemble de tous les points Pareto optimaux de S .

En utilisant ces définitions, l'efficacité des différents profils de stratégie peuvent être comparée. Dans un jeu de n joueurs, les paiements obtenus par les joueurs avec un profil de stratégie $x = (x_1, x_2, \dots, x_n)$ peuvent être représentés comme un point en R^n : $(\pi_1(x), \pi_2(x), \dots, \pi_n(x))$. Selon cette représentation, un profil de stratégie $x = (x_1, x_2, \dots, x_n) \in \prod_{i \in N} X_i$ est un profil *Pareto optimal* (ou un profil qui aboutit à un résultat $(\pi_1(x), \pi_2(x), \dots, \pi_n(x))$ Pareto optimal) s'il n'existe pas d'autre profil dans l'ensemble $\prod_{i \in N} X_i$ avec lequel le paiement d'au moins un joueur est plus élevé tandis que les paiements des autres joueurs ne sont pas moins élevés. Autrement dit, un profil de stratégie est Pareto optimal s'il n'est pas possible d'augmenter le paiement d'un joueur quelconque sans réduire celui d'un autre.

L'équilibre de Nash d'un jeu ne correspond pas nécessairement à un profil de stratégie Pareto optimal. Considérons l'exemple de la Figure 2.2 connu sous le nom du *dilemme du prisonnier*. Le dilemme du prisonnier est un jeu à somme non-nulle dans lequel les deux prisonniers J_1 et J_2 sont accusés de conspirer pour un crime. Le juge promet que si l'un d'eux avoue le crime, celui-ci sera libéré et le deuxième prisonnier obtiendra la peine maximale de 6 ans. Si les deux avouent, tous les deux seront condamnés à une peine de 5 ans. Si aucun prisonnier n'avoue, ils seront condamnés à une peine de seulement 1 an. Chaque joueur possède alors deux stratégies : avouer ou ne pas avouer.

J ₁ \ J ₂		
	<i>avouer</i>	<i>ne pas avouer</i>
<i>avouer</i>	$(\pi_1(a, a) = -5, \pi_2(a, a) = -5)$	$(\pi_1(a, na) = 0, \pi_2(a, na) = -6)$
<i>ne pas avouer</i>	$(\pi_1(na, a) = -6, \pi_2(na, a) = 0)$	$(\pi_1(na, na) = -1, \pi_2(na, na) = -1)$

Figure 2.2. Dilemme du prisonnier

Une stratégie $y \in X_i$ est dite strictement dominante si $\pi_i(y, x_{-i}) > \pi_i(z, x_{-i})$ pour tous $z \in X_i - \{y\}$ et tous $x_{-i} \in \prod_{j \in N - \{i\}} X_j$. S'il existe une stratégie strictement dominante pour un joueur, alors cette stratégie sera jouée par ce joueur à l'équilibre de Nash du jeu. Si tous les joueurs ont une stratégie strictement dominante, alors l'équilibre de Nash est unique et se trouve à la rencontre de ces stratégies.

Dans le jeu de dilemme du prisonnier, la stratégie « avouer » est strictement dominante sur « ne pas avouer » pour tous les deux prisonniers. Donc, il existe un équilibre unique en stratégies pures : $(avouer, avouer)$. Par contre, le profil de stratégie $(avouer, avouer)$ est Pareto dominé par le profil $(ne\ pas\ avouer, ne\ pas\ avouer)$. Ici, le profil $(ne\ pas\ avouer, ne\ pas\ avouer)$ est un profil Pareto optimal. À l'équilibre, chacun des prisonniers choisit d'avouer même s'ils gagneraient à coopérer, en choisissant de ne pas avouer.

Dans les chaînes logistiques décentralisées, on suppose généralement que chaque acteur maximise l'espérance mathématique de ses profits (ou minimise l'espérance mathématique de ses coûts). Un tel comportement caractérise un acteur *neutre au risque*. Les performances optimales sont obtenues si l'équilibre du système décentralisé correspond à la solution théorique du système centralisé, c'est-à-dire à la solution qui maximise la somme des fonctions de paiement des acteurs. Par la suite, nous nous concentrons sur la coordination des chaînes logistiques décentralisées où chaque acteur est neutre au risque. Nous citons Gan *et al.* (2004) pour plus d'information sur la coordination des chaînes logistiques lorsque les acteurs sont averses au risque. Considérons une chaîne logistique décentralisée avec n acteurs neutres au risque, où chaque acteur i maximise sa fonction de paiement $\pi_i(x)$. Soit $\pi_0(x) = \sum_{i \in N} \pi_i(x)$ la fonction de profit du système centralisé correspondant et $x^0 = \arg \max_x \pi_0(x)$ la solution optimale du système centralisé. Le profit maximal du système centralisé définit la borne supérieure du profit total du système décentralisé : $\sum_{i \in N} \pi_i(x^*) \leq \pi_0(x^0)$. Si l'équilibre du système décentralisé, $x^* = (x_1^*, x_2^*, \dots, x_n^*)$, ne correspond pas à la solution optimale du système centralisé, $x^0 = (x_1^0, x_2^0, \dots, x_n^0)$, alors la décentralisation des décisions conduit à une dégradation des performances pour l'ensemble de la chaîne : $\sum_{i \in N} \pi_i(x^*) < \pi_0(x^0)$.

Dans la littérature sur la gestion de chaînes logistiques, la majorité des travaux s'intéresse au fait qu'un équilibre de Nash ne maximise pas le gain total de la chaîne logistique, et cherche à déterminer des contrats qui peuvent améliorer l'efficacité de la chaîne. L'efficacité d'une chaîne logistique est définie comme le ratio du profit total du système décentralisé au profit maximal du système centralisé. Les contrats renforcent les engagements des partenaires par partage des risques, des profits ou des coûts, à travers des paiements de transfert entre les acteurs et motivent les acteurs à opter pour les actions optimales du système centralisé. Un contrat coordonne la chaîne logistique si son application positionne l'équilibre du système décentralisé à la solution optimale du système centralisé (Cachon, 2003). Autrement dit, un contrat qui coordonne la chaîne logistique permet d'obtenir 100 % d'efficacité. Le prix d'achat, les remises sur quantité et les subventions/pénalités sont des exemples de paramètres de contrat utilisés dans le but de coordonner les chaînes logistiques décentralisées. Cachon (2003) et Tsay *et al.* (1999) fournissent des revues de littérature sur la coordination des chaînes logistiques par contrats. Nous classifions cette littérature en trois catégories.

La première catégorie concerne les travaux qui analysent l'équilibre de la chaîne logistique décentralisée et cherchent à déterminer un contrat qui améliore son efficacité ou qui coordonnent la chaîne (Cachon et Zipkin, 1999 ; Caldentey et Wein, 2003). Le contrat proposé doit idéalement aboutir à un équilibre qui Pareto domine l'équilibre sans le contrat. S'il existe un contrat qui coordonne la chaîne logistique et qui aboutit à un équilibre qui Pareto domine l'équilibre sans le contrat, les acteurs accepteront l'application du contrat proposé.

Les travaux dans la deuxième catégorie ne définissent pas les règles du jeu entre les acteurs. L'équilibre du système et le contrat adopté par les acteurs sont des résultats d'un processus de négociation non précisé. Le but est de déterminer les contrats qui sont susceptibles d'être adoptés par les acteurs indépendamment des règles du processus de négociation et des pouvoirs de négociation des acteurs, c'est-à-dire les contrats Pareto optimaux. Un contrat est dit Pareto optimal s'il n'est pas possible d'augmenter le paiement d'un joueur quelconque sans réduire celui d'un autre en appliquant un autre contrat dans l'ensemble de contrats considéré (Cachon, 2003 ; Cachon 2004). Si l'ensemble de contrats considéré contient des contrats qui coordonnent la chaîne logistique, alors les contrats qui coordonnent la chaîne constituent la Pareto frontière de cet ensemble. Si les contrats dans la Pareto frontière permettent de partager le gain total de la chaîne arbitrairement entre les acteurs, les acteurs opteront idéalement pour un contrat dans la Pareto frontière et le gain de chaque acteur dépendra des règles du processus de négociation. La flexibilité du type de contrat proposé est un aspect important pour de telles analyses. Un type de contrat est dit flexible s'il permet de partager le gain total de la chaîne arbitrairement entre les acteurs en ajustant les valeurs de ses paramètres (Cachon, 2003). Ici, nous parlons d'un ensemble de contrats où chaque contrat spécifique de cet ensemble permet une division différente du gain total. Dans ce cadre, Jemai et Karaesmen (2007), Cachon (2003), Lariviere (1999) et Cachon (1999b) se focalisent sur un seul type de contrat à la fois. Ils déterminent les valeurs des paramètres qui coordonnent la chaîne logistique et montrent que l'ensemble de contrats obtenu permette de partager le gain total de la chaîne arbitrairement entre les acteurs. Cachon (2004) analyse différents types de contrats en même temps et cherche à déterminer la Pareto frontière de l'ensemble de contrats considéré.

La troisième catégorie concerne les travaux précisant la procédure de négociation entre les acteurs. La plupart des travaux que nous citons dans cette catégorie (Lariviere et Porteus, 2001 ; Dong et Rudi, 2004 ; Corbett et Tang, 1999 ; Corbett et de Groote, 2000 ; Corbett *et al.*, 2004 ; Cachon et Zhang, 2006 ; Cachon et Lariviere, 2001 ; Plambeck et Zenios, 2003) adoptent une approche principal-agent (Fudenberg et Tirole, 1991) dans laquelle le principal propose un contrat à son agent et l'agent peut soit accepter soit refuser cette proposition. Dans la littérature économique, le modèle principal-agent est souvent étudié en présence de sélection adverse (*adverse selection*), de signalisation ou de risque moral (*moral hazard*) (Fudenberg et Tirole, 1991). Dans le cas de sélection adverse, le principal est imparfaitement informé sur la nature (le type) de son agent. Dans le cas de signalisation, c'est l'agent qui est imparfaitement informé sur le type de son principal. Dans les deux cas, il s'agit donc d'étudier les conséquences d'une asymétrie d'information. Dans le cas de risque moral, le principal ne peut pas observer l'action de son agent. Le cas de risque moral est rarement traité dans la littérature sur la gestion de chaînes logistiques (Plambeck et Zenios, 2003). La section 2.4.1 expose les études avec information complète citées dans cette catégorie (Lariviere et Porteus, 2001 ; Dong et Rudi, 2004). La section 2.5 expose les études avec information asymétrique (Corbett et Tang, 1999 ; Corbett et de Groote, 2000 ; Corbett *et al.*, 2004 ; Cachon et Zhang, 2006 ; Cachon et Lariviere, 2001).

2.3.4. Jeux statiques et compétition dans les chaînes logistiques

Il existe plusieurs travaux qui analysent, dans le cadre de jeux statiques, la compétition entre les entreprises d'une même chaîne logistique. Cachon et Zipkin (1999) étudient une chaîne logistique à deux niveaux, constituée d'une part d'un détaillant ayant une demande finale exogène et d'autre part de son fournisseur. La demande finale qui s'étale sur plusieurs périodes est aléatoire et stationnaire. Les deux entreprises ont des délais d'approvisionnement fixes. Les acteurs possèdent des coûts de stockage locaux et ils partagent les coûts de rupture de stock du détaillant. Deux jeux statiques sont étudiés. Dans le premier jeu, nommé SE, la gestion des stocks dans les entreprises est accomplie suivant une politique de stock nominal du type échelon et la stratégie de chaque acteur $i = 1, 2$ détermine son niveau de stock nominal du type échelon $S_i \in X_i = [0, S]$ où S est une très grande constante. Dans le deuxième jeu, nommé SI, le stock de chaque étage est géré suivant une politique de stock nominal du type installation et la stratégie de chaque acteur $i = 1, 2$ est son niveau de stock nominal du type installation $\bar{S}_i$ où $S_i = \bar{S}_i + \bar{S}_2$ et $S_2 = \bar{S}_2$. Le critère de chaque acteur est de minimiser ses coûts moyens de stockage et de rupture sur un nombre infini de périodes. Les acteurs choisissent leur stratégie simultanément et, une fois que les stratégies sont choisies, chaque acteur s'engage à appliquer la stratégie adoptée pour chaque période. Les auteurs montrent que chaque jeu statique étudié possède un équilibre de Nash unique, qui ne correspond pas à la solution optimale du système centralisé. Les équilibres des jeux SE et SI sont différents. En termes de stock échelon, les acteurs installent les niveaux de stock nominaux plus élevés dans le jeu SI : $\bar{S}_1^* + \bar{S}_2^* > S_1^*$ et $\bar{S}_2^* > S_2^*$. Dans les deux cas, la compétition entre les entreprises diminue le niveau de stock nominal total de la chaîne en comparaison avec la solution optimale du système centralisé, noté (S_1^0, S_2^0) , qui peut être obtenue par l'algorithme de Clark et Scarf : $S_1^* < \bar{S}_1^* + \bar{S}_2^* < S_1^0 = \bar{S}_1^0 + \bar{S}_2^0$. Les auteurs proposent un contrat de coordination caractérisant des paiements de transfert linéaires entre les acteurs. L'application du contrat proposé ramène le système à ses performances optimales. Cachon (1999a) fournit une analyse détaillée des différents mécanismes de coordination proposés dans ce cadre.

Cachon (2001b) analyse la gestion compétitive des stocks dans un système de distribution constitué d'un fournisseur et de n détaillants identiques. La demande arrive chez les détaillants suivant un processus de Poisson. Les délais d'approvisionnement des entreprises sont fixes. La gestion des stocks chez les détaillants est accomplie suivant la politique (R_r, Q_r) et chez le fournisseur suivant la politique (R_w, Q_w) . Chaque acteur détermine son point de commande dans le but de minimiser ses coûts moyens par unité de temps. Dans ce jeu entre les détaillants et le fournisseur, il existe des équilibres multiples qui ne correspondent pas, en général, à la solution optimale du système centralisé. Les auteurs proposent plusieurs types de mécanismes de coordination.

Jemai et Karaesmen (2007) analysent une chaîne logistique à deux niveaux constituée d'un producteur et de son détaillant. La demande finale arrive chez le détaillant selon un processus de Poisson. Le producteur possède un système de fabrication à capacité limitée dans lequel les temps de fabrication successifs des produits sont indépendants et suivent tous une loi exponentielle. Les entreprises disposent de stocks locaux et contrôlent leurs niveaux de stocks nominaux. Le temps de transport entre les installations de stock des entreprises est supposé négligeable. Les auteurs montrent l'existence d'un équilibre de Nash. À l'équilibre, le niveau de stock nominal total du système décentralisé est toujours inférieur à la valeur optimale du niveau de stock nominal total du système centralisé. Un contrat de coordination qui définit un paiement de transfert linéaire entre les acteurs et qui coordonne la chaîne logistique décentralisée est proposé. Cachon (1999b) analyse un système similaire en supposant qu'une demande est perdue si elle n'est pas satisfaite.

Caldentey et Wein (2003) étudient les interactions entre un producteur fabriquant les produits finis et son détaillant. Le producteur fonctionne en mode de production à la commande et le détaillant dispose d'un stock de produits finis. Le système de fabrication du producteur est modélisé comme une file d'attente M/M/1. Les auteurs montrent qu'un équilibre de Nash unique existe quand le producteur contrôle sa capacité de production (son taux moyen de production) et le détaillant son niveau de stock nominal. L'équilibre de Nash ne correspond pas à la solution optimale du système centralisé. Ils proposent un mécanisme de coordination basé sur un paiement de transfert linéaire et montrent que le mécanisme proposé ramène l'équilibre du système décentralisé vers la solution optimale du système centralisé.

Dans la littérature, la compétition entre les entreprises d'un même niveau d'une chaîne logistique est souvent appelée la *compétition horizontale* tandis que la compétition entre les entreprises des différents niveaux est nommée la *compétition verticale*. Lippman et McCardle (1997) analysent la compétition horizontale entre les détaillants d'un système de distribution. Ils étudient un modèle duopole avec deux détaillants ainsi qu'un modèle oligopole avec n détaillants. Chaque détaillant affronte un problème de « marchand de journaux » et décide de son niveau de stock initial, qui peut aussi être interprété comme sa capacité de production. Les produits en stock sont ensuite utilisés pour satisfaire la demande qui survient sur une seule période. Les prix unitaires de vente sont exogènes et identiques. Par conséquent, la compétition n'est pas sur le prix de vente mais sur la disponibilité de produit. La demande agrégée du marché, qui est une variable aléatoire continue, est initialement attribuée aux détaillants selon une règle connue par tous les acteurs. Cette première allocation est indépendante des niveaux de stock initiaux des entreprises. Après la première allocation, si la quantité en stock chez un détaillant est inférieure à sa demande, alors une proportion de la demande qui n'est pas satisfaite est distribuée aux autres détaillants. C'est cette deuxième allocation qui crée la compétition entre les entreprises. Dans les modèles analysés, le critère de chaque acteur est de maximiser son profit moyen. Les auteurs montrent que, sous certaines conditions, chaque jeu statique étudié a un équilibre de Nash. Pour le modèle duopole, la somme des

niveaux de stock initiaux des entreprises en compétition est en général plus élevée que le niveau de stock initial d'une firme en situation de monopole. Mahajan et van Ryzin (2001) analysent un problème similaire avec n détaillants. Dans leur modèle, la demande du marché est le résultat des choix dynamiques de clients hétérogènes. Ils montrent que la compétition entre les entreprises augmente le niveau de stock initial dans le système.

Dans la littérature économique, il existe deux modèles de compétition classiques appliqués aux canaux de distribution : la compétition de Bertrand et la compétition de Cournot. Dans le modèle duopole de Bertrand classique, chaque entreprise $i = 1, 2$ décide de son prix de vente p_i et ensuite doit satisfaire sa demande réalisée qui dépend des prix de vente annoncés : $d_i(p_1, p_2)$. La compétition de Cournot porte sur la quantité de production. Chaque entreprise $i = 1, 2$ décide sa quantité de production q_i , et le prix de vente $p(q)$ dépend de la demande totale du marché $q = q_1 + q_2$. Dans la littérature sur la gestion de chaînes logistiques, il existe plusieurs travaux généralisant les modèles de Bertrand et de Cournot (Leng et Parlar, 2005).

Cachon et Harker (2002) présentent un modèle de compétition entre deux entreprises avec des économies d'échelle, c'est-à-dire que le coût par unité de demande est décroissant en quantité de demande pour chaque entreprise. Dans le canal de distribution considéré, chaque entreprise $i = 1, 2$ décide de son prix de vente p_i et la demande de chaque entreprise est ensuite réalisée en fonction des prix de vente fixés : $d_i(p_1, p_2)$. Deux jeux conformes à cette structure sont considérés. Dans le premier, les clients arrivent à l'entreprise i selon un processus de Poisson ayant le taux $d_i(p_1, p_2)$ et les deux entreprises fonctionnent comme des serveurs exponentiels. Ici, le gain obtenu par une vente réalisée est définie comme « $p_i - g_i$ » où g_i est l'espérance mathématique du temps de séjour des clients, qui est déterminée par le choix de capacité de production (ou de service) de l'entreprise. Dans le deuxième jeu, le taux de demande $d_i(p_1, p_2)$ est déterministe et chaque entreprise a un coût de commande fixe. La quantité de commande de chaque entreprise est déterminée par la formule bien connue de la *Quantité Économique de Commande* (QEC). Les auteurs analysent les situations où un équilibre de Nash en stratégies pures n'existe pas. Pour un modèle similaire de QEC, Bernstein et Federgruen (2003) analysent les compétitions de Bertrand et de Cournot entre n détaillants.

2.4. JEUX DYNAMIQUES

Les jeux statiques que nous avons analysés dans la section précédente ne permettent pas de représenter des modèles évoluant au cours du temps. Les *jeux dynamiques*, nommé aussi *jeux sous forme extensive*, tiennent compte de l'ordre dans lequel les joueurs choisissent leur stratégie. Ainsi, un jeu sous forme extensive peut être représenté comme un arbre de décision dans lequel le sommet initial et les sommets intermédiaires représentent les sommets de décision et les sommets terminaux représentent les résultats

possibles du jeu. Pour chaque joueur $i \in N$, un sommet avec le label $[i, k]$ est un sommet de décision contrôlé par le joueur i . Ici, k représente l'état d'information du joueur i concernant les stratégies adoptées dans les étapes précédentes. Quand un sommet de décision avec le label $[i, k]$ est le sommet racine, nous notons $k = 0$. Les arcs sortant d'un sommet avec le label $[i, k]$ représentent les stratégies disponibles du joueur i à ce coup. La Figure 2.3 montre la représentation d'un jeu à deux joueurs dans lequel les joueurs choisissent leurs stratégies successivement. Dans la première étape du jeu, le joueur 1 agit en premier et choisit une de ces deux stratégies disponibles, x_1 ou x_2 . Pour les sommets de décision contrôlés par le joueur 2, les états d'information $k = 1$ et $k = 2$ indiquent que le joueur 2 observe respectivement les choix x_1 et x_2 . Dans la deuxième étape du jeu, le joueur 2 choisit alors une stratégie parmi y_1 et y_2 en connaissant la stratégie adoptée par le joueur 1 dans la première étape du jeu. À la suite de ces deux choix successifs, on détermine le gain de chaque joueur.

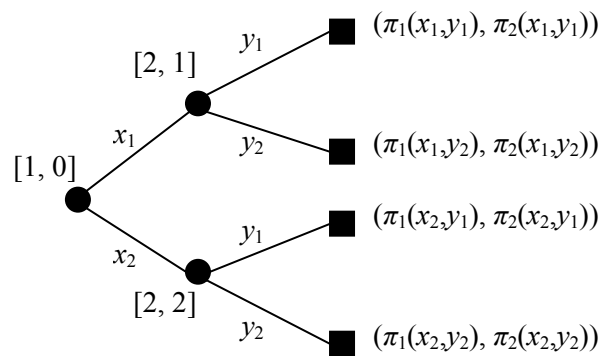


Figure 2.3. Représentation d'un jeu dynamique par un arbre de décision

Dans un jeu sous forme extensive, les sommets de décision contrôlés par le joueur i avec le même état d'information k constituent un *ensemble d'information*. Les sommets dans un même ensemble d'information ne sont pas distinguables par le joueur. Dans l'exemple de la Figure 2.3, chaque sommet est l'élément unique de son ensemble d'information. Donc, chaque joueur peut savoir à tout moment où il se situe dans l'arbre du jeu. Pour cet exemple, si l'on suppose que le joueur 2 choisit sa stratégie sans connaître la stratégie adoptée par le joueur 1, le jeu étudié devient équivalent à un jeu statique dans lequel les joueurs choisissent leurs stratégies simultanément. La Figure 2.4 montre la représentation de ce jeu statique à deux joueurs où l'état d'information $k = 3$ pour le joueur 2 indique que le joueur 2 ne connaît pas la stratégie adoptée par le joueur 1. Dans ce cas, les sommets de décision contrôlés par le joueur 2 constituent un ensemble d'information. Le joueur 2 ne peut pas savoir dans quel sommet il se trouve.

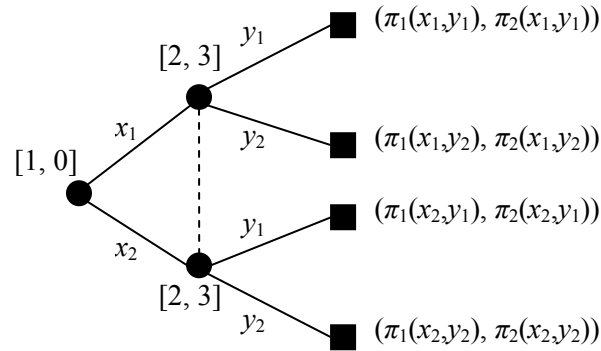


Figure 2.4. Représentation d'un jeu statique à deux joueurs avec un arbre de décision

Un jeu à *information parfaite* est un jeu sous forme extensive dont tous les ensembles d'information sont des singletons. Autrement dit, un jeu à information parfaite est un jeu qui se déroule en deux ou plusieurs étapes dans lesquelles les joueurs choisissent leur stratégie successivement (jamais simultanément). À chaque étape d'un jeu à information parfaite, le joueur qui détermine sa stratégie connaît les stratégies adoptées par lui-même et par les autres joueurs dans les étapes passées. Par exemple, le jeu d'échec est un jeu à information parfaite. Dans le cas contraire, le jeu est à *information imparfaite*.

Le concept de résolution, souvent utilisé dans la littérature sur la gestion de chaînes logistiques, est l'*équilibre parfait en sous-jeu* qui est un raffinement du concept d'équilibre de Nash pour les jeux sous forme extensive. Un sous-jeu d'un jeu sous forme extensive est un sous-arbre constitué d'un sommet racine qui est le seul élément de son ensemble d'information et de tous les sommets postérieurs de ce sommet racine. En outre, un sous-jeu doit contenir soit aucun, soit tous les éléments d'un ensemble d'information. L'équilibre parfait en sous-jeu d'un jeu sous forme extensive représente un équilibre de Nash pour chaque sous-jeu du jeu original. Pour les jeux à information parfaite ayant un nombre fini d'étapes, l'équilibre parfait en sous-jeu peut être déterminé par induction à rebours.

Par la suite, nous analysons les jeux sous forme extensive les plus souvent étudiés dans la littérature sur la gestion de chaînes logistiques, notamment les jeux de Stackelberg, les jeux répétés et les jeux stochastiques.

2.4.1. Jeux de Stackelberg

Les jeux sous forme extensive sont beaucoup moins traités dans les travaux s'intéressant aux chaînes logistiques, à l'exception des jeux en deux étapes. La notion de jeu en deux étapes a été introduite par H. von Stackelberg en 1934 comme une extension du modèle de Cournot. Dans un modèle duopole de Stackelberg, le premier joueur, appelé aussi le *meneur*, agit en premier et choisit une stratégie parmi l'ensemble de ses stratégies. Le deuxième joueur, le *suiveur*, observe ce choix et répond en adoptant une stratégie dans l'ensemble de ses stratégies possibles (voir l'exemple de la Figure 2.3). Similairement, un

modèle oligopole à n joueurs peut être défini comme un jeu à information parfaite dans lequel les joueurs choisissent leurs stratégies successivement dans un ordre prédéterminé. Par la suite, nous nous concentrons sur les jeux en deux étapes successives, appelé aussi *jeu de Stackelberg*.

L'*équilibre de Stackelberg* est l'aboutissement rationnel d'un jeu de Stackelberg. L'équilibre de Stackelberg est obtenu par induction à rebours, en déterminant d'abord la courbe de réaction du suiveur : $x_2^*(x_1) = \arg \max_{x_2 \in X_2} \pi_2(x_1, x_2)$. Dans le cas d'information complète, le meneur prédit correctement la courbe de réaction du suiveur et il tient compte du comportement du suiveur en intégrant la courbe de réaction de celui-ci à sa propre fonction de paiement : $\pi_1(x_1, x_2^*(x_1))$. Dans la première étape du jeu, le meneur opte pour une stratégie x_1^* qui maximise sa fonction de paiement $\pi_1(x_1, x_2^*(x_1))$. Dans la deuxième étape du jeu, en observant la stratégie adoptée par le meneur, le suiveur opte alors pour une stratégie x_2^* qui maximise $\pi_2(x_1^*, x_2)$. Selon ces réflexions, l'équilibre de Stackelberg peut être exprimé par les relations (Cachon et Netessine, 2004) :

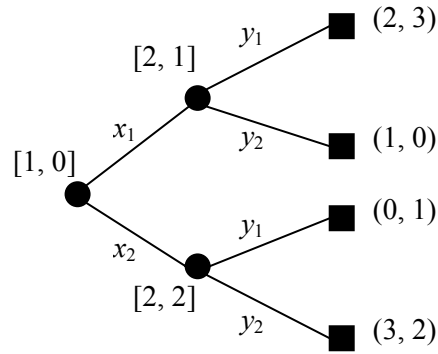
$$x_1^* \in \arg \max_{x_1 \in S_1} \pi_1(x_1, x_2^*(x_1)) \quad \text{et} \quad x_2^* \in x_2^*(x_1^*) \quad (2.3)$$

Le jeu de Stackelberg définit une relation de force non équitable en faveur du meneur car le meneur peut imposer l'équilibre qui lui convient en agissant en premier. Le gain du meneur à l'équilibre de Stackelberg est au moins aussi élevé que son gain à l'équilibre de Nash. Dans le pire des cas, le meneur opte pour un équilibre de Nash. Par conséquent, dans le cas d'information complète, un joueur préfère toujours de mener un jeu de Stackelberg que de participer à un jeu statique. La préférence du suiveur entre un jeu de Stackelberg et un jeu statique dépend de la structure spécifique du problème analysé.

Considérons le jeu statique (a.1) de la Figure 2.5. Il existe deux équilibres de Nash dans ce jeu, qui sont (x_1, y_1) et (x_2, y_2) . L'équilibre (x_1, y_1) favorise le joueur 2 et l'équilibre (x_2, y_2) favorise le joueur 1. Supposons maintenant que le joueur 1 est le meneur dans ce jeu (voir le jeu (a.2) dans la Figure 2.5). Le joueur 1 optera forcément pour x_2 en sachant que le suiveur choisira y_2 en observant son choix. L'équilibre de Stackelberg correspondant est alors (x_2, y_2) . Étant le meneur, le joueur 1 impose l'équilibre qui lui convient.

1 \ 2	y_1	y_2
x_1	(2, 3)	(1, 0)
x_2	(0, 1)	(3, 2)

(a.1) Jeu sous forme normale



(a.2) Jeu sous forme extensive

1 \ 2	y_1	y_2
x_1	(2, 1)	(3, 2)
x_2	(1, 4)	(4, 3)

(b) Jeu sous forme normale

Figure 2.5. Exemples des jeux à deux joueurs

Dans certaines situations, un joueur peut préférer être le suiveur qu'être le meneur dans un jeu de Stackelberg. Par exemple, le jeu statique (b) de la Figure 2.5 ne possède pas un équilibre de Nash. Si le joueur 1 est le meneur dans ce jeu, il impose l'équilibre (x_1, y_2) . Si le joueur 2 est le meneur, l'équilibre de Stackelberg résultant est (x_2, y_2) . Par conséquent, le joueur 2 préfère être le meneur mais le joueur 1 préfère être le suiveur.

L'existence d'un équilibre de Stackelberg est facile à montrer si les fonctions de paiement des acteurs sont continues. Néanmoins, démontrer l'unicité de l'équilibre peut être beaucoup plus difficile. Si la fonction de paiement du suiveur, $\pi_2(x_1, x_2)$, est strictement quasi-concave en x_2 et la fonction de paiement du meneur, $\pi_1(x_1, x_2^*(x_1))$, est strictement quasi-concave en x_1 , alors l'équilibre de Stackelberg unique (x_1^*, x_2^*) s'écrit de la manière suivante :

$$x_1^* = \arg \max_{x_1 \in S_1} \pi_1(x_1, x_2^*(x_1)) \quad \text{et} \quad x_2^* = x_2^*(x_1^*) \quad (2.4)$$

Les jeux en deux étapes successives sont bien indiqués pour modéliser les interactions entre les entreprises dans le cas où une entreprise est dominante. La situation de dominance est très fréquente dans les chaînes logistiques, par exemple pour un donneur d'ordre face à ses fournisseurs. Nous utilisons le concept de l'équilibre de Stackelberg dans le chapitre 4.

Parmi les travaux qui se sont intéressés aux jeux de Stackelberg, Lariviere et Porteus (2001) étudient une chaîne logistique à deux niveaux où le fournisseur (le meneur) décide de son prix de vente et le détaillant

(le suiveur), confronté à une demande aléatoire, détermine la quantité à commander qui maximise son espérance de profit. Le problème du détaillant a la même structure qu'un modèle de marchand de journaux. Le détaillant accepte un contrat (un prix de vente) proposé par le fournisseur si ce contrat lui permet d'obtenir au moins une espérance de profit nulle. Les fonctions de paiement du fournisseur, π_s , et du détaillant, π_r , s'écrivent comme suit :

$$\pi_s(w, q) = (w - c)q \quad (2.5)$$

$$\pi_r(w, q) = p[\min(D, q)] - wq \quad (2.6)$$

où p est le prix de vente du détaillant, w est le prix de vente du fournisseur, c est le coût unitaire de production et D est la variable aléatoire représentant la demande. L'espérance de profit total du système centralisé est $\pi_0(q) = p[\min(D, q)] - cq$. Puisque le fournisseur observe la meilleure réponse du détaillant, notée $q^*(w)$, la fonction de paiement du fournisseur peut être réécrite comme $\pi_s(w, q^*(w)) = (w - c)q^*(w)$ ou, de façon équivalente, comme $\pi_s(w^*(q), q) = (w^*(q) - c)q$ où $w(q) = q^{-1}(w)$. Donc, le fournisseur détermine la quantité de commande qui maximise ses profits et propose un prix de vente qui induit la quantité qu'il a choisi. Dans ce problème, l'équilibre de Stackelberg ne correspond pas à la solution optimale du système centralisé. Le prix de vente fixé par le fournisseur est toujours supérieur à son coût unitaire de production et la quantité de commande du détaillant est toujours inférieure à la quantité de commande optimale du système centralisé. Ce phénomène est connu dans la littérature sous le nom de *marginalisation double*. Dans la présentation classique de ce phénomène, le fournisseur décide de son prix de vente w et le détaillant fixe un prix p pour les produits qu'il vend sur le marché. La demande du marché, notée $q(p)$, est déterministe et décroissante en p . Dans ce cas, le fournisseur fixe un prix supérieur à son coût unitaire de production et le détaillant fixe un prix qui est supérieur au prix de vente optimal du système centralisé. Quand le prix de marché est exogène et la demande est aléatoire, le phénomène de marginalisation double apparaît sous la forme d'un niveau de stock insuffisant.

Les contrats de prix de vente sont souvent utilisés en pratique. La facilité d'application de ces contrats est leur point fort. Un contrat qui est difficile à appliquer peut engendrer des coûts administratifs trop élevés. Dans ces types de cas, les entreprises préfèrent souvent un contrat simple même s'il ne coordonne pas la chaîne logistique. Lariviere et Porteus (2001) analysent les effets de la variabilité de la demande sur les performances du système décentralisé. Ils montrent que, quand la demande est moins variable, la quantité de commande est moins sensible au prix de vente et donc le fournisseur fixe un prix plus élevé. Le gain du fournisseur augmente tandis que le gain du détaillant diminue et le gain total du système décentralisé est plus élevé. Dong et Rudi (2004) analysent un système similaire avec n détaillants. Les détaillants

peuvent livrer les produits entre eux dans le but de compenser la demande non satisfaite d'un détaillant par le stock d'inventures d'un autre. Dans ce jeu de Stackelberg entre le fournisseur et les détaillants, le fournisseur obtient, en général, le bénéfice de la livraison inter-détaillants en imposant un prix de vente plus élevé.

Le phénomène de marginalisation peut être évité en utilisant des contrats plus complexes, par exemple, un contrat de partage des revenus qui impose que le détaillant partage ses revenus provenant des ventes réalisées avec le fournisseur en plus le prix qu'il paie pour chaque produit qu'il achète (Cachon, 2003). Lariviere (1999), Tsay *et al.* (1999), Cachon (1999a), Cachon (2003) et Cachon (2004) analysent différents types de contrat qui coordonnent ce système décentralisé, c'est-à-dire qui permettent de ramener la quantité de commande du détaillant vers la quantité optimale de commande du système centralisé.

Corbett et Tang (1999) analysent le phénomène de marginalisation double avec une demande déterministe. Le fournisseur (le meneur) offre un contrat à son détaillant (le suiveur) et le détaillant décide de son prix de vente, noté p . La demande finale dépend du prix de vente annoncé : $q(p) = a - bp$ où $a \geq 0$ et $b \geq 0$ sont des paramètres connus. En d'autres termes, le détaillant détermine la quantité de commande q et le prix de marché est réalisé comme $p(q)$. Les auteurs analysent deux types de contrat : un contrat de prix de vente (w) et un contrat (w, L) constitué d'un paiement de transfert fixe L et d'un prix de vente w . En utilisant un contrat (w, L) , le problème d'optimisation du fournisseur s'écrit comme suit :

$$\max_{w, L} \pi_s(w, L, q^*) = (w - s)q^* - L \quad (2.7)$$

sous les contraintes

$$q^* = \arg \max_q \pi_r(w, L, q) = (p(q) - w - c)q + L \quad (2.8)$$

$$\pi_r(w, L, q^*) \geq \pi_r^{\min} \quad (2.9)$$

où s est le coût unitaire de production chez le fournisseur et c est le coût unitaire de production chez le détaillant. La contrainte de compatibilité d'incitation (2.8) définit la courbe de réaction du détaillant comme la stratégie maximisant sa fonction de paiement. La contrainte de rationalité individuelle (2.9) définit une borne inférieure pour le profit maximal du détaillant, car le détaillant accepte un contrat seulement si ce contrat lui permet d'obtenir un profit supérieur ou égal à son *profit de réservation* $\pi_r^{\min} > 0$. Le profit de réservation du détaillant exprime son opportunité dans le marché, c'est-à-dire le profit minimal qu'il peut obtenir en travaillant avec un autre fournisseur. Selon une hypothèse classique, le profit de réservation d'une entreprise est supposé nul si le marché est parfaitement compétitif (voir par

exemple Lariviere et Porteus (2001)). Dans le problème étudié, si le fournisseur propose un contrat de prix de vente (w), alors on rencontre le phénomène de marginalisation double. Le fournisseur obtient un tiers du profit total du système décentralisé. Les auteurs montrent que, en appliquant un contrat (w, L) , le fournisseur utilise le paiement fixe afin de laisser le détaillant avec son profit de réservation : $\pi_r(w, L^*, q^*) = \pi_r^{\min}$ pour $\forall w$ en fixant $L^* = \pi_r^{\min} - (p(q^*) - w - c)q^*$. Si $L^* < 0$, alors c'est le détaillant qui paie des frais de franchise à son fournisseur. Le prix de vente qui maximise le profit du fournisseur est $w^* = s$. Avec le contrat (w^*, L^*) , le fournisseur obtient le profit optimal du système centralisé moins le profit de réservation du fournisseur. Le contrat (w^*, L^*) coordonne alors cette chaîne logistique à deux niveaux.

Wang et Gerchak (2003) analysent un système d'assemblage confronté à une demande aléatoire. Dans le premier jeu analysé, l'assembleur (le meneur) décide un prix d'achat pour chaque composant et les fournisseurs (les suiveurs) décident leurs capacités de production. Dans le deuxième cas, les fournisseurs (les meneurs) décident leurs prix de vente et l'assembleur (le suiveur) décide sa capacité de production. Caldentey et Wein (2003) étudient un système de production/stockage avec un fournisseur qui décide de sa capacité de production et un détaillant qui décide de son niveau de stock nominal. Ils déterminent l'équilibre de Stackelberg en considérant d'abord le fournisseur comme le meneur et ensuite le détaillant comme le meneur. Cachon et Zipkin (1999) analysent un système dans lequel chaque acteur décide de son niveau de stock nominal. Ils étudient deux jeux de Stackelberg (1) avec le fournisseur comme meneur et (2) avec le détaillant comme meneur.

2.4.2. Jeux répétés

Un jeu répété se déroule en plusieurs périodes et, à chaque période, les joueurs jouent le même jeu. Autrement dit, l'ensemble de stratégies disponibles et la fonction de paiement de chaque joueur ainsi que les règles de jeu sont les mêmes pour chaque période. Dans un jeu répété, la fonction de paiement d'un joueur à une période donnée dépend des stratégies adoptées dans la période en question mais ne dépend pas des stratégies adoptées dans les périodes précédentes. Cette indépendance en temps est la raison pour laquelle les jeux répétés n'ont pas trouvé beaucoup d'applications dans le cadre de chaînes logistiques.

Notons que l'indépendance en temps n'implique pas que les joueurs ne prennent pas en compte les stratégies adoptées dans les périodes précédentes. En effet, un jeu répété peut avoir plusieurs équilibres et opter pour le même équilibre de Nash à chaque période et seulement l'un d'eux. Puisque dans un jeu répété les joueurs peuvent observer les aboutissements des jeux précédents, ils peuvent employer des *stratégies de menace* : un joueur choisi la même stratégie tant que son rival ne change pas sa stratégie et il change sa stratégie au moment où son rival change sa stratégie. Cette menace de choisir une autre stratégie relève d'un désir de coopération plutôt que d'une démarche conflictuelle et peut même induire

le meilleur aboutissement possible. Dans ce cadre, Debo et Sun (2005) analysent le modèle de Lariviere et Porteus (2001) dans un environnement répétitif et montrent que les performances optimales peuvent être obtenues avec un contrat de prix de vente si les facteurs d'actualisation des acteurs sont suffisamment élevés.

2.4.3. Jeux stochastiques

La dépendance en temps est une propriété qu'on retrouve souvent dans les chaînes logistiques par exemple avec les cumuls des stocks et des ruptures d'une période à l'autre. Un jeu stochastique est un type de jeu multi-périodes avec la dépendance en temps : la fonction de paiement d'un joueur à une période donnée dépend des stratégies adoptées dans la période en question ainsi que des stratégies adoptées dans les périodes précédentes.

La structure d'un jeu stochastique est essentiellement une combinaison d'un jeu statique avec un processus de décision Markovien : à chaque période, en plus des ensembles de stratégies et des fonctions de paiement, s'ajoute un mécanisme de transition $\Pr\{s'/s, x\}$, la probabilité de passer de l'état s à l'état s' en appliquant l'action x . La détermination des probabilités de transition repose typiquement sur le processus de la demande moyennant des hypothèses telles que des demandes indépendantes et identiquement distribuées. Sous ces hypothèses et s'il y a un seul décideur, on aboutit à une solution stationnaire telle qu'une politique d'approvisionnement et de stockage $(S-1, S)$, (R, Q) , etc. Néanmoins, dans le contexte d'un jeu stochastique, des équilibres non-stationnaires peuvent exister. Une approche standard est de supposer que la solution finale interdit de tels équilibres. Sous l'hypothèse que l'équilibre est stationnaire, on passe à un système équivalent de jeu statique et l'équilibre du jeu stochastique est obtenu comme une séquence d'équilibre de Nash. En adoptant cette approche, Cachon et Zipkin (1999) analysent une chaîne logistique à deux niveaux dans laquelle le fournisseur et le détaillant déterminent leurs niveaux de stock nominaux simultanément à chaque période sur un nombre infini de périodes. Netessine *et al.* (2006) analysent la compétition sur la disponibilité de produit entre deux détaillants. Les détaillants déterminent leurs niveaux de stock nominaux à chaque période. Bernstein et Federgruen (2004) étudient la compétition entre n détaillants quand la demande aléatoire d'un détaillant dépend de son prix de vente et de son niveau de service. Chaque détaillant décide de son prix de vente, de son niveau de service et de son niveau de stock nominal.

Van Mieghem (1999) analyse un jeu stochastique à deux périodes. Un producteur et un sous-traitant ayant accès à différents marchés déterminent leurs capacités de production dans la première période. Les demandes aléatoires des marchés sont ensuite réalisées. Dans la deuxième période, les acteurs déterminent leurs quantités de production et de vente avec l'option de sous-traitance. Les auteurs déterminent d'abord l'équilibre du sous-jeu de production et de sous-traitance et montrent ensuite qu'un équilibre parfait en sous-jeu existe. En utilisant une approche similaire, Van Mieghem et Dada (1999)

analysent la compétition entre n détaillants. Les détaillants déterminent leurs capacités de production simultanément dans la première période et leurs quantités de production sous les contraintes de capacités dans la deuxième période. Le prix de vente, noté $p = \varepsilon - Q$ où ε est une variable aléatoire, est ensuite réalisé comme une fonction de la quantité totale de production Q . Les auteurs montrent que le jeu étudié possède un équilibre parfait en sous-jeu.

2.5. JEUX AVEC INFORMATION ASYMÉTRIQUE

Dans les travaux précédents, les stratégies des différents acteurs ainsi que leurs fonctions d'utilité sont supposées connues de tous les joueurs. Cependant dans les chaînes logistiques, il est fréquent de trouver des entreprises bénéficiant d'une meilleure prévision que d'autres ou d'une connaissance supérieure des fonctions de paiement de leurs partenaires. La théorie des jeux traite ce type de situations dans le cadre des jeux à information asymétrique. Nous pouvons classer les jeux à information asymétrique étudiés dans la littérature en trois catégories : les jeux en deux étapes où le suiveur manque d'information, les jeux en deux étapes où le meneur manque d'information et les jeux statiques à information asymétrique.

Dans la première catégorie, Cachon et Lariviere (2001) s'intéressent à un modèle dans lequel le détaillant (le meneur), qui fait face à une demande aléatoire, propose un contrat à son fournisseur. Le contrat proposé par le détaillant est constitué du nombre d'engagements (la quantité minimale de commande), noté m , du nombre d'options, notée o (où $m + o$ détermine la quantité maximale de commande), et des paiements correspondants : le détaillant paie le prix w_m pour chaque unité d'engagement, le prix w_o pour chaque unité d'option et le prix d'achat w_e pour chaque unité réellement livrée par le fournisseur. Les paramètres du contrat, notamment m , o , w_m , w_o et w_e , traduisent les stratégies du détaillant. Ici, le paiement correspondant aux engagements et aux options, $w_m m + w_o o$, peut être interprété comme un paiement fixe de la part du détaillant. Si le fournisseur accepte ce contrat, il décide de la capacité de production qu'il va mettre en œuvre. Après la réalisation de la demande, le détaillant passe sa commande dont la quantité doit être entre les quantités minimale et maximale déclarées. Les auteurs montrent qu'il n'est jamais profitable de déclarer une quantité minimale de commande. Dans le cas d'information complète, si le fournisseur est obligé de choisir une capacité de production égale à la quantité maximale de commande, alors le détaillant propose un contrat ayant les paramètres o , w_o et w_e . Le détaillant gagne le profit total du système centralisé en laissant le fournisseur avec un profit nul. Si le fournisseur est libre de choisir une capacité de production inférieure à la quantité maximale de commande, alors le détaillant offre un contrat ayant le prix d'achat w_e comme seul paramètre. Dans ce cas, l'équilibre du système décentralisé ne correspond pas à la solution optimale du système centralisé.

Dans le cas d'information asymétrique, Cachon et Lariviere (2001) définissent la demande du détaillant comme $D_\theta = \theta X$ où X est une variable aléatoire avec la distribution de probabilité $f(x)$ et θ est un paramètre ayant deux valeurs possibles : H (demande élevée) et L (demande faible). Le détaillant connaît

la valeur du paramètre θ et la distribution de probabilité $f(x)$ mais le fournisseur connaît seulement la distribution de probabilité $f(x)$. Le paramètre θ représente le type du détaillant. Avant d'observer le contrat proposé par le détaillant, le fournisseur attribue la probabilité ρ que la valeur exacte du paramètre θ soit H et la probabilité $1 - \rho$ que la valeur exacte du paramètre θ soit L . Le fournisseur sait qu'un détaillant du type L a tendance à transmettre des prévisions fausses pour obtenir une capacité de production plus élevée. Les auteurs se focalisent sur les équilibres distinguant un détaillant du type H d'un détaillant du type L . Autrement dit, dans les équilibres considérés, un détaillant du type H offre seulement un contrat du type H tandis qu'un détaillant du type L offre seulement un contrat du type L . Par conséquent, en observant le type du contrat offert, le fournisseur apprend la valeur réelle du paramètre θ et actualise la probabilité ρ soit à 1 soit à 0. Les auteurs montrent que, dans le cas où le fournisseur est obligé de choisir une capacité de production égale à la quantité maximale de commande, un détaillant du type H peut signaler son type au détaillant sans frais à travers le contrat qu'il propose. Par contre, quand le fournisseur est libre de choisir une capacité de production inférieure à la quantité maximale de commande, le fournisseur est obligé de payer des frais supplémentaire, par exemple en offrant un prix d'achat unitaire plus élevé ou un paiement de transfert fixe additionnel en comparaison avec le cas d'information complète, afin de pouvoir signaler son type à son fournisseur. Une autre solution est de déclarer une quantité minimale de commande qui induit des frais supplémentaire seulement si la demande réalisée est inférieure à la quantité déclarée.

Dans la deuxième catégorie, Corbett et Tang (1999) généralisent le problème (2.7) – (2.9) en supposant que le coût unitaire de production chez le détaillant est une variable aléatoire pour le fournisseur. Le fournisseur (le meneur) ne connaît pas le coût unitaire de production actuel chez le détaillant (le suiveur), mais possède une distribution de probabilité à priori $f(x)$ définie sur l'intervalle $[c_{\min}, c_{\max}]$. Les auteurs étudient l'application des contrats (w) , (w, L) et $(w(q), L(q))$. En appliquant le contrat $(w(q), L(q))$, le prix de vente et le paiement de transfert fixe dépendent de la quantité de commande optée par le détaillant. La quantité de commande optimale du détaillant, notée $q(c)$, dépend de son coût de production. Donc, le fournisseur peut reformuler ce contrat comme $(w(c), L(c))$. En effet, puisque le fournisseur ne connaît pas le coût de production réel, il offre un menu de contrats $\{w(x), L(x), q(x)\}$, appelé aussi *contrat de mécanisme*. En choisissant un contrat $(w(x), L(x))$, le détaillant annonce un coût de production x et fixe sa quantité de commande à $q(x)$. Par contre, le détaillant peut avoir intérêt à annoncer un coût de production faux : $x \neq c$. Selon le *principe de révélation*, il existe un contrat de mécanisme avec lequel le détaillant révèle son vrai coût : $x = c$. En d'autres termes, le fournisseur, étant le meneur, peut concevoir un contrat de mécanisme avec lequel la révélation de son vrai type est l'option la plus profitable pour le détaillant. Pour ce cas, le problème d'optimisation du fournisseur s'écrit comme suit :

$$\max_{w(\cdot), L(\cdot)} \int_{c_{\min}}^{c_{\max}} ((w(x) - s)q(x) - L(x))f(x)dx \quad (2.10)$$

sous les contraintes

$$q(x) = \arg \max_q \pi_r(w(x), L(x), q) = (p(q) - w(x) - c)q + L(x) \quad (2.11)$$

$$c = \arg \max_x \pi_r(w(x), L(x), q(x)) = (p(q(x)) - w(x) - c)q(x) + L(x) \quad \forall c \in [c_{\min}, c_{\max}] \quad (2.12)$$

$$\pi_r(w(c), L(c), q(c)) \geq \pi_r^{\min} \quad \forall c \in [c_{\min}, c_{\max}] \quad (2.13)$$

La contrainte (2.12) indique que le contrat de mécanisme $\{w(x), L(x), q(x)\}$ force le détaillant à déclarer son vrai type. Les auteurs déterminent les valeurs optimales du contrat de mécanisme $\{w(x), L(x), q(x)\}$. Contrairement à la première catégorie, la révélation d'information ne peut jamais être obtenue sans frais quand le meneur manque d'information (Cachon et Netessine, 2004). Le fournisseur est toujours obligé de verser un paiement supplémentaire, appelé *rente d'information*, à son détaillant afin de pouvoir apprendre son vrai type.

Il existe plusieurs travaux utilisant le principe de révélation. Corbett *et al.* (2004) généralisent le modèle étudié par Corbett et Tang (1999) en considérant que le fournisseur aussi a un profit de réservation positif. Corbett et de Groote (2000) analysent une chaîne logistique à deux niveaux avec une demande finale déterministe et un coût de commande fixe à chaque niveau. Le fournisseur (le meneur) ne connaît pas le coût unitaire de stockage réel de son détaillant, noté h , et propose un menu de contrats $\{P(x), q(x)\}$ où $P(x)$ est le paiement de transfert, $q(x)$ est la quantité de commande du détaillant et x est la variable aléatoire représentant le coût unitaire de stockage du détaillant. Cachon et Zhang (2006) étudient un système de production/stockage avec un fournisseur (le suiveur) qui décide de son taux moyen de production, supposé exponentiel, noté μ , et un détaillant (le meneur) qui décide de son niveau de stock nominal. Le coût encouru par le fournisseur est « $b\mu$ ». Le détaillant ne connaît pas le coût de capacité b et propose un menu de contrat $\{R(x), \mu(x)\}$ à son fournisseur, où $R(x)$ est le prix d'achat unitaire.

Dans la troisième catégorie, Cachon et Lariviere (1999) analysent une chaîne logistique à deux niveaux constituée d'un fournisseur et de n détaillants. Dans la première période du jeu, le fournisseur décide de sa capacité de production et annonce ainsi un mécanisme d'allocation de capacité. Dans la deuxième période, en connaissant les stratégies adoptées par le fournisseur, chaque détaillant observe sa demande et détermine sa quantité de commande. Si la somme des quantités de commande dépasse la capacité installée chez le fournisseur, le fournisseur utilise le mécanisme d'allocation annoncé afin d'allouer la capacité disponible entre les détaillants. Le jeu de la deuxième période est un jeu statique à information asymétrique : les détaillants prennent leurs décisions simultanément et chaque détaillant observe seulement sa propre demande. Soit $\theta_i \in \Theta_i$ le type (la demande) du détaillant i où Θ_i est l'ensemble de

types disponibles du détaillant i et $\mu_i(\theta_{i-1})$ la distribution de probabilité cumulative de la variable aléatoire θ_{i-1} étant donné le type du détaillant i . La distribution de probabilité $\mu_i(\theta_{i-1})$ est connue de tous les joueurs. Dans ce jeu statique à information asymétrique, chaque détaillant i maximise $Z_i(x_i(\theta_i), x_{-i}(\theta_{i-1})) = \int_{\prod_{j \neq i} \Theta_j} \pi_i(x_i(\theta_i), x_{-i}(\theta_{i-1})) d\mu_i(\theta_{i-1})$. Un profil de stratégie $(x_1^*(\theta_1), x_2^*(\theta_2), \dots, x_n^*(\theta_n))$ est un *équilibre de Nash Bayesian* si $Z_i(x_i^*(\theta_i), x_{-i}^*(\theta_{-i})) \geq Z_i(x_i, x_{-i}^*(\theta_{-i}))$ pour $\forall x_i \in X_i(\theta_i)$ et $i = 1, \dots, n$. Les auteurs montrent qu'il existe un équilibre de Nash Bayesian dans ce jeu. Le fournisseur peut concevoir un mécanisme d'allocation de capacité avec lequel chaque détaillant commande exactement ce dont il a besoin. Par contre, la chaîne logistique est, en général, plus performante avec un mécanisme qui laisse les détaillants gonfler leurs quantités de commande.

2.6. JEUX COOPÉRATIFS

La théorie des jeux coopératifs se focalise sur la valeur créée par les joueurs et la façon que les joueurs décident pour partager cette valeur. Dans un jeu coopératif, aucun joueur ne possède de pouvoir à priori et tous les joueurs sont considérés comme des négociateurs actifs. Cette théorie nous permet de modéliser les résultats des processus de négociation complexes et de répondre les questions plus générales comme « Quelle est la pouvoir d'un acteur face à la compétition ? » (Brandenburge et Stuart, 2007). Dans la littérature sur la gestion de chaînes logistiques, les travaux traitant les jeux coopératifs sont très rares mais notons que c'est un domaine de recherche qui devient plus en plus populaire.

La littérature sur la gestion de chaînes logistiques utilise seulement les concepts liées aux *jeux coopératifs à utilité transférable*, appelés aussi *jeux de coalition à utilité transférable* ou *jeux sous forme caractéristique*. Selon l'hypothèse d'utilité transférable, les joueurs peuvent librement partager une commodité (un montant d'argent) entre eux à travers des paiements de transfert. Un jeu sous forme caractéristique est un tuple $\langle N, v \rangle$ où N est l'ensemble de joueurs avec $\text{card}(N) = n$ et $v : P(N) \rightarrow R$ est la *fonction caractéristique*. Les joueurs sont libres de former des *coalitions* qui sont bénéfiques pour eux. Ici, $P(N)$ représente l'ensemble de coalitions où chaque coalition $S \in P(N)$ est un sous-ensemble non-vide de joueurs : $P(N) = \{S \mid S \subseteq N \text{ et } S \neq \emptyset\}$. La coalition $S = N$ formée par tous les joueurs dans le jeu est appelée *grande coalition*. Pour chaque coalition $S \in P(N)$, la fonction caractéristique $v(S)$ spécifie la valeur maximale que les membres de la coalition S peuvent créer en utilisant les ressources de tous les joueurs dans cette coalition. Les joueurs peuvent partager la valeur maximale de la coalition qu'ils forment sans restriction. La théorie des jeux coopératifs se focalise sur les *actions jointes* des joueurs formant une coalition. L'ensemble d'actions jointes d'une coalition S est constitué de toutes les divisions possibles de la valeur $v(S)$ entre les membres de cette coalition. Le résultat d'un jeu sous forme caractéristique est alors la spécification de la coalition formée et l'action jointe de cette coalition.

Pour un jeu sous forme caractéristique $\langle N, v \rangle$, une *allocation d'utilité* est un vecteur $(\pi_i)_{i \in N} = (\pi_1, \pi_2, \dots, \pi_n)$ qui spécifie comment la valeur maximale d'une coalition est partagée entre les joueurs. Ici, la composante π_i représente l'utilité du joueur $i = 1, \dots, n$. Une allocation d'utilité $(\pi_i)_{i \in N}$ est *amissible* pour une coalition S si $\sum_{i \in S} \pi_i \leq v(S)$, c'est-à-dire si les joueurs formant la coalition S peuvent obtenir les allocations précisées par $(\pi_i)_{i \in N}$ en partageant $v(S)$ entre eux. Une allocation d'utilité $(\pi_i)_{i \in N}$ est dite *efficace* si $\sum_{i \in N} \pi_i = v(N)$. Une coalition S est *améliorante* par rapport à une allocation d'utilité $(\pi_i)_{i \in N}$ si $v(S) > \sum_{i \in S} \pi_i$. Si une coalition S est améliorante par rapport à $(\pi_i)_{i \in N}$, alors il existe une allocation d'utilité $(\varphi_i)_{i \in N}$ admissible pour la coalition S telle que $\varphi_i > \pi_i$ pour $\forall i \in S$.

Le concept de résolution le plus souvent utilisé dans la théorie des jeux coopératifs est le *cœur* du jeu. Une allocation d'utilité $(\pi_i)_{i \in N}$ est dans le cœur d'un jeu sous forme caractéristique $\langle N, v \rangle$ si $(\pi_i)_{i \in N}$ est une allocation d'utilité efficace et s'il n'y a aucune coalition améliorante par rapport à $(\pi_i)_{i \in N}$. Autrement dit, une allocation d'utilité $(\pi_i)_{i \in N}$ est dans le cœur si

$$\sum_{i \in N} \pi_i = v(N), \quad (2.14)$$

$$\sum_{i \in S} \pi_i \geq v(S) \quad \forall S \subseteq N. \quad (2.15)$$

La condition (2.14) (*condition d'efficacité*) assure que la valeur maximale de la coalition est distribuée entre les joueurs. La condition (2.15) (*condition de participation*) indique qu'il n'y a pas de coalition S et d'allocation d'utilité $(\varphi_i)_{i \in N}$ admissible pour la coalition S telle que $\varphi_i > \pi_i \quad \forall i \in S$. Une allocation d'utilité dans le cœur est individuellement rationnelle pour chaque joueur $i \in N$: $\pi_i = x(\{i\}) \geq v(\{i\})$. En outre, une allocation d'utilité dans le cœur satisfait le principe de contribution marginal pour chaque joueur $i \in N$: $\pi_i \leq v(N) - v(N \setminus \{i\})$.

Le cœur d'un jeu sous forme caractéristique contient toutes les allocations d'utilité satisfaisant les conditions (2.14) et (2.15). Si le cœur est un singleton, aucun joueur n'aura intérêt à dévier unilatéralement de la grande coalition et les joueurs opteront pour l'allocation d'utilité qui constitue le cœur. Par contre, le cœur peut être vide ou trop large. Dans le cas où le cœur est vide, il devient difficile de prédire les coalitions qui seront formées et les allocations d'utilités qui seront adoptées par les joueurs. Quand le cœur n'est pas vide mais trop large, le gain de chaque joueur reste indéterminé et dépend encore des règles du processus de négociation résiduel et des pouvoirs de négociation des joueurs.

En termes d'applications spécifiques dans la littérature sur la gestion de chaînes logistiques, Hartman *et al.* (2000) analysent une chaîne logistique constituée de n détaillants où chaque détaillant est face à une demande aléatoire. Le problème d'optimisation de chaque détaillant a la même forme qu'un modèle de

marchand de journaux. Les auteurs étudient un jeu de centralisation de stocks dans lequel les détaillants ont l'option de centraliser leurs stocks et peuvent partager les bénéfices résultant de cette mise en commun de stocks (*inventory pooling*). Les auteurs montrent que le cœur de ce jeu sous forme caractéristique n'est pas vide sous certaines conditions sur les distributions de probabilité des demandes. Muller *et al.* (2002) relâchent ces restrictions et montrent que le cœur n'est jamais vide indépendamment des distributions de probabilité des demandes. Hartman et Dror (2003) analysent un problème similaire dans le but de déterminer les corrélations entre les demandes des détaillants qui maximisent les bénéfices obtenus en centralisent les stocks. Hartman et Dror (2005) généralisent le modèle de centralisation de stocks en considérant n détaillants non identiques en termes de coût de stockage et de rupture et étudient ainsi le cas où l'allocation des bénéfices résultant de la mise en commun de stocks est faite après que les demandes soient réalisées chez les détaillants.

Dans les cas où le cœur d'un jeu sous forme caractéristique est vide ou trop large, l'utilisation de ce concept de résolution comme théorie prédictive devient énigmatique. Il existe plusieurs concepts de résolution autres que le cœur utilisés dans la théorie des jeux coopératifs : l'ensemble stable, l'ensemble de négociation, le *nucleus*, le *kernel* et la valeur de Shapley (détaillés dans Osborne et Rubinstein (1994) et Myerson (1991)). La valeur de Shapley est un concept de résolution qui attribue une allocation unique à chaque joueur et par conséquent qui aboutit à un résultat unique. Par contre, la valeur de Shapley ainsi que les autres concepts de résolution n'ont pas trouvé d'application dans la littérature sur la gestion de chaînes logistiques.

Dans la théorie des jeux coopératifs, les actions spécifiques des acteurs individuels ne sont pas spécifiées. Réciproquement, la théorie des jeux non-coopératifs se focalise sur les actions spécifiques des acteurs individuels. Plus récemment, une forme hybride de jeu qui mélange les aspects des jeux non-coopératifs et des jeux coopératifs a été proposée sous le nom de *jeu biforme* (Brandenburge et Stuart, 2007). Un jeu biforme peut être interprété comme un jeu non-coopératif ayant comme résultats des jeux coopératifs. Dans un jeu biforme, chaque joueur $i \in N$ possède un ensemble de stratégies disponibles, noté X_i . Le jeu se déroule en deux étapes. Dans la première étape, les joueurs choisissent leur stratégie d'une manière non-coopérative. Un profil de stratégie $x = (x_1, x_2, \dots, x_n)$ opté par les joueurs dans la première étape détermine la fonction caractéristique $v(x) : P(N) \rightarrow R$ du jeu coopératif de la deuxième étape. Puisque le cœur de ce jeu coopératif est rarement un singleton, on attribue, en général, un index de confiance α_i pour chaque joueur $i \in N$. Soit $\pi_i^{\min}$ l'utilité minimale et $\pi_i^{\max}$ l'utilité maximale du joueur i dans le cœur. Si le cœur n'est pas vide, le joueur i espère gagner $\alpha_i \pi_i^{\max} + (1 - \alpha_i) \pi_i^{\min}$ dans le jeu coopératif de la deuxième étape. En attribuant une allocation spécifique à chaque joueur, le jeu de la première étape peut être analysé comme un jeu non-coopératif classique, par exemple comme un jeu statique. Nous citons Brandenburge et Stuart (2007) pour plus de détails sur les jeux biformes.

En comparaison avec les jeux coopératifs, les jeux biformes sont beaucoup plus traités dans la littérature sur la gestion de chaînes logistiques. Anupindi *et al.* (2001) analysent un système constitué de n détaillants. Les détaillants possèdent des installations de stock locales ainsi que des entrepôts centralisés. Dans la première étape du jeu (étape non-coopérative), chaque détaillant détermine son niveau de stock local et réserve un niveau de stock dans chaque entrepôt centralisé. Dans la deuxième étape du jeu (étape coopérative), après les réalisations des demandes, les détaillants peuvent former des coalitions et livrer les produits entre les différentes installations de stock appartenant à cette coalition dans le but de compenser la demande non satisfaite d'un détaillant par le stock invendu d'un autre. Les auteurs montrent qu'en général, le cœur de ce jeu biforme n'est pas vide et déterminent une allocation d'utilité qui est dans le cœur. Plambeck et Taylor (2005) analysent un modèle constitué de deux détaillants. Dans la première étape, chaque détaillant détermine sa capacité de production et son investissement sur l'innovation qui influence sa demande potentielle. Dans la deuxième étape, après les réalisations des demandes, les détaillants négocient sur l'allocation de la capacité totale du système. Chatain et Zemsky (2007) analysent les relations entre les entreprises qui veulent externaliser la gestion de certaines de leurs fonctions et les fournisseurs qui sont les prestataires potentiels. Dans la première étape du jeu biforme, les fournisseurs décident de rentrer ou non dans le marché. La deuxième étape est constituée des négociations entre les fournisseurs et les entreprises dans le marché. Wong *et al.* (2007) utilisent un jeu biforme dans le contexte de la mise en commun de pièces de rechange. Les entreprises possèdent des stocks des pièces de rechange utilisées dans le cas d'une panne dans le système de production. Dans la première étape du jeu, chaque entreprise décide de son niveau de stock nominal pour le stock local des pièces de rechange. Dans la deuxième étape, les entreprises ont l'option de livrer les pièces de rechange entre elles et négocient donc sur l'allocation des coûts résultants.

2.7. CONCLUSIONS

La gestion de chaînes logistiques est encore un vaste domaine de recherche, et les concepts de la théorie des jeux peuvent apporter des éléments importants pour répondre à ses diverses interrogations. Au long de ce chapitre, nous avons présenté quelques concepts de cette théorie et ses interactions avec la gestion de chaînes logistiques. Nous avons cité divers travaux réalisés en les classifiant par type de jeu étudié : jeux statiques, jeux de Stackelberg, jeux répétés, jeux stochastiques, jeux de signalisation, jeux de mécanisme, jeux statiques à information asymétrique, jeux coopératifs.

Un des problèmes qu'on peut rencontrer dans les chaînes logistiques décentralisées est que la gestion compétitive des stocks dans les systèmes de stockage multi-étages conduit, en général, à une diminution de la quantité totale de stock en comparaison avec celle d'une chaîne logistique centralisée. Dans le chapitre 4, nous quantifions ce phénomène pour une chaîne logistique décentralisée à deux étages de production/stockage. Nous analysons ce modèle décentralisé à l'aide des outils de la théorie des jeux introduits dans ce chapitre. Nous nous intéressons en particulier à des jeux de Stackelberg. Nous

proposons l'utilisation d'un contrat de coordination et nous montrons à travers l'équilibre de Stackelberg que l'utilisation de ce contrat ramène le système à ses performances optimales.

CHAPITRE III

3. Politique d'approvisionnement dans un système à plusieurs fournisseurs

3.1. INTRODUCTION

Les aléas liés aux délais d'approvisionnement des composants majeurs constituent un problème commun de la plupart des industries. Pour un composant majeur donné, établir des relations de longue durée avec le fournisseur le plus performant du marché en termes de délai de livraison peut être une stratégie pour assurer un approvisionnement satisfaisant en quantité et en délai. Néanmoins, les responsables des achats essaient souvent d'éviter la dépendance à un seul fournisseur à cause des multiples risques associés et ils favorisent les stratégies d'approvisionnements basés sur des relations commerciales avec plusieurs fournisseurs. Plus particulièrement dans un environnement aléatoire, le délai de livraison effectif des commandes et les coûts de stockage et de rupture peuvent être réduits en adoptant une stratégie qui favorise l'éclatement des commandes entre plusieurs fournisseurs.

Parmi les critères les plus importants pour choisir ses fournisseurs, on peut citer le prix d'achat (le prix net, les rabais, les conditions de paiement), la qualité, et le niveau de service du fournisseur (le délai de livraison, la variabilité du délai de livraison, la fiabilité, la flexibilité). Il est clair qu'en présence d'économies d'échelle à travers les coûts de commandes, les rabais sur quantité et les coûts de transport, effectuer les commandes d'approvisionnement chez plusieurs fournisseurs peut être plus coûteux. Par contre, dans la plupart des cas pratiques, les gains virtuels associés aux économies d'échelle peuvent être compensés par les gains sur les coûts de stockage et de rupture. En outre, les problèmes liés aux comportements opportunistes et aux asymétries d'information par rapport aux coûts vrais de fabrication peuvent être surmontés en adoptant une stratégie multi-fournisseurs qui intensifie la compétition entre les fournisseurs.

Il existe des exemples d'applications réelles qui favorisent les stratégies multi-fournisseurs dans le but de réduire le temps de service des clients finaux. Sun Micro Systems attribue les commandes en puces de mémoire à plusieurs fournisseurs en utilisant un système de scorecard. Un autre exemple significatif est le cas des producteurs Japonais de véhicules qui travaillent souvent avec deux fournisseurs pour les

composants majeurs (Cachon et Zhang, 2004). Malgré les applications réelles et les avantages potentiels, les études théoriques sur l'implémentation des stratégies multi-fournisseurs sont peu fréquentes dans la littérature. Nous analysons dans ce chapitre les stratégies multi-fournisseurs comme un des éléments de base en gestion de chaînes logistiques.

Nous étudions plus particulièrement la problématique d'approvisionnement d'un producteur ayant une demande externe et aléatoire d'un produit. Le producteur transforme un produit intermédiaire acheté en un produit vendu aux clients en un temps négligeable et dispose d'un stock à partir duquel ses clients vont être servis. Pour le produit intermédiaire utilisé, nous supposons qu'un nombre total de n fournisseurs homogènes en termes de prix et de qualité sont disponibles dans le marché. Les fournisseurs ont des capacités de fabrication limitées et des temps de fabrication aléatoires. Ils diffèrent en termes de variabilité du temps de fabrication, c'est-à-dire en termes de variabilité du délai de livraison. À chaque déclenchement de réapprovisionnement du stock, le producteur est libre de passer la commande à un fournisseur différent (Figure 3.1). La question que se pose est alors : doit-il effectuer toutes les commandes chez le fournisseur le plus performant du marché ou acheminer les commandes chez différents fournisseurs? Et s'il opte pour une stratégie multi-fournisseurs, quels sont les fournisseurs à qui il doit adresser les commandes ?

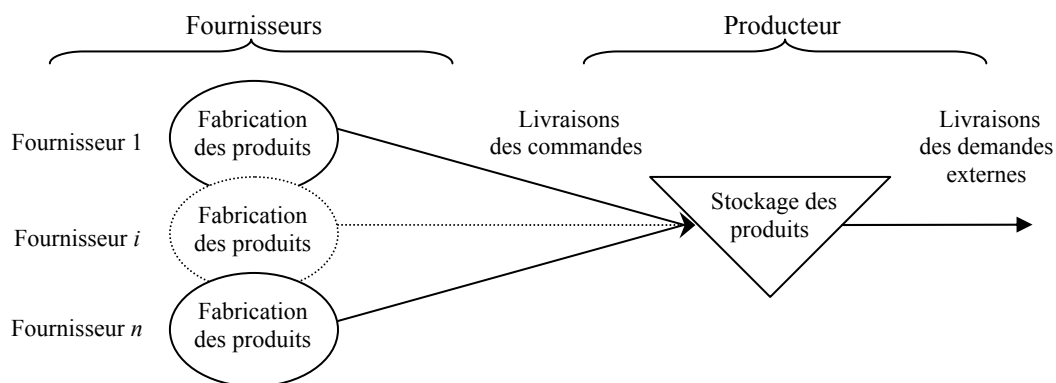


Figure 3.1. La chaîne logistique à deux étages

Les travaux analysant les stratégies multi-fournisseurs en présence des délais de livraison aléatoires se focalisent souvent sur le fractionnement de la quantité de commande entre plusieurs fournisseurs. Minner (2003) et Thomas et Tyworth (2006) fournissent des revues de littérature sur cette problématique. Dans ce cadre, les études analysent les effets des stratégies multi-fournisseurs sur les mesures statistiques de performances comme les délais moyens effectifs ou bien elles adoptent une démarche basée sur un critère économique.

Guo et Ganeshan (1995) analysent les effets de fractionner la quantité de commande Q entre n fournisseurs ayant les délais de livraison aléatoires $L_1, \dots, L_n$. Pour chaque commande de quantité Q , une commande de quantité Q/n est passée à chacun des n fournisseurs en même temps. Les auteurs supposent

que les délais de livraison des fournisseurs suivent la même loi de distribution (uniforme ou exponentielle) avec des caractéristiques identiques. En utilisant l'ensemble ordonné des délais de livraison, $L_{(1)} \leq L_{(2)} \leq \dots \leq L_{(n)}$, les délais de livraison effectifs sont définis comme le temps d'inter-arrivées des commandes successives : $T_1 = L_{(1)} = \min(L_1, \dots, L_n)$ et $T_r = L_{(r)} - L_{(r-1)}$ pour $r = 2, \dots, n$. Ils montrent que l'espérance et la variance du premier délai d'approvisionnement effectif, $E[T_1]$ et $Var[T_1]$, sont décroissantes de façon monotone en n . En outre, $E[T_r] \leq E[L]$ et $Var[T_r] \leq Var[L]$ où L est la variable aléatoire représentant le délai de livraison d'un seul fournisseur. Étant données des limites supérieures de $E[T_1]$ et de $Var[T_1]$, les auteurs déterminent le nombre satisfaisant de fournisseurs. Kelle et Miller (2001) analysent le cas de deux fournisseurs ayant des délais de livraison aléatoires, avec $E[L_1] \geq E[L_2]$. La demande étant aléatoire, ils approximent la demande pendant le délai d'approvisionnement comme une variable aléatoire suivant une loi de distribution exponentielle. Ils montrent que la probabilité de rupture de stock est réduite en fractionnant la quantité de commande entre les fournisseurs si $E[L_1] / E[L_2]$ est inférieur à un seuil déterminé. Pour un problème similaire, Fong *et al.* (2000) fournit les expressions analytiques de diverses mesures statistiques.

En prenant comme critère la minimisation de la somme des coûts de commande, de stockage et de rupture, Ramasesh *et al.* (2001) analysent les effets de fractionner la quantité de commande entre deux fournisseurs ayant des délais de livraison aléatoires indépendants et identiquement distribués (uniformes ou exponentiels). En appliquant la politique (R, Q) , la quantité de commande Q est fractionnée également entre les fournisseurs à chaque déclenchement de réapprovisionnement du stock. Le taux de la demande est supposé constant. Les valeurs optimales du point de commande R et de la quantité de commande Q sont obtenues par une recherche numérique pour les cas mono-fournisseur et bi-fournisseurs. Les auteurs montrent que la stratégie bi-fournisseurs diminue les coûts de stockage et de rupture. Mohebbi et Posner (1998) analysent un problème similaire avec des délais de livraison suivant des lois exponentielles non-identiques et des demandes arrivant selon un processus de Poisson composé. Ils supposent que les demandes non satisfaites sont entièrement perdues. La quantité de commande est partagée entre les fournisseurs selon une fraction $x \in (0,1)$ qui est une variable décision additionnelle. Les analyses numériques effectuées montrent que, en comparaison avec la stratégie mono-fournisseur, la stratégie bi-fournisseurs diminue les coûts opérationnels sauf si l'un des fournisseurs est beaucoup moins performant. Sedarage *et al.* (1999) étudient le cas de n fournisseurs ayant les délais non-identiques. Ils montrent par des analyses numériques qu'il existe un nombre optimal de fournisseurs qui minimise les coûts opérationnels.

Dans cette étude, nous supposons que les demandes arrivent chez le producteur selon un processus de Poisson et que les délais de livraison des fournisseurs suivent des lois exponentielles non-identiques. Les coûts de commande étant négligeables, le producteur applique une politique de stock nominal qui génère une commande d'une unité chaque fois qu'une demande unitaire arrive. Donc, la quantité de commande

ne peut pas être fractionnée entre les fournisseurs mais chaque commande peut être affectée à un fournisseur différent.

Dans le cas où le producteur ne dispose pas d'un stock de produit fini, c'est-à-dire dans le cas de production à la commande, le problème étudié devient proche du problème de répartition de charges entre plusieurs serveurs (processeurs, machines, etc.), dont le but est de déterminer la politique de répartition de trafic qui optimise des mesures de performances du système (le temps moyen d'attente des clients, etc.) (voir par exemple Liu et Righter (1998)). Les politiques de répartition de trafic utilisées sont classées par le niveau d'information nécessaire à leur fonctionnement. À part les caractéristiques de base du système comme le taux d'arrivée et les taux de service, la mise en place des politiques statiques en temps réel nécessite l'information sur les arrivées du système. Les politiques dynamiques peuvent être appliquées quand l'information sur l'état réel du système, comme les charges réelles des serveurs, est disponible. Le principe de la plus courte file d'attente est un exemple des politiques dynamiques utilisées. Pour le cas des serveurs homogènes dont les temps de service suivent des lois exponentielles, affecter le client au serveur ayant la plus courte file d'attente est prouvé optimal pour la plupart des modèles étudiés. En général, les politiques dynamiques mènent à un niveau de performances plus élevé que les politiques statiques. Par contre, le fait de disposer d'une information totale sur l'état réel du système n'est pas toujours réaliste. L'obtention de cette information peut être coûteuse ou peut consommer beaucoup de temps. En outre, dans le cas de serveurs hétérogènes, les performances des politiques dynamiques sont difficilement quantifiables. De l'autre côté, les performances des politiques statiques sont plus faciles à analyser. En outre, les politiques statiques sont bien adaptées pour la phase de conception d'un système de service car elles fournissent des limites de performances pour les systèmes contrôlés dynamiquement.

Dans cette étude, nous nous concentrons sur les politiques statiques. Nous supposons que le producteur a accès à l'information sur les taux moyens d'arrivées et des services. Le producteur applique une politique d'acheminement de commande probabiliste (voir par exemple Combé et Boxma (1994)). Autrement dit, chaque commande est affectée à un des fournisseurs selon des probabilités fixées à l'avance. La politique d'acheminement de commande proposée est une politique statique car chaque fois qu'une commande doit être affectée à un des fournisseurs le producteur n'utilise aucune information sur l'état réel du système d'approvisionnement. Le producteur décide des probabilités d'affectation des commandes une fois pour toutes au début de la période de fonctionnement du système. En prenant comme critère la minimisation des coûts de stockage et de rupture, les décisions du producteur sont alors le niveau de stock nominal et les probabilités d'affectation des commandes aux fournisseurs.

Benjaafar *et al.* (2004) analysent un problème similaire d'allocation de demande entre plusieurs fournisseurs en supposant l'existence des produits multiples. Chaque fournisseur est capable de fabriquer chaque produit avec un coût associé. Les auteurs utilisent une recherche numérique pour trouver les valeurs optimales du niveau de stock nominal et des probabilités d'affectation de commande. Pour le cas

de deux fournisseurs homogènes, ils montrent qu'une stratégie bi-fournisseurs diminue les coûts de stockage et de rupture. Dans cette étude, nous montrons qu'une solution approximative du problème étudié peut être obtenue en utilisant la solution optimale du problème dans le cas de production à la commande.

Dans la deuxième section, nous modélisons le système d'approvisionnement dans le cas de production pour stock et nous fournissons l'expression analytique de l'espérance de la somme des coûts de stockage et de rupture en fonction des variables de décision. La troisième section est consacrée au cas de la production à la commande dans laquelle nous déterminons les valeurs optimales des probabilités d'affectation des commandes aux fournisseurs en supposant que le niveau de stock nominal est nul. Dans la quatrième section, nous proposons une heuristique pour résoudre le problème d'optimisation du cas de production pour stock. Nous terminons ce chapitre par les analyses numériques et les conclusions.

3.2. MODÉLISATION DU PROBLÈME DANS LE CAS DE PRODUCTION POUR STOCK

Nous supposons que la demande arrive chaque fois chez le producteur en quantité unitaire et selon un processus de Poisson ayant le taux $\lambda > 0$. Notons que cette hypothèse est moins restrictive qu'il n'y paraît, dans la mesure où elle couvre aussi le cas assez fréquent de lots de fabrication de taille fixe et non fractionnable. Quand la demande arrive, elle est satisfaite s'il y a des produits dans le stock, sinon elle est retardée. La gestion de stock est accomplie suivant la politique de stock nominal $(S - 1, S)$. À l'état initial, le stock contient le niveau de stock nominal « S », avec $S \geq 0$. Le réapprovisionnement du stock est déclenché lorsque la position de stock devient « $S - 1$ », c'est-à-dire chaque fois qu'une demande arrive. À chaque déclenchement, le producteur effectue une commande unitaire chez un fournisseur extérieur.

Chaque fournisseur fabrique un produit dans une durée suivant une loi probabiliste exponentielle en traitant les commandes selon l'ordre FIFO. Soit μ_i le taux de distribution exponentielle du temps de fabrication pour le fournisseur $i = 1, \dots, n$ avec $\mu_i \neq \mu_j$ pour $j \in \{1, \dots, i-1, i+1, \dots, n\}$. Nous supposons que le temps de transport d'un produit entre les étages est négligeable. Quand le processus de fabrication d'une unité est terminé, l'unité prend sa place dans le stock de sortie du producteur s'il n'y a pas de ruptures de stock, sinon elle est utilisée pour satisfaire les demandes en attente en suivant l'ordre FIFO.

En appliquant la politique $(S - 1, S)$, nous supposons que lorsqu'une demande arrive, le producteur effectue une commande chez le fournisseur i avec la probabilité α_i où $0 \leq \alpha_i \leq 1$ et $\sum_{i=1}^n \alpha_i = 1$. Le processus d'acheminement de commande proposé est un processus d'acheminement de Bernoulli (*Bernoulli routing process*). Soit $\{N(t), t \geq 0\}$ le processus d'arrivée des demandes ayant le taux λ où $N(t)$ représente le nombre de demandes arrivées chez le producteur dans l'intervalle de temps $(0, t]$. Le processus d'acheminement de Bernoulli ne change pas les caractéristiques du processus d'arrivée (Ross,

2000) (voir l'annexe A, propriété A.1) : le processus $\{N_i(t), t \geq 0\}$ où $N_i(t)$ représente le nombre de commandes arrivées chez le fournisseur $i = 1, \dots, n$ dans l'intervalle de temps $(0, t]$ est un processus de Poisson ayant le taux $\alpha_i \lambda$. En outre, les processus de Poisson $\{N_i(t), t \geq 0\}$, $i = 1, \dots, n$, sont mutuellement indépendants avec $N(t) = \sum_{i=1}^n N_i(t)$.

Donc, nous pouvons modéliser le système de fabrication de chaque fournisseur « i » comme une file d'attente M/M/1 avec le taux d'arrivée $\alpha_i \lambda$ et le taux de service μ_i . Chaque file peut être étudiée de manière indépendante. Nous avons par conséquent un réseau ouvert de files d'attente avec n files M/M/1 en parallèle (Figure 3.2).

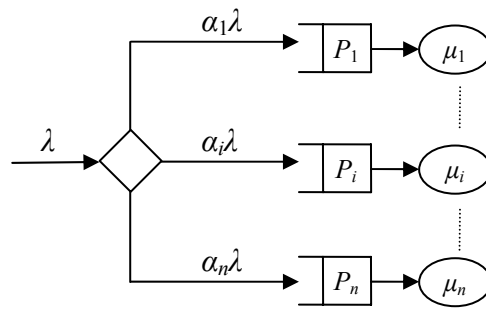


Figure 3.2. Le réseau ouvert de files d'attente avec n files en parallèle

Pour chaque fournisseur $i = 1, \dots, n$, l'évolution du système est décrite en utilisant des variables d'état en régime permanent :

P_i : le nombre de commandes en attente de fabrication

K_i : le nombre de commandes en attente de fabrication et en fabrication (P_i plus l'unité éventuellement en cours de fabrication)

Soit $P_{k_i} = \Pr\{K_i = k_i\}$ la probabilité stationnaire d'avoir k_i commandes dans le système de fabrication du fournisseur i . La probabilité d'avoir k_i commandes dans le système de fabrication du fournisseur i ayant le taux d'arrivée $\alpha_i \lambda$ et le taux de service μ_i est :

$$P_{k_i} = \left(\frac{\alpha_i \lambda}{\mu_i} \right)^{k_i} \left(1 - \frac{\alpha_i \lambda}{\mu_i} \right) \quad i = 1, \dots, n \quad (3.1)$$

La condition nécessaire et suffisante pour l'existence de la probabilité P_{k_i} en régime permanent est $\alpha_i \rho_i < 1$ où $\rho_i = \lambda / \mu_i$ et $\alpha_i \rho_i$ est le taux d'utilisation de la file M/M/1 ($\alpha_i \lambda, \mu_i$). Ensuite, l'état du réseau ouvert de files d'attente s'exprime par le vecteur $\mathbf{K} = (K_1, \dots, K_n)$. Sachant que le nombre de commandes

dans chaque file est indépendant de celui des autres, la probabilité d'être dans l'état $(k_1, \dots, k_n)$ en régime permanent est déterminée par le produit (Baskett *et al.*, 1975) :

$$\Pr\{K_1 = k_1, K_2 = k_2, \dots, K_n = k_n\} = P_{k_1} \times P_{k_2} \times \dots \times P_{k_n} = \prod_{i=1}^n (\alpha_i \rho_i)^{k_i} (1 - \alpha_i \rho_i) \quad (3.2)$$

Chez le producteur, chaque unité dans le stock induit un coût de stockage « h » et chaque demande retardée induit un coût de rupture de stock « b » par unité de temps. Le producteur décide de son niveau de stock nominal « S » et des paramètres de Bernoulli « $\alpha_1, \alpha_2, \dots, \alpha_n$ » dans le but de minimiser l'espérance de la somme des coûts de stockage et de rupture par unité de temps. Ici, les valeurs optimales des paramètres de Bernoulli expriment les probabilités optimales d'affectation des commandes aux fournisseurs en régime permanent.

L'évolution du nombre de produits dans le stock de sortie du producteur est décrite en utilisant les variables d'état en régime permanent suivantes :

V : le nombre de commandes qui ne sont pas encore livrées

I : le niveau de stock possédé

B : le niveau de rupture de stock

Selon la politique de stock nominal, la position de stock (= *commandes attendues* + *quantité en stock* – *demandes retardées*) reste constante au niveau de stock nominal S . Donc, $S = V + I - B$ en régime permanent. Par conséquent, $I = [S - V]^+$ et $B = [V - S]^+$, où $[x]^+$ dénote $\max\{x, 0\}$.

Dans le cas de n fournisseurs, le nombre de commandes qui ne sont pas encore livrées est le nombre total de commandes dans le réseau ouvert qui comprend les files d'attente des commandes chez les fournisseurs (Figure 3.3). En régime permanent, nous pouvons écrire :

$$V = K_1 + K_2 + \dots + K_n$$

La probabilité d'avoir v commandes dans le réseau ouvert de files d'attente s'écrit

$$P_v = \Pr\{K_1 + K_2 + \dots + K_n = v\}. \quad (3.3)$$

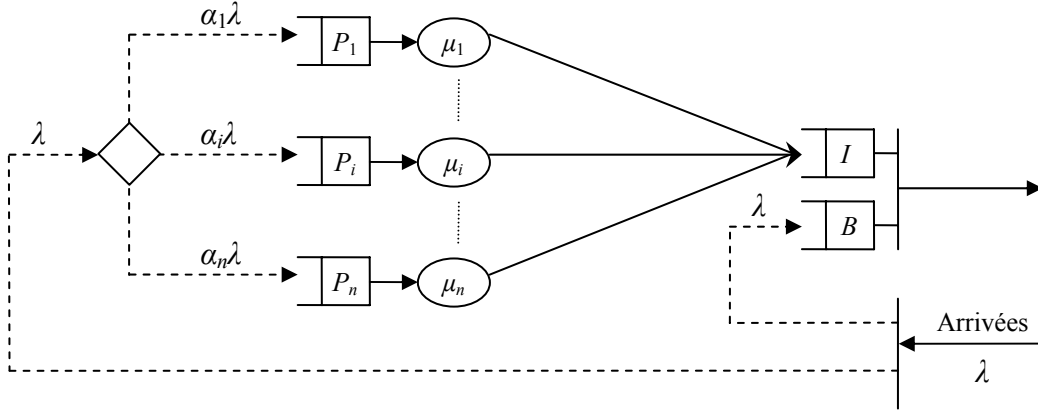


Figure 3.3. Représentation du système par les files d'attente

Ensuite, les mesures de performances du producteur, qui sont le niveau moyen de stock possédé $E[I]$ et le niveau moyen de rupture de stock $E[B]$, sont calculées par les équations :

$$E[I] = \sum_{v=0}^S (S - v) P_v \quad (3.4)$$

$$E[B] = \sum_{v=S}^{\infty} (v - S) P_v \quad (3.5)$$

La fonction de coût du producteur s'écrit alors

$$C(S, \alpha_1, \alpha_2, \dots, \alpha_n) = hE[I] + bE[B] = h \sum_{v=0}^S (S - v) P_v + b \sum_{v=S}^{\infty} (v - S) P_v. \quad (3.6)$$

Notons que P_v est une fonction des paramètres de Bernoulli « $\alpha_1, \alpha_2, \dots, \alpha_n$ ». En supposant $\alpha_i \rho_i \neq \alpha_j \rho_j$ pour $\forall i \neq j$, nous pouvons obtenir l'expression mathématique de la fonction de distribution de probabilité P_v en utilisant les propriétés des fonctions de génération de probabilité (Annexe B) :

$$P_v = \sum_{i=1}^n \left(\left(\prod_{\substack{j=1 \\ j \neq i}}^n \frac{1 - \alpha_j \rho_j}{\alpha_i \rho_i - \alpha_j \rho_j} \right) (\alpha_i \rho_i)^{n+v-1} (1 - \alpha_i \rho_i) \right) \quad (3.7)$$

Le cas $\alpha_i \rho_i = \alpha_j \rho_j$ pour $\exists i \neq j$ ne peut pas être exclu a priori. Toutefois, dans les analyses numériques, la valeur $\alpha_i \rho_i$ peut être représentée approximativement en remplaçant $\alpha_i \rho_i = \alpha_j \rho_j$ par $\alpha_i \rho_i = \alpha_j \rho_j + \varepsilon$ où $\varepsilon \ll \max_{i=1, \dots, n} \rho_i$.

Soit $E[V]$ le nombre moyen de commandes dans le réseau ouvert de files d'attente :

$$E[V] = \sum_{v=0}^{\infty} vP_v = \sum_{i=1}^n \frac{\alpha_i \rho_i}{1 - \alpha_i \rho_i} = \sum_{i=1}^n \frac{\alpha_i \lambda}{\mu_i - \alpha_i \lambda} \quad (3.8)$$

En utilisant la relation $S = E[V] + E[I] - E[B]$ et l'expression (3.7), nous pouvons écrire la fonction de coût du producteur comme suit :

$$\begin{aligned} C(S, \alpha_1, \dots, \alpha_n) &= (h + b)E[I] + b(E[V] - S) \\ &= (h + b) \sum_{i=1}^n \left(\left(\prod_{\substack{j=1 \\ j \neq i}}^n \frac{1 - \alpha_j \rho_j}{\alpha_i \rho_i - \alpha_j \rho_j} \right) \left(S(\alpha_i \rho_i)^{n-1} - \frac{(\alpha_i \rho_i)^n (1 - \alpha_i^S \rho_i^S)}{1 - \alpha_i \rho_i} \right) \right) + b \left(\sum_{i=1}^n \frac{\alpha_i \rho_i}{1 - \alpha_i \rho_i} - S \right) \end{aligned} \quad (3.9)$$

La fonction de coût (3.6) s'écrit également

$$\begin{aligned} C(S, \alpha_1, \dots, \alpha_n) &= h(S - E[V]) + (h + b)E[B] \\ &= h \left(S - \sum_{i=1}^n \frac{\alpha_i \rho_i}{1 - \alpha_i \rho_i} \right) + (h + b) \sum_{i=1}^n \left(\left(\prod_{\substack{j=1 \\ j \neq i}}^n \frac{1 - \alpha_j \rho_j}{\alpha_i \rho_i - \alpha_j \rho_j} \right) \frac{(\alpha_i \rho_i)^{S+n}}{1 - \alpha_i \rho_i} \right). \end{aligned}$$

La fonction de coût (3.9) n'est pas nécessairement convexe par rapport aux variables de décisions « $S, \alpha_1, \alpha_2, \dots, \alpha_n$ ». Par conséquent, la minimisation de la fonction de coût (3.9) sous les contraintes de l'admissibilité des paramètres de Bernoulli est un problème d'optimisation difficile.

3.3. POLITIQUE D'APPROVISIONNEMENT OPTIMAL DANS LE CAS DE PRODUCTION À LA COMMANDE

Dans le cas de production à la commande, le producteur est supposé installer un niveau de stock nominal nul, $S = 0$. La fonction de coût du producteur pour le cas de production à la commande exprime le coût moyen de rupture qui s'écrit

$$C(\alpha_1, \alpha_2, \dots, \alpha_n) = bE[V] = b \sum_{i=1}^n \frac{\alpha_i \lambda}{\mu_i - \alpha_i \lambda}. \quad (3.10)$$

Les variables de décision du producteur sont les paramètres de Bernoulli « $\alpha_1, \alpha_2, \dots, \alpha_n$ ». Selon la formule de LITTLE (Kleinrock, 1975), les valeurs optimales des paramètres de Bernoulli peuvent être obtenues en minimisant le temps moyen de séjour des commandes dans le réseau ouvert de files d'attente, $E[W] = E[V] / \lambda$. Notons que le temps moyen de séjour des commandes dans le réseau ouvert de files d'attente exprime ainsi le délai moyen d'approvisionnement du producteur.

Sans perte de généralité, les n fournisseurs peuvent être numérotés dans l'ordre décroissant de leurs taux de service : $\mu_1 > \mu_2 > \dots > \mu_n > 0$. Ensuite, le problème d'optimisation du producteur se formule de la manière suivante :

$$\Pi_1: \min_{\alpha_1, \alpha_2, \dots, \alpha_n} \sum_{i=1}^n \frac{\alpha_i}{\mu_i - \alpha_i \lambda} \quad (3.11)$$

sous les contraintes

$$0 \leq \alpha_i \leq 1 \quad i = 1, \dots, n \quad (3.12)$$

$$\sum_{i=1}^n \alpha_i = 1 \quad (3.13)$$

$$\frac{\alpha_i \lambda}{\mu_i} < 1 \quad i = 1, \dots, n \quad (3.14)$$

Les contraintes (3.12) et (3.13) expriment les relations nécessaires pour que les paramètres de Bernoulli soient admissibles. La stabilité du réseau ouvert de files d'attente nécessite les restrictions (3.14). En outre, la condition nécessaire et suffisante pour l'existence d'un ensemble des paramètres de Bernoulli $(\alpha_i, i = 1, \dots, n)$ qui satisfait les contraintes (3.12), (3.13) et (3.14) est

$$\lambda < \sum_{i=1}^n \mu_i. \quad (3.15)$$

Nous supposons que la condition (3.15) est satisfaite par les données du problème.

Toutes les contraintes sont linéaires et sous ces contraintes, la fonction objectif $E[W]$ est convexe car la satisfaction de la contrainte (3.14) implique $\partial^2 E[W] / \partial \alpha_i^2 > 0$ où

$$\frac{\partial^2 E[W]}{\partial \alpha_i^2} = \frac{\partial}{\partial \alpha_i} \left(\frac{\mu_i}{(\mu_i - \alpha_i \lambda)^2} \right) = \frac{2\mu_i \lambda}{(\mu_i - \alpha_i \lambda)^3}. \quad (3.16)$$

Par conséquent, le problème Π_1 est un problème d'optimisation convexe avec une solution unique.

3.3.1. Résolution du problème relaxé

Considérons tout d'abord un problème relaxé qui comprend le critère (3.11) et la contrainte (3.13). Soit Π_2 le problème relaxé correspondant. Le Lagrangien du problème Π_2 s'écrit

$$L(\alpha_1, \dots, \alpha_n, u) = \left(\sum_{i=1}^n \frac{\alpha_i}{\mu_i - \alpha_i \lambda} \right) - u \left(\sum_{i=1}^n \alpha_i - 1 \right). \quad (3.17)$$

où u est le paramètre de Lagrange associé à la contrainte (3.13). Soit α_i^* la valeur optimale du paramètre de Bernoulli α_i pour $i = 1, \dots, n$. Alors, la solution optimale du problème Π_2 satisfait les conditions suivantes :

$$\frac{\partial L}{\partial \alpha_i} = \frac{\mu_i}{(\mu_i - \alpha_i^* \lambda)^2} - u = 0 \quad i = 1, \dots, n \quad (3.18)$$

$$\sum_{i=1}^n \alpha_i^* = 1 \quad (3.19)$$

Pour quel que soit $i \in \{1, \dots, n\}$ et quel que soit $j \in \{1, \dots, n\}$, les conditions (3.18) impliquent

$$\frac{\mu_i}{(\mu_i - \alpha_i^* \lambda)^2} = \frac{\mu_j}{(\mu_j - \alpha_j^* \lambda)^2}. \quad (3.20)$$

Nous pouvons réécrire l'équation (3.20) de la manière suivante :

$$\mu_j - \alpha_j^* \lambda = \frac{\sqrt{\mu_j} (\mu_i - \alpha_i^* \lambda)}{\sqrt{\mu_i}} \quad (3.21)$$

En sommant sur j les deux termes de l'équation (3.21), nous obtenons

$$\sum_{j=1}^n (\mu_j - \alpha_j^* \lambda) = \sum_{j=1}^n \frac{\sqrt{\mu_j} (\mu_i - \alpha_i^* \lambda)}{\sqrt{\mu_i}}. \quad (3.22)$$

Sous la condition (3.19), l'équation (3.22) devient

$$\sum_{j=1}^n \mu_j - \lambda = \frac{\mu_i - \alpha_i^* \lambda}{\sqrt{\mu_i}} \sum_{j=1}^n \sqrt{\mu_j}. \quad (3.23)$$

Donc, nous pouvons écrire la propriété suivante.

Propriété 3.1 : Les valeurs optimales des paramètres de Bernoulli pour le problème relaxé sont définies par les équations :

$$\alpha_i^* = \frac{1}{\lambda} (\mu_i - \tau_n \sqrt{\mu_i}) \quad i = 1, \dots, n \quad (3.24)$$

$$\text{où } \tau_n = \frac{\sum_{j=1}^n \mu_j - \lambda}{\sum_{j=1}^n \sqrt{\mu_j}}$$

Notons que les valeurs optimales des paramètres de Bernoulli, $(\alpha_i^*)_{i=1}^n$, définies par les relations (3.24) satisfont la contrainte (3.13) par construction :

$$\sum_{i=1}^n \alpha_i^* = \sum_{i=1}^n \frac{1}{\lambda} (\mu_i - \tau_n \sqrt{\mu_i}) = 1$$

Si les valeurs optimales des paramètres de Bernoulli $(\alpha_i^*)_{i=1}^n$ satisfont les contraintes (3.12) et (3.14), alors le minimum du problème Π_2 est admissible et par conséquent optimal pour le problème Π_1 . En utilisant (3.24), nous pouvons récrire les contraintes (3.14) comme suit :

$$\mu_i - \tau_n \sqrt{\mu_i} < \mu_i \quad i = 1, \dots, n \quad (3.25)$$

D'après la condition (3.15), $\tau_n > 0$. Donc, les contraintes (3.25) sont satisfaites naturellement. Les contraintes (3.12) peuvent être remplacées par

$$0 \leq \mu_i - \tau_n \sqrt{\mu_i} \leq \lambda \quad i = 1, \dots, n. \quad (3.26)$$

L'inégalité $\mu_i - \tau_n \sqrt{\mu_i} \leq \lambda$ est satisfaite car $\sum_{i=1}^n \alpha_i^* = 1$. L'inégalité $\mu_i - \tau_n \sqrt{\mu_i} \geq 0$ devient

$$\mu_i \geq \tau_n^2 \quad i = 1, \dots, n. \quad (3.27)$$

3.3.2. Problème de choix restrictif

Selon la règle de numération de fournisseur retenue, $\mu_1 > \mu_2 > \dots > \mu_n > 0$. Par conséquent, si l'inégalité (3.27) est violée par i_0 où $1 \leq i_0 \leq n$, alors elle est aussi violée par $i = i_0 + 1, \dots, n$. Pour ce cas, nous pouvons formuler un problème de choix restrictif où $\alpha_i^* = 0$ est imposé pour $i = i_0, \dots, n$. Afin de montrer la pertinence du problème de choix restrictif, nous définissons le paramètre τ_m pour $m = 1, \dots, n+1$ en supposant $\mu_{n+1} = 0$:

$$\tau_m = \frac{\sum_{j=1}^m \mu_j - \lambda}{\sum_{j=1}^m \sqrt{\mu_j}} \quad (3.28)$$

L'évolution de ce paramètre est basée sur les lemmes suivants.

Lemme 3.1 : *Pour les valeurs positives des paramètres τ_m et τ_{m+1} ($m \leq n-1$), l'évolution de τ_m suit les règles ci-dessous :*

$$(1) \begin{cases} \tau_{m+1} < \tau_m \\ \Leftrightarrow \mu_{m+1} < \tau_m^2 \\ \Leftrightarrow \mu_{m+1} < \tau_{m+1}^2 \end{cases},$$

$$(2) \begin{cases} \tau_{m+1} = \tau_m \\ \Leftrightarrow \mu_{m+1} = \tau_m^2 \\ \Leftrightarrow \mu_{m+1} = \tau_{m+1}^2 \end{cases},$$

$$(3) \begin{cases} \tau_{m+1} > \tau_m \\ \Leftrightarrow \mu_{m+1} > \tau_m^2 \\ \Leftrightarrow \mu_{m+1} > \tau_{m+1}^2 \end{cases}.$$

Lemme 3.2 : *Le paramètre τ_m croît avec m pour $1 \leq m \leq m^*$ et décroît de manière monotone avec m pour $m^* < m < n$. La valeur maximale du paramètre τ_m est obtenue alors pour m^* ($1 \leq m^* \leq n$) qui est l'indice unique satisfaisant les relations suivantes :*

$$\sum_{i=1}^{m^*} \mu_i > \lambda \quad (3.29)$$

$$\mu_{m^*+1} \leq \tau_{m^*}^2 < \mu_{m^*} \quad (3.30)$$

Les démonstrations des lemmes 3.1 et 3.2 se trouvent dans l'annexe B.

La propriété suivante décrit la solution optimale du problème Π_1 .

Propriété 3.2 : *Considérons l'indice m^* ($1 \leq m^* \leq n$) qui satisfait les relations (3.29) et (3.30). Sous la condition que la relation (3.15) soit satisfaite, les valeurs optimales des paramètres de Bernoulli pour le problème Π_1 sont :*

$$\alpha_i^* = \begin{cases} \frac{1}{\lambda} (\mu_i - \tau_{m^*} \sqrt{\mu_i}) & \text{pour } i = 1, \dots, m^* \\ 0 & \text{pour } i = m^* + 1, \dots, n \end{cases} \quad (3.31)$$

La démonstration de la propriété 3.2 se trouve dans l'annexe B.

Dans le cas où $\mu_n > \tau_n^2$, l'indice unique qui satisfait les relations (3.29) et (3.30) est $m^* = n$. Donc, la solution optimale du problème Π_1 définie par la propriété 3.2 correspond à la solution optimale du problème Π_2 définie par la propriété 3.1. Pour ce cas, le producteur continue à exercer des relations commerciales avec chacun des n fournisseurs. Autrement dit, à chaque déclenchement de réapprovisionnement du stock, le producteur effectue la commande chez le fournisseur $i = 1, \dots, n$ avec une probabilité non-nulle.

Dans le cas contraire, le producteur détermine un seuil d'acceptation selon les taux de service des fournisseurs. Le seuil d'acceptation du producteur est exprimé à travers l'indice m^* défini par le lemme 3.2. Le producteur cesse d'exercer des relations commerciales avec les fournisseurs ayant un taux de service plus faible que μ_{m^*} . À chaque déclenchement de réapprovisionnement du stock, le producteur effectue la commande chez le fournisseur $i = 1, \dots, n$ avec une probabilité non nulle si $\mu_i \geq \mu_{m^*}$. Notons que cette règle de sélection est valide si les fournisseurs sont hétérogènes. Dans le cas où les taux de fabrication des fournisseurs sont égales, $\mu = \mu_1 = \mu_2 = \dots = \mu_n$, la solution optimale du problème est $\alpha^* = \alpha_1^* = \alpha_2^* = \dots = \alpha_n^* = 1/n$ selon la propriété 3.1.

3.4. MÉTHODE APPROXIMATIVE DANS LE CAS DE PRODUCTION POUR STOCK

Dans le cas de production pour stock, le producteur installe un niveau de stock nominal non-négatif, $S \geq 0$. Le niveau de stock nominal « S » devient alors une variable de décision additionnelle du producteur en comparaison avec le cas de production à la commande. Notons que le producteur peut aussi décider d'installer un niveau de stock nominal nul. Donc, le cas de production pour stock correspond à un cas général qui inclut aussi le cas de production à la commande.

Le problème d'optimisation du producteur, noté Π_3 , et de trouver les valeurs optimales du niveau de stock nominal « S » et des paramètres de Bernoulli « $\alpha_1, \alpha_2, \dots, \alpha_n$ » qui minimisent la fonction de coût (3.9) sous les contraintes (3.12), (3.13) et (3.14). Étant donné la complexité de la fonction (3.9), nous proposons une méthode approximative pour résoudre le problème Π_3 . La méthode de résolution approximative proposée calcule les valeurs des variables de décisions en deux étapes, chaque étape correspondant à un problème d'optimisation :

- Dans la première étape, le niveau de stock nominal est supposé d'être égal à zéro, $S = 0$. Les seules variables de décision sont les paramètres de Bernoulli « $\alpha_1, \alpha_2, \dots, \alpha_n$ ». Les hypothèses de cette étape sont identiques aux hypothèses du cas de production à la commande. Le problème d'optimisation correspondant est alors le problème Π_1 étudié dans la section 3.3.

- Dans la deuxième étape, les valeurs des paramètres de Bernoulli sont supposées données. La seule variable de décision est le niveau de stock nominal « S ». Nous appelons ce problème le problème d'optimisation Π_4 . Les valeurs optimales des paramètres de Bernoulli pour le problème Π_1 obtenues dans la première étape sont les données du problème Π_4 .

3.4.1. Calcul des paramètres de Bernoulli

En supposant $S = 0$, les valeurs des paramètres de Bernoulli peuvent être calculées d'après les résultats exposés dans la section 3.3. L'indice m^* est obtenu selon le lemme 3.2. En suite, les valeurs des paramètres de Bernoulli peuvent être calculées en utilisant l'expression (3.31).

3.4.2. Calcul du niveau de stock nominal

Le problème d'optimisation Π_4 consiste à trouver la valeur optimale du niveau de stock nominal « S » qui minimise la fonction $C(S)$ définie par l'équation (3.9), dont les paramètres de Bernoulli sont les résultats du problème Π_1 . La valeur optimale de S pour le problème Π_4 peut être calculée en utilisant la formule de la fraction critique de la version discrète du problème de vendeur de journaux (voir par exemple Veatch et Wein (1996)).

Soit $G(S) = C(S+1) - C(S)$ l'accroissement de la fonction de coût (3.9). La fonction $G(S)$ s'écrit

$$G(S) = (h + b) \Pr\{V \leq S\} - b. \quad (3.32)$$

La fonction de probabilité cumulative $F(S) = \Pr\{V \leq S\}$ du nombre de commandes dans le réseau ouvert de files d'attente est croissante et positive par définition. Par conséquent, la fonction $G(S)$ est croissante en S pour $S \geq 0$. Soit S^* la valeur optimale de S qui minimise la fonction $C(S)$. Les conditions nécessaires et suffisantes pour que S^* soit la valeur optimale sont donc les suivantes :

$$C(S^*) \leq C(S^* - 1) \Rightarrow G(S^* - 1) \leq 0 \quad (3.33)$$

$$C(S^*) < C(S^* + 1) \Rightarrow G(S^*) > 0 \quad (3.34)$$

En utilisant (3.32), les conditions d'optimalités (3.33) et (3.34) peuvent être réécrites comme suit :

$$F(S - 1) \leq \frac{b}{h + b} < F(S) \quad (3.35)$$

La politique d'approvisionnement optimale $\alpha^*(m^*) = (\alpha_i^*(m^*))_{i=1}^n$ implique $\alpha_i^* = 0$ pour $i = m^* + 1, \dots, n$. Donc, le nombre total des commandes dans le réseau ouvert des files d'attente est égal au nombre total

des commandes dans les files d'attente de premiers m^* fournisseurs. En utilisant (3.7) et (3.35), la valeur optimale du niveau de stock nominal pour le problème Π_4 est la valeur minimale de S qui satisfait l'inégalité $F(S) > b/(h+b)$:

$$S^* = \max \left\{ S : \sum_{v=0}^{S-1} P_v \leq \frac{b}{h+b} \right\} \quad (3.36)$$

où

$$\sum_{v=0}^S P_v = \sum_{i=1}^{m^*} \left(\left(\prod_{\substack{j=1 \\ j \neq i}}^{m^*} \frac{1 - \alpha_j \rho_j}{\alpha_i \rho_i - \alpha_j \rho_j} \right) (\alpha_i \rho_i)^{m^*-1} (1 - \alpha_i^{S+1} \rho_i^{S+1}) \right). \quad (3.37)$$

La valeur optimale du stock nominal se définit ainsi comme

$$S^* = \lfloor \hat{S} \rfloor$$

où $\lfloor \hat{S} \rfloor$ est le plus grand nombre entier inférieur ou égal à la valeur $\hat{S}$ qui satisfait l'égalité

$$\sum_{v=0}^{\hat{S}-1} P_v = \frac{b}{h+b}.$$

3.5. ANALYSES NUMÉRIQUES

La méthode approximative proposée se base sur deux simplifications. La première consiste à remplacer le problème d'optimisation Π_3 ayant comme variables de décision « $S, \alpha_1, \alpha_2, \dots, \alpha_n$ » en deux problèmes d'optimisation indépendants : le problème d'optimisation avec les variables de décision « $\alpha_1, \alpha_2, \dots, \alpha_n$ » et le problème d'optimisation avec la variable de décision « S ». La deuxième simplification est de résoudre le problème d'optimisation ayant les variables de décision « $\alpha_1, \alpha_2, \dots, \alpha_n$ » pour une valeur imposée de S , $S = 0$. Notons que la méthode approximative proposée ne peut pas être appliquée itérativement en actualisant la valeur de S dans le problème Π_1 . Par conséquent, la qualité de la solution obtenue ne peut pas être assurée.

Nous évaluons la méthode approximative proposée dans le cas de deux fournisseurs. Pour ce cas, la solution optimale du problème Π_3 peut être obtenue par une recherche numérique dans la région admissible. Nous considérons un problème avec $\lambda = 1$, $h = 1$ et $b = 1000$ et nous comparons les solutions obtenues pour les valeurs différents des taux moyen de fabrication (Tableau 3.1).

Tableau 3.1. Analyses numériques dans le cas de deux fournisseurs

	Taux de service	Nombre de fournisseurs	α_1	α_2	S	Coût total
1	$\mu_1 = 1.25$ $\mu_2 = 0.5$	Solution optimale avec 1 fournisseur	1,000	0,000	30,000	30,957
		Solution optimale avec 2 fournisseurs	0,738	0,262	15,000	14,492
		Solution approximative avec 2 fournisseurs	0,791	0,209	16,000	15,501
2	$\mu_1 = 1.25$ $\mu_2 = 0.6$	Solution optimale avec 1 fournisseur	1,000	0,000	30,000	30,957
		Solution optimale avec 2 fournisseurs	0,698	0,302	14,000	13,275
		Solution approximative avec 2 fournisseurs	0,748	0,252	14,000	13,969
3	$\mu_1 = 1.25$ $\mu_2 = 0.7$	Solution optimale avec 1 fournisseur	1,000	0,000	30,000	30,957
		Solution optimale avec 2 fournisseurs	0,661	0,339	13,000	12,290
		Solution approximative avec 2 fournisseurs	0,707	0,293	13,000	12,727
4	$\mu_1 = 1.25$ $\mu_2 = 0.8$	Solution optimale avec 1 fournisseur	1,000	0,000	30,000	30,957
		Solution optimale avec 2 fournisseurs	0,627	0,373	12,000	11,449
		Solution approximative avec 2 fournisseurs	0,667	0,333	12,000	11,726
5	$\mu_1 = 1.25$ $\mu_2 = 0.9$	Solution optimale avec 1 fournisseur	1,000	0,000	30,000	30,957
		Solution optimale avec 2 fournisseurs	0,596	0,404	11,000	10,738
		Solution approximative avec 2 fournisseurs	0,628	0,372	11,000	10,908
6	$\mu_1 = 1.25$ $\mu_2 = 1$	Solution optimale avec 1 fournisseur	1,000	0,000	30,000	30,957
		Solution optimale avec 2 fournisseurs	0,566	0,434	10,000	10,198
		Solution approximative avec 2 fournisseurs	0,590	0,410	10,000	10,292
7	$\mu_1 = 2$ $\mu_2 = 0.5$	Solution optimale avec 1 fournisseur	1,000	0,000	9,000	9,955
		Solution optimale avec 2 fournisseurs	0,851	0,149	9,000	8,618
		Solution approximative avec 2 fournisseurs	1,000	0,000	9,000	9,955
8	$\mu_1 = 2$ $\mu_2 = 0.6$	Solution optimale avec 1 fournisseur	1,000	0,000	9,000	9,955
		Solution optimale avec 2 fournisseurs	0,820	0,180	8,000	8,281
		Solution approximative avec 2 fournisseurs	0,966	0,034	9,000	9,438
9	$\mu_1 = 2$ $\mu_2 = 0.7$	Solution optimale avec 1 fournisseur	1,000	0,000	9,000	9,955
		Solution optimale avec 2 fournisseurs	0,790	0,210	8,000	7,992
		Solution approximative avec 2 fournisseurs	0,932	0,068	9,000	9,051
10	$\mu_1 = 2$ $\mu_2 = 0.8$	Solution optimale avec 1 fournisseur	1,000	0,000	9,000	9,955
		Solution optimale avec 2 fournisseurs	0,762	0,238	8,000	7,781
		Solution approximative avec 2 fournisseurs	0,897	0,103	8,000	8,673
11	$\mu_1 = 2$ $\mu_2 = 0.9$	Solution optimale avec 1 fournisseur	1,000	0,000	9,000	9,955
		Solution optimale avec 2 fournisseurs	0,735	0,265	8,000	7,628
		Solution approximative avec 2 fournisseurs	0,863	0,137	8,000	8,256
12	$\mu_1 = 2$ $\mu_2 = 1$	Solution optimale avec 1 fournisseur	1,000	0,000	9,000	9,955
		Solution optimale avec 2 fournisseurs	0,709	0,291	7,000	7,357
		Solution approximative avec 2 fournisseurs	0,828	0,172	8,000	7,953

Les analyses numériques montrent les avantages économiques de la stratégie bi-fournisseurs même dans le cas où un des fournisseurs est beaucoup moins performant. Acheminer les commandes chez les deux fournisseurs à la place d'acheminer toutes les commandes chez le fournisseur le plus performant du marché diminue le niveau de stock nominal et les coûts de stockage et de rupture. La méthode approximative donne des résultats satisfaisants pour les valeurs de la fonction objectif. La déviation moyenne entre les résultats de la méthode approximative et les solutions optimales est moins de 8 %. Par contre, la méthode approximative surestime la valeur du paramètre de Bernoulli associé au fournisseur le plus performant avec une déviation moyenne plus de 12 %.

3.6. CONCLUSIONS

Dans le but d'évaluer les diminutions possibles des coûts de stockage et de rupture, nous avons analysé l'application d'une stratégie multi-fournisseurs dans une chaîne logistique à deux niveaux ayant les demandes et les délais de livraison aléatoires. Nous avons étudié le problème d'approvisionnement d'un producteur ayant une demande externe et aléatoire d'un produit. Le producteur applique une politique de stock nominal afin de gérer son stock à partir duquel ses clients vont être servis. Nous avons supposé que le producteur peut envoyer chaque commande d'approvisionnement de stock à un fournisseur différent disponible dans le marché. Les fournisseurs sont homogènes en termes de prix et de qualité mais hétérogènes en termes de variabilité du délai de livraison. Nous avons modélisé chaque fournisseur comme un serveur exponentiel ayant un taux moyen de service différent. Les fournisseurs ont des capacités de fabrication limitées dans le sens où ils ne peuvent fabriquer qu'un produit à la fois.

Nous avons supposé que les commandes de réapprovisionnement sont allouées parmi les fournisseurs disponibles dans le marché selon un processus d'acheminement de Bernoulli qui définit une probabilité fixe d'affectation de commande pour chaque fournisseur. À chaque déclenchement de réapprovisionnement du stock, le producteur peut ensuite utiliser les probabilités d'affectation de commande afin de choisir le fournisseur auquel la commande doit être affectée.

Considérons un problème dans lequel le stock de sortie du système est distribué entre les fournisseurs, c'est-à-dire qu'il existe un stock de sortie localisé chez chaque fournisseur et le producteur qui produit à la commande achemine les demandes des clients finaux aux fournisseurs. Pour un tel problème, les positions de stock des fournisseurs peuvent être contrôlées indépendamment par le système d'information. C'est aussi le cas d'un problème d'assemblage dans lequel chaque fournisseur fabrique un produit différent. Il y a une interdépendance entre les stocks des produits intermédiaires et de produit fini, mais le producteur peut contrôler les positions de stock des produits intermédiaires et de produit fini indépendamment. Dans cette étude, nous avons supposé que le stock de sortie du système est mis en commun (*inventory pooling*) chez le producteur. Nous avons étudié un système de stockage dans un seul

endroit et d'un seul type de produit. Cette structure centralisée du système ne permet pas de décomposer le problème de gestion de stock étudié.

Dans la deuxième section, nous avons fourni les expressions analytiques des mesures de performances du producteur qui sont le niveau moyen de stock possédé et le niveau moyen de rupture de stock. Nous avons exprimé la somme des coûts moyens de stockage et de rupture en fonction des variables de décisions : le niveau de stock nominal et les probabilités d'affectation.

Les valeurs optimales des probabilités d'affectation de commande sont obtenues pour le cas de production à la commande, c'est-à-dire pour le cas où le coût unitaire de stockage est suffisamment élevé pour que le producteur installe un niveau de stock nominal nul. L'étude menée montre qu'acheminer les commandes chez plusieurs fournisseurs est plus profitable qu'acheminer toutes les commandes chez un seul fournisseur. La stratégie multi-fournisseurs diminue le délai moyen d'approvisionnement du producteur.

Dans le cas de production pour stock, le problème d'optimisation consiste à trouver les valeurs du niveau de stock nominal et des probabilités d'affectation qui minimisent les coûts moyens de stockage et de rupture. Puisque la fonction objectif n'est pas nécessairement convexe par rapport aux variables de décisions, nous avons proposé une méthode approximative pour résoudre ce problème. Nous avons validé cette technique par des analyses numériques. Ces analyses montrent que, en comparaison avec une stratégie mono-fournisseur, une stratégie multi-fournisseurs diminue le niveau de stock nominal nécessaire. En gardant moins de produit dans le stock, le producteur diminue la somme des coûts de stockage et de rupture de stock.

CHAPITRE IV

4. Optimisation des décisions dans une chaîne logistique décentralisée à deux étages de production/stockage

4.1. INTRODUCTION

La date de disponibilité des produits est un attribut de la concurrence qui joue à la fois sur la rapidité de mise sur le marché de produits nouveaux et sur celle de livraison de commandes de produits existants. Dans le cadre des réseaux d'entreprises, la date de disponibilité d'un produit vendu aux clients dépend de l'efficacité des différentes entreprises qui contribuent à sa création : les entreprises fabriquant des composants, des produits semi-finis ou des produits finis. Les délais de livraison de commandes en produits finis peuvent aussi être provoqués par des délais de livraison imprévus des composants ou des produits semi-finis. Afin d'assurer un approvisionnement satisfaisant en quantité et en délai, les entreprises entrent dans des relations contractuelles de longue durée avec leurs fournisseurs. Les caractéristiques d'une telle relation sont déterminantes pour la satisfaction des clients finaux. Dans ce chapitre, nous étudions les interactions entre les entreprises adjacentes des chaînes logistiques et les caractéristiques spécifiques des relations contractuelles permettant aux entreprises d'améliorer les performances du fonctionnement global, tout en limitant les risques encourus par chacun des partenaires.

Nous analysons un maillon élémentaire à deux niveaux d'une chaîne logistique, composé de deux entreprises adjacentes. Par commodité de terminologie, l'entreprise aval sera appelée « le producteur » et l'entreprise amont « le fournisseur ». Par leurs activités de transformation et/ou d'assemblage, ces deux entreprises contribuent à l'augmentation de la valeur des produits de la filière. Le fournisseur transforme une matière première en un produit intermédiaire qui est ensuite utilisé dans le processus de fabrication d'un produit fini chez le producteur. En outre, les deux entreprises disposent de stocks qui leur permettent de contrôler leurs productions et leurs livraisons. Pour le réapprovisionnement de son stock, chaque entreprise lance des ordres de fabrication internes et des commandes externes de produit intermédiaire ou de matière première à l'entreprise en amont.

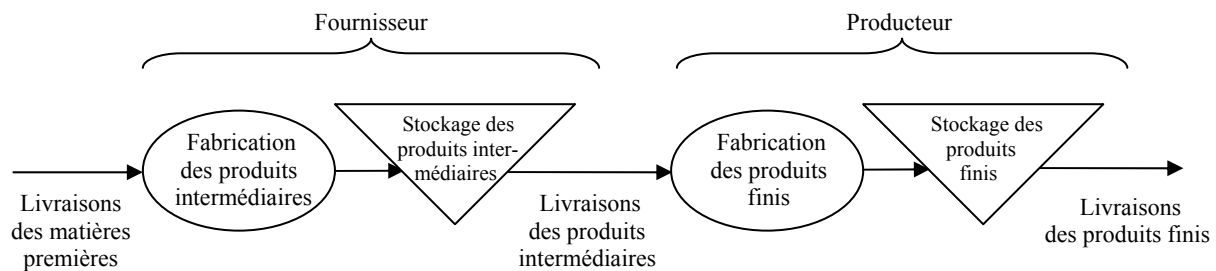


Figure 4.1. La chaîne logistique à deux étages de production/stockage

Dans cette chaîne logistique à deux étages de production/stockage, le producteur vend un produit fini au marché de consommateurs. La demande du marché final, les temps de fabrication des produits finis chez le producteur et des produits intermédiaires chez le fournisseur sont aléatoires. À cause des aléas en interne et externe, toutes les demandes du marché ne peuvent pas être satisfaites immédiatement. En outre, le délai de livraison des produits finis dépend non seulement du niveau d'utilisation des stocks et des performances de fabrication du producteur mais aussi du délai de livraison des produits intermédiaires. Autrement dit, le niveau de service du système résulte des performances de fabrication et des décisions d'approvisionnement des stocks de deux entreprises.

Nous supposons que chaque entreprise est une entité individuelle qui vise à optimiser sa politique d'approvisionnement par rapport à ses propres critères économiques. Dans cette chaîne logistique décentralisée, les décisions distribuées des partenaires peuvent être moins efficaces qu'un mécanisme maximisant les performances globales de la chaîne. Un tel mécanisme déterminerait les niveaux d'utilisation des ressources de stockage des entreprises dans le but d'obtenir un niveau de service satisfaisant. Dans le système décentralisé, l'existence et l'utilisation de ressources de stockage génèrent des coûts pour le fournisseur mais peuvent assurer au producteur un approvisionnement satisfaisant en quantité et en délai. Par contre, le fournisseur n'a pas un intérêt direct concernant le niveau de service de la chaîne. Le système décentralisé nécessite alors des mécanismes de coordination motivant le fournisseur à opter pour un niveau d'utilisation de ressources de stockage qui permet d'obtenir des niveaux satisfaisants pour le délai de livraison des produits intermédiaires et par conséquent pour le délai de livraison des produits finis. Dans ce chapitre, nous proposons un contrat de coordination qui permet d'amener les performances du système décentralisé vers les performances théoriques du modèle centralisé.

Cachon (1999a), Cachon et Zipkin (1999), Lee et Whang (1999) et Porteus (2000) étudient la problématique de gestion des stocks dans les chaînes logistiques décentralisées. Cachon et Zipkin (1999) analysent une chaîne logistique à deux niveaux, constituée d'une part d'un détaillant ayant une demande aléatoire stationnaire et d'autre part de son fournisseur. Les deux entreprises ont des délais d'approvisionnement certains. Clark et Scarf (2004) montrent qu'une politique de stock nominal est optimale pour ce système. Dans les jeux étudiés, la stratégie de chaque acteur détermine son niveau de

stock nominal. Cachon et Zipkin (1999) montrent que, même quand le fournisseur partage les coûts de rupture de stock du détaillant, l'équilibre de Nash ne correspond pas à la solution optimale du système centralisé. Ils proposent un contrat de coordination définissant un paiement de transfert linéaire. Pour un système similaire, Lee et Whang (1999) proposent un paiement de transfert non-linéaire qui correspond aux paiements utilisés par Clark et Scarf (2004). Porteus (2000) propose l'utilisation des jetons de responsabilité. Le fournisseur envoie chez le détaillant un jeton de responsabilité chaque fois qu'une commande de détaillant ne peut pas être satisfaite à cause d'une rupture de stock. Du point de vue du détaillant, le jeton envoyé correspond à un produit réel. Les mécanismes de coordination proposés par Lee et Whang (1999) et Porteus (2000) sont similaires car ils compensent les pertes du détaillant résultant des ruptures de stock du fournisseur. Le détaillant choisit alors son niveau de stock nominal sans prendre en compte les délais de livraisons du fournisseur. Les trois mécanismes cités coordonnent la chaîne logistique, c'est-à-dire que leur utilisation ramène l'équilibre du système décentralisé vers la solution optimale du système centralisé, qui peut être obtenue en utilisant l'algorithme de Clark et Scarf (2004).

Dans ce chapitre, nous généralisons les études citées en considérant que le système de fabrication de chaque étage est à capacité limitée et que les temps de fabrication des produits chez les entreprises sont des variables aléatoires. Nous proposons un contrat de coordination qui définit le prix d'achat des produits intermédiaires en fonction du délai de livraison observé du fournisseur. L'application de ce contrat impose une pénalité chez le fournisseur pour les livraisons retardées de produits intermédiaires. Nous montrons qu'en utilisant ce contrat, les pertes du producteur résultant des ruptures de stock du fournisseur ne sont pas totalement compensées et que le producteur détermine son niveau de stock nominal en prenant en compte les délais de livraison du fournisseur. Nous montrons ainsi que le contrat proposé coordonne la chaîne logistique analysée.

Par la suite, nous supposons que les commandes des clients finaux arrivent chez le producteur selon un processus de Poisson et que les temps de fabrication des produits dans les entreprises sont des variables aléatoires suivant des lois exponentielles. Nous pouvons mentionner plusieurs travaux qui analysent les systèmes avec demande et temps de fabrication aléatoires en s'appuyant sur la théorie des files d'attente. Caldentey et Wein (2003) étudient les interactions entre un producteur fabriquant les produits finis et son détaillant. Le système de fabrication du producteur fonctionne comme une file d'attente $M/M/1$. Le détaillant dispose d'un stock de produits finis. Les auteurs montrent qu'un équilibre de Nash existe quand le producteur contrôle son taux de production et le détaillant son niveau de stock nominal. Jemai (2003) et Jemai et Karaesmen (2007) analysent un système dans lequel les deux entreprises, le producteur et le détaillant, disposent de stocks locaux et contrôlent leurs niveaux de stocks nominaux. Le producteur possède un système de fabrication à capacité limitée. Le temps de transport entre les installations de stock des entreprises est supposé négligeable. Les auteurs montrent qu'un équilibre de Nash existe sous

différentes hypothèses. Cachon (1999b) propose différents mécanismes de coordination pour un système similaire en supposant qu'une demande est perdue si elle n'est pas satisfaite.

Gupta et Weerawat (2006) s'intéressent aux interactions entre un producteur de produit fini et son fournisseur de produit intermédiaire en supposant que le producteur ne dispose pas de stocks de sortie. Le système de fabrication de chaque entreprise est à capacité limitée. Les auteurs proposent des contrats de partage des revenus qui peuvent coordonner cette chaîne logistique. Gupta *et al.* (2004) généralisent cette étude à un système dans lequel le producteur dispose d'un stock de produits finis. Ils proposent des contrats de partage des revenus qui imposent une pénalité chez le fournisseur pour les demandes retardées du producteur. Ils démontrent ainsi que la valeur optimale du niveau de stock nominal du producteur peut être obtenue avec une recherche numérique dans un intervalle fermé.

Dans ce chapitre, nous proposons un contrat qui impose une pénalité chez le fournisseur seulement pour les ruptures de stock de produits intermédiaires. Le fournisseur n'est pas concerné par les ruptures de stock de produits finis qui peuvent aussi être provoquées par des temps de fabrication longs ou par un niveau de stock nominal insuffisant du producteur. Contrairement à Gupta *et al.* (2004), nous montrons que les niveaux de stocks nominaux des entreprises et les paramètres du contrat proposé peuvent être optimisés simultanément.

Après cette introduction, nous exposons le modèle de pilotage de flux étudié. La troisième section est consacrée aux calculs des mesures des performances du système analysé. Dans la quatrième section, nous définissons un jeu de Stackelberg dans lequel le producteur est le meneur et le fournisseur est le suiveur. Nous déterminons ensuite l'équilibre de Stackelberg du système décentralisé. Nous utilisons la solution optimale du système centralisé pour montrer l'efficacité du contrat proposé. Nous déterminons la solution optimale du système centralisé dans la quatrième section. Nous terminons ce chapitre par des conclusions.

4.2. MODÈLE DE PILOTAGE DE FLUX

Dans le système analysé, chaque entreprise dispose d'un stock de sortie à partir duquel ses demandes vont être servies. Quand la demande arrive, elle est satisfaite s'il y a des produits en stock, sinon elle est retardée. La gestion de stock dans les entreprises est accomplie suivant la politique de stock nominal. À l'état initial, le stock de l'entreprise de niveau $i = 1, 2$ contient le niveau de stock nominal « S_i » qui détermine le niveau maximal du stock. Le réapprovisionnement du stock est déclenché lorsque la position de stock devient inférieure au niveau S_i , c'est-à-dire chaque fois qu'une demande arrive. Nous supposons que la demande arrive chez le producteur à chaque fois en quantité unitaire et selon un processus de Poisson ayant le taux λ . Pour simplifier les notations, nous supposons que le producteur utilise un produit intermédiaire délivré par le fournisseur pour fabriquer un produit fini. De la même façon, le fournisseur utilise une unité de matière première pour fabriquer un produit intermédiaire. Les entreprises ne disposent

pas de stocks d'entrée. Selon cette dépendance, à chaque niveau $i = 1, 2$, l'arrivée d'une demande unitaire déclenche, en interne, un ordre de fabrication unitaire et en même temps une commande unitaire pour l'entreprise de niveau $i - 1$. Ainsi, l'arrivée d'une demande chez le producteur génère simultanément une demande unitaire pour chaque entreprise de niveau antérieur. Par hypothèse, l'entreprise de niveau « 0 » est un stock infini de matière première et les temps de transport des produits entre les étages sont négligeables. Sous ces hypothèses, les matières premières sont toujours disponibles. Par conséquent, le fournisseur reçoit une matière première chaque fois qu'une demande arrive.

En appliquant la politique de stock nominal $(S_i - 1, S_i)$, le système de fabrication fonctionne si le niveau de stock est inférieur au niveau S_i dans le but de le ramener à son niveau maximal. Les temps de fabrication du produit intermédiaire chez le fournisseur et du produit fini chez le producteur sont des variables aléatoires à distribution de probabilité exponentielle ayant respectivement les taux μ_1 et μ_2 ($\mu_1 \neq \mu_2$). Le taux d'utilisation $\rho_i = \lambda / \mu_i$ de chaque entreprise $i = 1, 2$ satisfait la condition de stabilité $\rho_i < 1$. Quand le processus de fabrication d'une unité est terminé, l'unité prend sa place dans le stock de sortie correspondant s'il n'y a pas de ruptures de stock, ou bien elle est utilisée pour satisfaire les commandes retardées en suivant l'ordre FIFO. Lorsqu'une unité est envoyée d'un étage à l'autre en réponse à une demande, l'unité rejoint une file d'attente à capacité illimitée devant le serveur suivant. Selon ces hypothèses, les produits intermédiaires dans le système sont partagés entre le stock de sortie du fournisseur et la file d'attente du producteur. Dans la Figure 4.2, la répartition physique des produits intermédiaires entre les entreprises est représentée par deux stations de synchronisation consécutives.

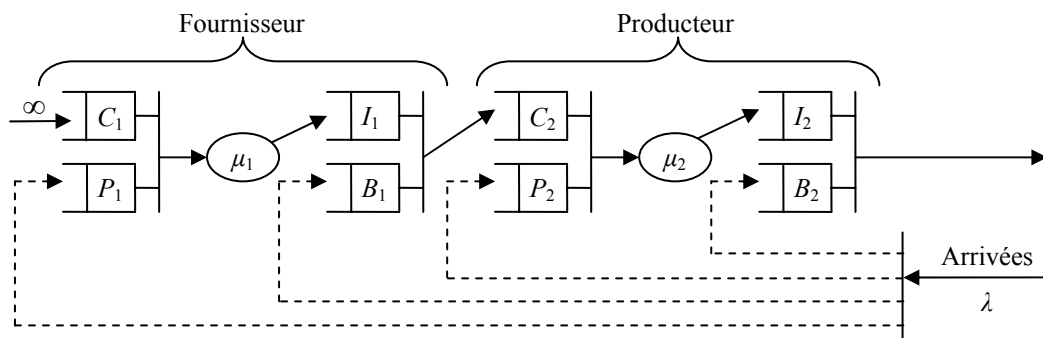


Figure 4.2. Représentation du système par les files d'attente

Pour chaque niveau $i = 1, 2$, l'évolution du système est décrite en utilisant des variables d'état en régime permanent :

- P_i : le nombre de commandes en attente de fabrication au niveau i
- K_i : le nombre de commandes en attente de fabrication et en fabrication au niveau i (P_i plus l'unité éventuellement en cours de fabrication)
- C_i : le nombre d'unités délivrées par l'entreprise de niveau $i - 1$ et en attente de fabrication au niveau i

Selon la loi de LITTLE (Kleinrock, 1975), des relations de dépendance similaires peuvent être établies pour les temps de séjour des commandes et les délais de livraison des entreprises. Soit L_i le temps de séjour des commandes et W_i le temps de séjour des produits dans le système de fabrication de niveau $i = 1, 2$. Soit D_i le délai de livraison de l'entreprise de niveau $i = 1, 2$. Le temps de séjour des commandes dans le système de fabrication du producteur « L_2 » est défini comme la somme du délai provoqué par le fournisseur « D_1 » et du temps de séjour des produits intermédiaires dans le système de fabrication du producteur « W_2 »: $L_2 = D_1 + W_2$. Sous l'hypothèse que l'entreprise de niveau « 0 » est un stock infini de matière première, $L_1 = W_1$ chez le fournisseur.

Notons qu'une politique de stock nominal n'est pas nécessairement optimale pour le système analysé en tenant compte du nombre d'en-cours, c'est-à-dire du nombre total de produits intermédiaires dans le système (Zipkin, 2000). Veatch et Wein (1994) analysent un système où chaque unité de produit fini dans le stock induit un coût de stockage $h > 1$, chaque demande retardée de produit fini induit un coût de rupture de stock b et chaque unité de produit intermédiaire dans le système induit un coût de stockage égal à 1 par unité de temps. Le critère à minimiser est le coût moyen actualisé sur un intervalle de temps infini. Les auteurs montrent qu'une politique de stock nominal n'est jamais optimale pour le système analysé. Ils utilisent la programmation dynamique afin de déterminer la meilleure politique de pilotage de flux et montrent que la politique optimale peut être une politique compliquée et difficile à appliquer. Un des avantages des politiques de stock nominal est leur facilité d'implémentation. Pour cette raison, les politiques de stock nominal sont souvent préférées pour le pilotage des systèmes à structure linéaire.

En utilisant la même approche que celle utilisée par Veatch et Wein (1994), Karesmen et Dallery (2000) comparent les performances des politiques de stock nominal, des politiques Kanban et des politiques Kanban généralisé. Les analyses numériques effectuées montrent que la meilleure politique de stock nominal est moins performante que la meilleure politique Kanban ou que la meilleure politique Kanban généralisé dans le cas où le producteur a un taux de fabrication plus faible. La meilleure politique Kanban généralisé est en générale plus performante que la meilleure politique de stock nominal ou que la meilleure politique Kanban. Zipkin (2000) compare aussi les performances de ces trois politiques de pilotage de flux par des analyses numériques et montre que les différences entre les coûts optimaux sont légères. En outre, les analyses numériques effectuées par Karesmen et Dallery (2000) et Zipkin (2000) montrent que les différences entre les niveaux de stocks nominaux de la meilleure politique de stock nominal et de la meilleure politique Kanban généralisé sont au plus 1. Pour des systèmes mono-étage de production/stockage, Liberopoulos et Dallery (2002) conjecturent que le niveau de stock nominal de la meilleure politique Kanban généralisé est égal au niveau de stock nominal de la meilleure politique de stock nominal. Selon ces résultats, même si la meilleure politique de stock nominal n'est pas nécessairement optimale pour le système analysé, elle est plutôt performante. La meilleure politique de

stock nominal peut aussi être utilisée afin de paramétrer une politique Kanban généralisé de façon approchée.

4.3. CALCULS DES MESURES DE PERFORMANCES

Dans cette étude, nous utilisons les niveaux moyens de stock possédé et les niveaux moyens de rupture de stock des entreprises comme des mesures de performances du système analysé. Soit $P_{k_i} = \Pr\{K_i = k_i\}$ la probabilité d'avoir k_i commandes dans le système de fabrication de niveau $i = 1, 2$. En utilisant cette définition, le niveau moyen de stock possédé $E[I_i]$ et le niveau moyen de rupture de stock $E[B_i]$ sont calculés par les équations :

$$E[I_i] = \sum_{k_i=0}^{S_i} (S_i - k_i) P_{k_i} \quad (4.1)$$

$$E[B_i] = \sum_{k_i=S_i}^{\infty} (k_i - S_i) P_{k_i} \quad (4.2)$$

Dans la suite, nous analysons le système de fabrication de chaque niveau $i = 1, 2$ afin de déterminer la probabilité P_{k_i} et les mesures de performances $E[I_i]$ et $E[B_i]$.

4.3.1. Le système de fabrication de produits intermédiaires

Sous l'hypothèse que les matières premières sont toujours disponibles, l'arrivée des matières premières chez le fournisseur suit exactement la demande finale. C'est donc un processus de Poisson ayant le taux λ . Ainsi, nous pouvons modéliser le système de fabrication chez le fournisseur comme une file d'attente M/M/1 avec taux d'utilisation $\rho_1 = \lambda / \mu_1$. Sous ces conditions, le nombre de commandes dans le système de fabrication « K_1 » est égal au nombre de clients dans la file M/M/1 (λ, μ_1).

La probabilité d'avoir k_1 commandes dans le système de fabrication est définie par :

$$P_{k_1} = \rho_1^{k_1} (1 - \rho_1) \quad (4.3)$$

En utilisant les équations (4.1), (4.2) et (4.3), le niveau moyen de stock possédé $E[I_1]$ et le niveau moyen de rupture de stock $E[B_1]$ sont obtenus comme suit :

$$E[I_1] = S_1 - \frac{\rho_1(1 - \rho_1^{S_1})}{1 - \rho_1} \quad (4.4)$$

$$E[B_1] = \frac{\rho_1^{S_1+1}}{1-\rho_1} \quad (4.5)$$

La probabilité d'avoir une rupture de stock chez le fournisseur est $\Pr\{I_1 = 0\} = \Pr\{K_1 \geq S_1\} = \rho_1^{S_1}$. Notons que le temps de séjour des produits dans le système de fabrication du fournisseur « W_1 » est égal au temps de séjour dans la file M/M/1 (λ, μ_1). Nous savons que le temps de séjour dans une file M/M/1 (λ, μ_i) est une variable aléatoire qui suit une loi exponentielle ayant le taux $\mu_i - \lambda$. Le délai de livraison du fournisseur « D_1 » est nul si $I_1 > 0$ et égal à W_1 si $I_1 = 0$. Dans la Figure 4.4, le délai de livraison du fournisseur est représenté par une phase (voir l'annexe C) ayant le taux exponentiel $\mu_1 - \lambda$. Le niveau moyen du délai de livraison du fournisseur s'écrit alors :

$$E[D_1] = \rho_1^{S_1} E[W_1] = \frac{\rho_1^{S_1}}{\mu_1 - \lambda}. \quad (4.6)$$

L'expression (4.6) satisfait la relation $E[D_1] = E[B_1] / \lambda$.

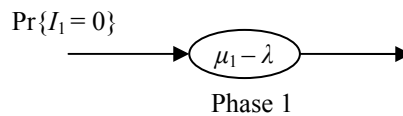


Figure 4.4. Représentation de D_1 par une phase

4.3.2. Le système de fabrication de produits finis

L'analyse exacte de ce système de fabrication est possible seulement si le fournisseur produit à la commande ou s'il possède un stock infini de produits intermédiaires. Dans le cas où le fournisseur produit à la commande avec $S_1 = 0$, le système à deux étages fonctionne comme un réseau de files d'attente en tandem où l'arrivée des produits intermédiaires dans la file d'attente de fabrication du producteur est un processus de Poisson ayant le taux λ . Par conséquent, le nombre d'unités dans le système de fabrication du producteur « N_2 » est égal au nombre de clients dans une file M/M/1 avec taux d'utilisation $\rho_2 = \lambda / \mu_2$. De la même façon, dans le cas limite où $S_1 \rightarrow \infty$, le système de fabrication de chaque niveau fonctionne comme une file d'attente M/M/1.

Dans le cas général où $0 < S_1 < \infty$, l'arrivée des produits intermédiaires chez le producteur n'est plus un processus de Poisson. Étant donnée l'arrivée d'un produit intermédiaire à l'instant t , le temps pour la prochaine arrivée est égal au temps d'inter-arrivées des demandes si $I_1(t) > 0$, au temps de fabrication d'un produit intermédiaire si $B_1(t) > 0$, et au maximum des deux si $I_1(t) - B_1(t) = 0$ (Buzacott *et al.*, 1991). Selon ces dépendances, les temps d'inter-arrivées des produits intermédiaires sont successivement corrélés. Autrement dit, il existe des dépendances probabilistes entre W_2 et W_1 ainsi que N_2 et N_1 (Lee et

Zipkin, 1992). Par conséquent, le système de fabrication du producteur ne peut pas être modélisé comme une file d'attente élémentaire. L'évaluation exacte de ce système de fabrication nécessite l'analyse de la chaîne de Markov en temps continu correspondante, en s'appuyant par exemple sur des outils de simulation.

Lee et Zipkin (1992) et Buzacott *et al.* (1991) proposent des méthodes approximatives (appelées respectivement LZ et BPS¹) pour l'évaluation analytique des performances des systèmes de stock nominal à n étages avec $S_i \geq 0$ pour $i = 1, \dots, n$. Les méthodes LZ et BPS¹ se basent sur l'hypothèse que l'arrivée des produits à chaque niveau $i > 1$, est un processus de Poisson ayant le taux λ . Sous cette hypothèse, le système de fabrication de chaque niveau i fonctionne comme une file d'attente M/M/1 (λ, μ_i) en isolation. En utilisant cette hypothèse, Buzacott *et al.* déterminent l'espérance du délai de livraison $E[D_i]$ de chaque niveau i . Lee et Zipkin déterminent les niveaux moyens des variables d'état en régime permanent $E[X_i]$, $X_i = K_i, I_i, B_i$, $i = 1, \dots, n$. Basées sur la même hypothèse, les deux méthodes donnent les mêmes résultats analytiques. Duri *et al.* (2000) montrent l'équivalence de ces deux méthodes et proposent des extensions de la méthode LZ pour des systèmes de stock nominal avec des temps de fabrication suivant des lois de type phase. Lee et Zipkin (1995) et Wang et Su (2007) généralisent la méthode LZ pour des systèmes de stock nominal plus complexes.

Buzacott *et al.* proposent aussi une méthode alternative (appelée BPS²) pour un système à deux étages. Ils déterminent la distribution de probabilité du temps d'inter-arrivées des produits à deuxième niveau et traitent le système de fabrication de deuxième niveau comme une file d'attente G/M/1. Notons que la méthode BPS² est encore une méthode approximative car elle méconnaît le fait que les temps d'inter-arrivées des produits à deuxième étage sont successivement corrélés. Gupta et Selvaraju (2006) proposent une méthode approximative (appelée GS) qui se base sur des améliorations de la méthode BPS² ainsi que des extensions pour un système à n étages.

Les analyses numériques de Gupta et Selvaraju montrent que, dans un système à deux étages, les erreurs moyennes absolues obtenues par les méthodes LZ (ou BPS¹), BPS² et GS sont respectivement 1.97 %, 1.4 % et 0.71 % pour les valeurs de $E[K_2]$ et 4.93 %, 4.65 % et 2.79 % pour les valeurs de $E[B_2]$. Les méthodes BPS² et GS peuvent sous-estimer les valeurs réelles de $E[K_2]$ et $E[B_2]$. Néanmoins, la méthode LZ (ou BPS¹) fournit toujours des bornes supérieures pour $E[K_2]$ et $E[B_2]$.

Liu *et al.* (2004) analysent un système dans lequel les temps d'inter-arrivées des demandes ainsi que les temps de fabrication des produits suivent des lois générales. Ils proposent une méthode approximative pour caractériser le processus d'arrivée des produits à chaque niveau i . Afin de déterminer $E[K_i]$, ils supposent que N_i est indépendant de B_{i-1} . Autrement dit, ils traitent le système de fabrication de niveau i comme une file d'attente G/G/1. Pour un système où les temps d'inter-arrivées des demandes et les temps

de fabrication des produits suivent des lois exponentielles, leur méthode détermine que l'arrivée des produits à chaque niveau est un processus de Poisson et donne donc les mêmes résultats analytiques que la méthode LZ.

Dans cette étude, nous adoptons la méthode approximative LZ afin de développer des résultats analytiques. Lee et Zipkin utilisent la méthode analytique proposée par Sovorons et Zipkin (1991) pour les systèmes de stock nominal avec temps de séjour des produits dans les étages indépendants et exogènes. La méthode de Sovorons et Zipkin (1991) se base sur les propriétés des lois de type phase (Neuts, 1994) et utilise des représentations du type matrice-exponentielle. Dans la suite, la méthode LZ est utilisée pour calculer approximativement P_{k_2} en utilisant des fonctions analytiques des distributions de densité de probabilité à la place des représentations matrice-exponentielle. L'application de la méthode LZ pour un système à n étages est détaillée dans l'annexe C.

4.3.2.1. La méthode approximative LZ

Sous l'hypothèse que l'arrivée des produits intermédiaires chez le producteur est un processus de Poisson ayant le taux λ , le temps de séjour des produits dans le système de fabrication du producteur « W_2 » est indépendant du temps de séjour des produits dans le système de fabrication du fournisseur « W_1 » et est égal au temps de séjour dans une file M/M/1 (λ, μ_2). Rappelons que le temps de séjour des commandes dans le système de fabrication du producteur est défini par $L_2 = D_1 + W_2$. L'hypothèse d'indépendance des temps de séjour des produits dans les systèmes de fabrication consécutifs nous permet d'écrire la fonction de densité de probabilité $f_{L_2}(t)$ comme un produit de convolution,

$$f_{L_2}(t) = f_{D_1+W_2}(t) = (f_{D_1} * f_{W_2})(t) = \int_0^t f_{D_1}(t-s)f_{W_2}(s) ds$$

où « $*$ » est l'opérateur de convolution. Par la suite, le temps de séjour des commandes dans le système de fabrication suit une loi de type phase dont la représentation graphique est donnée dans la Figure 4.5. Selon cette représentation, l'attente provoquée par une rupture de stock du produit intermédiaire est une opération supplémentaire chez le producteur. Cette opération est exécutée qu'avec une probabilité égale à la probabilité de rupture de stock chez le fournisseur.

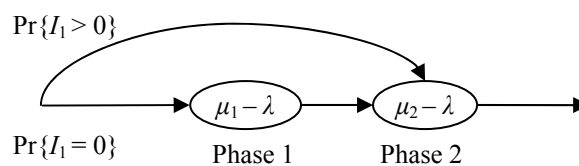


Figure 4.5. Représentation de L_2 par une loi de type phase

Nous pouvons réécrire la fonction de densité de probabilité $f_{L_2}(t)$ comme

$$f_{L_2}(t) = \rho_1^{S_1} f_{W_1+W_2}(t) + (1 - \rho_1^{S_1}) f_{W_2}(t)$$

$$\text{où } f_{W_1+W_2}(t) = (f_{W_1} * f_{W_2})(t) = \int_0^t f_{W_1}(t-s) f_{W_2}(s) ds.$$

En supposant $\rho_1 \neq \rho_2$, la fonction de densité de probabilité hypo-exponentielle $f_{W_1+W_2}(t)$ est

$$f_{W_1+W_2}(t) = \frac{\rho_1(1-\rho_2)}{\rho_1-\rho_2}(\mu_1-\lambda)e^{-(\mu_1-\lambda)t} - \frac{\rho_2(1-\rho_1)}{\rho_1-\rho_2}(\mu_2-\lambda)e^{-(\mu_2-\lambda)t}.$$

En développant, la fonction $f_{L_2}(t)$ est obtenue comme suit :

$$f_{L_2}(t) = \left(\frac{\rho_1^{S_1+1}(1-\rho_2)}{\rho_1-\rho_2} \right) (\mu_1-\lambda)e^{-(\mu_1-\lambda)t} + \left(1 - \frac{\rho_1^{S_1+1}(1-\rho_2)}{\rho_1-\rho_2} \right) (\mu_2-\lambda)e^{-(\mu_2-\lambda)t}$$

Le nombre de commandes dans le système de fabrication « K_2 » étant égal au nombre de demandes arrivant pendant le temps de séjour « L_2 », nous pouvons obtenir la probabilité d'avoir k_2 demandes dans le système de fabrication :

$$P_{k_2} = \left(\frac{\rho_1^{S_1+1}(1-\rho_2)}{\rho_1-\rho_2} \right) \rho_1^{k_2} (1-\rho_1) + \left(1 - \frac{\rho_1^{S_1+1}(1-\rho_2)}{\rho_1-\rho_2} \right) \rho_2^{k_2} (1-\rho_2)$$

A partir de cette expression, nous calculons le niveau moyen de stock et le niveau moyen de rupture de stock :

$$E[I_2] = S_2 - \frac{\rho_1^{S_1+1}}{1-\rho_1} - \frac{\rho_2}{1-\rho_2} + \left(\frac{\rho_1^{S_1+1}(1-\rho_2)}{\rho_1-\rho_2} \right) \frac{\rho_1^{S_2+1}}{1-\rho_1} + \left(1 - \frac{\rho_1^{S_1+1}(1-\rho_2)}{\rho_1-\rho_2} \right) \frac{\rho_2^{S_2+1}}{1-\rho_2} \quad (4.7)$$

$$E[B_2] = \left(\frac{\rho_1^{S_1+1}(1-\rho_2)}{\rho_1-\rho_2} \right) \frac{\rho_1^{S_2+1}}{1-\rho_1} + \left(1 - \frac{\rho_1^{S_1+1}(1-\rho_2)}{\rho_1-\rho_2} \right) \frac{\rho_2^{S_2+1}}{1-\rho_2} \quad (4.8)$$

La probabilité d'avoir une rupture de stock chez le producteur est

$$\Pr\{I_2 = 0\} = \Pr\{K_2 \geq S_2\} = \left(\frac{\rho_1^{S_1+1}(1-\rho_2)}{\rho_1-\rho_2} \right) \rho_1^{S_2} + \left(1 - \frac{\rho_1^{S_1+1}(1-\rho_2)}{\rho_1-\rho_2} \right) \rho_2^{S_2}.$$

Lee et Zipkin déterminent la fonction de densité de probabilité du délai de livraison du producteur en utilisant la propriété donnée par Sovorons et Zipkin (1991) : le nombre de commandes retardées « B_2 » est égal au nombre de demandes arrivant pendant le délai de livraison « D_2 ». En combinant cette propriété avec les propriétés des lois de type phase, le délai de livraison du producteur suit une loi de type phase (Figure 4.6) avec la fonction de densité de probabilité

$$f_{D_2}(t) = \Pr\{K_1 \geq S_1 + S_2\}f_{W_1+W_2}(t) + (\Pr\{K_2 \geq S_2\} - \Pr\{K_1 \geq S_1 + S_2\})f_{W_2}(t)$$

où $\Pr\{K_1 \geq S_1 + S_2\} = \rho_1^{S_1+S_2}$.

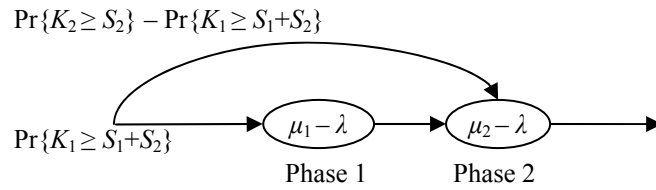


Figure 4.6. Représentation de D_2 par une loi de type phase

Le niveau moyen du délai de livraison du fournisseur est obtenu comme suit :

$$E[D_2] = \left(\frac{\rho_1^{S_1+1}(1-\rho_2)}{\rho_1-\rho_2} \right) \frac{\rho_1^{S_2}}{\mu_1-\lambda} + \left(1 - \frac{\rho_1^{S_1+1}(1-\rho_2)}{\rho_1-\rho_2} \right) \frac{\rho_2^{S_2}}{\mu_2-\lambda} \quad (4.9)$$

L'expression (4.9) satisfait la relation $E[D_2] = E[B_2] / \lambda$. Notons que la méthode approximative LZ est exacte dans le cas où $S_1 = 0$.

4.4. SYSTÈME DÉCENTRALISÉ ET ÉQUILIBRE DE STACKELBERG

Dans le système décentralisé étudié, chaque entreprise décide de son niveau de stock nominal $S_i \geq 0$ dans le but de maximiser son profit moyen. À chaque niveau $i = 1, 2$, chaque unité produite induit un coût de production c_i . Le prix de vente d'un produit fini est p_2 ($p_2 \geq c_1 + c_2$). Nous supposons que le prix d'achat d'une matière première est inclus dans le coût unitaire de production c_1 . Puisque chaque entreprise dispose de stocks pour servir ses demandes, les coûts de stockage sont encourus à chaque niveau $i = 1, 2$: chaque unité dans le stock de niveau i induit un coût de stockage h_i par unité de temps. Cependant, les coûts de rupture n'apparaissent que chez le producteur, pour les livraisons retardées de produits finis : chaque demande de produit fini retardée induit un coût de rupture de stock b_2 par unité de temps.

D'après une hypothèse générale en théorie de gestion des stocks, le coût de stockage d'un produit augmente avec la valeur ajoutée. Nous supposons donc $h_2 > h_1$. Nous pourrions aussi supposer que le fournisseur est chargé d'un coût unitaire de stockage h_1 pour chaque produit intermédiaire présent dans le

système de fabrication du producteur. Par contre, la considération de ces coûts additionnels ne changerait pas les calculs car, en utilisant la méthode LZ, l'espérance « $h_1 E[N_2]$ » est une constante. Nous excluons ce type des coûts pour simplifier la présentation.

Dans le cas d'un système centralisé, le problème consiste à trouver les niveaux de stocks nominaux qui maximisent l'espérance de profit de la chaîne logistique globale. L'espérance de profit du système centralisé s'écrit

$$\pi_0(S_1, S_2) = \lambda(p_2 - c_1 - c_2) - h_1 E[I_1] - h_2 E[I_2] - b_2 E[B_2]. \quad (4.10)$$

Donc, le niveau de stock nominal du premier étage, S_1 , du système centralisé est déterminé en prenant en compte les coûts de rupture de produits finis. Par contre, dans le système décentralisé, le fournisseur n'a pas intérêt à installer un niveau de stock nominal positif. Une approche possible pour motiver le fournisseur à installer un niveau de stock nominal positif est de partager les coûts de rupture de stock de produits finis entre le fournisseur et le producteur, avec une fraction $x \in [0,1]$ qui est supposée exogène. En utilisant cette approche, chaque demande de produit fini retardée induit un coût « $x b_2$ » chez le fournisseur et « $(1-x) b_2$ » chez le producteur (Caldentey et Wein, 2003; Cachon et Zipkin, 1999). Ces coûts de rupture peuvent être interprétés comme les pénalités de bonne volonté qui n'ont pas une influence égale sur les entreprises. Dans cette étude, l'objectif est de définir un mécanisme de coordination qui se base plutôt sur les coûts locaux de rupture de stock représentant les pénalités de livraison retardée.

Nous définissons un jeu de Stackelberg où le producteur (le meneur) propose un contrat fixant son prix d'achat. Nous appelons ce contrat (p_1, b_1) . Le contrat (p_1, b_1) définit le prix d'achat d'un produit intermédiaire comme « $P_1(D_1) = p_1 - b_1 D_1$ » où D_1 est le délai de livraison observé du fournisseur. Le producteur propose donc un prix ajusté en fonction du délai de livraison observé du fournisseur en imposant une pénalité pour les livraisons retardées de produits intermédiaires. Les paramètres de ce contrat, notamment $p_1 \geq 0$ et $b_1 \geq 0$, sont des variables de décision additionnelles du producteur. Le fournisseur reçoit le contrat proposé et décide de son niveau de stock nominal S_1 . Notons enfin, selon une hypothèse classique de cas limite pour un marché parfaitement compétitif, que le fournisseur n'accepte ce contrat que s'il lui permet d'obtenir un profit supérieur ou égal à zéro. Cette hypothèse n'est pas limitative dans la mesure où les résultats peuvent être facilement transposés au cas où le seuil d'acceptation du contrat n'est plus nul et correspond à une valeur de profit minimale.

Le contrat (p_1, b_1) définit un paiement de transfert entre les acteurs. Le producteur paie le prix $p_1 - b_1 D_1$ pour chaque produit intermédiaire qu'il achète à son fournisseur. L'espérance mathématique du paiement de transfert correspondant s'écrit

$$E[T(S_1, p_1, b_1)] = E[\lambda (p_1 - b_1 D_1)]. \quad (4.11)$$

En utilisant la relation $E[D_1] = E[B_1] / \lambda$, nous pouvons réécrire l'espérance de paiement de transfert comme suit :

$$E[T(S_1, p_1, b_1)] = E[\lambda p_1 - b_1 B_1] \quad (4.12)$$

Par conséquent, il existe deux techniques possibles pour appliquer le contrat (p_1, b_1) : soit le producteur propose un prix ajusté en fonction du délai de livraison observé du fournisseur, soit il propose un prix d'achat constant et il impose une pénalité pour les ruptures de stock du fournisseur. Dans le deuxième cas, p_1 représente le prix d'achat d'un produit intermédiaire et b_1 représente la pénalité de rupture de stock par unité de temps pour une commande de produit intermédiaire retardée. Le producteur paie alors le prix p_1 pour chaque produit intermédiaire qu'il achète à son fournisseur et le fournisseur paie une pénalité de rupture de stock au producteur. Par la suite, afin de simplifier la présentation, nous supposons que le producteur applique le contrat (p_1, b_1) en proposant un prix d'achat constant p_1 et une pénalité de rupture de stock b_1 .

En utilisant le contrat (p_1, b_1) , les espérances de profit du fournisseur et du producteur s'écrivent :

$$\pi_1(S_1, p_1, b_1) = E[T(S_1, p_1, b_1)] - \lambda c_1 - h_1 E[I_1] \quad (4.13)$$

$$\pi_2(S_1, S_2, p_1, b_1) = \lambda(p_2 - c_2) - h_2 E[I_2] - b_2 E[B_2] - E[T(S_1, p_1, b_1)] \quad (4.14)$$

En utilisant (4.12), nous obtenons

$$\pi_1(S_1, p_1, b_1) = \lambda(p_1 - c_1) - h_1 E[I_1] - b_1 E[B_1], \quad (4.15)$$

$$\pi_2(S_1, S_2, p_1, b_1) = \lambda(p_2 - p_1 - c_2) + b_1 E[B_1] - h_2 E[I_2] - b_2 E[B_2]. \quad (4.16)$$

Dans la suite, nous supposons que S_i est une variable de décision continue non-négative pour chaque entreprise de niveau $i = 1, 2$. Cette hypothèse méconnaît la restriction de S_i à l'ensemble de nombres entiers non-négatifs, qui est plus adaptée aux applications industrielles. En contrepartie, cette hypothèse nous permet de développer des résultats analytiques qui peuvent être utilisés pour approximer les valeurs optimales entières.

4.4.1. Problème d'optimisation du fournisseur

Le problème d'optimisation du fournisseur consiste à trouver la valeur optimale de S_1 qui maximise sa fonction d'utilité. En utilisant les expressions (4.4), (4.5) et (4.15), la fonction d'utilité du fournisseur $\pi_1(S_1) = \pi_1(S_1, p_1, b_1)$ s'écrit de la manière suivante :

$$\pi_1(S_1) = \lambda(p_1 - c_1) - h_1 \left(S_1 - \frac{\rho_1}{1 - \rho_1} \right) - (h_1 + b_1) \frac{\rho_1^{S_1+1}}{1 - \rho_1} \quad (4.17)$$

La dérivée première de la fonction $\pi_1(S_1)$ s'écrit

$$\pi_1'(S_1) = -h_1 - (h_1 + b_1) \frac{\rho_1^{S_1+1}}{1 - \rho_1} \ln \rho_1.$$

La dérivée deuxième est obtenue comme

$$\pi_1''(S_1) = -(h_1 + b_1) \frac{\rho_1^{S_1+1}}{1 - \rho_1} (\ln \rho_1)^2.$$

Soit $b_1^{\min}$ la valeur de b_1 qui satisfait $\pi_1'(0) = 0$:

$$b_1^{\min} = -h_1 \left(1 + \frac{1 - \rho_1}{\rho_1 \ln \rho_1} \right) \quad (4.18)$$

Il est clair que $\pi_1''(S_1) < 0$. Par conséquent, la fonction $\pi_1(S_1)$ est concave pour $S_1 \geq 0$. En outre, $\pi_1'(S_1) < 0$ si $b_1 < b_1^{\min}$. Donc, la fonction $\pi_1(S_1)$ est décroissante pour $S_1 \geq 0$ dans le cas où $b_1 < b_1^{\min}$. En utilisant ces observations, nous définissons la meilleure stratégie du fournisseur en réponse à un contrat proposé par le producteur avec la propriété suivante.

Propriété 4.1 : *Étant donné le contrat (p_1, b_1) , la valeur optimale du niveau de stock nominal qui maximise la fonction d'utilité du fournisseur $\pi_1(S_1)$ est*

$$S_1^*(b_1) = \begin{cases} \frac{\ln(\alpha(b_1))}{\ln \rho_1} & \text{si } b_1 > b_1^{\min} \\ 0 & \text{si } b_1 \leq b_1^{\min} \end{cases} \quad (4.19)$$

où

$$\alpha(b_1) = \frac{-h_1(1 - \rho_1)}{(h_1 + b_1)\rho_1 \ln \rho_1}. \quad (4.20)$$

Le seul paramètre du contrat qui a une influence sur la meilleure stratégie du fournisseur est la pénalité de rupture de stock b_1 . La valeur $b_1^{\min}$ est la valeur minimale de la pénalité b_1 pour laquelle le fournisseur a intérêt à installer un niveau de stock nominal positif.

Pour montrer cette propriété, soit $\varphi_1(x) = (1-x)/x \ln x$. La fonction $\varphi_1(x)$ est croissante pour $x \in (0,1)$, car $\ln x \leq x-1 \quad \forall x > 0$:

$$\varphi_1'(x) = \frac{x - \ln x - 1}{(x \ln x)^2} > 0$$

Notons aussi que $\lim_{x \rightarrow 1^-} \varphi_1(x) = -1$. Par conséquent, $b_1^{\min} > 0$ et la relation $b_1 > b_1^{\min}$ est opérante.

Nous observons que la meilleure stratégie du fournisseur $S_1^*(b_1)$ est croissante en b_1 pour $b_1 > b_1^{\min}$. Donc, le producteur peut augmenter le niveau de stock nominal du fournisseur en imposant une pénalité $b_1 > b_1^{\min}$. En outre, la condition de non-négativité $S_1^*(b_1) \geq 0$ est naturellement satisfaite car $\alpha(b_1^{\min}) = 1$.

4.4.2. Problème d'optimisation du producteur

En utilisant les expressions (4.7), (4.8) et (4.16), nous calculons la fonction d'utilité du producteur comme suit :

$$\begin{aligned} \pi_2(S_1, S_2, p_1, b_1) = & \lambda(p_2 - p_1 - c_2) - h_2 \left(S_2 - \frac{\rho_2}{1 - \rho_2} - \frac{\rho_1^{S_1+1}}{1 - \rho_1} \right) + b_1 \frac{\rho_1^{S_1+1}}{1 - \rho_1} \\ & - (h_2 + b_2) \left(\left(\frac{\rho_1^{S_1+1}(1 - \rho_2)}{\rho_1 - \rho_2} \right) \frac{\rho_1^{S_2+1}}{1 - \rho_1} + \left(1 - \frac{\rho_1^{S_1+1}(1 - \rho_2)}{\rho_1 - \rho_2} \right) \frac{\rho_2^{S_2+1}}{1 - \rho_2} \right) \end{aligned} \quad (4.21)$$

Le problème d'optimisation du producteur se formule de la manière suivante :

$$\Pi_1: \quad \max_{p_1, b_1, S_2} \pi_2(S_1^*, S_2, p_1, b_1) \quad (4.22)$$

sous les contraintes

$$S_1^* = \arg \max_{S_1} \pi_1(S_1, p_1, b_1) \quad (4.23)$$

$$\pi_1(S_1^*, p_1, b_1) \geq 0. \quad (4.24)$$

Le problème du producteur consiste à trouver les valeurs optimales des paramètres du contrat (p_1, b_1) et du niveau de stock nominal S_2 qui maximisent son profit moyen et qui satisfont les contraintes (4.23) et (4.24). La contrainte de compatibilité d'incitation (4.23) définit la meilleure stratégie du fournisseur comme la stratégie maximisant sa fonction d'utilité. Rappelons que le fournisseur accepte un contrat si ce

contrat lui permet d'obtenir un profit non-négatif. La contrainte de rationalité individuelle (4.24) assure l'acceptation du contrat par le fournisseur en définissant une borne inférieure pour le profit maximal du fournisseur.

Le producteur est le premier joueur dans ce jeu de Stackelberg. Donc, le producteur peut prédire la meilleure stratégie du fournisseur et il peut utiliser cette connaissance dans son problème d'optimisation. En connaissant la meilleure stratégie du fournisseur (4.19), le problème d'optimisation Π_1 devient :

$$\Pi_2: \max_{p_1, b_1, S_2} \pi_2(S_1^*(b_1), S_2, p_1, b_1) \quad (4.25)$$

sous la contrainte

$$\pi_1(S_1^*(b_1), p_1, b_1) \geq 0. \quad (4.26)$$

4.4.2.1. Les valeurs optimales des paramètres du contrat

Le Lagrangien du problème Π_2 s'écrit

$$L(S_2, p_1, b_1, u) = \pi_2(S_1^*(b_1), S_2, p_1, b_1) + u \pi_1(S_1^*(b_1), p_1, b_1)$$

où u est le paramètre de Lagrange correspondant à la contrainte (4.26). Nous calculons la dérivée partielle du Lagrangien par rapport à p_1 : $\partial L / \partial p_1 = -\lambda(1-u)$, et nous obtenons la valeur optimale du paramètre de Lagrange $u^* = 1$ qui satisfait la condition d'optimalité $\partial L / \partial p_1 = 0$. Donc, d'après la condition nécessaire d'optimalité de Karush-Kuhn-Tucker, le profit du fournisseur est égal à zéro à l'optimum, ce qui indique que, la valeur de p_1 qui maximise $\pi_2(S_2, p_1, b_1) = \pi_2(S_1^*(b_1), S_2, p_1, b_1)$ est $p_1^*(b_1)$, qui satisfait $\pi_1(S_1^*(b_1), p_1^*(b_1), b_1) = 0$. La propriété suivante définit la valeur optimale du prix d'achat proposé par le producteur.

Propriété 4.2 : *Étant donnés b_1 et S_2 , la valeur optimale du prix d'achat p_1 qui maximise $\pi_2(S_2, p_1, b_1)$ s'écrit comme suit :*

$$p_1^*(b_1) = c_1 + \frac{h_1}{\lambda} \left(S_1^*(b_1) - \frac{\rho_1}{1 - \rho_1} \right) + \frac{h_1 + b_1}{\lambda} \left(\frac{\rho_1^{S_1^*(b_1)+1}}{1 - \rho_1} \right) \quad (4.27)$$

Analysons la valeur optimale de p_1 en utilisant les équations (4.13) et (4.14). D'après l'expression (4.13), la condition $\pi_1(S_1, p_1, b_1) = 0$ implique

$$E[T(S_1, p_1, b_1)] = \lambda p_1 - b_1 E[B_1] = \lambda c_1 + h_1 E[I_1]. \quad (4.28)$$

Autrement dit, le producteur propose un prix p_1 qui compense la totalité des coûts opérationnels et des coûts de rupture du fournisseur quelle que soit la pénalité de rupture b_1 . Cette proposition laisse une marge de profit nulle pour le fournisseur. Le paiement de transfert correspondant à cette proposition est la somme des coûts opérationnels du fournisseur. Le producteur encourt alors les coûts opérationnels du fournisseur. En utilisant (4.28), nous pouvons réécrire l'expression (4.16) comme suit :

$$\pi_2(S_1, S_2, p_1, b_1) = \lambda(p_2 - c_1 - c_2) - h_1 E[I_1] - h_2 E[I_2] - b_2 E[B_2] \quad (4.29)$$

Par comparaison avec l'expression (4.10), l'expression (4.29) montre que, en appliquant le contrat (p_1, b_1) , le producteur obtient le profit total de la chaîne logistique à deux-étages de production/stockage.

Par la suite, nous cherchons la valeur optimale de la pénalité de rupture de stock b_1 qui sert à contrôler le niveau de stock nominal du fournisseur S_1 . Afin de faciliter les formulations, nous utilisons $\pi_2(S_2, b_1)$ pour indiquer $\pi_2(S_1^*(b_1), S_2, p_1^*(b_1), b_1)$:

$$\begin{aligned} \pi_2(S_2, b_1) = & \lambda(p_2 - c_1 - c_2) - h_1 \left(S_1^*(b_1) - \frac{\rho_1}{1 - \rho_1} + \frac{\rho_1^{S_1^*(b_1)+1}}{1 - \rho_1} \right) - h_2 \left(S_2 - \frac{\rho_2}{1 - \rho_2} - \frac{\rho_1^{S_1^*(b_1)+1}}{1 - \rho_1} \right) \\ & - (h_2 + b_2) \left(\left(\frac{\rho_1^{S_1^*(b_1)+1}(1 - \rho_2)}{\rho_1 - \rho_2} \right) \frac{\rho_1^{S_2+1}}{1 - \rho_1} + \left(1 - \frac{\rho_1^{S_1^*(b_1)+1}(1 - \rho_2)}{\rho_1 - \rho_2} \right) \frac{\rho_2^{S_2+1}}{1 - \rho_2} \right) \end{aligned}$$

La dérivée partielle de la fonction $\pi_2(S_2, b_1)$ par rapport à b_1 pour $b_1 > b_1^{\min}$ s'écrit

$$\frac{\partial \pi_2(S_2, b_1)}{\partial b_1} = \frac{h_1(b_1 + h_2 - (h_2 + b_2)\tau_0(S_2))}{(h_1 + b_1)^2 \ln \rho_1} \quad (4.30)$$

où

$$\tau_0(S_2) = \frac{(1 - \rho_1)(1 - \rho_2)}{\rho_1 - \rho_2} \left(\frac{\rho_1^{S_2+1}}{1 - \rho_1} - \frac{\rho_2^{S_2+1}}{1 - \rho_2} \right). \quad (4.31)$$

Lemme 4.1 : $\tau_0(S_2)$ est positive et décroissante pour $S_2 \geq 0$ avec $\tau_0(S_2) : [0, \infty) \rightarrow (0, 1]$.

Démonstration : Soit $\varphi_2(x) = x^{S_2+1} \ln x / (1 - x)$. La fonction $\varphi_2(x)$ est décroissante pour $x \in (0, 1)$ si $S_2 \geq 0$:

$$\varphi_2'(x) = \frac{x^{S_2}(S_2(1 - x) + 1) \ln x}{(1 - x)^2} + \frac{x^{S_2}}{1 - x} = \frac{x^{S_2}((S_2(1 - x) + 1) \ln x - x + 1)}{(1 - x)^2} < 0$$

Par conséquent,

$$\tau'_0(S_2) = \frac{(1-\rho_1)(1-\rho_2)}{\rho_1-\rho_2} \left(\frac{\rho_1^{S_2+1} \ln \rho_1}{1-\rho_1} - \frac{\rho_2^{S_2+1} \ln \rho_2}{1-\rho_2} \right) < 0$$

En outre, $\tau_0(0) = 1$ et $\lim_{S_2 \rightarrow \infty} \tau_0(S_2) = 0$. $\square$

La propriété suivante caractérise la valeur optimale de la pénalité b_1 .

Propriété 4.3 : *Étant donné S_2 , la fonction d'utilité $\pi_2(S_2, b_1)$ est strictement quasi-concave pour $b_1 > b_1^{\min}$ et constante pour $b_1 \leq b_1^{\min}$. La valeur optimale de la pénalité de rupture b_1 qui maximise $\pi_2(S_2, b_1)$ est*

$$b_1^*(S_2) = \begin{cases} b_1^{\text{opt}}(S_2) & \text{si } \tau_0(S_2) > (h_2 + b_1^{\min})/(h_2 + b_2) \\ b_1^{\min} & \text{si } \tau_0(S_2) \leq (h_2 + b_1^{\min})/(h_2 + b_2) \end{cases} \quad (4.32)$$

où $b_1^{\text{opt}}(S_2) = (h_2 + b_2)\tau_0(S_2) - h_2$.

Démonstration : Si $\tau_0(S_2) > (h_2 + b_1^{\min})/(h_2 + b_2)$, $b_1^{\text{opt}} > b_1^{\min}$. Alors, (4.30) indique que $\pi_2(S_2, b_1)$ est strictement croissante pour $b_1^{\min} < b_1 < b_1^{\text{opt}}$ et strictement décroissante pour $b_1 > b_1^{\text{opt}}$. Si $\tau_0(S_2) \leq (h_2 + b_1^{\min})/(h_2 + b_2)$, (4.30) indique que $\pi_2(S_2, b_1)$ est strictement décroissante pour $b_1 > b_1^{\min}$. $\square$

Le lemme 4.1 montre que, si $b_2 \geq b_1^{\min}$, il existe une valeur $S_2^{\max}$ qui satisfait

$$\tau_0(S_2^{\max}) = \frac{h_2 + b_1^{\min}}{h_2 + b_2}. \quad (4.33)$$

Dans le cas où $b_2 < b_1^{\min}$, $\tau_0(S_2) < (h_2 + b_1^{\min})/(h_2 + b_2)$ et $b_1^*(S_2) = b_1^{\min}$. En effet, le coût de rupture du producteur est suffisamment bas pour que le producteur n'ait pas besoin de forcer le fournisseur à installer un niveau de stock nominal positif. Donc, le producteur propose la valeur minimale de la pénalité de rupture $b_1^{\min}$. En recevant le contrat $(p_1^*(b_1^{\min}), b_1^{\min})$, le fournisseur installe $S_1^*(b_1^{\min}) = 0$.

Dans le cas où $b_2 \geq b_1^{\min}$, la proposition du producteur dépend de son niveau de stock nominal S_2 . Si $S_2 \geq S_2^{\max}$, $\tau_0(S_2) \leq (h_2 + b_1^{\min})/(h_2 + b_2)$ et $b_1^*(S_2) = b_1^{\min}$. Autrement dit, en prenant en compte les taux de fabrication et les coûts opérationnels du système, le niveau de stock nominal du producteur est

suffisant pour satisfaire les demandes sans que le fournisseur installe un niveau de stock nominal positif. Dans ce cas, le producteur offre le contrat $(p_1^*(b_1^{\min}), b_1^{\min})$ et le fournisseur installe $S_1^*(b_1^{\min}) = 0$. Si $S_2 < S_2^{\max}$, $\tau_0(S_2) > (h_2 + b_1^{\min})/(h_2 + b_2)$. Par conséquent, le producteur offre le contrat $(p_1^*(b_1^{\text{opt}}), b_1^{\text{opt}})$ et impose un niveau de stock nominal positif pour le fournisseur. En recevant ce contrat, le producteur installe un niveau de stock nominal positif, $S_1^*(b_1^{\text{opt}}) > 0$. Nous observons ainsi que la fonction $b_1^{\text{opt}}(S_2)$ est décroissante pour $S_2 \geq 0$. Puisque le producteur couvre les coûts opérationnels du fournisseur (voir (4.29)), le producteur prend en compte les coûts de stockage du fournisseur quand il détermine la pénalité de rupture à proposer. Dans le cas où $b_2 \geq b_1^{\min}$, le producteur ne propose jamais une pénalité de rupture plus élevée que b_2 , $b_1^{\min} \leq b_1^*(S_2) \leq b_2$.

Dans tous les cas, en offrant le contrat $(p_1^*(b_1^*), b_1^*)$, le producteur obtient le profit total du système centralisé, en laissant le fournisseur avec un profit nul. Rappelons que la valeur de la pénalité de rupture b_1 sert à contrôler le niveau de stock nominal du fournisseur et a une influence indirecte sur les grandeurs du paiement de transfert et du profit maximal du producteur seulement si $b_1 > b_1^{\min}$. Pour $b_1 \leq b_1^{\min}$, la fonction d'utilité du producteur est constante en b_1 . Par conséquent, le contrat $(p_1^*(b_1^*), b_1^*)$ peut aussi être appliqué en définissant $b_1^*(S_2)$ comme suit :

$$b_1^*(S_2) = \begin{cases} b_1^{\text{opt}}(S_2) & \text{si } \tau_0(S_2) > (h_2 + b_1^{\min})/(h_2 + b_2) \\ 0 & \text{si } \tau_0(S_2) \leq (h_2 + b_1^{\min})/(h_2 + b_2) \end{cases}$$

Pour cette application, le producteur offre la pénalité $b_1^*(S_2) = 0$ si $b_2 < b_1^{\min}$ ou si $b_2 \geq b_1^{\min}$ et $S_2 \geq S_2^{\max}$. En utilisant (4.27), le contrat correspondant peut être écrit comme $(c_1, 0)$. En recevant le contrat $(c_1, 0)$, le fournisseur installe $S_1^*(0) = 0$.

4.4.2.2. La valeur optimale du niveau de stock nominal

Le problème d'optimisation restant consiste à trouver la valeur optimale du niveau de stock nominal S_2 qui maximise $\pi_2(S_2, b_1^*(S_2))$. Soit $\alpha(S_2) = \alpha(b_1^{\text{opt}}(S_2))$ et $\pi_2(S_2) = \pi_2(S_2, b_1^*(S_2))$ avec la dérivée première :

$$\pi_2'(S_2) = \begin{cases} -h_2 - (h_2 + b_2)\tau_1(S_2) & \text{si } \tau_0(S_2) > (h_2 + b_1^{\min})/(h_2 + b_2) \\ -h_2 - (h_2 + b_2)\tau_2(S_2) & \text{si } \tau_0(S_2) \leq (h_2 + b_1^{\min})/(h_2 + b_2) \end{cases}$$

où

$$\tau_1(S_2) = \left(\frac{\alpha(S_2)\rho_1(1-\rho_2)}{\rho_1-\rho_2} \right) \frac{\rho_1^{S_2+1} \ln \rho_1}{1-\rho_1} + \left(1 - \frac{\alpha(S_2)\rho_1(1-\rho_2)}{\rho_1-\rho_2} \right) \frac{\rho_2^{S_2+1} \ln \rho_2}{1-\rho_2}, \quad (4.34)$$

$$\tau_2(S_2) = \left(\frac{\rho_1(1-\rho_2)}{\rho_1-\rho_2} \right) \frac{\rho_1^{S_2+1} \ln \rho_1}{1-\rho_1} + \left(1 - \frac{\rho_1(1-\rho_2)}{\rho_1-\rho_2} \right) \frac{\rho_2^{S_2+1} \ln \rho_2}{1-\rho_2}. \quad (4.35)$$

Notons que la fonction $\tau_1(S_2)$ est définie pour $S_2 \in [0, S_2^{\max})$.

Lemme 4.2: $\tau_2(S_2)$ est croissante pour $S_2 \geq 0$.

Démonstration: La fonction $\varphi_3(x) = 1/\varphi_1(x)$ est décroissante pour $x \in (0,1)$. Par conséquent, dans le cas

où $\rho_i > \rho_j$, $\frac{\rho_i \ln \rho_i}{1-\rho_i} < \frac{\rho_j \ln \rho_j}{1-\rho_j}$ qui nous donne $\frac{\rho_j(1-\rho_i) \ln \rho_j}{\rho_i(1-\rho_j) \ln \rho_i} < 1$ et $\left(\frac{\rho_j(1-\rho_i) \ln \rho_j}{\rho_i(1-\rho_j) \ln \rho_i} \right)^2 < 1$. Donc,

pour $S_2 \geq 0$,

$$\tau_2'(S_2) = \frac{(1-\rho_2)\rho_1^2\rho_2^{S_2}(\ln \rho_1)^2}{(\rho_1-\rho_2)(1-\rho_1)} \left(\left(\frac{\rho_1}{\rho_2} \right)^{S_2} - \left(\frac{\rho_2(1-\rho_1) \ln \rho_2}{\rho_1(1-\rho_2) \ln \rho_1} \right)^2 \right) > 0. \square$$

En utilisant le lemme 4.2, la propriété suivante définit la valeur optimale de S_2 pour le cas $b_2 < b_1^{\min}$.

Propriété 4.3 : Dans le cas où $b_2 < b_1^{\min}$, la fonction d'utilité du producteur $\pi_2(S_2)$ est concave pour $S_2 \geq 0$. La valeur optimale du niveau de stock nominal qui maximise $\pi_2(S_2)$ est alors $S_2^* = S_2^{\text{opt}_2}$ qui satisfait

$$\tau_2(S_2^{\text{opt}_2}) = \frac{-h_2}{h_2 + b_2}.$$

La fonction $\tau_1(S_2)$ n'est pas nécessairement croissante pour $S_2 \in [0, S_2^{\max})$. Par contre, les analyses numériques effectuées montrent que la fonction d'utilité $\pi_2(S_2)$ est strictement quasi-concave dans le cas où $b_2 \geq b_1^{\min}$. Notons que, pour $S_2 < S_2^{\max}$, $\alpha(S_2) < 1$ et par conséquent $\tau_1(S_2) > \tau_2(S_2)$ avec $\lim_{S_2 \rightarrow S_2^{\max}} \tau_1(S_2) = \tau_2(S_2^{\max})$. En utilisant ces observations, la propriété suivante définit la valeur optimale du niveau de stock nominal S_2 et les conditions d'optimalité nécessaires correspondantes.

Propriété 4.4 : Dans le cas où $b_2 \geq b_1^{\min}$, la fonction d'utilité du producteur $\pi_2(S_2)$ est strictement quasi-concave pour $S_2 \geq 0$ si $\tau_1(S_2) = -h_2/(h_2 + b_2)$ a une solution unique quand

$\tau_2(S_2^{\max}) > -h_2/(h_2 + b_2)$ et $\tau_1(S_2)$ est croissante en S_2 quand $\tau_2(S_2^{\max}) \leq -h_2/(h_2 + b_2)$. Sous ces conditions, la valeur optimale du niveau de stock nominal qui maximise $\pi_2(S_2)$ est

$$S_2^* = \begin{cases} S_2^{\text{opt}_1} & \text{si } \tau_2(S_2^{\max}) > -h_2/(h_2 + b_2) \\ S_2^{\text{opt}_2} & \text{si } \tau_2(S_2^{\max}) \leq -h_2/(h_2 + b_2) \end{cases}$$

$$\text{où } \tau_1(S_2^{\text{opt}_1}) = \frac{-h_2}{h_2 + b_2} \text{ et } \tau_2(S_2^{\text{opt}_2}) = \frac{-h_2}{h_2 + b_2}.$$

La Figure 4.7 montre l'évolution des niveaux de stocks nominaux S_1^* et S_2^* avec les différentes valeurs de ρ_1 et h_1 . Pour chaque problème, $b_1 > b_1^{\min}$ et les conditions de quasi-concavité de la propriété 4.4 sont satisfaites. Pour $h_1 = 0.6$, le fournisseur installe un niveau de stock nominal positif, $S_1^* > 0$, si le taux de fabrication du fournisseur est faible en comparaison avec le taux de fabrication du producteur, $\mu_1 < \mu_2$. Néanmoins, pour $h_1 = 0.2$, le producteur incite le fournisseur à installer un niveau de stock nominal positif même quand le taux de fabrication du fournisseur est plus élevé que le taux de fabrication du producteur, $\mu_1 > \mu_2$. Puisque le producteur compense les coûts de stockage du fournisseur en proposant le contrat $(p_1^*(b_1^*), b_1^*)$, le producteur force le fournisseur à installer un niveau de stock nominal élevé si le taux de fabrication et le coût unitaire de stockage du fournisseur sont faibles. Dans le cas où le coût unitaire de stockage du fournisseur est élevé, le producteur compense un faible taux de fabrication chez le fournisseur en augmentant le niveau de stock nominal S_2^* .

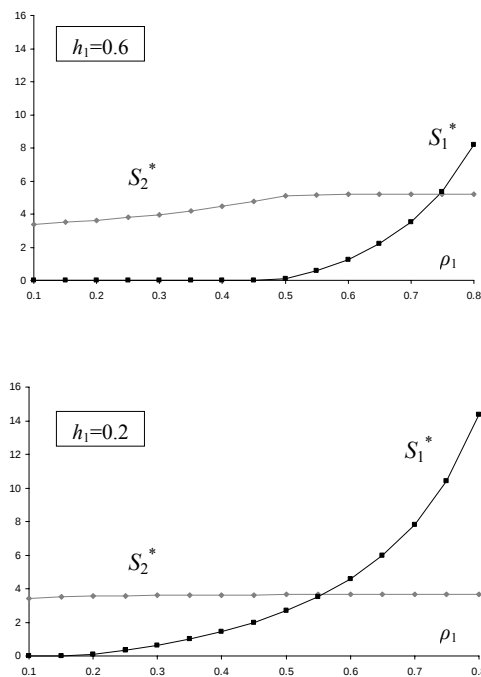


Figure 4.7. Les niveaux de stocks nominaux S_2^* et $S_1^*(b_1^*(S_2^*))$ pour $h_2 = 0.8$, $b_2 = 10$, $\rho_2 = 0.5$

4.5. PROBLÈME CENTRALISÉ ET PERFORMANCE DU CONTRAT

Les performances optimales d'une chaîne logistique nécessitent l'exécution d'un ensemble précis d'actions. Les actions conduisant à l'obtention des performances optimales sont définies comme les décisions optimisant les objectifs de la chaîne logistique globale. Dans cette section, nous déterminons les actions optimales de la chaîne logistique analysée afin d'étudier l'efficacité du contrat proposé.

Le profit moyen de la chaîne logistique centralisé (4.10) est la somme des fonctions d'utilité du fournisseur (4.17) et du producteur (4.21), $\pi_0(S_1, S_2) = \pi_1(S_1, p_1, b_1) + \pi_2(S_1, S_2, p_1, b_1)$. Les actions déterminant les performances du système sont les niveaux de stocks nominaux S_1 et S_2 .

Gallego et Zipkin (1999) et Zipkin (2000) montrent que si la méthode LZ est utilisée pour déterminer approximativement les mesures de performance d'un système centralisé à n étages, alors l'algorithme de Clark et Scarf (2004) peut être utilisé pour calculer la meilleure politique de stock nominal du type échelon. Nous savons aussi que chaque politique de stock nominal du type échelon est équivalente à une politique de stock nominal du type installation. Par conséquent, les valeurs optimales des niveaux de stocks nominaux S_1 et S_2 peuvent être calculées en utilisant l'algorithme de Clark et Scarf. Par contre, les résultats de l'algorithme de Clark et Scarf ne sont pas analytiquement comparables avec l'équilibre de Stackelberg obtenu dans la section précédente. Afin de pouvoir comparer l'équilibre de Stackelberg avec la solution optimale de la chaîne logistique, nous calculons d'abord la valeur du niveau de stock nominal S_1 qui maximise $\pi_0(S_1, S_2)$.

La dérivée partielle de la fonction $\pi_0(S_1, S_2)$ par rapport à S_1 s'écrit

$$\frac{\partial \pi_0(S_1, S_2)}{\partial S_1} = -h_1 - (h_1 - h_2 + (h_2 + b_2)\tau_0(S_2)) \frac{\rho_1^{S_1+1} \ln \rho_1}{1 - \rho_1}.$$

En suite,

$$\frac{\partial \pi_0(S_1, S_2)}{\partial S_1} = -(h_1 - h_2 + (h_2 + b_2)\tau_0(S_2)) \frac{\rho_1^{S_1+1} (\ln \rho_1)^2}{1 - \rho_1}.$$

Propriété 4.5 : *Étant donné S_2 , la fonction d'utilité $\pi_0(S_1, S_2)$ est concave pour $S_1 \geq 0$ si $\tau_0(S_2) > (h_2 + b_1^{\min})/(h_2 + b_2)$ et décroissante pour $S_1 > 0$ si $\tau_0(S_2) \leq (h_2 + b_1^{\min})/(h_2 + b_2)$. La valeur optimale du niveau de stock nominal S_1 qui maximise $\pi_0(S_1, S_2)$ est alors*

$$S_1^0(S_2) = \begin{cases} \ln(\alpha(S_2)) / \ln \rho_1 & \text{if } \tau_0(S_2) > (h_2 + b_1^{\min})/(h_2 + b_2) \\ 0 & \text{if } \tau_0(S_2) \leq (h_2 + b_1^{\min})/(h_2 + b_2) \end{cases}.$$

Nous observons que $S_1^0(S_2) = S_1^*(b_1^*(S_2))$. Par conséquent,

$$\pi_0(S_1^0(S_2), S_2) = \pi_2(S_1^*(b_1^*(S_2)), p_1^*(b_1^*(S_2)), b_1^*(S_2), S_2).$$

Notons que, dans la section précédente, nous avons utilisé $\pi_2(S_2)$ pour indiquer $\pi_2(S_1^*(b_1^*(S_2)), p_1^*(b_1^*(S_2)), b_1^*(S_2), S_2)$. Ensuite, la valeur optimale du niveau de stock nominal S_2 qui maximise $\pi_0(S_1^0(S_2), S_2)$ peut être obtenue par les propriétés 4.3 et 4.4, $S_2^0 = S_2^*$. L'équilibre du système décentralisé correspond alors à la solution optimale du système centralisé :

$$S_1^0(S_2^0) = S_1^*(b_1^*(S_2^*))$$

$$S_2^0 = S_2^*$$

Autrement dit, l'utilisation du contrat (p_1, b_1) ramène les performances du système décentralisé au niveau de celles du système centralisé. En utilisant le contrat (p_1, b_1) , la fonction d'utilité du producteur devient identique à l'espérance de profit du système centralisé (voir (4.29) et (4.10)). Donc, le producteur opte pour les valeurs optimales des niveaux de stocks nominaux du système centralisé. Il utilise la valeur de la pénalité de rupture $b_1^*(S_2^0)$ pour inciter le fournisseur à installer $S_1^0(S_2)$. En outre, l'espérance de profit total du système décentralisé dans l'équilibre de Stackelberg est égale à l'espérance de profit maximal du système centralisé :

$$\pi_0^*(S_1^0, S_2^0) = \pi_1^*(S_1^*, p_1^*, b_1^*) + \pi_2^*(S_1^*, S_2^*, p_1^*, b_1^*)$$

Puisque le profit du fournisseur est nul dans l'équilibre de Stackelberg, $\pi_1^*(S_1^*, p_1^*, b_1^*) = 0$, le producteur obtient le profit total du système centralisé, $\pi_2^*(S_1^*, S_2^*, p_1^*, b_1^*) = \pi_0^*(S_1^0, S_2^0)$.

Propriété 4.6 : Dans le jeu de Stackelberg entre le producteur et le fournisseur, le producteur offre le contrat $(p_1^*(b_1^*), b_1^*)$ maximisant son espérance de profit. En recevant ce contrat, le fournisseur installe le niveau de stock nominal $S_1^*(b_1^*)$. Ce contrat permet au producteur d'obtenir l'espérance de profit total du système centralisé, en laissant le fournisseur avec une espérance de profit nulle.

Un contrat est considéré efficace s'il coordonne la chaîne logistique, c'est-à-dire si son application positionne l'équilibre du système décentralisé à la solution optimale du système centralisé. Nous pouvons conclure que le contrat $(p_1^*(b_1^*), b_1^*)$ est efficace car il coordonne la chaîne logistique étudiée. Un autre aspect important est la facilité d'application d'un contrat. Un contrat qui est difficile à appliquer peut engendrer des coûts administratifs trop élevés. Les contrats simples sont souvent préférés même s'ils

n'optimisent pas les performances de la chaîne logistique. Nous arguons que le contrat proposé dans cette étude est un contrat simple car le producteur peut appliquer ce contrat en proposant un prix d'achat unitaire qui dépende du délai de livraison observé du fournisseur, $P_1(D_1) = p_1^*(b_1^*) - b_1^* D_1$ (voir (4.11)).

4.5.1. Prise en compte le pouvoir de négociation du fournisseur

Dans le jeu de Stackelberg étudié, nous avons supposé que le producteur a une situation dominante comme celle d'un donneur d'ordre. Étant le meneur, le producteur impose l'équilibre qui lui convient en agissant en premier. Le contrat proposé lui permet d'obtenir le profit total du système centralisé. Les résultats obtenus peuvent être généralisés pour un système dans lequel le fournisseur a un certain pouvoir de négociation. L'approche standard pour modéliser le pouvoir de négociation d'une entreprise est de supposer que l'entreprise a un profit de réservation positif et exogène (Cachon, 2003 ; Corbett et Tang, 1999 ; Corbett *et al.* 2004). Le profit de réservation exprime l'opportunité de l'entreprise dans le marché et détermine le seuil d'acceptation d'un contrat. Autrement dit, l'entreprise accepte un contrat si ce contrat lui permet d'obtenir un profit supérieur ou égal à son profit de réservation. Le pouvoir de négociation d'une entreprise est supposé croissant avec le niveau de son profit de réservation.

Supposons que le profit de réservation du fournisseur est $\pi_1^{\min} > 0$. Pour ce cas, le problème d'optimisation du producteur Π_1 peut être réécrit en remplaçant la contrainte (4.24) par

$$\pi_1(S_1^*, p_1, b_1) \geq \pi_1^{\min}.$$

Les valeurs élevées de $\pi_1^{\min}$ diminuent le pouvoir de négociation du producteur et peuvent même générer des profits négatifs pour le producteur. Pour aligner les pouvoirs de négociation des entreprise nous supposons ainsi que le profit minimal acceptable du producteur est $\pi_2^{\min} > 0$. Si $\pi_1^{\min} + \pi_2^{\min} \leq \pi_0^*(S_1^0, S_2^0)$ alors nous pouvons montrer qu'il existe un équilibre de Stackelberg dans lequel le producteur offre le contrat $(p_1^{**}(b_1^*), b_1^*)$ où $p_1^{**}(b_1^*)$ est la valeur qui satisfait $\pi_1(S_1^*(b_1^*), p_1^{**}(b_1^*), b_1^*) = \pi_1^{\min}$:

$$p_1^{**}(b_1^*) = p_1^*(b_1^*) + \frac{\pi_1^{\min}}{\lambda}$$

En utilisant les équations (4.13) et (4.14), le paiement de transfert correspondant s'écrit comme la somme des coûts opérationnels du fournisseur et de son profit de réservation :

$$E[T(S_1, p_1, b_1)] = \lambda c_1 + h_1 E[I_1] + \pi_1^{\min}.$$

En recevant le contrat $(p_1^{**}(b_1^*), b_1^*)$, le fournisseur installe $S_1^*(b_1^*)$. Le producteur installe S_2^* . Donc, le contrat $(p_1^{**}(b_1^*), b_1^*)$ coordonne la chaîne logistique. Dans l'équilibre de Stackelberg, les espérances de profit des entreprises s'écrivent :

$$\pi_1^*(S_1^*, p_1^*, b_1^*) = \pi_1^{\min}$$

$$\pi_2^*(S_1^*, S_2^*, p_1^*, b_1^*) = \pi_0^*(S_1^0, S_2^0) - \pi_1^{\min}$$

L'exigence du fournisseur pour obtenir au moins son profit de réservation en contractant avec le producteur oblige le producteur à partager le profit total du système centralisé avec le fournisseur. Le producteur n'accepte pas une relation contractuelle si $\pi_0^*(S_1^0, S_2^0) - \pi_1^{\min} < \pi_2^{\min}$.

4.6. CONCLUSIONS

Dans ce chapitre, nous avons étudié un modèle dynamique de flux reflétant les aspects aléatoires des demandes et des délais de livraisons dans les chaînes logistiques. L'évaluation exacte des mesures de performances du système analysé est possible seulement dans des situations particulières. Nous avons étudié le cas général avec $S_i \geq 0$ pour $i = 1, 2$ et nous avons déterminé les expressions analytiques des mesures de performances du système en nous basant sur la méthode approximative proposée par Lee et Zipkin (1992).

Nous avons considéré que chaque entreprise de la chaîne est une entité individuelle qui est principalement concernée par l'optimisation de ses propres objectifs. Cette focalisation locale conduit souvent à une dégradation des performances pour l'ensemble de la chaîne. Nous avons défini un jeu de Stackelberg entre les acteurs en supposant que le producteur a une situation dominante. Ensuite, nous avons montré que les performances optimales peuvent être obtenues en contractant sur un ensemble de paiements de transfert tel que l'objectif de chaque entreprise devient aligné avec l'objectif de la chaîne logistique.

Le contrat proposé est constitué de deux paramètres : le prix d'achat d'un produit intermédiaire et la pénalité de rupture de stock pour une commande de produit intermédiaire retardée. Le producteur peut aussi appliquer ce contrat en proposant à son fournisseur un prix unitaire d'achat qui diminue en fonction du délai de livraison observé. Contrairement aux études menées, le contrat proposé impose une pénalité pour le fournisseur seulement pour les livraisons retardées de produits intermédiaires et non pour les livraisons retardées de produits finis aux clients. Nous avons montré que le producteur peut simultanément optimiser les paramètres du contrat et son niveau de stock nominal. Nous avons déterminé analytiquement les niveaux de stocks nominaux des entreprises dans l'équilibre de Stackelberg.

Le paiement de transfert entre les acteurs est constitué d'une part d'un paiement qui dépende de la quantité de produit intermédiaires achetés et d'autre part d'un paiement de pénalité qui dépende du niveau moyen de rupture de stock du fournisseur. Puisque le taux moyen d'arrivée des demandes chez le producteur est fixe, le paiement correspondant aux achats est un paiement fixe : le producteur paie « λp_1 » à son fournisseur. L'existence d'un paiement fixe entre les acteurs permet au producteur de laisser le fournisseur avec son profit minimal acceptable : $\lambda p_1 = \lambda c_1 + h_1 E[I_1] + b_1 E[B_1] + \pi_1^{\min}$, alors $\pi_1(S_1, p_1, b_1) = \pi_1^{\min}$. Ainsi, la fonction objectif du producteur devient alignée avec la fonction objectif de la chaîne logistique centralisée : $\pi_2(S_1, S_2, p_1, b_1) = \pi_0(S_1, S_2) - \pi_1^{\min}$. Par conséquent, l'application du contrat proposé pour les valeurs optimales de ses paramètres élève les performances globales du système décentralisé au niveau de celles du système centralisé.

Le contrat proposé peut aussi être interprété comme un contrat de partage des revenus (Cachon et Lariviere, 2005). Pour chaque produit fini vendu sur le marché, le producteur paie une portion des revenus obtenus par cette vente à son fournisseur. Pour chaque produit fini vendu, la portion du fournisseur est alors $p_1 - b_1 E[D_1]$.

CONCLUSIONS GÉNÉRALES ET PERSPECTIVES

Dans le cadre de cette thèse, nous nous sommes intéressés aux implémentations de la politique de stock nominal dans les systèmes de production/stockage mono-étages et multi-étages. Nous avons commencé nos travaux par analyser l'application d'une stratégie multi-fournisseurs dans une chaîne logistique à deux niveaux. Nous avons étudié le problème de gestion des stock d'un producteur ayant une demande externe aléatoire d'un produit en supposant que la demande arrive en quantité unitaire et selon un processus de Poisson ayant le taux λ . Le producteur applique une politique de stock nominal $(S - 1, S)$ afin de gérer son stock à partir duquel ses clients vont être servis. Selon la politique de stock nominal, le producteur effectue une commande d'approvisionnement d'une unité chaque fois qu'une demande arrive. Nous avons analysé le cas où le producteur a l'option d'envoyer chaque commande d'approvisionnement à un fournisseur différent disponible dans le marché. Chaque fournisseur disponible fonctionne dans un mode de production à la commande avec un système de fabrication à capacité limitée, dans le sens où une unité de produit peut être fabriquée à la fois. Pour chaque fournisseur $i = 1, \dots, n$, le temps de fabrication est une variable aléatoire qui suit une loi exponentielle ayant le taux μ_i où les taux moyen de fabrication des fournisseurs se diffèrent : $\mu_i \neq \mu_j$ si $i \neq j$.

Pour ce système de production/stockage, nous avons étudié l'application d'une politique d'acheminement de commande probabiliste, nommé processus d'acheminement de Bernoulli, qui définit une probabilité d'affectation de commande α_i pour chaque fournisseur $i = 1, \dots, n$. À chaque déclenchement de réapprovisionnement du stock, le producteur détermine alors le fournisseur à qui il va affecter la commande selon les probabilités d'affectation de commande $(\alpha_i)_{i=1}^n$ fixées à l'avance. La probabilité d'affectation de commande α_i peut aussi être interprétée comme la fraction des demandes qui sont satisfaites par le fournisseur i . En appliquant un processus d'acheminement de Bernoulli, le système d'approvisionnement du producteur est un réseau ouvert de files d'attente constitué de n files d'attente M/M/1 en parallèle : M/M/1($\alpha_i \lambda, \mu_i$), $i = 1, \dots, n$. Nous avons fourni l'expression analytique de l'espérance de la somme des coûts moyens de stockage et de rupture en fonction des variables de décision du producteur, notamment le niveau de stock nominal S et les probabilités d'affectation $(\alpha_i)_{i=1}^n$.

Nous avons d'abord analysé le cas où le système fonctionne dans une mode de production à la commande avec $S = 0$. La fonction objectif du producteur dans ce cas est son délai moyen d'approvisionnement.

Nous avons déterminé les valeurs optimales des probabilités d'affectation de commande, notées $(\alpha_i^*)_{i=1}^n$, qui minimisent le délai moyen d'approvisionnement du producteur. L'étude menée montre qu'acheminer les commandes chez plusieurs fournisseurs à la place d'acheminer toutes les commandes chez le fournisseur le plus performant du marché diminue le délai moyen d'approvisionnement du producteur. La politique optimale d'acheminement de commande $(\alpha_i^*)_{i=1}^n$ définit un critère de sélection de fournisseur basé sur les performances des fournisseurs disponibles, ce critère est le *taux moyen de fabrication minimal* acceptable. Les fournisseurs ayant un taux de fabrication supérieur ou égal au taux moyen de fabrication minimal prédéfini sont jugés performants et des probabilités d'affectation de commande non-nulles leur sont attribuées. En revanche, la politique optimale attribue des probabilités d'affectation de commande nulles aux fournisseurs ayant un taux de fabrication inférieur au taux moyen de fabrication minimal prédéfini. Dans le cas où chaque fournisseur a le même taux moyen de fabrication, la politique optimale d'acheminement de commande est $(\alpha_i^* = 1/n)_{i=1}^n$.

Pour le cas de production pour stock avec $S \geq 0$, nous avons proposé de déterminer le niveau de stock nominal du producteur en supposant que les valeurs des probabilités d'affectation de commande sont données par $(\alpha_i^*)_{i=1}^n$. Les analyses numériques effectuées pour le cas de deux fournisseurs montrent que les différences entre les coûts optimaux et les coûts obtenus par la méthode approximative proposée sont légères. Le niveau de stock nominal obtenu par la méthode approximative ne diffère de sa valeur optimale que dans quatre problèmes parmi douze analysés où la différence est au plus 1. Par contre, la méthode approximative surestime la probabilité d'affectation de commande associée au fournisseur le plus performant. Le résultat le plus significatif obtenu des analyses numériques effectuées est que, en comparaison avec le cas où toutes les commandes d'approvisionnement doivent être affectées au fournisseur le plus performant du marché, opter pour une stratégie multi-fournisseurs diminue le niveau de stock nominal et les coûts moyens de stockage et de rupture du producteur.

L'application d'une politique d'acheminement de commande probabiliste en respectant les valeurs réelles de ces paramètres peut être difficile. Par contre, une telle politique permet d'approximer le comportement d'une entreprise qui voudrait contrôler les charges de travail de ses installations de production. Le modèle présenté peut être un outil pendant la phase de conception d'un réseau logistique, par exemple en déterminant les fournisseurs partenaires et le nombre de fournisseurs qui améliore les performances de la chaîne logistique. Le fait que les coûts moyens de stockage et de rupture sont diminués en optant pour une stratégie multi-fournisseurs peut être interprété comme une forme de mise en commun de risques. Dans la suite de ces travaux, notre défi à court terme sera de déterminer analytiquement les effets d'une stratégie multi-fournisseurs sur les niveaux moyens de stock possédés et de rupture de stock et les variances du délai d'approvisionnement et du nombre de commandes attendues. Nous voulons ainsi généraliser les résultats obtenus en considérant des distributions de probabilité plus générales. À long terme, il semble

intéressant d'analyser les effets d'une stratégie multi-fournisseurs dans un système multi-produits où chaque fournisseur est capable de fabriquer plusieurs produits. L'analyse de la politique optimale de gestion de flux en présence des produits multiples est complexe et mérite d'être étudiée en profondeur dans des travaux futurs. Une autre perspective est d'introduire les concepts liés à la gestion de flux de retours concernant des produits ou des emballages qui peuvent avoir pour objet la gestion du service après vente, la réutilisation des composants ou des emballages ou le recyclage.

Suite aux études effectuées sur les stratégies multi-fournisseurs, nous nous sommes focalisés sur la coordination des chaînes logistiques décentralisées. Nous avons analysé une chaîne logistique à deux niveaux gérés par deux acteurs différents : un producteur de produit fini et son fournisseur de produit intermédiaire. Chaque entreprise de niveau $i = 1$ (fournisseur), 2 (producteur) dispose d'un stock de sortie et d'un système de fabrication à capacité limitée qui approvisionne ce stock. Nous parlons alors d'une chaîne logistique à deux étages de production/stockage. Nous avons supposé que la demande du producteur suit un processus de Poisson ayant le taux λ . La demande du fournisseur est constituée des commandes d'approvisionnement effectuées par le producteur. Les temps de fabrication du produit intermédiaire chez le fournisseur et du produit fini chez le producteur sont des variables aléatoires à distribution de probabilité exponentielle ayant respectivement les taux μ_1 et μ_2 où $\mu_1 \neq \mu_2$. Afin de simplifier les notations, nous avons considéré que le producteur nécessite d'une unité de produit intermédiaire afin de fabriquer une unité de produit fini.

La gestion de stock dans chaque entreprise est accomplie suivant une politique de stock nominal du type installation: $(S_i - 1, S_i)$, $i = 1, 2$. En appliquant une politique de stock nominal à chaque étage, l'arrivée d'une demande finale déclenche une demande unitaire et en même temps un ordre de fabrication unitaire pour chaque entreprise de niveau $i = 1, 2$. Ainsi, en présence des matières nécessaires, le système de fabrication de l'étage i est en marche quand le niveau de stock de l'étage est inférieur au niveau de stock nominal S_i . Les mesures de performances du système analysé sont les niveaux moyens de stock possédés et les niveaux moyens de rupture de stock des entreprises. L'analyse exacte des mesures de performances est possible seulement dans des cas particuliers, notamment quand $S_1 = 0$ ou quand $S_1 \rightarrow \infty$. Afin d'obtenir des expressions analytiques des mesures de performances pour le cas général où $0 < S_1 < \infty$, nous avons adopté une méthode approximative issue de la littérature (Lee et Zipkin, 1992) qui se base sur les propriétés des lois de type phase.

Dans cette chaîne logistique décentralisée, chaque entreprise i est une entité individuelle qui décide de son niveau de stock nominal $S_i \geq 0$ dans le but de maximiser son profit moyen. Les performances du producteur vis-à-vis des clients dépendent non seulement de son niveau de stock nominal et de son taux de fabrication mais aussi des performances de son fournisseur : quand tous les produits intermédiaires dans le système sont consommés, une pénurie de produit intermédiaire est provoquée chez le producteur et cette pénurie bloque le système de fabrication de produits finis en présence des demandes. Le niveau de

stock nominal du fournisseur influence donc les performances globales du système. Par contre, le fournisseur n'a pas un intérêt direct concernant le niveau de service de la chaîne. Le système décentralisé nécessite alors des mécanismes de coordination motivant le fournisseur à opter pour un niveau de stock nominal qui permet d'obtenir des niveaux satisfaisants pour le délai de livraison de produits intermédiaires et par conséquent pour le délai de livraison de produits finis.

Afin d'analyser les interactions entre les deux entreprises, nous avons défini un jeu de Stackelberg en supposant que le producteur a une situation dominante. Nous avons proposé l'utilisation d'un contrat de coordination ayant deux paramètres : le prix d'achat d'un produit intermédiaire, $p_1 \geq 0$, et la pénalité que le fournisseur paie à son producteur pour chaque livraison retardée de produit intermédiaire par unité de temps, $b_1 \geq 0$. Le producteur étant le meneur dans ce jeu de Stackelberg, son problème d'optimisation consiste à trouver les valeurs optimales des paramètres du contrat d'une part, et de son niveau de stock nominal d'autre part qui maximisent son profit moyen et qui satisfont la contrainte de compatibilité d'incitation et la contrainte de rationalité individuelle. Dans cette forme, le modèle analysé peut être interprété comme un modèle « principal-agent » sans la présence de sélection adverse, de signalisation ou de risque moral.

Nous avons montré que la meilleure stratégie du fournisseur qui détermine son niveau de stock optimal dépend seulement de la pénalité de rupture de stock imposée par le producteur. Elle est donc croissante en b_1 . Le producteur peut concevoir un contrat (p_1^*, b_1^*) qui minimise son profit moyen et qui est acceptable par le fournisseur. En outre, nous avons montré que le producteur peut simultanément optimiser les paramètres du contrat et son niveau de stock nominal et nous avons déterminé analytiquement les niveaux de stocks nominaux des entreprises dans l'équilibre de Stackelberg.

L'application du contrat proposé pour les valeurs optimales de ses paramètres coordonne la chaîne logistique, c'est-à-dire élève les performances globales du système décentralisé au niveau de celles du système centralisé. La valeur optimale de la pénalité b_1^* est utilisée dans le but de contrôler le niveau de stock nominal du fournisseur. Le producteur détermine la valeur optimale de la pénalité b_1^* pour inciter le fournisseur à installer la valeur optimale du niveau de stock nominal S_1^* qui maximise l'espérance de profit total du système centralisé. Le paiement fixe λp_1^* correspondant à la valeur optimale du prix d'achat p_1^* compense les coûts opérationnels du fournisseur (coûts de production, de stockage et de ruptures) et permet au producteur de laisser son fournisseur avec son profit de réservation. Ainsi, le producteur obtient la différence entre l'espérance de profit total du système centralisé et le profit de réservation du fournisseur. Le producteur opte lui-même pour la valeur optimale du niveau de stock nominal S_2^* qui maximise l'espérance de profit total du système centralisé. L'espérance de profit total du

système décentralisé dans l'équilibre de Stackelberg est donc égale à l'espérance de profit maximal du système centralisé.

L'application du contrat proposé n'est pas limitée aux cas où les temps d'inter-arrivées des demandes et les temps de service ont une distribution exponentielle. Dans un système de stock nominal, le niveau de stock nominal est toujours croissant avec la pénalité de rupture de stock. Donc, le producteur peut toujours concevoir un contrat avec la valeur appropriée de la pénalité de rupture qui impose la valeur optimale du niveau de stock nominal du système centralisé pour le premier étage. En outre, les coûts opérationnels du fournisseur peuvent toujours être compensés à travers un paiement fixe. Le producteur peut donc toujours obtenir la différence entre l'espérance de profit total du système centralisé et le profit de réservation du fournisseur et laisser le fournisseur avec son profit de réservation. Une extension naturelle de cette étude est alors d'analyser les valeurs optimales des niveaux de stocks nominaux des entreprises et des paramètres du contrat pour les cas de temps d'inter-arrivées des demandes et de temps de service suivant des distributions de probabilité plus générales. Nous privilégierons, à court terme, les analyses des cas où les temps de fabrication ou les temps d'inter-arrivées des demandes suivent des lois de type phase qui permettent d'approximer de près des distributions de probabilité générales.

Le type de contrat proposé peut être utilisé dans une chaîne logistique à n étages. Considérons un système à n étages de production/stockage. Le coût de rupture de stock, b_n , et le prix de vente, p_n , de l'entreprise n sont exogènes. Examinons un jeu dynamique à information parfaite dans lequel les joueurs choisissent leurs stratégies successivement. L'entreprise n propose un contrat (p_{n-1}, b_{n-1}) à l'entreprise $n - 1$ et décide aussi de son niveau de stock nominal S_n . L'entreprise $n - 1$ observe le contrat proposé et propose lui-même un contrat (p_{n-2}, b_{n-2}) à l'entreprise $n - 2$ et détermine son niveau de stock nominal S_{n-1} . De la même façon, chaque joueur $i = n - 2, \dots, 2$ détermine ses stratégies (p_{i-1}, b_{i-1}) et S_i en observant les stratégies adoptées par le joueur $i - 1$. L'entreprise de niveau « 0 » étant un stock infini, l'entreprise de niveau 1 détermine seulement son niveau de stock nominal S_1 en recevant le contrat (p_1, b_1) proposé par l'entreprise de niveau 2. L'équilibre de ce jeu dynamique peut être déterminé par induction à rebours. Puisque chaque entreprise $i = 2, \dots, n$ fait un transferts d'argent fixe, λp_{i-1} , à son fournisseur, chaque entreprise $i < n$ obtient son profit de réservation, $\pi_i^{\min}$, et l'entreprise n obtient le profit total du système décentralisé moins $\sum_{i=1}^n \pi_i^{\min}$. Notons que, en utilisant la méthode approximative LZ, les niveaux moyens de stock possédé et de rupture de stock d'une entreprise i ne dépend pas seulement des niveaux de stock nominaux S_i et S_{i-1} mais dépend aussi des niveaux de stock nominaux S_j pour $j = i - 2, \dots, 1$. Selon cette dépendance, l'application d'un contrat (p_i, b_i) dans le but de contrôler l'étage i ne permet pas de coordonner tous les n étages, car chaque entreprise $i = 2, \dots, n$ détermine la pénalité de rupture b_{i-1} et donc le niveau de stock S_{i-1} dans le but de maximiser le profit total des i premiers étages. Autrement dit, un

étage $i = 2, \dots, n$ peut contrôler le niveau de stock nominal de l'étage $i - 1$ en appliquant un contrat (p_i, b_i) mais ne peut pas contrôler les niveaux de stock nominaux des entreprises en amont de l'étage $i - 1$. Dans la continuité de nos travaux, il serait intéressant d'étudier la dégradation des performances liée à la décentralisation dans les systèmes à n étages et d'analyser des différents types de contrat qui peuvent coordonner la chaîne logistique dans son ensemble.

À plus long terme, notre défi sera d'analyser les systèmes similaires de distribution constitués d'un fournisseur et de plusieurs producteurs où le fournisseur traitera les commandes de chaque producteur comme une classe de demandes différente. Dans ce contexte, il semble intéressant d'étudier la compétition entre les producteurs. Encore à long terme, nous voulons aussi étudier des modèles similaires qui considèrent les concepts d'information avancée sur les demandes futures et évaluer ses effets quand les informations entre acteurs sont asymétriques.

ANNEXE A

A. Processus Stochastiques

A.1. LOIS DE PROBABILITÉ ET PROCESSUS STOCHASTIQUES

Dans cette section, nous présentons les définitions des lois de probabilité et des processus stochastiques utilisés lors de cette thèse.

A.1.1. La loi géométrique

Une variable aléatoire X prenant des valeurs discrètes et non-négatives suit une loi géométrique ayant le paramètre $0 < p < 1$ si sa fonction de distribution de probabilité est donnée par :

$$\Pr\{X = n\} = (1 - p)p^n, \quad n = 0, 1, \dots$$

La variable aléatoire X a la moyenne $E[X] = \frac{p}{1 - p}$ et la variance $Var[X] = \frac{p}{(1 - p)^2}$.

A.1.2. La loi de Poisson

Une variable aléatoire X prenant des valeurs discrètes et non-négatives suit une loi de Poisson ayant le paramètre $\lambda > 0$ si sa fonction de distribution de probabilité est donnée par :

$$\Pr\{X = n\} = e^{-\lambda} \frac{(\lambda)^n}{n!}, \quad n = 0, 1, \dots$$

La variable aléatoire X a la moyenne $E[X] = \lambda$ et la variance $Var[X] = \lambda$.

A.1.3. La loi exponentielle

Une variable aléatoire X prenant des valeurs continues et non-négatives suit une loi exponentielle ayant le paramètre $\lambda > 0$ si sa fonction de distribution de probabilité est donnée par :

$$f_X(t) = \lambda e^{-\lambda t}, \quad t \geq 0.$$

La variable aléatoire X a la moyenne $E[X] = 1/\lambda$ et la variance $Var[X] = 1/\lambda^2$.

Une variable aléatoire continue X a la *propriété sans mémoire* si $\Pr\{X > s + t \mid X > t\} = \Pr\{X > s\}$ pour tous $s, t \geq 0$. La loi exponentielle est la seule loi continue qui possède la propriété sans mémoire.

A.1.4. Processus stochastique

Un processus stochastique à temps continu $\{X(t), t \geq 0\}$ est une collection des variables aléatoires $X(t)$ pour $t \geq 0$. L'index t représente le temps continu et $X(t)$ représente l'état du processus à l'instant t .

Un processus stochastique à temps discret $\{X_n, n \geq 0\}$ est défini de la même façon où l'index $n = 0, 1, 2, \dots$ représente le temps discret.

Un processus stochastique est à état discret si la variable d'état prend des valeurs discrètes et est à état continu si la variable d'état prend des valeurs continues.

A.1.5. Chaîne de Markov à temps discret

Une chaîne de Markov à temps discret $\{X_n, n \geq 0\}$ est un processus stochastique à temps discret et à état discret où

$$\Pr\{X_{n+1} = j \mid X_n = i, X_{n-1} = i_{n-1}, \dots, X_1 = i_1, X_0 = i_0\} = p_{ij}$$

pour tous les états $i_0, i_1, \dots, i_{n-1}, i, j$ et tous $n \geq 0$.

La matrice des probabilités de transition $\mathbf{P} = (p_{ij})$ satisfait $\mathbf{P} \geq 0$ et $\mathbf{P}\mathbf{1} = \mathbf{1}$ où $\mathbf{1}$ est le vecteur colonne dont tous ces éléments sont égaux à 1.

Le vecteur $\boldsymbol{\pi}(n) = (\pi_i(n))$ où $\pi_i(n) = \Pr\{X_n = i\}$ décrit l'évolution du processus comme suit : $\boldsymbol{\pi}(n) = \boldsymbol{\pi}(n-1)\mathbf{P} = \boldsymbol{\pi}(n-2)\mathbf{P}^2 = \dots = \boldsymbol{\pi}(0)\mathbf{P}^n$.

Si une chaîne de Markov $\{X_n, n \geq 0\}$ est irréductible et positive récurrente (voir Ross (2000)), une distribution de probabilité stationnaire $\boldsymbol{\pi}$ existe telle qu'elle est la solution unique qui satisfait l'équation de balance $\boldsymbol{\pi} = \boldsymbol{\pi}\mathbf{P}$ et l'équation de normalisation $\boldsymbol{\pi}\mathbf{1} = 1$. Si une distribution de probabilité stationnaire $\boldsymbol{\pi}$ existe et les probabilités initiales satisfont $\boldsymbol{\pi}(0) = \boldsymbol{\pi}$, alors $\boldsymbol{\pi}(n) = \boldsymbol{\pi}$ pour tous $n \geq 0$. Si une chaîne de Markov $\{X_n, n \geq 0\}$ est aussi apériodique (voir Ross (2000)), la limite $\lim_{n \rightarrow \infty} \boldsymbol{\pi}(n)$ existe et est égale à $\boldsymbol{\pi}$.

A.1.6. Chaîne de Markov à temps continu

Une chaîne de Markov à temps continu $\{X(t), t \geq 0\}$ est un processus stochastique à temps continu et à état discret où

$$\Pr\{X(t+s) = j | X(s) = i, X(u) = x(u), 0 \leq u < s\} = \Pr\{X(t+s) = j | X(s) = i\}$$

pour tous les états $x(u), 0 \leq u < s, i, j$ et tous $s, t \geq 0$. En outre, $\Pr\{X(t+s) = j | X(s) = i\}$ est indépendant de s .

Une chaîne de Markov à temps continu est aussi définie comme un processus stochastique qui a comme propriété :

- (i) Chaque fois que le processus entre dans un état i , l'intervalle de temps que le processus reste dans cet état avant de faire une transition à un autre état suit une distribution de probabilité exponentielle ayant le taux v_i ;
- (ii) Quand le processus quitte l'état i , il entre dans un état j avec la probabilité p_{ij} où $\mathbf{P} = (p_{ij})$ satisfait $\mathbf{P} \geq 0$, $\mathbf{P}\mathbf{1} = \mathbf{1}$ et $p_{ii} = 0$ pour tous i .

Le générateur (la matrice des taux de transitions) $\mathbf{Q} = (q_{ij})$ où $q_{ij} = v_i p_{ij}$ pour $j \neq i$ et $q_{ij} = -v_i$ pour $j = i$ satisfait $\mathbf{Q}\mathbf{1} = \mathbf{0}$.

Définissons le vecteur $\boldsymbol{\pi}(t) = (\pi_i(t))$ où $\pi_i(t) = \Pr\{X(t) = i\}$. Si une chaîne de Markov $\{X(t), t \geq 0\}$ est irréductible et positive récurrente (voir Ross (2000)), la limite $\lim_{n \rightarrow \infty} \boldsymbol{\pi}(n)$ existe telle qu'elle est la solution unique $\boldsymbol{\pi} = \lim_{n \rightarrow \infty} \boldsymbol{\pi}(n)$ qui satisfait l'équation de balance $\boldsymbol{\pi}\mathbf{Q} = \mathbf{0}$ et l'équation de normalisation $\boldsymbol{\pi}\mathbf{1} = 1$.

A.1.7. Processus de comptage

Un processus stochastique à temps continu et à état discret $\{N(t), t \geq 0\}$ est un processus comptage si $N(t)$ représente le nombre d'arrivées dans un intervalle de temps $(0, t]$. Un processus comptage satisfait les conditions suivantes :

- (i) $N(t) \geq 0$;
- (ii) $N(t)$ prend des valeurs entières ;
- (iii) Si $s < t$, alors $N(s) < N(t)$;

- (iv) Si $s < t$, alors $N(t) - N(s)$ est le nombre d'arrivées dans l'intervalle de temps $(s, t]$.

Pour $s > 0$, soit $N_s(t) = N(t + s) - N(t)$ l'incrément du processus de comptage dans l'intervalle $(t, t + s]$. Un processus de comptage a des *incréments indépendants*, si les variables aléatoires $N_s(t)$ sont indépendantes pour tous $s, t \geq 0$. Un processus de comptage a des *incréments stationnaires*, si les incréments $N_s(t)$ suivent la même loi de probabilité pour tous $s, t \geq 0$. Un processus de comptage ayant des incréments stationnaires est appelé un processus de comptage stationnaire.

A.1.8. Processus de Poisson

Un processus de comptage $\{N(t), t \geq 0\}$ est un processus de Poisson ayant le taux λ si

- (i) $N(0) = 0$;
- (ii) Le processus a des incréments indépendants ;
- (iii) Le nombre d'arrivées dans tout intervalle de longueur t suit une loi de Poisson de moyen λt .
C'est-à-dire, pour tous $s, t \geq 0$:

$$\Pr\{N(s+t) - N(s) = n\} = e^{-\lambda t} \frac{(\lambda t)^n}{n!}.$$

La condition (iii) indique qu'un processus de Poisson a des incréments stationnaires. En outre, $\text{Var}[N(t)] = E[N(t)] = \lambda t$.

Considérons un processus de Poisson. Soit T_1 le temps d'occurrence de la première arrivée et, pour $k > 1$, T_k le temps écoulé entre $(k-1)^{\text{ième}}$ et $k^{\text{ième}}$ arrivées. En utilisant les conditions (ii) et (iii), on peut montrer que les variables aléatoires $T_k, k = 1, 2, \dots$, sont indépendantes et suivent toutes une loi exponentielle de moyen $1/\lambda$: pour $k = 1, 2, \dots$, $\Pr\{T_k > t\} = e^{-\lambda t}$.

Par conséquent, un processus de Poisson est aussi caractérisé par la loi de probabilité des temps d'inter-arrivées :

Un processus de comptage $\{N(t), t \geq 0\}$ est un processus de Poisson de taux λ si et seulement si les temps d'inter-arrivées $T_k, k = 1, 2, \dots$, sont indépendants et suivent tous une loi exponentielle de moyen $1/\lambda$.

Un *processus de renouvellement* est un processus de comptage ayant des temps d'inter-arrivées indépendants et identiquement distribués. La loi de probabilité des temps d'inter-arrivées peut être une loi quelconque. Un processus de Poissons est alors un processus de renouvellement ayant des temps d'inter-arrivées exponentiels.

Propriété A.1 : Considérons un processus de Poisson $\{N(t), t \geq 0\}$ ayant le taux λ . Supposons que chaque fois qu'un événement se produit, il est classifié comme un type d'événement i , $i = 1, \dots, k$, avec la probabilité p_i où $\sum_{i=1}^k p_i = 1$. Soit $N_i(t)$ le nombre d'événements du type i qui se sont produits dans l'intervalle $(0, t]$. Notons que $N(t) = \sum_{i=1}^k N_i(t)$. Chaque processus stochastiques $\{N_i(t), t \geq 0\}$, $i = 1, \dots, k$, est alors un processus de Poisson ayant le taux $p_i \lambda$. En outre, les processus de Poisson $\{N_i(t), t \geq 0\}$, $i = 1, \dots, k$, sont mutuellement indépendants (Ross, 2000).

Propriété A.2 : Considérons un ensemble $\{N_i(t), t \geq 0, i = 1, \dots, k\}$ où $\{N_i(t), t \geq 0\}$ est un processus de Poisson ayant le taux λ_i pour $i = 1, \dots, k$. Si les processus de Poisson $\{N_i(t), t \geq 0\}$, $i = 1, \dots, k$, sont indépendants, alors le processus stochastique $\{N(t), t \geq 0\}$ où $N(t) = \sum_{i=1}^k N_i(t)$ est un processus de Poisson ayant le taux $\lambda = \sum_{i=1}^k \lambda_i$ (Ross, 2000).

A.1.9. Processus de Poisson composé

Un processus stochastique $\{X(t), t \geq 0\}$ est un processus de Poisson composé s'il peut être représenté comme

$$X(t) = \sum_{i=1}^{N(t)} Y_i$$

où $\{N(t), t \geq 0\}$ est un processus de Poisson et $\{Y_i, i \geq 1\}$ est une famille des variables aléatoires indépendantes et identiquement distribuées qui sont aussi indépendantes du processus $\{N(t), t \geq 0\}$.

A.2. MODÈLES DE FILES D'ATTENTE

Un modèle de file d'attente est une description abstraite d'un système réel de file d'attente. Les principaux domaines d'application des modèles de file d'attente sont les systèmes manufacturiers, les systèmes de production et de stockage, les systèmes de communication et les systèmes d'information. Les modèles de file d'attente sont particulièrement utiles pour l'analyse des effets de congestion dans les systèmes à capacité limitée. Dans les systèmes de file d'attente, les aléas et les limites de capacité provoquent ensemble la congestion.

La configuration basique d'un système de file d'attente peut être décrite de la manière suivante (Figure A.1) : les clients arrivent à intervalle aléatoire pour un service exécuté par un serveur dans un temps aléatoire. Un client arrivant se place dans la file pour attendre son tour. Une fois le serveur libéré, le client entre en service et occupe le serveur pendant tout son temps de service. Puis le client libère le serveur et quitte le système.

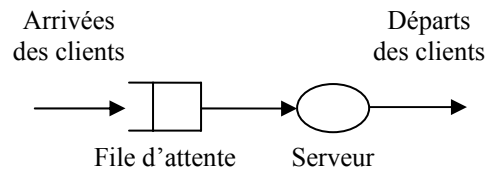


Figure A.1. Représentation graphique d'un système de file d'attente simple

A.2.1. Caractérisation des modèles de file d'attente

La définition d'un modèle de file d'attente nécessite principalement la caractérisation du processus d'arrivée, de la distribution du temps de service, du nombre de serveurs, de la capacité du système et de la discipline de service.

A.2.1.1. Processus d'arrivée

Le processus de Poisson est le processus stochastique le plus utilisé dans la théorie des files d'attente pour la modélisation du processus d'arrivée des clients. Pour $t \geq 0$, soit $N(t)$ la variable aléatoire indiquant le nombre d'arrivées dans un intervalle de temps $(0, t]$. Le processus stochastique d'arrivée $\{N(t), t \geq 0\}$ est alors un processus de comptage. Soit $N(s + t) - N(s)$ l'incrément d'un processus de comptage dans l'intervalle $(s, s + t]$. Un processus de Poisson est un processus de comptage $\{N(t), t \geq 0\}$ ayant des incréments indépendants et stationnaires.

Un processus de Poisson a un seul paramètre, $\lambda > 0$, appelé le taux d'arrivée. Pour tous $s, t \geq 0$, la variable aléatoire $N(s + t) - N(s)$ suit une loi de Poisson de moyenne λt . Pour $t = 1$, $\lambda = E[N(1)]$ exprime alors le nombre moyen d'arrivées par unité de temps, ce qui explique l'appellation « taux moyen d'arrivée » donnée à λ .

Dans les systèmes de production et de stockage, le processus de Poisson est largement utilisé pour modéliser la demande. Quand la demande est modélisée comme un processus de Poisson $\{D(t), t \geq 0\}$, les demandes sont supposées unitaires. $D(t)$ représente la demande cumulée dans un intervalle de temps $(0, t]$. Dans la plupart des situations pratiques, la demande se comporte vraiment comme un processus de Poisson. Une raison pour ceci est que les sources de la demande sont souvent de petites tailles et presque indépendants, par exemple quand les clients sont dispersés dans une région étendue. Un tel processus agrégé peut être modélisé convenablement par un processus de Poisson. La justification théorique de ce fait est le théorème de Palm-Khintchine. Soit $\{N(t) = N_1(t) + N_2(t) + N_m(t), t \geq 0\}$ la superposition de m processus de comptage indépendants et stationnaires. Selon le théorème de Palm-Khintchine, sous certaines conditions de régularité, la superposition $\{N(t) = N_1(t) + N_2(t) + N_m(t), t \geq 0\}$ se comporte approximativement comme un processus de Poisson lorsque m est assez grand. Ce théorème est souvent nommé le théorème central limite des processus de comptage.

Les temps d'inter-arrivées (les temps écoulés entre les arrivées successives) d'un processus de Poisson sont indépendants et identiquement distribués. Les temps d'inter-arrivées suivent tous une loi exponentielle de moyen $1/\lambda$.

Dans la théorie des files d'attente, le processus d'arrivée des clients est souvent modélisé par un processus de comptage ayant des temps d'inter-arrivées indépendants et identiquement distribués et le processus d'arrivée est caractérisé par la loi de probabilité des temps d'inter-arrivées. Dans le cas plus général, les temps d'inter-arrivées suivent une loi générale non-précisée de moyen $1/\lambda$.

Un autre élément important que l'on peut prendre en considération est les arrivées groupées (*batch arrivals*). Les arrivées groupées nécessitent la détermination de la loi de probabilité caractérisant les tailles de lots d'arrivées. Le processus de Poisson composé est un processus stochastique utilisé pour représenter les demandes ayant des quantités variables (*lumpy demand*). Quand la demande est représentée comme un processus de Poisson composé $\{D(t) = \sum_{i=1}^{N(t)} Y_i, t \geq 0\}$, les demandes non-unitaires arrivent suivant un processus de Poisson $\{N(t), t \geq 0\}$. La quantité de chaque demande est une variable aléatoire Y où Y_i représente la quantité de la $i^{\text{ème}}$ demande. Les quantités des demandes successives sont indépendantes et identiquement distribuées.

A.2.1.2. Distribution du temps de service

La caractérisation du temps de service est similaire à celle des arrivées. Soit X_i le temps de service de $i^{\text{ème}}$ client. Dans la théorie des files d'attente, on suppose souvent que les temps de service des clients, c'est-à-dire les variables aléatoires $X_i, i = 1, 2, \dots$, sont indépendantes et identiquement distribuées.

La distribution la plus souvent utilisée pour la caractérisation du temps de service X est la distribution de probabilité exponentielle. Le seul paramètre d'une distribution de probabilité exponentielle est $\mu > 0$ qui détermine sa moyenne $E[X] = 1/\mu$ et sa variance $Var[X] = 1/\mu^2$. Le paramètre μ exprime alors le « taux moyen de service ».

Toutes les durées réelles ne peuvent pas être représentées par une loi exponentielle. Les temps de service qui ne peuvent pas être décrits par une loi exponentielle sont souvent caractérisés en utilisant les généralisations de la loi exponentielle, comme la distribution d'Erlang, la distribution Gamma, la distribution hyper-exponentielle, la distribution hypo-exponentielle, et la distribution de Cox. Soit X la somme de n variables aléatoires suivant chacune une loi exponentielle de taux μ . La distribution de la variable aléatoire X est alors une distribution d'Erlang (ou Erlang- n) Notons que n est un nombre entier positif. La distribution Gamma est une généralisation de la distribution d'Erlang pour $n > 0$. La distribution hypo-exponentielle peut aussi être interprétée comme une généralisation de la distribution d'Erlang dans laquelle chaque variable aléatoire X_i suit une loi exponentielle de taux μ_i .

Les lois de type phase (Annexe C) permettent de caractériser des généralisations plus étendues. La famille de lois de type phase est constituée des distributions exponentielles et des distributions de toutes les sommes et les combinaisons finies des variables aléatoires exponentielles. Les distributions d'Erlang, hyper-exponentielle, hypo-exponentielle, et Cox sont des cas particuliers de cette famille de distributions. Chacune de ces distributions peut être caractérisée par une matrice analogue à celle d'une loi de type phase (Annexe C). L'utilisation des lois de type phase a un grand intérêt, car elles permettent de décrire des durées réelles plus complexes. Par exemple, le processus de fabrication d'un produit peut passer par plusieurs étapes de construction et de vérification ayant des temps opératoires suivant des lois exponentielles. En outre, des distributions de probabilité générales peuvent être approximées de près par des lois de type phase afin de tirer avantage des propriétés markoviens et des flexibilités analytiques des lois de type phase.

A.2.1.3. Nombre de serveurs

Le nombre de serveurs indique le nombre maximal d'exécutions en parallèle du même service. Dans un système de file d'attente multi-serveurs, les clients qui arrivent se placent dans une seule file d'attente. Chaque fois qu'un serveur est libéré, un client en attente dans la file entre en service. Les temps de service des serveurs sont généralement indépendants et identiquement distribués.

A.2.1.4. Capacité du système

Dans certains systèmes de file d'attente, des contraintes physiques ou organisationnelles peuvent exister et limitent la longueur maximale de la file. Dans ces types de cas, la capacité du système indique le nombre maximal de clients qui peuvent se retrouver dans le système (en attente de service et en service). Dans un système de production, cette capacité peut être liée à une limite de l'espace de stockage.

A.2.1.5. Discipline de service

La discipline de service détermine l'ordre dans lequel les clients sont sélectionnés pour le service. Les disciplines les plus utilisées sont : premier arrivé premier servi (*First In First Out* (FIFO), *First Come First Served* (FCFS)), dernier arrivé premier servi (*Last In First Out* (LIFO), *Last Come First Served* (LCFS)), sélection aléatoire, temps de service le plus court d'abord, règles de priorité préemptives (le service en cours d'exécution peut être interrompu) ou non-préemptives.

A.2.2. Notation des modèles de file d'attente

Dans la théorie des files d'attente, la notation de Kendall (premièrement proposée par D.G. Kendall en 1953) est un système standard pour décrire les caractéristiques essentielles d'un modèle de file d'attente. La notation de Kendall a la forme A/B/C/K/N/D où :

A : distribution des temps d'inter-arrivées
 B : distribution du temps de service
 C : nombre de serveurs
 K : capacité du système
 N : nombre de clients existant dans l'univers considéré
 D : discipline de service

Les différents symboles utilisés pour la caractérisation de la distribution des temps d'inter-arrivées et du temps de service sont : M (Markovien) pour une distribution exponentielle, G (Général) pour une distribution quelconque, PH pour une loi de type phase, E_k pour une distribution d'Erlang- k , D pour un temps déterministe, etc. Le symbole M^x est utilisé pour décrire un processus de poisson composé. Quand un système de file d'attente est caractérisé en utilisant seulement trois champs A/B/C, on sous-entend que le système est à capacité infinie, que la population est infinie et que la discipline de service est FIFO. Quelques exemples sont M/M/1, G/M/1, M/G/1, G/G/1, M/M/c,

A.2.3. Évaluation de performances

Soit $X(t)$ le nombre de clients dans un système de file d'attente (le nombre de clients en attente de service plus le nombre de clients en service) à l'instant $t \geq 0$. Sous certaines conditions, la distribution de $X(t)$ a une limite pour $t \rightarrow \infty$:

$$P_n = \lim_{t \rightarrow \infty} \Pr\{X(t) = n\}$$

L'existence de cette limite montre que, à long terme, le système atteint un régime permanent indépendant de son état initial. La limite P_n est interprétée comme la probabilité d'avoir exactement n clients dans le système en régime permanent. Quand les probabilités P_n , $n \geq 0$ existent, on dit que le processus stochastique $\{X(t), t \geq 0\}$ est ergodique. Pour la plupart des systèmes de files d'attente, la condition générale pour l'existence des probabilités P_n , $n \geq 0$, est la stabilité du système. Un système de file d'attente est dit stable si le nombre de clients dans le système ne peut pas augmenter jusqu'à l'infini. Si le processus stochastique $\{X(t), t \geq 0\}$ est régénératif P_n peut aussi être interprété comme la proportion du temps (à long terme) pendant laquelle le système a n clients. Notons qu'un système de file d'attente ayant un processus de renouvellement comme processus d'arrivée est régénératif.

Si les probabilités P_n , $n \geq 0$, existent, nous pouvons décrire l'état du système en régime permanent en utilisant la variable d'état :

X : le nombre de clients dans le système en régime permanent

Pour $n \geq 0$, la variable aléatoire X satisfait

$$P_n = \Pr\{X = n\} = \lim_{t \rightarrow \infty} \Pr\{X(t) = n\}.$$

Les autres quantités fondamentales utilisées afin d'analyser les performances d'un système sont :

- X_f : le nombre de clients en attente dans la file en régime permanent
- W : le temps de séjour des clients dans le système en régime permanent
(le temps d'attente plus le temps de service)
- W_f : le temps de séjour des clients dans la file d'attente en régime permanent

Dans la théorie des files d'attente, on s'intéresse plutôt aux mesures de performances espérées :

- $E[X]$: le nombre moyen de clients dans le système
- $E[X_f]$: le nombre moyen de clients en attente
- $E[W]$: le temps moyen de séjour des clients dans le système
- $E[W_f]$: le temps moyen de séjour des clients dans la file d'attente

Des relations entre ces mesures des performances sont désignées sous le nom de la loi de LITTLE (Kleinrock, 1975). Soit $N(t)$ le nombre de clients arrivés dans un intervalle de temps $(0, t]$. Le taux moyen d'arrivée $\lambda_a = \lim_{T \rightarrow \infty} N(T)/T$ exprime le nombre moyen de clients arrivés dans le système par unité de temps. La loi de LITTLE indique que

$$E[X] = \lambda_a E[W]. \quad (\text{A.1})$$

Si on applique la loi de LITTLE seulement à la file d'attente, on obtient

$$E[X_f] = \lambda_a E[W_f]. \quad (\text{A.2})$$

La loi de LITTLE est valide pour presque tous les systèmes de file d'attente indépendamment du processus d'arrivée, du nombre de serveurs, ou de la discipline de service.

La propriété PASTA (*Poisson Arrivals See Time Averages*) (Ross, 2000) est une propriété souvent utilisée pour l'analyse des systèmes de file d'attente du type $M/\cdot/\cdot$. Soit a_n la proportion des clients (à long terme) qui trouvent n clients dans le système quand ils arrivent. Selon la propriété PASTA, $a_n = P_n$ si et seulement si le processus d'arrivée est un processus de Poisson. En d'autres termes, la probabilité qu'un client quelconque trouve n clients dans le système en régime permanent est exactement la probabilité que le nombre de clients dans le système soit n en régime permanent. Pour la plupart des systèmes de file d'attente du type $M/\cdot/\cdot$, la propriété PASTA permet d'obtenir les mesures de performances du système sans déterminer la distribution de la variable d'état X .

A.2.4. La file M/M/1

Dans un système de file d'attente M/M/1 (λ, μ), les clients arrivent selon un processus de Poisson ayant le taux λ . Le temps de service suit une distribution de probabilité ayant le taux μ . La discipline de service est FIFO.

Soit $X(t)$ le nombre de clients dans le système à l'instant $t \geq 0$. Le processus stochastique $\{X(t), t \geq 0\}$ est une chaîne de Markov à temps continu ayant le diagramme de transition de la Figure A.2.

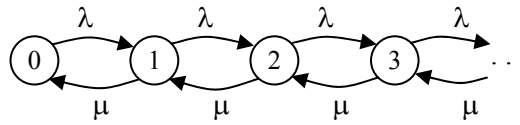


Figure A.2. Diagramme de transition d'un système de file d'attente M/M/1 (λ, μ)

Les probabilités $P_n, n \geq 0$, sont les solutions uniques des équations d'équilibre (*balance equations*) et de l'équation de normalisation du système. La probabilité que le processus soit dans l'état $n \geq 0$ en régime permanent est donnée par :

$$P_0 = \Pr\{X = 0\} = 1 - \frac{\lambda}{\mu} \quad (\text{A.3})$$

$$P_n = \Pr\{X = n\} = \left(\frac{\lambda}{\mu}\right)^n \left(1 - \frac{\lambda}{\mu}\right), \quad n \geq 1 \quad (\text{A.4})$$

La variable aléatoire X suit alors une distribution de probabilité géométrique. Notons qu'une condition nécessaire pour l'existence de la probabilité P_n pour $n = 0, \dots, \infty$ est $\lambda < \mu$. Autrement dit, le taux d'utilisation $\rho = \lambda / \mu$ qui exprime la proportion du temps pendant lequel le serveur est occupé doit satisfaire la condition $\rho < 1$. C'est la condition de stabilité du système. Quand $\lambda > \mu$, le nombre de clients dans le système augmente sans limite et donc les probabilités $P_n, n \geq 0$, n'existent pas (la chaîne de Markov n'est positive récurrente). Pour la plupart des systèmes de file d'attente ayant un seul serveur avec un temps moyen d'inter-arrivées $1/\lambda$ et un temps moyen de service $1/\mu$, la condition $\lambda / \mu < 1$ est aussi la condition de stabilité pour que les probabilités $P_n, n \geq 0$, existent. Pour un système ayant c serveurs, cette condition est décrite par $\lambda / c\mu < 1$.

Les mesures de performances du système sont obtenues en utilisant les expressions (A.3) et (A.4) et les relations (A.1) et (A.2) où $\lambda_a = \lambda$:

$$E[X] = \sum_{n=0}^{\infty} nP_n = \frac{\lambda}{\mu - \lambda}$$

$$E[W] = \frac{E[X]}{\lambda} = \frac{\lambda}{\mu - \lambda}$$

$$E[W_f] = E[W] - 1/\mu = \frac{1}{\mu(\mu - \lambda)}$$

$$E[X_f] = \lambda E[W_f] = \frac{\lambda^2}{\mu(\mu - \lambda)}$$

En utilisant la propriété sans mémoire de la distribution du temps de service et la propriété PASTA du processus d'arrivée, il est possible d'obtenir la distribution du temps de séjour : le temps de séjour d'un client quelconque dans un système de file d'attente M/M/1 (λ, μ) suit une distribution de probabilité exponentielle ayant le taux $\mu - \lambda$ (Ross, 2000).

La réversibilité du processus stochastique $\{X(t), t \geq 0\}$ indique que le processus de départ d'un système de file d'attente M/M/1 (λ, μ) est un processus de Poisson ayant le taux λ (Ross, 2000). Cette propriété est connue sous le nom du théorème de Burke : Pour un système de file d'attente M/M/1 (λ, μ), M/M/c (λ, μ) ou M/M/ ∞ (λ, μ) en régime permanent (1) le processus de départ est un processus de Poisson ayant le taux λ (2) le nombre de clients dans le système est indépendant des instants de départs antérieurs.

A.2.5. Réseaux des files d'attente

Un réseau des files d'attente est constitué de systèmes de file d'attente interconnectés. Les clients circulent entre les différents systèmes de files. Les réseaux des files d'attente sont bien adaptés pour la modélisation des interactions entre les différentes ressources d'un système de production, de stockage, de communication ou d'information.

Les réseaux des files d'attente sont principalement classés en trois types : réseaux ouverts, réseaux fermés et réseaux mixtes. Dans les réseaux de files d'attente ouverts, les clients qui entrent dans le réseau peuvent en sortir. Si un réseau ouvert est acyclique, les clients peuvent visiter un système de file d'attente seulement une fois avant de sortir du réseau. La représentation graphique d'un réseau ouvert (acyclique) de files d'attente constitué de deux systèmes de file d'attente en tandem est donnée dans la Figure A.3.

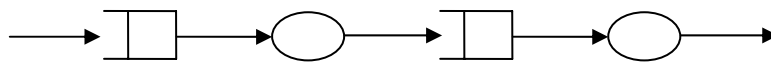


Figure A.3. Réseau ouvert (acyclique) des files d'attente en tandem

Dans les réseaux de files d'attente réentrants, les clients peuvent visiter un système de file d'attente plusieurs fois avant de sortir du réseau. Ce retour en arrière dans le réseau est appelé un feed-back. La Figure A.4 donne la représentation graphique d'un réseau ouvert des files d'attente avec feed-back où les

clients revisitent le premier système de file d'attente avec une probabilité p et passent au système de file d'attente suivant avec une probabilité $1 - p$.

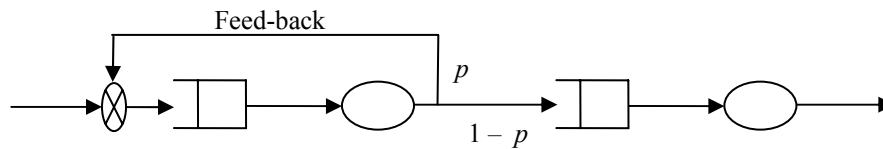


Figure A.4. Réseau ouvert des files d'attente en tandem avec feed-back

Dans les réseaux de files d'attente fermés, il n'y a jamais de nouveaux clients qui entrent le réseau et les clients existants n'en sortent jamais (Figure A.5).

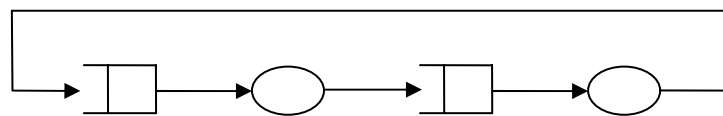


Figure A.5. Réseau fermé des files d'attente en tandem

S'il existe plusieurs classes de client, un réseau ouvert doit être ouvert pour chaque classe et un réseau fermé doit être fermé pour chaque classe. Un réseau des files d'attente qui est ouvert pour certaines classes et fermé pour d'autres est appelé un réseau mixte (Figure A.6).

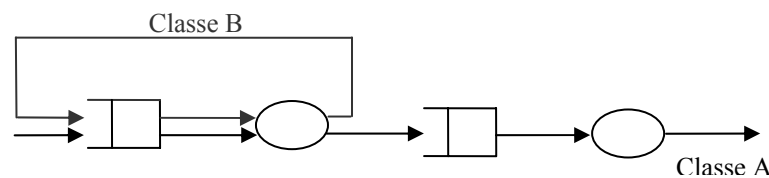


Figure A.6. Réseau mixte des files d'attente en tandem

Dans certaines situations, les clients peuvent être acheminés vers différents systèmes de files d'attente selon des probabilités prédéterminées (Figure A.7). Ces types de situations se rencontrent par exemple quand il existe des systèmes de files d'attente exécutant chacun un service différent et où les clients nécessitent les services disponibles selon des probabilités prédéterminées, ou quand il existe des systèmes de files d'attente exécutant tous le même service, où chaque client est acheminé vers un des systèmes de file d'attente selon des probabilités prédéterminées.

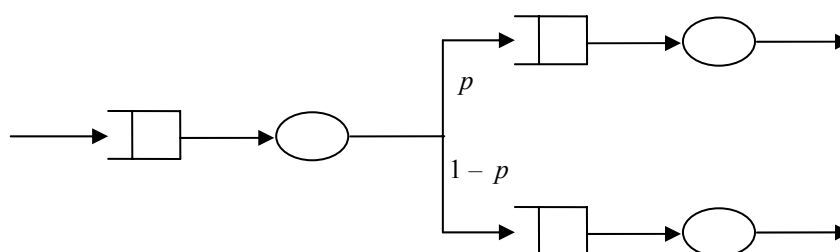


Figure A.7. Acheminement probabiliste dans un réseau ouvert des files d'attente

A.2.5.1. Réseaux à forme produit

Considérons un réseau ouvert acyclique de files d'attente en tandem constitué de k files d'attente à capacité illimitée. Il existe une seule classe de client. Chaque système de file d'attente i , $i = 1, \dots, k$, a un seul serveur dont la distribution du temps de service est exponentielle de taux μ_i . La discipline de service est FIFO. Les clients arrivent de l'extérieur dans le premier système de file d'attente selon un processus de Poisson ayant le taux λ . Une fois le client servi dans le système de file d'attente i , il entre dans le système de file d'attente $j = i + 1$, avec une probabilité égale à 1. Les clients partent du réseau des files d'attente après le service dans le système $j = k$.

Soit X_i le nombre de clients dans le système de file d'attente i en régime permanent. Le premier système de file d'attente est exactement un modèle M/M/1 (λ, μ_1). Selon le théorème de Burke, le processus de départ du premier système de file d'attente est un processus de Poisson ayant le taux λ . Par conséquent, le deuxième système de file d'attente est aussi un modèle M/M/1 (λ, μ_2). En continuant ce raisonnement pour $i = 3, \dots, k$, la probabilité d'avoir n_i clients dans le système de file d'attente i , $i = 1, \dots, k$, en régime permanent est

$$\Pr\{X_i = n_i\} = \left(\frac{\lambda}{\mu_i}\right)^{n_i} \left(1 - \frac{\lambda}{\mu_i}\right).$$

Supposons que la condition $\lambda < \mu_i$ soit satisfaite pour chaque $i = 1, \dots, k$. Selon le théorème de Burke, le nombre de clients dans un système de file d'attente i en régime permanent est indépendant des instants de départ antérieurs du système i . Puisque les instants des départ antérieurs constituent le processus d'arrivée du système de file d'attente $i + 1$, les variables d'état X_i et X_{i+1} sont mutuellement indépendantes pour $i = 1, \dots, k$. La probabilité que le réseau des files d'attente soit dans l'état $(n_1, \dots, n_k)$ en régime permanent est alors déterminé par le produit :

$$\Pr\{X_1 = n_1, \dots, X_k = n_k\} = \prod_{i=1}^k \left(\frac{\lambda}{\mu_i}\right)^{n_i} \left(1 - \frac{\lambda}{\mu_i}\right)$$

Ce résultat peut être généralisé pour les réseaux à acheminements probabilistes. Considérons un réseau ouvert des files d'attente constitué de k files d'attente à capacité illimitée ayant chacun un seul serveur avec une distribution du temps de service exponentielle. La discipline de service est FIFO. Les clients arrivent de l'extérieur dans le système de file d'attente i , $i = 1, \dots, k$, selon un processus de Poisson ayant le taux r_i (les processus de Poisson $i = 1, \dots, k$ sont indépendants). Une fois le client servi par le système de file d'attente i , le client rejoint le système de file d'attente $j = 1, \dots, k$, avec une probabilité P_{ij} où $\sum_{j=1}^k P_{ij} < 1$ et $1 - \sum_{j=1}^k P_{ij}$ représente la probabilité que le client parte du réseau des files d'attente après le service dans le système i . Notons que selon les propriétés A.1 et A.2, le processus d'arrivée de chaque

système de file d'attente est un processus de Poisson s'il n'y a pas de feed-back dans le réseau. Dans le cas contraire, les arrivées d'un système de files d'attente peuvent être dépendantes et le processus d'arrivée du système n'est pas un processus de Poisson. Par conséquent, les systèmes de file d'attente ne sont pas nécessairement des modèles M/M/1. En tous cas, la probabilité que le réseau des files d'attente soit dans l'état $(n_1, \dots, n_k)$ en régime permanent est exprimée par le produit

$$\Pr\{X_1 = n_1, \dots, X_k = n_k\} = \prod_{i=1}^k \left(\frac{\lambda_i}{\mu_i} \right)^{n_i} \left(1 - \frac{\lambda_i}{\mu_i} \right), \quad (\text{A.5})$$

où $\lambda_i = r_i + \sum_{j=1}^k \lambda_j P_{ji}$ doit satisfaire la condition $\lambda_i < \mu_i$ pour chaque $i = 1, \dots, k$. Cette résultat peut être vérifié en montrant que les probabilités (1.5) satisfont les équations de balance du réseau des files d'attente. Il est connu sous le nom du théorème de Jackson. Notons que le théorème de Jackson considère originalement des systèmes de files d'attente multi-serveurs.

La plupart des résultats connus sur les réseaux des files d'attente à forme produit sont fournis par Baskett *et al.* (1975). Selon la définition de Baskett *et al.*, un réseau des files d'attente est un réseau à forme produit si la probabilité que le réseau soit dans un état $(x_1, \dots, x_k)$ (la définition d'état peut être différente d'un système de file d'attente à l'autre) en régime permanent peut être exprimée sous la forme :

$$\Pr\{\mathbf{S} = (x_1, \dots, x_k)\} = C \times d(\mathbf{S}) \times \prod_{i=1}^k f_k(x_k)$$

où C est un constant de normalisation et $d(\mathbf{S})$ est une fonction de nombre totale de clients dans le réseau et $f_k(x_k)$ est une fonction qui dépend du type du système de file d'attente k . Baskett *et al.*, analysent des réseaux ouverts, fermés et mixtes avec des disciplines de service et des distributions du temps service variées. Ils ont caractérisé la catégorie des réseaux à forme produit.

A.2.5.2. Réseaux des files d'attente avec blocage

Si les systèmes de file d'attente dans un réseau sont à capacité limitée, une situation appelée *blocage* peut survenir. Le blocage désigne l'inactivité d'un système de file d'attente, en présence de clients en attente de service, à cause des limites de capacité des systèmes de file d'attente en aval. Dans la littérature, il existe plusieurs types de blocage. Nous citons à titre d'exemples le « blocage après service » et le « blocage avant service ». Dans le cas de blocage après service, un client qui vient d'être servi par un système de file d'attente i peut joindre la file d'attente du système $i + 1$ s'il y a une place disponible dans la file. Sinon, le client reste en attente de transfert et le serveur i reste bloqué jusqu'à ce qu'une place se libère dans le système $i + 1$. Dans le cas de blocage avant service, un client entre en service dans le système de file d'attente i seulement s'il y a une place disponible dans la file d'attente $i + 1$.

Les réseaux des files d'attente avec blocage ont des solutions à forme produit seulement dans des cas particuliers. Des techniques d'approximation analytiques ont été proposées afin de déterminer les probabilités d'état en régime permanent des réseaux des files d'attente avec blocage (voir par exemple Dallery et Frein (1993)).

Les lignes de production sont souvent modélisées par des réseaux des files d'attente en tandem avec blocage où chaque système de file d'attente correspond à un poste de travail. Papadopoulos et Heavey (1996) fournissent une revue de la littérature sur les modèles files d'attente des lignes de production. Pour une revue extensive de la littérature sur les applications de la théorie des files dans les systèmes manufacturiers, voir Rao *et al.* (1998).

A.2.5.3. Réseaux des files d'attente avec stations de synchronisation

Les réseaux des files d'attente permettent de représenter l'intégration synchronisée des flux. Considérons un système de files d'attente ayant des flux d'arrivée multiples, où chaque flux correspond à une entité différente. Les arrivées multiples du système peuvent être synchronisées en utilisant une station de synchronisation (Figure A.8). Un départ de la station de synchronisation vers le serveur ne se produit qu'en présence d'une entité dans chaque file et si le serveur est libre. Dans le cadre d'un système de production, chaque flux peut correspondre à un type de produit et le serveur peut être une machine exécutant une opération d'assemblage.

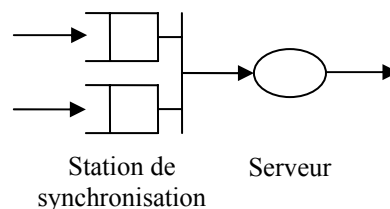


Figure A.8. Système de file d'attente avec synchronisation d'arrivées

Dans la littérature, les réseaux des files d'attente avec séparation de flux et intégration synchronisée de flux sont souvent appelés les réseaux *Fork/Join* avec synchronisation (Figure A.9). Dans le cadre d'un système de production, la séparation de flux peut correspondre par exemple à une opération de désassemblage.

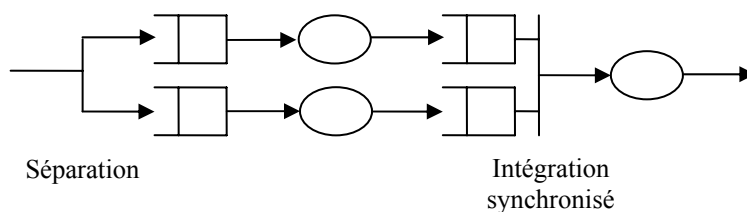


Figure A.9. Réseau Fork/Join avec synchronisation

A.3. FORMULATION DE LA FONCTION $C(S)$

Dans cette section, nous développons la fonction (1.1) présentée dans le chapitre 1. Soit $P_k = \Pr\{K = k\}$ la probabilité d'avoir k commandes dans la file d'attente M/M/1 (μ, λ) en régime permanent :

$$P_k = \rho^k (1 - \rho)$$

Les mesures de performances qui sont le niveau moyen de stock possédé $E[I]$ et le niveau moyen de rupture de stock $E[B]$ sont calculés par les équations

$$E[I] = S \sum_{k=0}^S P_k - \sum_{k=0}^S k P_k = S(1 - \rho) \sum_{k=0}^S \rho^k - (1 - \rho) \sum_{k=0}^S k \rho^k ,$$

$$E[B] = \sum_{k=S}^{\infty} k P_k - S \sum_{k=S}^{\infty} P_k = (1 - \rho) \sum_{k=S}^{\infty} k \rho^k - S(1 - \rho) \sum_{k=S}^{\infty} \rho^k ,$$

où

$$\sum_{k=0}^S \rho^k = \frac{1 - \rho^{S+1}}{1 - \rho} ,$$

$$\sum_{k=0}^S k \rho^k = \rho \sum_{k=0}^S k \rho^{k-1} = \rho \frac{d}{d\rho} \left(\sum_{k=0}^S \rho^k \right) = \rho \frac{d}{d\rho} \left(\frac{1 - \rho^{S+1}}{1 - \rho} \right) = -\frac{(S+1)\rho^{S+1}}{1 - \rho} + \frac{\rho(1 - \rho^{S+1})}{(1 - \rho)^2} ,$$

$$\sum_{k=S}^{\infty} k \rho^k = \rho \sum_{k=S}^{\infty} k \rho^{k-1} = \rho \frac{d}{d\rho} \left(\sum_{k=S}^{\infty} \rho^k \right) = \rho \frac{d}{d\rho} \left(\frac{\rho^S}{1 - \rho} \right) = \frac{S\rho^S}{1 - \rho} + \frac{\rho^{S+1}}{(1 - \rho)^2} ,$$

$$\sum_{k=S}^{\infty} \rho^k = \sum_{k=0}^{\infty} \rho^k - \sum_{k=0}^{S-1} \rho^k = \frac{1}{1 - \rho} - \frac{1 - \rho^S}{1 - \rho} = \frac{\rho^S}{1 - \rho} .$$

En utilisant ces relations, le niveau moyen de stock possédé $E[I]$ et le niveau moyen de rupture de stock $E[B]$ s'écrivent

$$E[I] = S - \frac{\rho(1 - \rho^S)}{1 - \rho} , \tag{A.6}$$

$$E[B] = \frac{\rho^{S+1}}{1 - \rho} . \tag{A.7}$$

La fonction de coût $C(S)$ est alors

$$C(S) = hE[I] + bE[B] = h \left(S - \frac{\rho(1-\rho^S)}{1-\rho} \right) + b \left(\frac{\rho^{S+1}}{1-\rho} \right). \quad (\text{A.8})$$

L'espérance mathématique du nombre de commandes dans la file d'attente M/M/1 (μ, λ) s'écrit

$$E[K] = \sum_{k=0}^{\infty} kP_k = \frac{\rho}{1-\rho} \quad (\text{A.9})$$

En utilisant la relation $S = E[K] + E[I] - E[B]$, le niveau moyen de rupture de stock $E[B]$ peut aussi être obtenu par les équations (A.6) et (A.9).

La fonction de coût (A.8) peut aussi être exprimée comme suit :

$$C(S) = (h+b)E[I] + b(E[K] - S) = (h+b) \left(S - \frac{\rho(1-\rho^S)}{1-\rho} \right) + b \left(\frac{\rho}{1-\rho} - S \right)$$

La fonction de coût (A.8) s'écrit également

$$C(S) = h(S - E[K]) + (h+b)E[B] = h \left(S - \frac{\rho}{1-\rho} \right) + (h+b) \left(\frac{\rho^{S+1}}{1-\rho} \right).$$

ANNEXE B

B. Démonstrations pour le modèle à plusieurs fournisseurs

B.1. FORMULATION DE LA FONCTION P_v

Dans cette section, nous approfondissons l'obtention de l'expression analytique de la fonction de distribution de probabilité P_v définie par l'équation (3.3). Nous présentons tout d'abord le concept de fonction de génération en théorie des probabilités. Ensuite, nous montrons l'obtention de la fonction (3.7) en utilisant les propriétés des fonctions de génération.

B.1.1. Fonctions de génération de probabilité

Soit X une variable aléatoire discrète et non-négative ayant la fonction de distribution de probabilité $P_x = \Pr\{X = x\}$. La fonction de génération de probabilité de la variable aléatoire X est définie comme

$$G_X(z) = \sum_x P_x z^x = E[z^X] \quad (\text{B.1})$$

pour $|z| \leq 1$. Puisque $\sum_x P_x = 1$, la somme $\sum_x P_x z^x$ est convergente pour $|z| \leq 1$.

La fonction de génération de probabilité $G_X(z)$ détermine la distribution de probabilité unique P_x , c'est-à-dire $G_X(z) = G_Y(z)$ si et seulement si $P_x = P_y$. En outre, $G_X(z)$ détermine les moments de la variable X à travers l'équation $G_X^{(r)}(1) = E[X(X-1)\dots(X-r+1)]$ où $G_X^{(r)}(z) = \frac{d^r}{dz^r} G_X(z)$. L'utilisation des fonctions de génération de probabilité facilite l'obtention de la fonction de distribution de probabilité et les moments quand on travaille avec les séquences, les sommes, ou les collections des variables aléatoires discrètes.

La fonction de génération de probabilité de la somme des variables aléatoires indépendantes est le produit des fonctions de génération de probabilités individuelles. Soit X et Y les variables aléatoires

indépendantes ayant les fonctions de génération de probabilité $G_X(z)$ et $G_Y(z)$ respectivement. La fonction de génération de probabilité $G_{X+Y}(z)$ de la variable $X + Y$ s'écrit

$$G_{X+Y}(z) = E[z^{X+Y}] = E[z^X \times z^Y] = E[z^X] \times E[z^Y] = G_X(z) \times G_Y(z). \quad (\text{B.2})$$

B.1.2. Obtention de la fonction P_v

Puisque la fonction de génération de probabilité détermine la distribution de probabilité de façon unique, nous pouvons obtenir la fonction de distribution de probabilité P_v en utilisant la fonction de génération de probabilité de la variable aléatoire V :

$$G_V(z) = G_{K_1+K_2+\dots+K_n}(z) = G_{K_1}(z) \times G_{K_2}(z) \times \dots \times G_{K_n}(z) \quad (\text{B.3})$$

Pour $i = 1, \dots, n$, la fonction de génération de probabilité de la variable aléatoire K_i s'écrit

$$G_{K_i}(z) = \frac{1 - \alpha_i \rho_i}{1 - \alpha_i \rho_i z}. \quad (\text{B.4})$$

Alors, la fonction de génération de probabilité $G_V(z)$ est en forme de produit :

$$G_V(z) = \prod_{i=1}^n \frac{1 - \alpha_i \rho_i}{1 - \alpha_i \rho_i z} \quad (\text{B.5})$$

En supposant $\alpha_i \rho_i \neq \alpha_j \rho_j$ pour $\forall i \neq j$, nous pouvons réécrire la fonction $G_V(z)$ en forme de somme (Kleinrock, 1975) :

$$G_V(z) = \sum_{i=1}^n \frac{A_i}{1 - \alpha_i \rho_i z} \quad (\text{B.6})$$

Calculons les coefficients $A_1, A_2, \dots, A_n$ qui satisfont l'égalité des équations (B.5) et (B.6). Le coefficient A_i pour $i = 1, \dots, n$ est la valeur qui satisfait l'égalité

$$(1 - \alpha_i \rho_i z) \left(\prod_{j=1}^n \frac{1 - \alpha_j \rho_j}{1 - \alpha_j \rho_j z} \right) = (1 - \alpha_i \rho_i z) \left(\sum_{j=1}^n \frac{A_j}{1 - \alpha_j \rho_j z} \right),$$

ou également l'égalité

$$(1 - \alpha_i \rho_i) \left(\prod_{\substack{j=1 \\ j \neq i}}^n \frac{1 - \alpha_j \rho_j}{1 - \alpha_j \rho_j z} \right) = A_i + (1 - \alpha_i \rho_i z) \left(\sum_{\substack{j=1 \\ j \neq i}}^n \frac{A_j}{1 - \alpha_j \rho_j z} \right)$$

pour $z = 1 / \alpha_i \rho_i$:

$$A_i = \left(\prod_{\substack{j=1 \\ j \neq i}}^n \frac{1 - \alpha_j \rho_j}{\alpha_i \rho_i - \alpha_j \rho_j} \right) (\alpha_i \rho_i)^{n-1} (1 - \alpha_i \rho_i) \quad (\text{B.7})$$

En utilisant l'équation (B.7), la fonction $G_V(z)$ s'écrit

$$G_V(z) = \sum_{i=1}^n \left(\left(\prod_{\substack{j=1 \\ j \neq i}}^n \frac{1 - \alpha_j \rho_j}{\alpha_i \rho_i - \alpha_j \rho_j} \right) \frac{(\alpha_i \rho_i)^{n-1} (1 - \alpha_i \rho_i)}{1 - \alpha_i \rho_i z} \right). \quad (\text{B.8})$$

Sachant que

$$\frac{1}{1 - \alpha_i \rho_i z} = \sum_{v=0}^{\infty} (\alpha_i \rho_i z)^v,$$

nous pouvons écrire la fonction $G_V(z)$ comme suit :

$$G_V(z) = \sum_{i=1}^n \left(\left(\prod_{\substack{j=1 \\ j \neq i}}^n \frac{1 - \alpha_j \rho_j}{\alpha_i \rho_i - \alpha_j \rho_j} \right) (\alpha_i \rho_i)^{n-1} (1 - \alpha_i \rho_i) \sum_{v=0}^{\infty} (\alpha_i \rho_i z)^v \right) \quad (\text{B.9})$$

Selon la définition (B.1), la fonction de distribution de probabilité P_v est alors

$$P_v = \sum_{i=1}^n \left(\left(\prod_{\substack{j=1 \\ j \neq i}}^n \frac{1 - \alpha_j \rho_j}{\alpha_i \rho_i - \alpha_j \rho_j} \right) (\alpha_i \rho_i)^{n+v-1} (1 - \alpha_i \rho_i) \right). \quad (\text{B.10})$$

B.2. DÉMONSTRATION DU LEMME 3.1

En utilisant la définition (3.28) du paramètre τ_m , nous pouvons écrire

$$\tau_m \sum_{i=1}^{m+1} \sqrt{\mu_i} - \tau_m \sqrt{\mu_{m+1}} = \tau_m \sum_{i=1}^m \sqrt{\mu_i} = \sum_{i=1}^{m+1} \mu_i - \mu_{m+1} - \lambda, \quad (\text{B.11})$$

$$\tau_{m+1} \sum_{i=1}^m \sqrt{\mu_i} + \tau_{m+1} \sqrt{\mu_{m+1}} = \tau_{m+1} \sum_{i=1}^{m+1} \sqrt{\mu_i} = \sum_{i=1}^m \mu_i + \mu_{m+1} - \lambda. \quad (\text{B.12})$$

Les relations (B.11) et (B.12) nous donnent

$$\mu_{m+1} - \tau_m \sqrt{\mu_{m+1}} = (\tau_{m+1} - \tau_m) \sum_{i=1}^{m+1} \sqrt{\mu_i}, \quad (\text{B.13})$$

$$\mu_{m+1} - \tau_{m+1} \sqrt{\mu_{m+1}} = (\tau_{m+1} - \tau_m) \sum_{i=1}^m \sqrt{\mu_i}. \quad (\text{B.14})$$

Les règles du lemme 3.1 sont les résultats directs des relations (B.13) et (B.14).

B.3. DÉMONSTRATION DU LEMME 3.2

Nous exposons la démonstration en deux étapes. Dans la première étape, nous montrons l'existence de l'indice m^* qui satisfait les relations (3.29) et (3.30). Dans la deuxième étape, nous montrons l'unicité de l'indice m^* et que l'évolution du paramètre τ_m suit le lemme 3.2.

B.3.1. Existence de l'indice m^*

Quel que soit ensemble des paramètres $(\lambda, \mu_1, \dots, \mu_n)$, la condition (3.15) implique $\tau_n > 0$. Soit k le plus petit indice qui satisfait $\sum_{i=1}^k \mu_i > \lambda$. En écrivant $m + 1 = k$ dans l'équation (B.14), nous obtenons

$$\mu_k - \tau_k \sqrt{\mu_k} = (\tau_k - \tau_{k-1}) \sum_{i=1}^{k-1} \sqrt{\mu_i}. \quad (\text{B.15})$$

Sachant que $\tau_k > 0$ et que $\tau_{k-1} \leq 0$, l'équation (B.15) implique $\mu_k > \tau_k^2$. Si $\mu_{k+1} \leq \tau_k^2$, alors $m^* = k$. Sinon, la relation $\mu_{k+1} > \tau_k^2$ implique $\mu_{k+1} > \tau_{k+1}^2$ et $\tau_{k+1} > \tau_k$ selon la troisième règle du lemme 3.1. Ensuite, si $\mu_{k+2} \leq \tau_{k+1}^2$, alors $m^* = k + 1$. Sinon, le processus itératif continue pour les indices $m = k + 2, \dots, n$. En continuant ce processus itératif, supposons que $\mu_{n-1} > \tau_{n-1}^2$ et $\mu_n > \tau_{n-1}^2$ pour $m = n - 1$. La relation $\mu_n > \tau_{n-1}^2$ implique $\mu_n > \tau_n^2$ et $\tau_n > \tau_{n-1}$. Nous savons que $\mu_{n+1} \leq \tau_n^2$ car $\mu_{n+1} = 0$. Alors $m^* = n$. Donc, il existe un indice m^* tel que $1 \leq m^* \leq n$ qui satisfait les relations (3.29) et (3.30).

Notons que dans le cas où $k = n$, l'équation (B.15) implique $\mu_n > \tau_n^2$ et par conséquent $m^* = n$.

B.3.2. Unicité de l'indice m^* et évolution du paramètre τ_m

Pour les indices $m = m^*+1, \dots, n$, nous avons la relation $\sum_{i=1}^m \mu_i > \sum_{i=1}^{m^*} \mu_i > \lambda$ qui implique $\tau_m > 0$. Sachant que $\mu_{m^*+1} \leq \tau_{m^*}^2$, nous pouvons écrire $\tau_{m^*+1} \leq \tau_{m^*}$ et $\mu_{m^*+1} \leq \tau_{m^*+1}^2$ en utilisant le lemme 3.1. Donc, la relation $\mu_{m^*+2} < \mu_{m^*+1} \leq \tau_{m^*+1}^2$ implique $\tau_{m^*+2} < \tau_{m^*+1}$. En appliquant le même raisonnement pour $m = m^*+3, \dots, n-1$, nous pouvons montrer que la valeur du paramètre τ_m décroît de manière monotone avec m pour $m^* < m < n$. Autrement dit, $\tau_m > \tau_{m+1}$ pour $m = m^*+1, \dots, n-1$.

Selon les équations (B.13) et (B.14), l'évolution du paramètre τ_m doit être croissante avec m pour $1 \leq m \leq k-1$, car $\mu_i > 0$ pour $i = 1, \dots, n$. Sachant que $\tau_m > 0$ pour $m = k, \dots, m^*$ et que $\mu_{m^*} > \tau_{m^*}^2$, nous pouvons écrire $\tau_{m^*} > \tau_{m^*-1}$ et $\mu_{m^*} > \tau_{m^*-1}^2$ selon le lemme 3.1. Par conséquent, la relation $\mu_{m^*-1} > \mu_{m^*} > \tau_{m^*-1}^2$ implique $\tau_{m^*-1} > \tau_{m^*-2}$. Nous pouvons appliquer le même raisonnement pour $m = k, \dots, m^*-3$. Donc, la valeur du paramètre τ_m est croissante avec m pour $1 \leq m \leq m^*$ et finalement la valeur maximale du paramètre τ_m est obtenue pour m^* . En outre, puisque nous avons $\mu_{m+1} < \mu_m \leq \tau_m^2$ pour $m = m^*+1, \dots, n$ et $\mu_m > \mu_{m+1} > \tau_m^2$ pour $m = k, \dots, m^*-1$, l'indice m^* est l'indice unique qui satisfait les relations (3.29) et (3.30).

B.4. DÉMONSTRATION DE LA PROPRIÉTÉ 3.2

Nous allons exposer la démonstration en deux étapes. Dans la première étape, nous montrons que la politique $\alpha^*(m^*) = (\alpha_i^*(m^*))_{i=1}^n$ définie par la propriété 3.2 est admissible. Dans la deuxième étape, nous montrons que la politique $\alpha^*(m^*)$ est optimale.

B.4.1. Admissibilité de la politique $\alpha^*(m^*)$

L'ensemble des paramètres de Bernoulli $\alpha^*(m^*) = (\alpha_i^*(m^*))_{i=1}^n$ défini par (3.31) satisfait la contrainte (3.13) :

$$\sum_{i=1}^n \alpha_i^* = \sum_{i=1}^{m^*} \frac{1}{\lambda} (\mu_i - \tau_{m^*} \sqrt{\mu_i}) = \frac{1}{\lambda} \left(\sum_{i=1}^{m^*} \mu_i - \tau_{m^*} \sum_{i=1}^{m^*} \sqrt{\mu_i} \right) = 1 \quad (\text{B.16})$$

Si $m^* = 1$, l'ensemble $(\alpha_1^* = 1, \alpha_i^* = 0 \text{ pour } i = 1, \dots, n)$ est admissible. Si $m^* > 1$, puisque $\mu_i > \mu_{m^*}$ pour $i = 1, \dots, m^*-1$ et $\mu_{m^*} > \tau_{m^*}^2$, nous pouvons écrire $\mu_i > \mu_{m^*} > \tau_{m^*}^2$ qui implique $\mu_i > \tau_{m^*} \sqrt{\mu_i}$ pour

$i = 1, \dots, m^*$. Selon la définition (3.31), nous pouvons conclure que $\alpha_i^* > 0$ pour $i = 1, \dots, m^*$. De plus, les contraintes (3.14) sont satisfaites car $\tau_{m^*} > 0$. Donc, la propriété 3.2 définit une solution admissible.

B.4.2. Optimalité de la politique $\alpha^*(m^*)$

D'après la propriété 3.1, la politique $\alpha^*(m^*)$ est optimale pour $m^* = n$. Si $m^* < n$, sachant que le problème Π_1 est un problème d'optimisation convexe, il suffit de montrer que l'ensemble des paramètres de Bernoulli $(\alpha_1^*, \alpha_2^*, \dots, \alpha_{m^*}^*, 0, \dots, 0)$ constitue un optimum local.

Considérons une variation $\partial\alpha_i$, pour $i = 1, \dots, n$, telle que $(\alpha_i + \partial\alpha_i)_{i=1}^n$ soit admissible. L'admissibilité de la solution $(\alpha_i + \partial\alpha_i)_{i=1}^n$ nécessite en particulier $\alpha_i + \partial\alpha_i \geq 0$ pour $i = 1, \dots, m^*$, $\partial\alpha_j \geq 0$ pour $j = m^* + 1, \dots, n$ et

$$\sum_{i=1}^{m^*} \partial\alpha_i + \sum_{j=m^*+1}^n \partial\alpha_j = 0. \quad (\text{B.17})$$

La variation de la fonction objectif $E[W]$ s'écrit

$$\partial E[W] = \sum_{i=1}^{m^*} \left(\frac{\alpha_i^* + \partial\alpha_i}{\mu_i - (\alpha_i^* + \partial\alpha_i)\lambda} - \frac{\alpha_i^*}{\mu_i - \alpha_i^*\lambda} \right) + \sum_{j=m^*+1}^n \left(\frac{\partial\alpha_j}{\mu_j - \partial\alpha_j\lambda} \right)$$

ou également

$$\partial E[W] = \sum_{i=1}^{m^*} \left(\frac{\partial\alpha_i \mu_i}{(\mu_i - (\alpha_i^* + \partial\alpha_i)\lambda)(\mu_i - \alpha_i^*\lambda)} \right) + \sum_{j=m^*+1}^n \left(\frac{\partial\alpha_j}{\mu_j - \partial\alpha_j\lambda} \right)$$

Supposons que $\partial\alpha_j > 0$ pour au moins un $j \geq m^* + 1$. Nous pouvons écrire alors l'inégalité suivante :

$$\partial E[W] > \sum_{i=1}^{m^*} \left(\frac{\partial\alpha_i \mu_i}{(\mu_i - \alpha_i^*\lambda)^2} \right) + \sum_{j=m^*+1}^n \left(\frac{\partial\alpha_j}{\mu_j} \right) = \sum_{i=1}^{m^*} \left(\frac{\partial\alpha_i}{\tau_{m^*}^2} \right) + \sum_{j=m^*+1}^n \left(\frac{\partial\alpha_j}{\mu_j} \right) = \frac{1}{\tau_{m^*}^2} \sum_{i=1}^{m^*} \partial\alpha_i + \sum_{j=m^*+1}^n \left(\frac{\partial\alpha_j}{\mu_j} \right)$$

Sachant que $\mu_j \leq \tau_{m^*}^2$ pour $j = m^* + 1, \dots, n$ et en utilisant l'équation (B.17), nous pouvons écrire

$$\partial E[W] > \frac{1}{\tau_{m^*}^2} \left(\sum_{i=1}^{m^*} \partial\alpha_i + \sum_{j=m^*+1}^n \partial\alpha_j \right) = 0.$$

Donc, la politique $\alpha^*(m^*)$ est optimale parmi les politiques ayant $\alpha_j^* > 0$ pour $j = m^*+1, \dots, n$. Notons que la politique $\alpha^*(m^*)$ est optimale parmi les politiques ayant $\alpha_j^* = 0$ pour $j = m^*+1, \dots, n$ par construction. La propriété 3.2 définit alors la solution optimale du problème Π_1 .

ANNEXE C

C. Lois de type phase et la méthode LZ

C.1. LOIS DE TYPE PHASE

La loi de probabilité du temps d'absorption dans une chaîne de Markov absorbante est dite de type phase. Dans la suite, nous exposons les définitions des lois de type phase à temps continu et à temps discret.

C.1.1. Lois de type phase à temps continu

Soit $\{X(t), t \geq 0\}$ une chaîne de Markov à temps continu ayant l'ensemble fini d'états $\{1, \dots, m+1\}$. Les états $1, \dots, m$ sont transitoires et l'état $m+1$ est absorbant. Le générateur S (la matrice des taux de transitions) d'une telle chaîne de Markov absorbante peut être décomposé comme suit :

$$S = \begin{bmatrix} G & \omega \\ \mathbf{0} & 0 \end{bmatrix}$$

où G est une matrice $m \times m$ et ω est un vecteur colonne de dimension m . $\mathbf{0}$ est le vecteur dont tous les éléments sont égaux à 0. La matrice G satisfait $G_{ii} < 0$ pour $1 \leq i \leq m$ et $G_{ij} \geq 0$ pour $i \neq j$. En outre, $G\mathbf{1} + \omega = \mathbf{0}$ où $\mathbf{1}$ est le vecteur colonne dont tous ces éléments sont égaux à 1. Le vecteur $\gamma = [\gamma_1 \dots \gamma_{m+1}]$ définit la distribution de probabilité initiale et satisfait $\gamma\mathbf{1} + \gamma_{m+1} = 1$. Soit $Z = \inf\{t \geq 0 : X(t) = m+1\}$ le temps d'absorption du processus de Markov dans l'état $m+1$. La variable aléatoire Z suit une loi de type phase à temps continu (*continuous phase type* ou CPH) d'ordre m avec les paramètres γ et G . On écrit $Z \sim \text{CPH}(\gamma, G)$. La fonction de distribution de probabilité de la variable aléatoire Z s'écrit

$$f_Z(t) = \gamma \exp(Gt) \omega$$

où $\exp(Gt) = e^{-Gt}$ est une fonction matrice exponentielle. La fonction de distribution cumulative et les moments de la variable aléatoire Z sont

$$F_Z(t) = 1 - \gamma \exp(Gt) \mathbf{1},$$

$$E[Z^n] = (-1)^n n! \gamma G^{-n} \mathbf{1}.$$

Les états transitoires $1, \dots, m$ de la chaîne de Markov correspondent aux phases du temps d'absorption. L'évolution du temps d'absorption est représentée par un graphe de services exponentiels qui est constitué de m phases où le temps de séjour dans la phase i suit une loi exponentielle de taux $-\lambda_{ii}$.

La famille de lois de type phase est constituée des distributions exponentielles et des distributions de toutes les sommes et les combinaisons finies de variables aléatoires exponentielles. Comme cas particuliers de cette famille de distributions, nous citons la distribution hypo-exponentielle, hyper-exponentielle, et Cox. La distribution hypo-exponentielle (ou Erlang généralisée) (Figure C.1) est la distribution de la somme de n variables aléatoires suivant des lois exponentielles. La distribution hypo-exponentielle peut être représentée par une loi de type phase ayant les paramètres

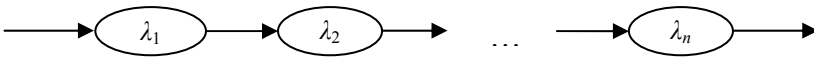
$$\mathbf{G} = \begin{bmatrix} -\lambda_1 & \lambda_1 & & & \\ & -\lambda_2 & \lambda_2 & & \\ & & \dots & & \\ & 0 & & -\lambda_{n-1} & \lambda_{n-1} \\ & & & & -\lambda_n \end{bmatrix}, \quad \boldsymbol{\omega} = \begin{bmatrix} 0 \\ 0 \\ \dots \\ 0 \\ \lambda_n \end{bmatrix}, \quad \boldsymbol{\gamma} = [1 \quad 0 \quad 0 \quad \dots \quad 0].$$


Figure C.1. Représentation de la distribution hypo-exponentielle

La distribution hyper-exponentielle (Figure C.2) est la distribution d'une combinaison finie de n variables aléatoires suivant des lois exponentielles. La distribution hyper-exponentielle peut être représentée par une loi de type phase ayant les paramètres

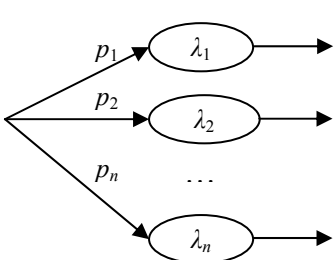
$$\mathbf{G} = \begin{bmatrix} -\lambda_1 & & & & \\ & -\lambda_2 & & & \\ & & \dots & & \\ & 0 & & -\lambda_{n-1} & \\ & & & & -\lambda_n \end{bmatrix}, \quad \boldsymbol{\omega} = \begin{bmatrix} \lambda_1 \\ \lambda_2 \\ \dots \\ \lambda_{n-1} \\ \lambda_n \end{bmatrix}, \quad \boldsymbol{\gamma} = [p_1 \quad p_2 \quad \dots \quad p_{n-1} \quad p_n].$$


Figure C.2. Représentation de la distribution hyper-exponentielle

Un exemple plus général est la distribution de Cox (Figure C.3) qui est une distribution d'Erlang généralisée dans laquelle, à chaque phase, il existe une possibilité de rentrer à l'état absorbant sans passer par les phases suivantes. Les paramètres définissant la distribution de type phase correspondante sont

$$\mathbf{G} = \begin{bmatrix} -\lambda_1 & p_1\lambda_1 & & & \\ & -\lambda_2 & p_2\lambda_2 & & 0 \\ & & \dots & & \\ 0 & & & -\lambda_{n-1} & p_{n-1}\lambda_{n-1} \\ & & & & -\lambda_n \end{bmatrix}, \quad \boldsymbol{\omega} = \begin{bmatrix} (1-p_1)\lambda_1 \\ (1-p_2)\lambda_2 \\ \dots \\ (1-p_{n-1})\lambda_{n-1} \\ \lambda_n \end{bmatrix}, \quad \boldsymbol{\gamma} = [1 \quad 0 \quad 0 \quad \dots \quad 0].$$

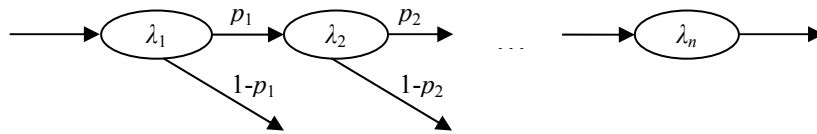


Figure C.3. Représentation de la distribution de Cox

C.1.2. Lois de type phase à temps discret

La définition de la loi de type phase à temps discret est similaire à la définition de CPH. Soit $\{X_n, n \geq 0\}$ une chaîne de Markov à temps discret ayant l'ensemble fini d'états $\{1, \dots, m+1\}$. Les états $1, \dots, m$ sont transitoires et l'état $m+1$ est absorbant. La matrice des probabilités de transition $\mathbf{T}$ de la chaîne de Markov $\{X_n, n \geq 0\}$ peut être décomposée comme suit :

$$\mathbf{T} = \begin{bmatrix} \mathbf{Q} & \boldsymbol{\eta} \\ \mathbf{0} & 1 \end{bmatrix}$$

où $\mathbf{Q}$ est une matrice $m \times m$ et $\boldsymbol{\eta}$ est un vecteur colonne de dimension m avec $\mathbf{Q}\mathbf{1} + \boldsymbol{\eta} = \mathbf{1}$. Le vecteur $\boldsymbol{\alpha} = [\alpha_1 \dots \alpha_{m+1}]$ définit la distribution de probabilité initiale et satisfait $\boldsymbol{\alpha}\mathbf{1} + \alpha_{m+1} = 1$. Soit $\theta = \min\{n \geq 0 : X_n = m+1\}$ le temps d'absorption du processus de Markov dans l'état $m+1$. La variable aléatoire θ suit une loi de type phase à temps discret (*discrete phase type* ou DPH) d'ordre m avec les paramètres $\boldsymbol{\alpha}$ et $\mathbf{Q}$. On écrit $\theta \sim \text{DPH}(\boldsymbol{\alpha}, \mathbf{Q})$. La fonction de distribution de probabilité, la fonction de distribution cumulative, et les moments de la variable aléatoires θ sont

$$f_\theta(k) = \boldsymbol{\alpha}\mathbf{Q}^{k-1}\boldsymbol{\eta},$$

$$F_\theta(k) = 1 - \boldsymbol{\alpha}\mathbf{Q}^k\mathbf{1},$$

$$E[\theta(\theta-1)\dots(\theta-n+1)] = n!\boldsymbol{\alpha}(\mathbf{I} - \mathbf{Q})^{-n}\mathbf{Q}^{n-1}\mathbf{1},$$

où $\mathbf{I}$ est la matrice d'identité.

La famille des lois de type phase à temps discret est constituée des distributions géométriques et des généralisations des distributions géométriques. Comme les lois de type phase à temps continu, elles ont des applications réelles et elles sont aussi utilisées pour approximer les processus non-markoviens.

C.2. MÉTHODE APPROXIMATIVE DE LEE ET ZIPKIN

Lee et Zipkin (1992) proposent une méthode approximative (appelée LZ) pour l'évaluation analytique des performances des systèmes de stock nominal à n étages avec $S_i \geq 0$ pour $i = 1, \dots, n$. La méthode approximative LZ se base sur l'hypothèse que l'arrivée des produits à chaque niveau $i > 1$ est un processus de Poisson ayant le taux λ . Selon cette hypothèse, le temps de séjour des produits dans le système de fabrication de niveau i , « W_i », suit une loi exponentielle de taux $\nu_i = \mu_i - \lambda$. Supposons que $L_i \sim \text{CPH}(\gamma_i, \mathbf{G}_i)$. Selon la propriété donnée par Sovorons et Zipkin (1991), le nombre de commandes dans le système de fabrication du niveau i , « K_i », est égal au nombre de demandes arrivant pendant le temps de séjour « L_i ». D'après les propriétés des lois de type phase (Neuts, 1994), K_i et $B_i = [K_i - S_i]^+$ suivent alors des lois de type phase à temps discret : $K_i \sim \text{DPH}(\alpha_i, \mathbf{Q}_i)$ et $B_i \sim \text{DPH}(\alpha_i \mathbf{Q}_i^{S_i}, \mathbf{Q}_i)$ avec les paramètres

$$\mathbf{Q}_i = \lambda(\lambda \mathbf{I} - \mathbf{G}_i)^{-1}, \quad (\text{C.1})$$

$$\alpha_i = \gamma_i \mathbf{Q}_i. \quad (\text{C.2})$$

Selon Sovorons et Zipkin (1991), le nombre de commandes retardées du niveau i , B_i , est égal au nombre de demandes arrivant pendant le délai de livraison D_i . En combinant cette propriété avec les propriétés des lois de type phase (Neuts, 1994), $D_i \sim \text{CPH}(\alpha_i \mathbf{Q}_i^{S_i}, \mathbf{G}_i)$. Selon la définition $L_{i+1} = D_i + W_{i+1}$, la fonction de densité de probabilité de L_{i+1} est un produit de convolution de deux lois de type phase. Nous pouvons écrire $f_{L_{i+1}}(t) = (f_{D_i} * f_{W_{i+1}})(t)$. Selon les propriétés des lois de type phase (Neuts, 1994), L_{i+1} suit alors une loi de type phase à temps continu : $L_{i+1} \sim \text{CPH}(\gamma_{i+1}, \mathbf{G}_{i+1})$ où

$$\gamma_{i+1} = [\gamma_i \mathbf{Q}_i^{S_i}, (1 - \gamma_i \mathbf{Q}_i^{S_i} \mathbf{1})], \quad (\text{C.3})$$

$$\mathbf{G}_{i+1} = \begin{bmatrix} \mathbf{G}_i & -\mathbf{G}_i \mathbf{1} \\ 0 & -\nu_{i+1} \end{bmatrix}. \quad (\text{C.4})$$

En développant, Lee et Zipkin réécrivent la matrice (C.4) comme suit :

$$\mathbf{G}_{i+1} = \begin{bmatrix} -v_1 & v_1 & & & \\ & -v_2 & v_2 & 0 & \\ & & \dots & & \\ & 0 & & -v_i & v_i \\ & & & & -v_{i+1} \end{bmatrix} \quad (\text{C.5})$$

Puisqu'il n'y'a pas de délai avant le premier serveur, L_1 suit la même loi de distribution que W_1 . En commençant par $\gamma_1 = [1]$, Lee et Zipkin calculent (C.3) récursivement pour $i = 1, \dots, n$.

En utilisant les propriétés des lois de type phase à temps discret, ils calculent les mesures de performances du système comme suit :

$$\Pr\{K_i > S_i\} = \alpha_i \mathbf{Q}_i^{S_i} \mathbf{1}$$

$$E[K_i] = \alpha_i (\mathbf{I} - \mathbf{Q}_i)^{-1} \mathbf{1}$$

$$E[B_i] = \alpha_i \mathbf{Q}_i^{S_i} (\mathbf{I} - \mathbf{Q}_i)^{-1} \mathbf{1}$$

$$E[I_i] = S_i - E[K_i] + E[B_i] = S_i - \alpha_i (\mathbf{I} - \mathbf{Q}_i)^{-1} \mathbf{1} + \alpha_i \mathbf{Q}_i^{S_i} (\mathbf{I} - \mathbf{Q}_i)^{-1} \mathbf{1}$$

La méthode approximative LZ est exacte quand $S_i = 0$ pour $i = 1, \dots, n$. Dans ce cas, $\gamma_n = [1 \ 0 \ 0 \dots 0]$ et le délai de livraison de niveau n est la somme de n variables aléatoires indépendantes suivant des lois exponentielles. Le délai de livraison de niveau n est alors le temps de séjour dans le réseau de files d'attente en tandem correspondant.

RÉFÉRENCES

- Anupindi, R., Bassok, Y. et E. Zemel, 2001. A general framework for the study of decentralized distribution systems. *Manufacturing and Service Operations Management*, vol. 3, no. 4, pp. 349-368.
- ARDA**, Y. et J.C. HENNET, 2007a. Inventory control in a decentralized two-stage make-to-stock queueing system. *International Journal of Systems Science*, Special Issue: Production Planning and Inventory Control, accepté pour publication en Juillet 2007.
- ARDA**, Y. et J.C. HENNET, 2006a. Inventory control in a multi-supplier system. *International Journal of Production Economics*, vol. 104, no. 2, pp. 249-259.
- ARDA**, Y. et J.C. HENNET, 2006b. Inventory control in a decentralized two-stage make-to-stock queueing System. 12th IFAC Symposium on Information Control Problems in Manufacturing, INCOM'06, 17-19 Mai 2006, Saint-Etienne (France).
- ARDA**, Y. et J.C. HENNET, 2006c. Coordination de chaînes logistiques, une approche par la théorie des jeux. 6^e Conférence Francophone de Modélisation et Simulation, MOSIM'06, 3-5 Avril 2006, Rabat (Maroc).
- ARDA**, Y., 2005. Impact des délais d'approvisionnement sur les performances de chaînes logistiques décentralisées. École Doctorale Systèmes, 6^e Congrès des Doctorants, 19-20 Mai 2005, Toulouse (France).
- ARDA**, Y. et J.C. HENNET, 2005. Supply chain coordination through contract negotiation. 44th IEEE Conference on Decision and Control, European Control Conference, ECC'05, 12-15 Décembre 2005, Seville (Spain).
- ARDA**, Y. et J.C. HENNET, 2004a. Optimizing the ordering policy in a supply chain. 11th IFAC Symposium on Information Control Problems in Manufacturing, INCOM'04, 5-7 Avril 2004, Salvador (Brazil).
- ARDA**, Y. et J.C. HENNET, 2004b. Inventory control in a multi-supplier system. 13th International Working Seminar on Production Economics, IGLS'04, 16-20 Février 2004, Igls/Innsbruck (Autriche), vol. 2, pp. 51-62.
- Axsater, S., 2003. Supply chain operations: Serial and distribution inventory systems. *Handbooks in Operations Research and Management Science, Volume 11: Supply Chain Management: Design, Coordination and Operation*, de Kok, A.G. et S.C. Graves, Ed., Elsevier.
- Axsater, S., 2000a. *Inventory Control*, Kluwer Academic Publishers.
- Axsater, S., 2000b. Exact analysis of continuous review (R, Q) policies in two-echelon inventory systems with compound Poisson demand. *Operations research*, vol. 48, no. 5, pp. 686-696.
- Axsater, S., 1993. Exact and approximate evaluation of batch-ordering policies for two-level inventory systems. *Operations Research*, vol. 41, no. 4, pp. 777-785.

- Axsater, S. et L. Juntti, 1996. Comparison of echelon stock and installation stock policies for two-level inventory systems. *International Journal of Production Economics*, vol. 45, pp. 303-310.
- Axsater, S. et K. Rosling, 1999. Ranking of generalised multi-stage KANBAN policies. *European Journal of Operational Research*, vol. 113, pp. 560-567.
- Axsater, S. et K. Rosling, 1993. Installation vs. echelon stock policies for multilevel inventory control. *Management Science*, vol. 49, no. 10, pp. 1274-1993.
- Babai, M.Z., 2005. Politiques de pilotage de flux dans les chaînes logistiques : impact de l'utilisation des prévisions sur la gestion des stocks. Thèse de doctorat, *Laboratoire Génie Industriel, École Central Paris*.
- Baskett, F., Chandy, K.M., Muntz R.R. et F.G. Palacios, 1975. Open, closed and mixed networks of queues with different classes of customers. *Journal of the Association for Computing Machinery*, vol. 22, no. 2, pp. 248-260.
- Baynat, B., Buzacott, J.A. et Y. Dallery, 2002. Multiproduct kanban-like control systems. *International Journal of Production Research*, vol. 40, no. 16, pp. 4225-4255.
- Baynat, B., Dallery, Y., Di Mascolo, M. et Y. Frein, 2001. A multi-class approximation technique for the analysis of kanban-like control systems. *International Journal of Production Research*, vol. 39, no. 2, pp. 307-328.
- Belvaux, G. et L.A. Wolsey, 2001. Modelling practical lot-sizing problems as mixed-integer programs. *Management Science*, vol. 47, no. 7, pp. 993-1007.
- Benjaafar, S., Mohsen, E.H. et F. Véricourt, 2004. Demand allocation in multiple-product multiple-facility make-to-stock systems. *Management Science*, vol. 50, no. 10, pp. 1431-1448.
- Berkley, B.J., 1992. A review of the Kanban production control research literature. *Production and Operations Management*, vol. 1, no. 4, pp. 393-411.
- Bernstein, F. et A. Federgruen, 2004. A general equilibrium model for industries with price and service competition. *Operations Research*, vol. 52, no. 6, pp. 868-886.
- Bernstein, F. et A. Federgruen, 2003. Pricing and replenishment strategies in a distribution system with competing retailers. *Operations Research*, vol. 51, no. 3, pp. 409-426.
- Billington, P.J., McClain, J.O. et L.J. Thomas, 1983. Mathematical programming approaches to capacity-constrained MRP systems: review, formulation and problem reduction. *Management Science*, vol. 29, no. 10, pp. 1126-1141.
- Brandenburger, A. et H. Stuart, 2007. Biform games. *Management Science*, vol. 55, no. 4, pp. 537-549.
- Buzacott, J.A., 1997. Continuous time distributed decentralized MRP. *Production Planning & Control*, vol. 8, no. 1, pp. 62-71.
- Buzacott, J.A., Price, S.M. et J.G. Shanthikumar, 1991. Service level in multistage MRP and base stock controlled production systems. *New Directions for Operations Research in a Manufacturing System*, Fandel, G., Gullledge, T. et A. Jones, Ed., Springer, pp. 445-463.

- Buzacott, J.A. et J.G. Shanthikumar, 1994. Safety stock versus safety time in MRP controlled production systems. *Management Science*, vol. 40, no. 12, pp. 1678-1689.
- Cachon, G.P., 2004. The allocation of inventory risk in a supply chain: push, pull, and advance-purchase discount contracts. *Management Science*, vol. 50, no. 2, pp. 222-238.
- Cachon, G.P., 2003. Supply chain coordination with contracts. *Handbooks in Operations Research and Management Science: Supply Chain Management: Design, Coordination and Operation*, de Kok, A.G. et S.C. Graves, Ed., Elsevier, pp. 229-340.
- Cachon, G.P., 2001a. Exact evaluation of batch-ordering inventory policies in two-echelon supply chains with periodic review. *Operations Research*, vol. 49, no. 1, pp. 79-98.
- Cachon, G.P., 2001b. Stock wars: inventory competition in a two-echelon supply chain with multiple retailers. *Operations Research*, vol. 49, no. 5, pp. 658-674.
- Cachon, G.P., 1999a. Competitive supply chain inventory management. *Quantitative Models for Supply Chain Management*, Tayur, S., Ganeshan, R. et M. Magazine, Ed., Kluwer Academic Publishers, pp. 111-146.
- Cachon, G.P., 1999b. Competitive and cooperative inventory management in a two-echelon supply chain with lost sales. Rapport technique, *The Wharton School, University of Pennsylvania* (opim.wharton.upenn.edu/~cachon).
- Cachon, G.P. et P.T. Harker, 2002. Competition and outsourcing with scale economies. *Management Science*, vol. 48, no. 10, pp. 1314-1333.
- Cachon, G.P. et M.A. Lariviere, 2005. Supply chain coordination with revenue-sharing contracts: strengths and limitations. *Management Science*, vol. 51, no. 1, pp. 30-44.
- Cachon, G.P. et M.A. Lariviere, 2001. Contracting to assure supply: how to share demand forecasts in a supply chain. *Management Science*, vol. 47, no. 5, pp. 629-646.
- Cachon, G.P. et M.A. Lariviere, 1999. Capacity choice and allocation: strategic behaviour and supply chain performance. *Management Science*, vol. 45, no. 8, pp. 1091-1108.
- Cachon, G.P. et S. Netessine, 2004. Game theory in supply chain analysis. *Handbook of Quantitative Supply Chain Analysis: Modeling in the eBusiness Era*, Smichi-Levi, D., Wu, S.D. et Z.J. Shen, Ed., Kluwer Academic Publishers, pp. 13-66.
- Cachon, G.P. et F. Zhang, 2006. Procuring fast delivery: sole sourcing with information asymmetry. *Management Science*, vol. 52, no. 6, pp. 881-896.
- Cachon, G.P. et P.H. Zipkin, 1999. Competitive and cooperative inventory policies in a two-stage supply chain. *Management Science*, vol. 45, no. 7, pp. 936-953.
- Caldentey, R. et L.M. Wein, 2003. Analysis of a decentralized production-inventory system. *Manufacturing and Service Operations Management*, vol. 5, no. 1, pp. 1-17.
- Cassandras, C.G., 1993. *Discrete Event Systems: Modeling and performance analysis*, Aksen Associates incorporated publishers.

- Chatain, O. et P. Zemsky, 2007. The horizontal scope of the firm: organizational tradeoffs vs. buyer-supplier relationships. *Management Science*, vol. 53, no. 4, pp. 550-565.
- Chaouiya, C. et Y. Dallery, 1997. Petri net models of pull control systems for assembly manufacturing systems. Proceedings of the Second International Workshop on Manufacturing and Petri Nets, International Conference on Application and Theory of Petri Nets, Di Cesare, F., M., Silva et R. Valette, ed., Toulouse, France : CNRS/LAAS, pp. 85-103.
- Chen, F. et Y.S. Zheng, 1997. One-warehouse multiretailer systems with centralized stock information. *Operations Research*, vol. 45, no. 2, pp. 275-287.
- Chen, F. et Y.S. Zheng, 1994a. Evaluating echelon stock (R, nQ) policies in serial production/inventory systems with stochastic demand. *Management Science*, vol. 40, no. 10, pp. 1262-1275.
- Chen, F. et Y.S. Zheng, 1994b. Lower bounds for multi-echelon stochastic inventory systems. *Management Science*, vol. 40, no. 11, pp. 1426-1443.
- Chen, F., 2000. Optimal policies for multi-echelon inventory problems with batch ordering. *Operations Research*, vol. 48, no. 3, pp. 376-389.
- Chen, F., Drezner, Z., Ryan, J.K. et D. Smichi-Levi, 2000. Quantifying the Bullwhip effect in a simple supply chain: The impact of forecasting, lead times, and information. *Management Science*, vol. 46, no. 3, pp. 436-443.
- Chopra, S. et P. Meindl, 2007. *Supply Chain Management: Strategy, Planning, and Operation*. 3^e édition, Pearson Prentice Hall.
- Clark, A.J. et H. Scarf, 1960 (2004). Optimal policies for a multi-echelon inventory problem. *Management Science*, vol. 6 (vol. 50), no. 4 (no. 12, Supplement), pp. 475-490 (pp. 1782-1790).
- Combé, M.B. et O.J. Boxma, 1994. Optimization of static traffic allocation policies. *Theoretical Computer Science A*, vol. 125, pp. 17-43.
- Corbett, C.J. et X. de Groote, 2000. A supplier's optimal discount policy under asymmetric information. *Management Science*, vol. 46, no. 3, pp. 444-450.
- Corbett, C.J. et C.S. Tang, 1999. Designing supply contracts: contract type and information asymmetry. *Quantitative Models for Supply Chain Management*, Tayur, S., Ganeshan, R. et M. Magazine, Ed., Kluwer Academic Publishers, pp. 269-298.
- Corbett, C.J., Zhou, D. et C.S. Tang, 2004. Designing supply contracts: contract type and information asymmetry. *Management Science*, vol. 50, no. 4, pp. 550-559.
- Dallery, Y. et G. Liberopoulos, 2000. Extended kanban control systems: combining kanban and base stock. *IIE Transactions*, vol. 32, pp. 369-386.
- Dallery, Y. et F. Yannick, 1993. On decomposition methods for tandem queueing networks with blocking. *Operations Research*, vol. 41, no. 2, pp. 386-399.
- Debo, L.G. et J. Sun. 2005. Repeatedly selling to the newsvendor in fluctuating markets. Rapport technique, Carnegie Mellon University, (<http://www.andrew.cmu.edu/user/jionsg/RepeatedNewsvendor.pdf>).

- Di Mascolo, M., Frein, Y. et Y. Dallery, 1996. An analytical method for performance evaluation of Kanban controlled production systems. *Operations Research*, vol. 44, no. 1, pp. 50-64.
- Di Mascolo, M., Frein, Y. Dallery, Y. et R. David, 1991. A unified modeling of kanban systems using Petri Nets. *International Journal of Flexible Manufacturing Systems*, vol. 3, no. 3/4, pp. 275-307.
- Dolgui, A. et M.A. Ould-Louly, 2002. A model for supply planning under lead time uncertainty. *International Journal of Production Economics*, vol. 78, pp. 145-152.
- Dong, L. et N. Rudi, 2004. Who benefits from transshipment? Endogenous vs. exogenous wholesale prices. *Management Science*, vol. 50, no. 5, pp. 645-657.
- Duri, C., Frein, Y. et M. Di Mascolo, 2000. Performance evaluation and design of base stock systems. *European Journal of Operational Research*, vol. 127, pp. 172-188.
- Federgruen A. et Y.S. Zheng, 1992. An efficient algorithm for computing an optimal (R, Q) policy in continuous review stochastic inventory systems. *Operations Research*, vol. 40, pp. 808-813.
- Federgruen A. et P. Zipkin, 1984. Computational issues in an infinite-horizon, multi-echelon inventory model. *Operations Research*, vol. 32, no. 4, pp. 818-836.
- Fudenberg, D. et J. Tirole, 1991. *Game Theory*, The MIT Press.
- Fong, D.K.H., Gempesaw, V.M. et J.K. Ord, 2000. Analysis of a dual sourcing inventory model with normal unit demand and Erlang mixture lead times. *European Journal of Operational Research*, vol. 120, pp. 97-107.
- Frein, Y., Di Mascolo, M. et Y. Dallery, 1995. On the design of generalized kanban control systems. *International Journal of Operations & Production Management*, vol. 15, no. 9, pp. 158-184.
- Gallego, G. et P. Zipkin, 1999. Stock positioning and performance estimation in serial production-transportation systems. *Manufacturing and Service Operations Management*, vol. 1, no. 1, pp. 77-88.
- Gan, X., Sethi, S.P. et H. Yan, 2004. Coordination of supply chains with risk-averse agents. *Production and Operations Management*, vol. 13, no. 2, pp. 135-149.
- Giard, V. (2003), *Gestion de la Production et des Flux*, 3^e édition, Economica.
- Guide, V.D.R. et R. Srivastava, 2000. A review of techniques for buffering against uncertainty with MRP systems. *Production Planning & Control*, vol. 11, no. 3, pp. 223-233.
- Guo, Y. et R. Ganeshan, 1995. Are more suppliers better? *Journal of Operational Research Society*, vol. 46, pp. 892-895.
- Gupta, D. et N. Selvaraju, 2006. Performance evaluation and stock allocation in capacitated serial supply systems. *Manufacturing and Service Operations Management*, vol. 8, no. 2, pp. 169-191.
- Gupta, D. et W. Weerawat, 2006. Supplier-manufacturer coordination in capacitated two-stage supply chains. *European Journal of Operational Research*, vol. 175, pp. 67-89.
- Gupta, D., Weerawat, W. et N. Selvaraju, 2004. Supply contracts for direct-to-market manufacturers. Rapport technique, University of Minnesota.

- Ha, A.Y., 1997a. Optimal dynamic scheduling policy for a make-to-stock production system. *Operations Research*, vol. 45, no. 1, pp. 42-53.
- Ha, A.Y., 1997b. Inventory rationing in a make-to-stock production system with several demand classes and lost sales. *Management Science*, vol. 43, no. 8, pp. 1093-1103.
- Ha, A.Y., 2000. Stock rationing in an $M/E_k/1$ make-to-stock queue. *Management Science*, vol. 46, no. 1, pp. 77-87.
- Hartman, B.C., et M. Dror, 2003. Optimizing centralized inventory operations in a cooperative game theory setting. *IIE Transactions*, vol. 35, pp. 243-257.
- Hartman, B.C., et M. Dror, 2005. Allocation of games from inventory centralization in newsvendor environments. *IIE Transactions*, vol. 37, pp. 93-107.
- Hartman, B.C., Dror, M. et M. Shaked, 2000. Cores of inventory centralization games. *Games and Economic Behavior*, vol. 31, pp. 26-49.
- HENNET, J.C. et Y. ARDA, 2007b. Supply chain coordination: a game theory approach. *Engineering Applications of Artificial Intelligence*, accepté pour publication en Août 2007.
- HENNET, J.C., 2003. A bimodal scheme for multi-stage production and inventory control. *Automatica*, vol. 39, pp. 793-805.
- Jemai, Z., 2003. Modèles stochastiques pour l'aide au pilotage des chaînes logistiques : L'impact de la décentralisation. Thèse de doctorat, Laboratoire Génie Industriel, École Central Paris.
- Jemai, Z. et F. Karaesmen, 2007. Decentralized inventory control in a two-stage capacitated supply chain. *IIE Transactions*, vol. 39, pp. 501-512.
- Karaesmen, F. et Y. Dallery, 2000. A performance comparison of pull type control mechanisms for multi-stage manufacturing. *International Journal of Production Economics*, vol. 68, pp. 59-71.
- Karaesmen, F., Buzacott J.A. et Y. Dallery, 2002. Integrating advance order information in make-to-stock production systems. *IIE Transactions*, vol. 34, pp. 649-662.
- Kelle, P. et P.A. Miller, 2001. Stockout risk and order splitting. *International Journal of Production Economics*, vol. 71, pp. 407-415.
- Khouja, M., 1999. The single-period (news-vendor) problem: literature review and suggestions for future research. *Omega, International Journal of Management Science*, vol. 27, pp. 537-553.
- Kleinrock, L., 1975. *Queueing systems Volume 1: Theory*, John Wiley & Sons.
- Koh, S.C.L., Saad, S.M. et M.H. Jones, 2002. Uncertainty under MRP-planned manufacture: Review and categorization. *International Journal of Production Research*, vol. 40, no. 10, pp. 2399-2421.
- Lariviere, M.A., 1999. Supply chain contracting and coordination with stochastic demand. *Quantitative Models for Supply Chain Management*, Tayur, S., Ganeshan, R. et M. Magazine, Ed., Kluwer Academic Publishers, pp. 235-268.
- Lariviere, M.A. et E.L. Porteus, 2001. Selling to the newsvendor: an analysis of price-only contracts. *Manufacturing and Service Operations Management*, vol. 3, no. 4, pp. 293-305.

- Lee, H. et S. Whang, 1999. Decentralized multi-echelon supply chains: incentives and information. *Management Science*, vol. 45, no. 5, pp. 633-640.
- Lee, H.L., Padmanabhan, V. et S. Whang, 1997. Information Distortion in a Supply Chain: The Bullwhip effect. *Management Science*, vol. 43, no. 4, pp. 546-558.
- Lee, H.L., So, K.S. et C.S. Tang, 2000. The value of information sharing in a two-level supply chain. *Management Science*, vol. 46, no. 5, pp. 627-643.
- Lee, Y.J. et P. Zipkin, 1995. Processing networks with inventories: sequential refinement systems. *Operations Research*, vol. 43, no. 6, pp. 1025-1036.
- Lee, Y.J. et P. Zipkin, 1992. Tandem queues with planned inventories. *Operations Research*, vol. 40, no. 5, pp. 936-947.
- Leng, M. et M. Parlar, 2005. Game theoretic applications in supply chain management: a review. *Information Systems and Operations Research*, vol. 43, no. 3, pp. 187-230.
- Liberopoulos, G. et S. Koukoulialos, 2005. Tradeoffs between base stock levels, number of kanbans, and planned supply lead times in production/inventory systems with advanced demand information. *International Journal of Production Economics*, vol. 96, pp. 213-232.
- Liberopoulos, G. et Y. Dallery, 2002. Base stock versus WIP cap in single-stage make-to-stock production-inventory systems. *IIE Transactions*, vol. 34, pp. 627-636.
- Liberopoulos, G. et Y. Dallery, 2003. Comparative modelling of multi-stage production-inventory control policies with lot sizing. *International Journal of Production Research*, vol. 41, no. 6, pp. 1273-1298.
- Lippman, S.A. et K.F. McCardle (1997). The competitive newsboy. *Operations Research*, vol. 45, no. 1, pp. 54-65.
- Liu, L., Liu, X. et D. Yao, 2004. Analysis and optimization of a multistage inventory-queue system. *Management Science*, vol. 50, no. 3, pp. 365-380.
- Liu, Z. et R. Righter, 1998. Optimal load balancing on distributed homogenous unreliable processors. *Operations Research*, vol. 46, no. 4, pp. 563-573.
- Mahajan, S. et G. van Ryzin, 2001. Inventory competition under dynamic consumer choice. *Operations Research*, vol. 49, no. 5, pp. 649-657.
- Matta, A., Dallery, Y. et M. Di Mascolo, 2006. Analysis of assembly systems controlled with kanbans. *European Journal of Operational Research*, vol. 166, pp. 310-336.
- Minner, S., 2003. Multi-supplier inventory models in supply chain management: a review. *International Journal of Production Economics*, vol. 81-82, pp. 265-279.
- Mohebbi, E. et M.J.M. Posner, 1998. Sole versus dual sourcing in a continuous-review inventory system with lost sales. *Computers in Industrial Engineering*, vol. 34, pp. 321-336.
- Moore, K.E. et S.M. Gupta, 1999. Stochastic coloured Petri net (SCPN) models of traditional and flexible kanban systems. *International Journal of Production Research*, vol. 37, no. 9, pp. 2135-2158.

- Muller, A., Scarsini, M., et M. Shaked, 2002. The newsvendor game has a nonempty core. *Games and Economic Behavior*, no. 38, pp. 118-126.
- Myerson, R.B., 1991. *Game Theory: Analysis of Conflict*, Harvard University Press.
- Netessine, S. Rudi, N. et Y. Wang, 2006. Inventory competition and incentives to back-order. *IIE Transactions*, vol. 38, pp.883-902.
- Neuts, M.F., 1994. *Matrix-Geometric Solutions in Stochastic Models*, Dover Publications.
- Osborne, M.J. et A. Rubinstein, 1994. *A Course in Game Theory*, The MIT Press.
- Ould-Louly, M.A. et A. Dolgui, 2004. The MPS parameterization under lead time uncertainty. *International Journal of Production Economics*, vol. 90, pp. 369-376.
- Ozdamar L. et G. Barbarosoglu, 2000. An integrated Lagrangean relaxation-simulated annealing approach to the multi-level multi-item capacitated lot sizing problem. *International Journal Production Economics*, vol. 68, pp. 319-331.
- Papadopoulos, H.T. et C. Heavey, 1996. Queueing theory in manufacturing systems analysis and design: a classification of models for production and transfer lines. *European Journal of Operational Research*, vol. 92, pp. 1-27.
- Plambeck, E. et T.A. Taylor, 2005. Sell the plant ? The impact of contract manufacturing on innovation, capacity, and profitability. *Management Science*, vol. 51, no. 1, pp. 133-150.
- Plambeck, E. et S. Zenios, 2003. Incentive efficient control of a make-to-stock production system. *Operations Research*, vol. 51, no. 3, pp. 371-386.
- Porteus, E.L., 2000. Responsibility tokens in supply chain management. *Manufacturing and Service Operations Management*, vol. 2, no. 2, pp. 203-219.
- Ramasesh, R.V., Ord, J.K., Hayya, J.C. et A. Pan, 1991. Sole versus dual sourcing in stochastic lead-time (s, Q) inventory models. *Management Science*, vol. 37, no. 4, pp. 428-443.
- Rao, U.S., 2003. Properties of the periodic review (R, T) inventory control policy for stationary, stochastic demand. *Manufacturing and Service Operations Management*, vol. 5, no. 1, pp. 37-53.
- Rao, S.S, Gunasekaran, A., Goyal, S.K. et T. Martikainen, 1998. Waiting line model applications in manufacturing. *International Journal of Production Economics*, vol. 54, pp. 1-28.
- Ross, S.M., 2000. *Introduction to probability models*, Seventh Edition, USA: Harcourt Academic Press.
- Sedarage D., Fujiwara, O. et H.T. Luong, 1999. Determining optimal order splitting and reorder level for N -supplier inventory systems. *European Journal of Operational Research*, vol. 116, pp. 389-404.
- Shang, K.H. et J.S. Song, 2003. Newsvendor bounds and heuristic for optimal policies in serial supply chains. *Management Science*, Vol. 49, no. 5, pp. 618-638.
- Sherbrooke, C.C., 1968. METRIC: A multi-echelon technique for recoverable item control. *Operations Research*, vol.16, pp. 122-141.
- Smichi-Levi, D., Kaminsky, P. et E. Smichi-Levi, 2003. *Designing and Managing the Supply Chain: Concepts, Strategies and Case Studies*. 2^e édition, McGraw-Hill.

- Song, J.S., 1994. The effect of leadtime uncertainty in a simple stochastic inventory model. *Management Science*, vol. 40, no.5, pp. 603-613.
- Sovorons, A. et P. Zipkin, 1991. Evaluation of one-for-one replenishment policies for multi-echelon inventory systems. *Management Science*, vol. 37, no. 1, pp. 68-83.
- Sovorons, A. et P. Zipkin, 1988. Estimating the performance of multi-level inventory systems. *Operations Research*, vol. 36, no. 1, pp. 57-72.
- Spearman, M.L. et M.A. Zazanis, 1992. Push and pull production systems: Issues and comparisons. *Operations Research*, vol. 40, no.3, pp. 521-532;
- Stadtler, H. 2005. Supply Chain Management – an overview. *Supply Chain Management and Advanced Planning, Concepts: Models, Software and Case Studies*, Stadtler, H. et C. Kilger, Ed., 3^e édition, Springer, pp. 7-28.
- Thomas, J.D. et J. Tyworth, 2006. Pooling lead-time risk by order splitting: a critical review. *Transportation Research Part E*, vol. 42, pp. 245-257.
- Tsay, A.A., Nahmias, S. et N. Agrawal, 1999. Modeling supply chain contracts: a review. *Quantitative Models for Supply Chain Management*, Tayur, S., Ganeshan, R. et M. Magazine, Ed., Kluwer Academic Publishers, pp. 301-336.
- Van Houtum, G.J., 2006. Multi-echelon production/inventory systems: Optimal policies, heuristics, and algorithms. *Tutorials in Operations Research: Models, Methods, and Applications for Innovative Decision Making*, Johnson, M.P., Norman, B. et N. Secomandi, Ed., INFORMS.
- Van Houtum, G.J., Inderfurth, K. et W.H.M. Zijm, 1996. Materials coordination in stochastic multi-echelon systems. *European Journal of Operational Research*, vol. 95, pp. 1-23.
- Van Mieghem, J.A., 1999. Coordinating investment, production, and subcontracting. Price versus production postponement: capacity and competition. *Management Science*, vol. 45, no. 7, pp. 954-971.
- Van Mieghem, J.A. et M. Dada, 1999. Price versus production postponement: capacity and competition. *Management Science*, vol. 45, no. 12, pp. 1631-1649.
- Veatch, M.H. et L.M. Wein, 1994. Optimal control of a two-station tandem production/inventory system. *Operations Research*, vol. 42, no. 2, pp. 337-350.
- Veatch, M.H. et L.M. Wein, 1996. Scheduling a make-to-stock queue: index policies and hedging points. *Operations Research*, vol. 44, no. 4, pp. 634-647.
- de Vericourt, F., Karaesmen, F. et Y. Dallery, 2002. Optimal stock allocation for a capacitated supply system. *Operations Research*, vol. 48, no. 11, pp. 1486-1501.
- de Vericourt, F., Karaesmen, F. et Y. Dallery, 2000. Dynamic scheduling in a make-to-stock system: a partial characterization of optimal policies. *Operations Research*, vol. 48, no. 5, pp. 811-819.
- Vollman, T.E., Berry, W.L. et D.C. Whybark, 1997. *Manufacturing Planning and Control Systems*, 4^e édition, Irwin.

- Wang, F.F. et C.T. Su, 2007. Performance evaluation of a multi-echelon production, transportation and distribution system: a matrix-analytical approach. *European Journal of Operational Research*, vol. 176, pp. 1066-1083.
- Wang, Y. et Y. Gerchak, 2003. Capacity games in assembly systems with uncertain demand. *Manufacturing and Service Operations Management*, vol. 5, no. 3, pp. 252-267.
- Wong, H., van Oudheusden, D. et D. Cattrysse, 2007. Cost allocation in spare parts inventory pooling. *Transportation Research Part E*, vol. 43, pp. 370-386.
- Zheng, Y.S. et A. Federgruen, 1991. Finding optimal (s, S) policies is about as simple as evaluating a single policy. *Operations Research*, vol. 39, no. 4, pp. 654-664.
- Zheng, Y.S. et P. Zipkin, 1990. A queueing model to analyze the value of centralized inventory information. *Operations Research*, vol. 38, no. 2, pp. 296-307.
- Zipkin, P.H., 2000. *Foundations of Inventory Management*, McGraw-Hill.

TITRE EN ANGLAIS : Replenishment policies in multi-supplier systems and Decisions Optimization in decentralized supply chains

RÉSUMÉ EN ANGLAIS : Coordinating product flows between the partners of a supply chain is a difficult task because of random variations in demand and supply processes and the antagonistic nature of the individual economic objectives of the partners. This study concentrates on the management of inter-organizational product flows in supply chains. Two approaches are analyzed with the aim of improving performances of production/inventory systems controlled by base stock type product flow control policies. In the first approach, the effects of multi-supplier strategies on the performances of supply chains are studied. It is shown that a multi-supplier strategy decreases the expected replenishment delay and the expected inventory holding and shortage costs. The second approach deals with the deviations from the set of supply chain optimal actions due to the decentralisation of decision rights in a two-stage supply chain. In the game theory framework, the partners play a two-stage game of the Stackelberg type. A coordination contract is proposed and it is shown that the optimal supply chain performance is achievable using the proposed contract.

MOTS-CLÉS EN ANGLAIS : Supply Chain Management, Inventory Control, Stochastic Models, Queueing Theory, Game Theory.

AUTEUR : Yasemin ARDA

TITRE : Politiques d'approvisionnement dans les systèmes à plusieurs fournisseurs et Optimisation des décisions dans les chaînes logistiques décentralisées

DIRECTEUR DE THÈSE : Jean-Claude HENNET

DATE ET LIEU DE SOUTENANCE : 14 Janvier 2008 au LAAS-CNRS

RÉSUMÉ : La coordination des flux physiques au sein des chaînes logistiques est une tâche difficile à cause du caractère aléatoire des variations dues au marché et aux partenaires commerciaux et des antagonismes existants entre les objectifs économiques des partenaires. Les travaux développés dans cette thèse s'intègrent dans le cadre de pilotage de flux inter-organisationnelle dans les chaînes logistiques. Nous analysons deux approches ayant le but d'améliorer les performances des systèmes de production/stockage pilotés par des politiques de pilotage flux du type stock nominal. Dans la première approche, nous étudions les effets des stratégies multi-fournisseurs sur les performances des chaînes logistiques. Nous montrons que le délai moyen d'approvisionnement et les coûts moyens de stockage et de rupture de stock peuvent être réduits en optant pour une stratégie multi-fournisseurs. Dans la deuxième approche, nous analysons la dégradation de performances due à la décentralisation des décisions dans une chaîne logistique à deux niveaux en définissant un jeu de Stackelberg entre les partenaires. Nous proposons un contrat de coordination et montrons que le contrat proposé ramène les performances du système décentralisé vers les performances optimales du système centralisé.

MOTS-CLÉS : Gestion de Chaînes Logistiques, Gestion des Stocks, Modèles Stochastiques, Théorie des Files d'Attente, Théorie des Jeux.

DISCIPLINE ADMINISTRATIVE : Systèmes Industriels

INTITULÉ ET ADRESS DU LABORATOIRE : LAAS-CNRS
7, avenue du Colonel Roche
31077, Toulouse, France