

N° d'ordre : 2496

Année 1998

THÈSE

présentée à

l'UNIVERSITE DES SCIENCES ET TECHNOLOGIES DE LILLE

pour l'obtention du titre de

DOCTEUR

En Productique : Automatique et Informatique Industrielle

par

François BOUSSU



**SIMULATION DE LA FILIERE TEXTILE/HABILLEMENT/DISTRIBUTION :
REDUCTION DE LA COMPLEXITE EN VUE D'UNE MEILLEURE
PREVISION DES VENTES**

Soutenue le 12 janvier 1998 devant la commission d'examen composée de :

Messieurs	Christian VASSEUR	Président et Directeur de Recherche
	Joël FAVREL	Rapporteur
	Gérard GOVAERT	Rapporteur
	Michel HAPPIETTE	Codirecteur
	Jean-Marie CASTELAIN	Examineur
	Pascal YIM	Examineur
	Jean-Jacques DENIMAL	Examineur

**à ma femme,
à mes enfants,
à mes parents.**

REMERCIEMENTS

Le travail présenté dans ce mémoire a été réalisé au sein du laboratoire GEMTEX (GÉNIE et Matériaux TEXtiles) de l'ENSAIT (Ecole Nationale Supérieure des Arts et Industries Textiles).

Je tiens à remercier :

Monsieur le Professeur Christian VASSEUR qui m'a fait l'honneur de présider le jury et d'avoir accepté de diriger les travaux de recherche,

Monsieur Joël FAVREL, Professeur à l'INSA de Lyon, et monsieur Gérard GOVAERT, Professeur à l'UTC, pour avoir accepté d'être les rapporteurs de ces travaux,

Monsieur Jean-Marie CASTELAIN, Professeur et directeur de L'ENSAIT, monsieur Pascal YIM, Maître de conférences à l'Ecole Centrale de Lille, ainsi que monsieur Jean Jacques DENIMAL, Maître de conférences à l'université de Lille, pour avoir accepté de faire partie du jury;

Monsieur Michel HAPPIETTE, Maître de conférences à l'ENSAIT et co-directeur de ce mémoire, pour son soutien sans faille et sans limite,

Monsieur Besoa RABENASOLO, Maître de Conférences à l'ENSAIT, pour sa patience, sa disponibilité, et surtout pour sa contribution majeure à la réalisation de ces travaux,

Monsieur Xianyi ZENG, Maître de Conférences à l'ENSAIT, pour ses précieux conseils,

Sans oublier l'ensemble des membres du GEMTEX et de l'ENSAIT pour la disponibilité et la patience qu'ils ont su m'accorder.

RESUME

L'application d'une démarche logistique globale ou "Quick Response" au sein des entreprises de la filière Textile/Habillement/Distribution (THD) repose essentiellement sur l'échange d'informations entre partenaires, traduites sous forme de messages EDI (Echange de Données Informatisées).

Cependant, différentes contraintes de traitement de l'information freinent son utilisation au sein des systèmes de gestion des partenaires de l'industrie textile telles que :

- l'absence d'outils de modélisation des flux et de simulation des stratégies d'échanges commerciaux entre différents partenaires,
- la contrainte d'utilisation et de choix d'un modèle de prévision,
- la complexité du problème d'identification.

Notre contribution consiste à proposer différents modèles et méthodes pour répondre aux différents environnements. Un cadre de modélisation des flux intra et inter entreprises est présenté afin de simuler différentes stratégies d'approvisionnements. Les exemples de simulation mettent en évidence la capacité du simulateur à traiter une information permettant de fournir une solution envisageable tout en tenant compte de l'ensemble des contraintes initiales.

Une identification des méthodes et modèles de prévision adaptés à l'environnement de vente des articles textiles est également proposée. L'application de six méthodes de prévision, et leurs évaluations par des mesures différentes de l'erreur sur des données de vente réelles, a permis de mettre en valeur les capacités d'adaptation et de précision des méthodes de lissage utilisant une procédure d'autorégulation de leurs propres paramètres.

Enfin, la réduction du nombre de données à traiter tout en minimisant la perte d'information est abordée. Les méthodologies de classification proposées constituent des méthodes d'analyse des données de vente des articles textiles et fournissent l'essentiel de l'information pour l'identification d'un modèle de prévision adapté. L'utilisation d'un algorithme génétique de classification, dont la capacité réside à explorer l'ensemble des solutions, a permis d'atteindre la répartition optimale globale.

Mots clés

Modélisation, Simulation, Classification, Prévision, Algorithme génétique, Textile.

ABSTRACT

The use of the "Quick Response" logistic method for the Textile/Apparel/Distribution (TAD) Channel firms is mainly based on the data exchange between partners thanks to the EDI language.

However, different constraints, due to the information treatment, occur and limit its application between the partners of the TAD Channel. These constraints are :

- the lack of modelling and simulation tools of logistic flows between the different actors to check their trading available strategy,
- the choice and the use of the well adapted forecasting method,
- the huge number of information to treat (complexity of the time series identification).

Our approach tends to submit different methods and models with respect to the previous problems. A modelling scheme of the different flows inside and between firms helps to simulate several possible supply strategies. The simulation results put on the fore the simulator advantages as regard to the information treatment and the ability to find (or not) one of the possible solution with respect to all the initial constraints and conditions.

A literature review of the forecasting methods and models adapted to the short term horizon is also proposed. A comparison made on six well known forecasting methods, evaluated with several measures based on the forecasting error, shows that the more precise procedures are the methods which are able to integrate a self-adaptation function of their own parameters.

And, the complexity problem can be solved by the use of data analysis methods which help to reduce the global information by minimisation of its lost. Two clustering methods and a symbolic description of items are proposed to give the minimum and essential information for model identification. The use of a genetic clustering algorithm helps also to find the optimal solution.

Key words

Modelling, Simulation, Clustering, Forecasting, Genetic algorithm, Textile.

TABLE DES MATIERES

INTRODUCTION GÉNÉRALE.....	11
I. Evolution de la prise en compte de la satisfaction du client au sein des méthodes de production	11
II. Adoption d'une démarche logistique intégrée ou globale	13
III. La démarche "Quick Response"	13
IV. Contribution de notre approche	16
 CHAPITRE 1	 19
 MODÉLISATION ET SIMULATION DE LA FILIÈRE	
TEXTILE/HABILLEMENT/DISTRIBUTION.....	19
I. Introduction	19
II. Modélisation d'une entreprise et de sa filière	22
II.1. Représentation globale d'une entreprise.....	23
II.2. Les fonctions fondamentales	24
II.3. Représentation globale de la filière THD.....	26
II.4. Modélisation et simulation par Hypernets	27
II.4.1. Choix de l'outil (Tableau 1).	27
II.4.2. Représentation graphique	30
II.4.3. Déclarations	31
II.4.4. Modèle d'une entreprise.....	32
II.4.5. Modélisation de la filière	38
(1) Principe	38
(2) Options de stratégies d'approvisionnement.....	39
III. Simulation	41
III.1. Simulation d'une entreprise.....	41
III.1.1. Date de livraison au plus tôt.....	41
III.1.2. Date de production au plus tard.....	45
III.1.3. Date au plus tôt et au plus tard de livraison d'une commande.....	47
III.2. Simulation de la filière Textile/Habillement/Distribution.....	48
III.2.1. Choix d'un partenaire fournisseur avec critère de sélection.....	48
III.2.2. Complément de commandes	51
III.3. Les capacités du simulateur	52

III.3.1. Informations spatiales	52
III.3.2. Informations temporelles.....	53
III.3.3. Algorithmes d'optimisation	53
IV. Conclusion.....	54

CHAPITRE 2 55

MODÈLES ET MÉTHODES DE PRÉVISION DES VENTES..... 55

I. Introduction	55
II. Présentation des méthodes de prévision usuelles.....	57
II.1. Notations	58
II.2. Méthodes intuitives.....	58
II.3. Méthodes endogènes	59
II.3.1. Principe de décomposition d'une série chronologique.	59
(1) Notion de tendance dans le comportement de vente des articles textiles.....	60
(2) Notion de "saisonnalité" dans le comportement de vente des articles textiles.	60
II.3.2. La méthode de prévision par la moyenne mobile.....	61
(1) La prévision par la moyenne mobile simple	61
(2) La prévision par la moyenne mobile double.....	61
II.3.3. Les méthodes de prévision par lissage exponentiel	62
(1) Classification des méthodes de prévision par lissage exponentiel.....	62
(2) La prévision par lissage exponentiel simple.....	64
(3) La prévision par lissage exponentiel double	64
(4) Choix du coefficient de lissage α :	64
(5) Méthode de prévision utilisant les notions de tendance et de «saisonnalité»	66
(a) Modèle additif	66
(b) Modèle multiplicatif	67
(6) Procédure d'auto-régulation de la méthode de Holt-Winters [VRO 96][VRO 97].....	68
II.3.4. Méthode de prévision basée sur un modèle issu d'une analyse de série temporelle.	69
II.4. Méthodes exogènes	70
II.5. Méthodes non linéaires	71
III. Critères d'un modèle de prévision.....	71
III.1. Choix d'une méthode	72
III.2. Evaluation d'une méthode	73
IV. Comparaison entre les différentes méthodes de prévision et validation des modèles pour un horizon de prévision h fixé.....	76
V. Méthodes de prévision spécifiques appliquées à l'industrie textile.....	79
V.1. Contexte de la prévision	79

V.2. Présentation des différentes méthodes textiles existantes	81
V.2.1. Horizon à court terme	82
(1) <i>Influence de la mode</i>	82
(a) Analyse de Garcia-Cuevas et al. [GAR 81]	82
(b) Méthode MODEFA	82
(2) <i>Ventes d'articles textiles soumis à de fortes variations</i>	83
V.2.2. Horizon à moyen terme	83
(1) <i>Regroupement des coefficients saisonniers</i>	83
(2) <i>Introduction des variables marketing textile</i>	84
V.2.3. Horizon à long terme	86
(1) <i>Estimation des quantités nécessaires à la mise en place de collection textile</i>	86
(2) <i>Optimisation de l'achat de la matière première textile</i>	88
V.3. Analyse comparative des méthodes de prévision textiles pour un horizon de prévision	
h=1.....	89
VI. Conclusion.....	96

CHAPITRE 3 98

L'APPORT DE LA CLASSIFICATION POUR LA PRÉVISION. 98

I. Introduction	98
II. Proposition d'une approche de classification.....	102
II.1. Proposition des méthodes de classification	103
II.1.1. Classification ascendante hiérarchique	103
II.1.2. Classification par partition floue.....	103
II.2. Critères d'agréations	104
II.3. Mesure de ressemblance entre les articles.....	105
II.3.1. distance entre des données de type quantitatif.....	106
II.3.2. distance entre des données de type qualitatif.....	106
II.4. Présentation des algorithmes retenus.....	106
II.4.1. Classification ascendante hiérarchique par la méthode des voisins réciproques	106
II.4.2. Algorithme de classification des "c" classes moyennes floues.....	108
II.5. Critère de validité utilisé.....	110
II.5.1. Pour les partitions floues	111
II.5.2. Pour les partitions non floues	112
II.5.3. Formulation d'une contrainte de compacité	112
(1) <i>locale additive</i>	113
(2) <i>globale additive</i>	113
(3) <i>locale multiplicative</i>	114
(4) <i>globale multiplicative</i>	114

II.6. Description symbolique des articles textiles.....	115
II.6.1. Définitions des objets symboliques.....	115
II.6.2. Notations utilisées.....	117
II.6.3. Caractérisation des centres de classe	119
(1) pour la classification ascendante hiérarchique.....	119
(2) pour la classification par partition floue	120
II.6.4. Mesure de ressemblance entre les articles	121
(1) Pour les données de type quantitatif.....	121
(2) Pour les données de type qualitatif.....	121
(a) pour la classification ascendante hiérarchique	122
(b) pour la classification par partition floue	123
III. Applications textiles.....	123
III.1. Application de la méthode de classification ascendante hiérarchique	125
III.1.1. Application aux allures de vente.....	125
III.1.2. Application aux variables qualitatives	133
III.2. Application de la méthode de classification par partition floue.....	134
III.2.1. Application aux allures de ventes	134
III.3. Comparaison des résultats obtenus.....	139
IV. Mesure de l'influence des paramètres exogènes sur les données de vente ..	140
IV.1. Procédure de transformation des variables exogènes.	141
IV.2. Procédure de corrélation entre variables exogènes et allures de vente.....	141
IV.3. Application au paramètre exogène température	142
IV.4. Avantages et inconvénients de la méthode.....	146
V. Conclusion.....	147

CHAPITRE 4 149

AMÉLIORATION DE L'APPROCHE DE CLASSIFICATION PAR DES PROCÉDURES GÉNÉTIQUES..... 149

I. Principe de la méthodologie proposée.....	149
I.1. Principaux codages utilisés.....	150
I.1.1. Par affectation [JON 91].....	150
I.1.2. Par permutation [JON 91].....	151
I.1.3. Par permutation avec heuristique de décodage [JON 91].....	151
I.1.4. Comparaison	152
I.2. Opérateurs de recombinaison	152
I.3. Opérateurs de mutation	155
I.4. Algorithme génétique de classification	158

I.5. Choix des valeurs du couple (P_c et P_m).....	159
II. Applications des stratégies adoptées	160
II.1. Procédures sur l'opérateur de mutation	162
II.1.1. sans heuristique	162
II.1.2. avec heuristiques.....	165
(1) heuristique 1	165
(2) heuristique 2	166
(3) heuristique 3	169
II.1.3. Récapitulatif des simulations effectuées sur l'algorithme de classification	172
II.2. Application aux données réelles de vente textile.....	174
III. Conclusion.....	176
CONCLUSION GÉNÉRALE	177
ANNEXE 1	180
ANNEXE 2	186
ANNEXE 3	189
BIBLIOGRAPHIE.....	191

Introduction générale

I. Evolution de la prise en compte de la satisfaction du client au sein des méthodes de production

Depuis la révolution industrielle de la fin du 19^{ème} siècle, les différentes approches de gestion des ressources de production au sein des entreprises ont évolué. De nos jours, la prise en compte de l'information client est un des facteurs principaux de l'adaptation de la production et de la distribution aux différents environnements économiques [SIN 96]. L'application au début du 20^{ème} siècle du "management" scientifique au service de la production de masse, selon le modèle de Taylor [TAY 27], ignorait complètement les exigences des clients. Puis, au début des années 70, l'utilisation du contrôle qualité pour les opérations de production a permis d'intégrer les remarques des clients en vue d'améliorer la qualité des produits finis. Au début des années 80, la philosophie d'amélioration de la qualité totale (Total Quality Management), incluse dans le modèle japonais "Lean Production" [WOM 92], traduite dans la pratique par l'application du "juste à temps" [SRI 94], positionne clairement l'exigence client comme la source principale d'informations. Une seconde approche d'origine américaine, le "Quick Response" [DEF 94], du début des années 80, permet également de répondre rapidement à la demande du client. Au cours des années 70 et 80, les premières utilisations des technologies de l'information comme outils de réponse au système de gestion ont rendu des résultats peu significatifs pour beaucoup d'industriels, particulièrement en terme de retour sur investissement et de niveau de productivité [THU 92]. En effet, la non communication de l'information traitée a conduit à une mauvaise utilisation de ces technologies dans l'élaboration des stratégies

de vente [HAC 90] et plus particulièrement dans les techniques de gestion des ressources humaines et d'allocations des moyens [LOV 94]. Enfin, l'émergence au début des années 90 de divers modèles tels que : "Mass Customization" [PIN 93], "Business Process Redesign" [CRO 94], "Consumer Response" [KSA 94], "Agile Virtual Enterprise" [GOL 95], "Holonc Enterprise" [HUG 95], pose les bases d'un nouveau concept de production, sans fournir pour l'instant de résultats concrets d'applications. Le principal objectif est d'être adaptable au changement rapide dans un contexte de compétition à l'échelle mondiale tout en respectant les contraintes imposées par la demande client. L'ensemble des caractéristiques de ces méthodes sont englobées dans les principes d'application d'une démarche logistique intégrée.

L'évolution croissante de la prise en compte des exigences des consommateurs en terme de prix, de délai de livraison et de qualité des articles textiles, oblige les sociétés de distribution et de production à bouleverser leurs organisations internes et à avoir recours aux nouvelles technologies de communication et d'information telles que l'Echange de Données Informatisées (EDI) [ROB 91] ou l'outil réseau Internet [QUE 96]. De ce fait, une évolution de la communication entre les sociétés de distribution et de production a également été initiée [DAU 94]. L'adoption des technologies de communication et d'information au sein des systèmes de gestion doit répondre aux différentes contraintes de production tout en permettant de créer des gains de productivité et des profits financiers.

En d'autres termes, la gestion des flux qui doit être assurée par les producteurs à tous les niveaux de la chaîne de distribution, oblige à une réorganisation et à une flexibilité de la production. Dans cette optique d'optimisation des flux, la mise en place d'une démarche logistique globale permet ainsi d'édifier un cadre structurel d'organisation des différentes connexions inter et intra-entreprises reliant les différents points de décision. Cette démarche, appliquée dans le cadre d'un partenariat [DER 94], permet d'aboutir à un avantage compétitif pour l'ensemble des entreprises de la chaîne de distribution, en minimisant les coûts de stockage, en optimisant les ventes, en répondant exactement à la demande des clients en terme d'articles et de quantités nécessaires, et en évitant les pertes financières dues au solde des articles invendus.

II. Adoption d'une démarche logistique intégrée ou globale

Pour le secteur de l'industrie du textile et de l'habillement français, une source d'inspiration possible peut provenir des techniques de production des sociétés du Sentier¹ "qui pratiquaient le flux tendu avant même l'invention du concept" [BRU 94]. L'environnement économique actuel et les effets induits par une compétition intense entre les distributeurs obligent les producteurs à s'orienter dans la voie de l'artisanat industriel [BOU 94]. Cette orientation peut être traduite par l'application d'une démarche logistique globale au sein de l'entreprise. En effet, la fonction logistique dans l'entreprise est d'assurer au moindre coût la coordination de l'offre et de la demande tout en maintenant un niveau de qualité [MAT 87]. La démarche logistique intégrée, par delà les outils de gestion qu'elle fournit à la direction générale, est porteuse d'innovations organisationnelles et de réactivité. Ces innovations restent à être explorées par le dirigeant pour adapter son entreprise à un environnement économique de plus en plus dépourvu de repères [PON 93]. Dans un contexte de mondialisation des marchés, l'adoption d'une démarche "Quick Response", définie comme cas particulier d'une démarche logistique intégrée, permet d'apporter une "réponse rapide aux consommateurs" en s'appuyant sur une communication de l'information entre partenaires. L'utilisation de ce concept n'est pas encore réellement présente au sein des entreprises textiles européennes principalement à cause du bouleversement culturel que celui-ci impose.

III. La démarche "Quick Response"

Son objectif est d'optimiser les flux matières, "informationnels" et "décisionnels" entre les différents partenaires de la filière Textile/Habillement/Distribution. La mise en place d'une démarche "Quick Response" peut s'opérer en plusieurs étapes suivant un ordre croissant proportionnel au degré de confiance entre partenaires. L'un des objectifs principaux consiste à éliminer les points de stockages intermédiaires [ROB 94] et de combiner une production flexible avec une grande variété de collections d'articles en

¹ Il s'agit d'un quartier de Paris regroupant un ensemble de détaillants d'articles textiles organisés de façon à répondre rapidement à des délais de livraison très courts.

appliquant des méthodes "justes à temps" aux différentes étapes de production à travers la chaîne de distribution [CON 85]. La mise en œuvre d'autres messages EDI, tel que le document EDI de planification "Planning Schedule Document", offre au fabricant la possibilité de communiquer ses besoins à ses fournisseurs suivant un horizon de planification convenu entre partenaires. Ce message établit une simple relation entre un partenaire distributeur et un partenaire producteur. Mais, dans le cadre d'une optimisation des critères inclus dans les différentes stratégies d'approvisionnement entre un distributeur et plusieurs de ses fournisseurs, un outil de simulation et de modélisation des différents liens entre partenaires permet d'analyser et de traiter efficacement chacun des messages de planification [BOU 97a][BOU 97b]. Enfin, l'amélioration des systèmes internes de prévision permet de disposer d'une planification conjointe et de développer un modèle de gestion des stocks de type "juste à temps" entre les différents partenaires [SHI 86]. Chaque acteur peut ainsi se rapprocher de l'acte de vente, contribuer à la "mise en équation" complète et dynamique de la filière et à l'élaboration d'un système commun de prévision. Le choix des modèles de prévision, adéquats à s'appliquer dans un environnement économique mondial et dans le cadre d'une démarche "Quick Response", constitue une prise de décision importante pour l'industriel. Ainsi, un état des méthodes de prévision existantes, abordé dans le chapitre concernant les différentes méthodologies de prévision, apporte un ensemble de recommandations pour effectuer le choix d'un modèle approprié. Un ensemble d'applications de modèles de prévision dans l'industrie textile apporte un éclairage supplémentaire sur les différentes contraintes économiques et spécifiques d'utilisation des procédures prévisionnelles.

Cependant, l'application du concept "Quick Response" au sein des industries textiles américaines fait apparaître des dysfonctionnements, révélés par une étude réalisée en collaboration avec des industriels et des distributeurs textiles [HUN 90][HUN 95]. Le retard de la mise en œuvre de la démarche a pu être quantifié par des résultats d'analyse fournis par des cabinets de consultants [KSA 88][AAC 89], et par des simulations d'après des modèles stochastiques de gestion des ressources de production mesurant différents critères de performance tels que : le retour sur investissement, le taux de service [NUT 91][HUN 92][HUN 96]. Une première explication provient d'au moins trois des caractéristiques originales propres à la filière textile/habillement :

- *Les problèmes liés à la gestion d'une immense quantité de produits textiles différents proposés aux consommateurs* : Un distributeur peut proposer entre un et un million et demi d'articles différents trois fois par an. Cet ordre de grandeur est bien plus important que dans l'alimentaire ou dans n'importe quelle autre industrie. De ce fait, la maîtrise du développement conjoint de nouveaux produits dans le contexte d'une intégration virtuelle sollicite les systèmes de gestion des stocks et de prévision du partenariat. L'enjeu est en effet d'extraire l'information stratégique d'immenses bases de données constituées au jour le jour. Disposer d'une capacité de traitement et d'analyse de telles bases de données constitue donc à moyen terme un véritable avantage compétitif. Quelques expériences sont en cours, notamment le développement récent de la méthode RADAR [HUT 94]. La réduction du nombre des informations à traiter peut également être résolue par une approche de regroupement des allures de vente des articles textiles [HAP 96a] suivant des objectifs de classification précis.
- *L'effet envahissant de la mode* : La durée de vie moyenne du vêtement est d'autant plus faible qu'un plus grand nombre d'autres produits subit l'influence de la mode [RHO 97]. De ce fait, la constitution d'historiques de ventes sur des périodes suffisantes, afin d'effectuer une identification du modèle de prévision adéquat, est compromise pour la plus grande majorité des articles. L'attribution d'un modèle de prévision pour les articles nouveaux d'une collection peut être effectué par le biais d'un transfert des connaissances d'une saison à l'autre. Cette procédure de transfert fait l'objet d'une approche abordée dans le chapitre concernant la méthodologie et les critères de classification.
- *La structure même de la filière* : Contrastant fortement avec les secteurs de l'automobile ou de l'électronique, par exemple, dans lesquels les fabricants du produit fini occupent une position dominante face aux sous-traitants et fournisseurs de pièces détachées et aux distributeurs, l'industrie de l'habillement, composée de petites entreprises, a peu d'influence sur son amont c'est à dire le textile, et sur son aval c'est à dire le commerce de détail.

Ainsi, les industriels de l'habillement qui réussiront l'intégration des nouvelles technologies de communication et d'information seront ceux qui utiliseront les données des points de vente pour anticiper la demande saisonnière, estimer les modifications des

projections d'achat, et maintenir les stocks au minimum, tout en tenant leurs fournisseurs textiles informés de leurs propres besoins. L'identification des besoins propres reposent sur une connaissance globale de l'environnement économique et des différents comportements de vente des articles textiles. Parmi l'ensemble des méthodes disponibles d'analyse, une approche de classification, en fonction des différents critères d'agrégation définissant l'objectif à atteindre, suivie d'une procédure d'évaluation des résultats de « partitionnements » obtenus, permet d'obtenir des résultats d'analyse spécifiques au problème posé. La formulation différente des critères d'agrégation permet, par exemple :

- d'établir une base de connaissances sur les caractéristiques propres des articles, ou endogènes, influant sur les allures de vente des produits, par l'utilisation d'un critère d'agrégation basé sur une représentation symbolique des objets [HAP 97],
- de distinguer les influences, plus ou moins prononcées, de phénomènes extérieurs aux caractéristiques propres des articles agissant sur l'environnement économique, traduites dans une procédure de classification des influences des paramètres exogènes.

Un besoin de logiciels élaborés pour répondre à différents objectifs d'analyse de données et de profondes modifications dans la gestion des entreprises d'habillement s'avère donc nécessaire, ainsi qu'une orientation vers des techniques de production flexibles et rapides. Dans cette optique, le producteur peut réellement ajouter une valeur à son produit en répondant plus efficacement à la demande du consommateur.

IV. Contribution de notre approche

Dans l'état actuel d'une démarche "Quick Response", différentes contraintes de traitement de l'information freinent son utilisation au sein des systèmes de gestion des partenaires de l'industrie textile :

- l'absence d'outils de modélisation des flux et de simulation des stratégies d'échanges commerciaux entre différents partenaires empêche toute évaluation objective et comparative selon des critères de performances,
- la contrainte d'utilisation et de choix d'un modèle de prévision,
- la taille du problème d'identification due au nombre important des données à traiter.

Dans le chapitre 1, une première approche de notre travail consiste à fournir un modèle complet de la filière afin d'analyser et de simuler différentes stratégies de planification prévisionnelle entre partenaires pour optimiser un critère relatif à un objectif précis. Cet outil de simulation, couplé au système de gestion interne de l'entreprise, repose sur une modélisation simple des principales fonctions de l'entreprise. La standardisation de ce modèle permet à chaque acteur de la filière de l'adapter à ses propres contraintes. L'utilisation de cet outil de simulation doit permettre, entre autre, de véhiculer des messages EDI de planification au sein de la filière et d'apporter une aide à la décision sur les politiques de réservation des ressources temporelles, humaines, financières et matières. Une amélioration de la précision des stratégies prévisionnelles repose sur l'identification et l'utilisation du modèle de prévision le plus adapté aux contraintes économiques ou segments de marché spécifiques.

C'est pourquoi dans le chapitre 2, après une synthèse bibliographique générale sur les méthodes de prévision les plus couramment utilisées par les industriels, une seconde approche va consister à les identifier en fonction des différents objectifs prévisionnels de l'industrie textile. En effet, suivant l'environnement économique et le type de partenariat établi entre client et fournisseur, différentes méthodologies d'optimisation des ressources financières et matérielles sont disponibles pour chacun des acteurs de la filière textile. De même, le problème de choix du modèle adapté à la situation économique est abordé sous l'angle de l'exploitation de comparaisons des résultats de simulations entre différents modèles et suivant un panorama de critères d'évaluation des méthodes de prévision. Quelques exemples de méthodologies prévisionnelles appliquées à différents environnements textiles apportent un éclairage complémentaire sur les contraintes et conditions spécifiques d'utilisation d'un modèle pour la vente d'articles textiles. Une application personnelle, à partir de 6 méthodes de prévisions évaluées par 5 mesures de l'erreur différentes sur des données de vente d'un confectionneur en relation avec un distributeur dans le cadre d'une démarche « Quick Response », permet d'identifier les méthodes de prévision, parmi les plus couramment utilisées, les mieux adaptées et les plus précises à ce contexte économique spécifique.

Dans le chapitre 3, une troisième approche aborde le problème lié à la réduction de la dimension de la procédure d'identification des modèles de prévision sur les séries temporelles des ventes d'articles textiles. Son objectif est d'extraire parmi l'ensemble des

historiques de vente des articles, la quantité nécessaire et suffisante d'informations à partir de méthodes de classification, afin de l'intégrer au modèle de prévision adéquat. Le caractère global ou local des informations extraites s'apprécie en fonction du degré de finesse souhaité pour l'analyse des séries temporelles. Ceci est rendu possible par l'utilisation de différents types de critères d'agrégation. Une étude sur les critères de validité des méthodes de classification est également abordée afin de choisir le nombre juste suffisant d'informations utiles. Deux applications de méthodologie, traduisant différents objectifs de classification en fonction de la formulation différente des critères d'agrégation, sont également présentées pour fournir des outils d'aide à l'analyse des données. Les résultats obtenus peuvent être enrichis par une mesure de l'influence des paramètres exogènes sur les données de vente.

Enfin chapitre 4, une amélioration des valeurs initiales des paramètres de notre approche de classification est rendue possible par l'utilisation de procédures génétiques. Cette méthode basée sur la recherche d'une valeur minimale d'une fonction objective, traduisant l'objectif d'agrégation, repose sur l'utilisation d'une représentation par codage des différentes partitions envisageables. Chacune des partitions disponibles, regroupées au sein d'une population, subit des procédures de transformation sur les valeurs de codage, dont le caractère peut être assimilé à des procédures de transformations génétiques, afin de faire évoluer l'ensemble des partitions existantes. L'intérêt de l'utilisation des méthodes de classification par des procédures de transformations génétiques réside dans l'exploration des différentes solutions de classifications possibles et l'identification rapide de l'optimum global relatif à la formulation de la fonction objective choisie.

Chapitre 1

Modélisation et simulation de la filière Textile/Habillement/Distribution

I. Introduction

L'application de la démarche "Quick Response" entre différents partenaires de la filière Textile/Habillement/Distribution (THD) nécessite une gestion fiable, rapide et rigoureuse de l'information. L'utilisation du langage EDI permet de véhiculer rapidement un message compréhensible pour tous les partenaires. Ainsi, l'information disponible, contenue dans les messages EDI, doit être traitée et interprétée par chaque entreprise de la filière pour être ensuite intégrée dans leur propre système de gestion (Figure 1). La mise à jour des informations permet d'établir une comparaison entre l'état réel et prévisionnel de l'entreprise. Elle est suivie d'une évaluation des performances logistiques traduite par différents critères tels que : le taux de service, la qualité de l'information, le taux de flexibilité, le taux de réactivité, le taux de productivité, etc.. L'écart constaté conduit à mettre en place des actions d'amélioration du système de gestion de l'entreprise et à formuler de nouveaux objectifs prévisionnels. L'utilisation d'un simulateur standard pour chaque entreprise de la filière permet de tester différentes stratégies logistiques entre partenaires afin d'optimiser les critères relatifs aux objectifs fixés. Cet outil peut être connecté au système de gestion des données de l'entreprise et exploiter l'ensemble des

capacités d'échange de l'information en interne, via les réseaux locaux, et en externe, via les supports réseaux de type Internet (cf. Figure 2). De plus, l'automatisation du traitement de l'information, reçue par les différents partenaires en fonction des critères à optimiser, peut être rendue possible et permet au décideur de l'entreprise de limiter son intervention aux sources de conflit plus complexes. L'approche de simulations proposées repose sur une modélisation simplifiée d'une entreprise dont l'un des objectifs est d'identifier les principaux points de traitement de l'information et leurs différentes relations. Les spécificités de modélisation [LEN 93] pour l'élaboration d'un modèle impliquent :

- l'établissement d'une liaison entre les modèles abstraits statiques et les modèles de simulation permettant une démarche de modélisation cohérente,
- l'augmentation de la puissance de modélisation et de diffusion des modèles,
- la prise en compte, de façon distincte, des aspects physiques, "informationnels" et "décisionnels",
- l'augmentation de la capacité de communication des modèles pour faciliter leur utilisation.

Face à toutes ces contraintes, une première représentation globale d'une entreprise, issue d'une démarche analytique, permet de mettre en place à un niveau hiérarchique supérieur les différentes fonctions fondamentales de l'entreprise ainsi que les principaux flux.

Le principe de la modélisation modulaire est de décomposer les systèmes en éléments plus simples afin de surmonter les problèmes liés à l'établissement et au traitement de modèles de grande dimension. Une représentation modulaire de l'entreprise, quelque soit le type de procédé de fabrication, permet d'aboutir à une architecture en réseau de la filière textile par connexion les uns aux autres des modules de représentation de l'entreprise. La spécificité de l'architecture en réseau de modules identiques permet d'atteindre un seuil de simplicité et de description de la compréhension de la modélisation.

Pour répondre à ces différentes contraintes, une identification des différentes sources principales de traitement de l'information est abordée dans la première partie de ce chapitre. Dans le cadre d'une application d'une démarche "Quick Response", la représentation globale d'une entreprise et de la filière THD permet de modéliser les liens

entre ces différentes fonctions. Le problème du choix d'un outil de modélisation et de simulation parmi l'ensemble des nombreuses applications existantes est abordé dans la deuxième partie. Enfin, des applications de simulations relatives à différentes stratégies sont exposées dans le troisième chapitre afin de montrer la possibilité d'une telle approche.

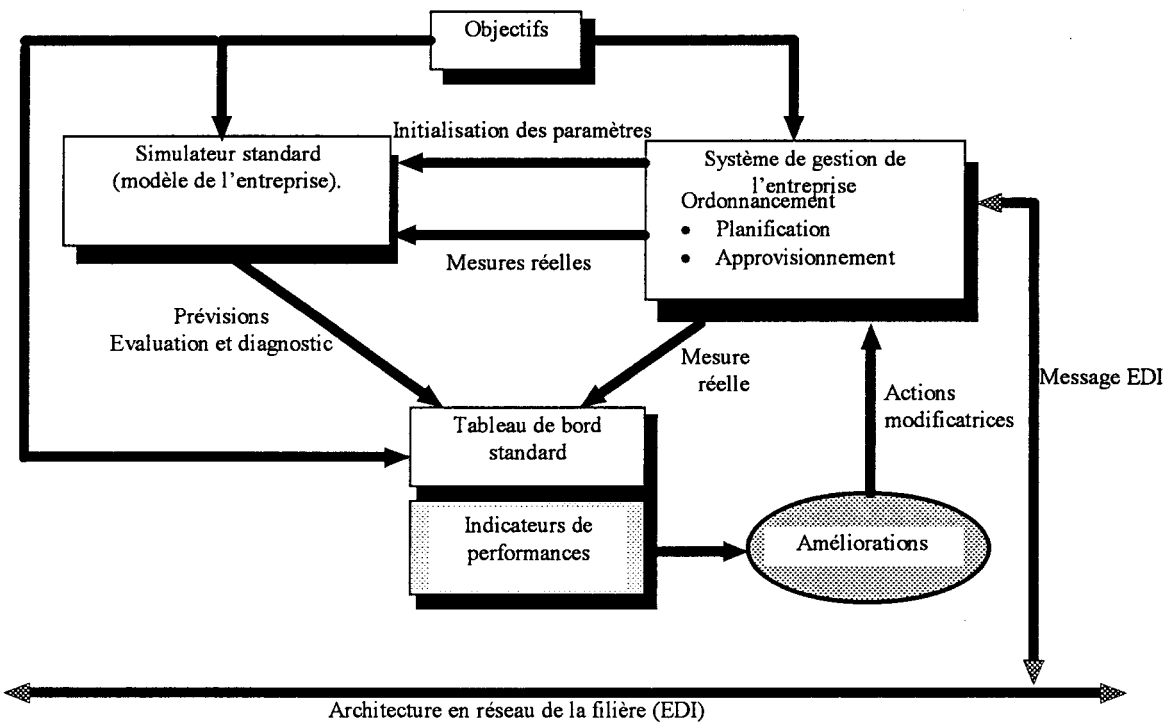


Figure 1. **Architecture générale d'une stratégie d'amélioration des performances**

Position du simulateur au sein des entreprises de la filière

La station EDITH se compose d'un ordinateur de type PC équipé d'un logiciel de codage des informations et de décodage des messages EDI selon la norme internationale des Nations Unies. Sa fonction principale consiste à traduire les informations venant de la base de données sous forme de message EDI, et réciproquement.

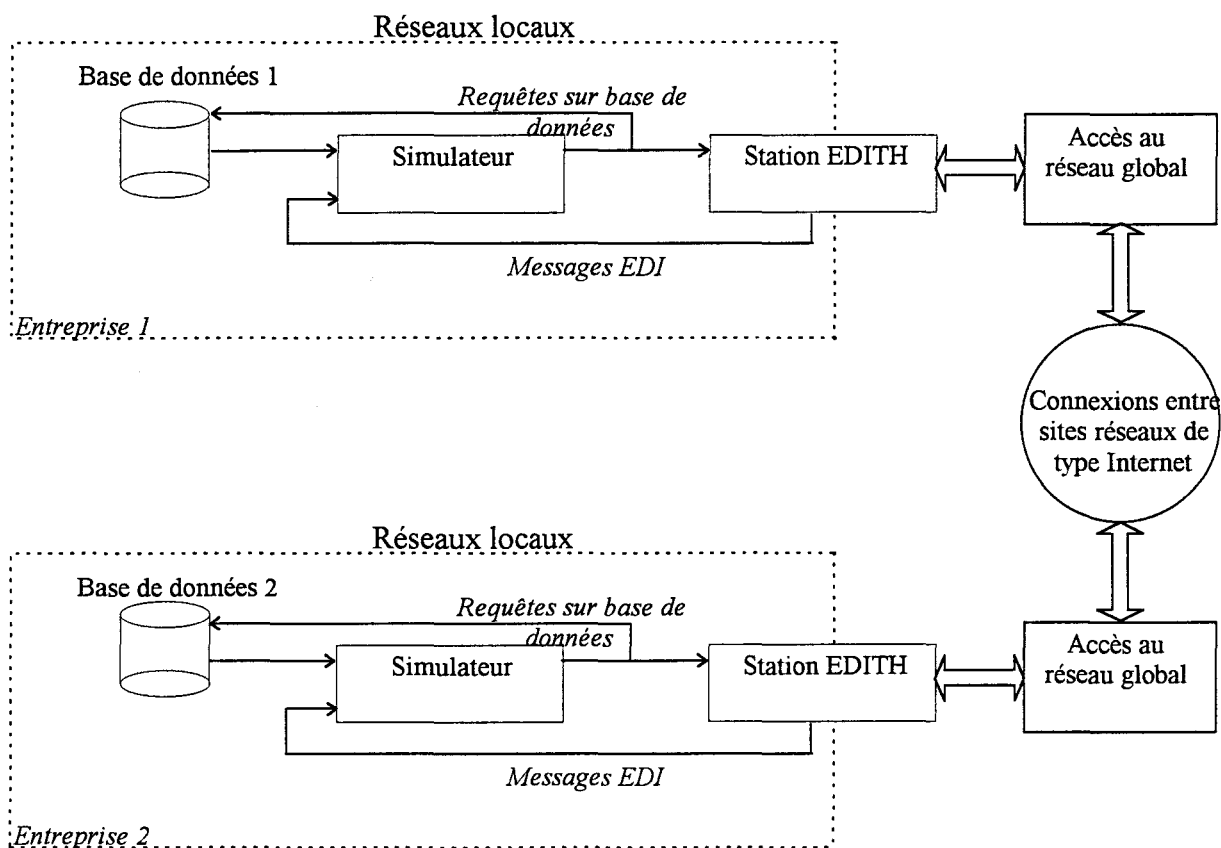


Figure 2. Position du simulateur au sein de la filière.

II. Modélisation d'une entreprise et de sa filière

Pour comprendre et donc pour donner du sens à un système complexe, une modélisation doit être effectuée pour construire son intelligibilité (compréhension) [MOI 88]. La proposition de la représentation globale de l'entreprise permet essentiellement de comprendre l'évolution du traitement de l'information aux différents points principaux de fonctionnement. Une représentation simplifiée, du type analyse systémique, permet, dans

un premier temps, de regrouper un certain nombre de fonctions de l'entreprise sous trois grandes fonctions majeures connectées chacune par deux grands types de flux. De plus, l'adaptation du modèle aux contraintes de chacun des industriels de la filière THD permet de standardiser la représentation globale de l'entreprise.

II.1. Représentation globale d'une entreprise

Dans cette optique de simplification et de standardisation, la représentation globale (Figure 3) de l'activité d'une entreprise repose sur les trois fonctions fondamentales :

OGF : Organisation de la Gestion des Flux

P : Production

V : Vente

qui agissent directement sur les deux types de flux :

F.I. : Flux d'Informations

F.M. : Flux de Matière

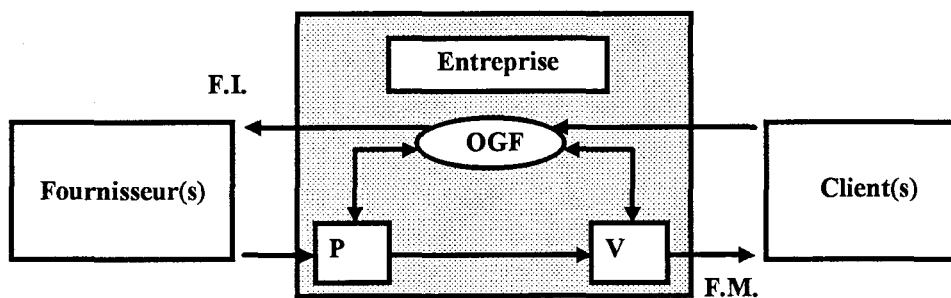


Figure 3. Représentation globale d'une entreprise

II.2. Les fonctions fondamentales

La fonction OGF.

L'Organisation de la Gestion des Flux, permet de collecter l'ensemble des informations provenant de la fonction V (Vente) de l'entreprise et des clients situés en aval, de traiter les informations vers la fonction P (Production) de l'entreprise et vers les fournisseurs situés en amont. Elle est la fonction principale qui englobe l'ensemble des paramètres internes et externes de l'entreprise pour optimiser la gestion des flux informations (F. I.) et matières (F. M.) suivant des critères définis par la stratégie commerciale de l'entreprise.

La fonction P.

La Production permet de transformer un ensemble de composants, venant du fournisseur amont ou disponibles en stock au sein de l'entreprise, en un ensemble de produits finis gérés ensuite par la fonction Vente de l'entreprise. Pour un producteur, la fonction P comporte le stock de composants, les en-cours de fabrication, le stock de produits finis et l'ensemble des flux logistiques assurant leur cohésion au plus juste. Pour un distributeur, la fonction P comprend l'ensemble des flux physiques logistiques nécessaires pour assurer correctement l'acheminement des produits finis.

La fonction V.

La Vente connecte le flux matière issu de l'entreprise vers le client aval (distributeur ou producteur) et génère un flux d'informations nécessaires à l'organisation de la gestion des flux physiques. La fonction V englobe l'ensemble des opérations et informations nécessaires au conditionnement et au transfert des produits finis vers le client.

Le Flux d'Information F.I.

Il caractérise l'ensemble des données nécessaires à la bonne gestion des flux des matières à l'intérieur et à l'extérieur des entreprises de la filière. La tendance actuelle s'oriente vers une utilisation de plus en plus importante d'un Echange de Données

Informatisées par secteur d'activité (EDITEX² pour le textile). L'utilisation d'un format commun de l'information permet d'assurer la compréhension du message délivré par l'ensemble des différents utilisateurs et d'éviter les opérations de traduction supplémentaires.

Le Flux des Matières F.M.

Il caractérise l'ensemble des bons (qualité) produits nécessaires (quantité) au bon moment (délai) au bon prix (coût) pour satisfaire la demande du client aval et répondant à l'ensemble des informations auxquelles ils sont associés.

A chaque acteur de la filière THD peut être associé une représentation globale de l'entreprise. Ainsi, par extension, il suffit d'associer à chacun des acteurs un module de représentation et de les connecter les uns aux autres afin de configurer une architecture en réseau de la filière (Figure 4). D'autres représentations de la filière THD existent [DOR 91]; certaines représentations n'exploitent que le flux des matières pour connecter les entreprises les unes aux autres, d'autres prennent trop en compte les spécificités des processus de fabrication de chaque acteur de la filière ce qui complique la compréhension de l'architecture en réseau. L'adoption d'une même représentation globale de l'entreprise par chacun des acteurs justifie l'intérêt d'une architecture réseau de la filière THD. En effet, si chacun des acteurs possède son propre simulateur couplé à son système de gestion interne, l'évaluation des différentes stratégies peut s'appliquer d'un bout à l'autre de la chaîne de distribution sans aucune opération de traduction ou de traitement supplémentaire de l'information.

² EDITEX : Association pour la promotion de l'EDI (Echange de Données Informatisées) dans la filière Textile/Habillement/Distribution

II.3. Représentation globale de la filière THD

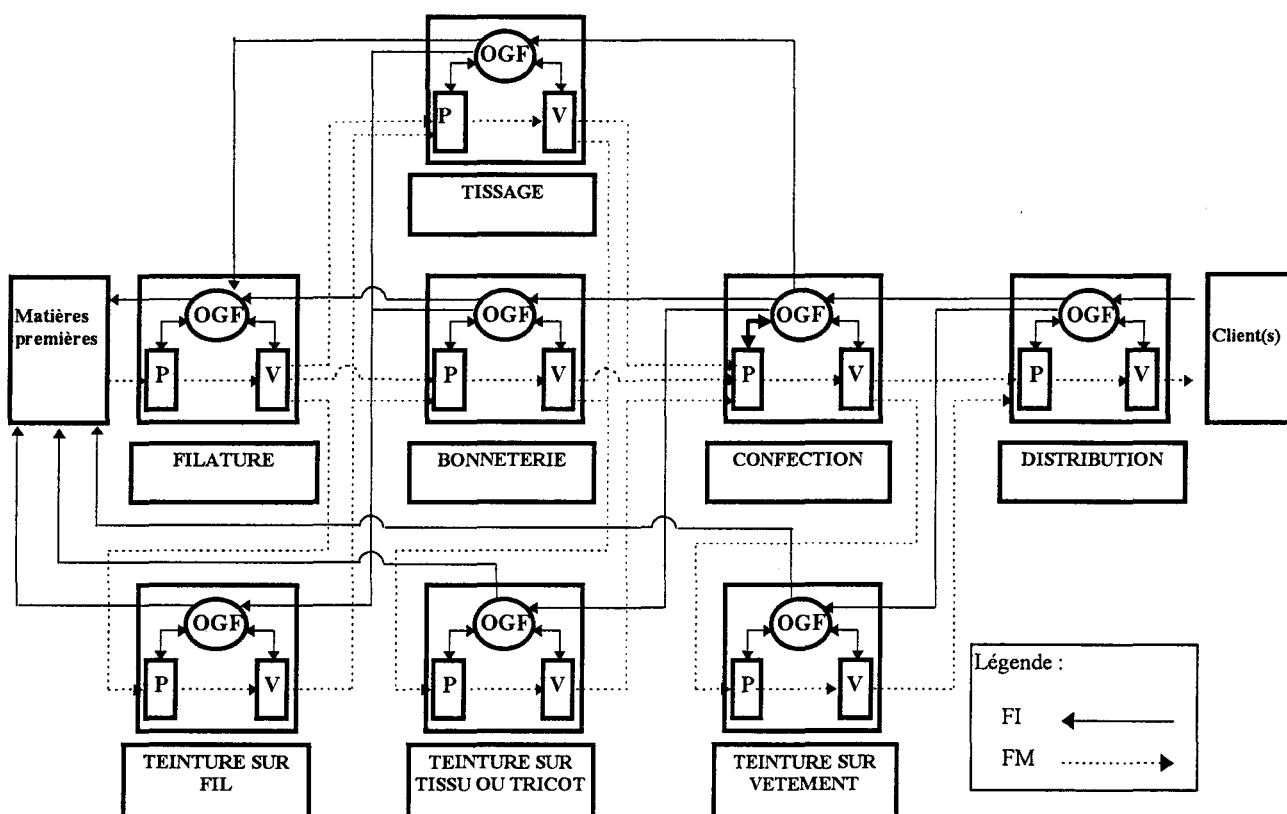


Figure 4. Représentation globale de la filière
Textile/Habillement/Distribution

Le schéma fait apparaître six classes d'acteurs qui participent à la transformation et/ou à l'obtention d'une valeur ajoutée du produit : les distributeurs, les confectionneurs, les tisserands, les tricoteurs, les filateurs, et les teinturiers sur fil, tissu, tricot, ou vêtement.

II.4. Modélisation et simulation par Hypernets

II.4.1. Choix de l'outil (Tableau 1).

De façon générale, les capacités et les possibilités de tout outil de modélisation et de simulation peuvent se résumer à des spécificités globales [LEN 95].

D'une part, la méthode de modélisation doit inclure :

- des méthodes et des réalisations de références élaborées pour les éléments composant le système de production et les relations entre ces éléments,
- un code de modélisation (langage) permettant de représenter de façon formelle le système,
- une approche structurée décrivant les différentes phases du travail de modélisation, d'évaluation, de réalisation et de documentation du système.

D'autre part, l'environnement de simulation doit :

- permettre de décrire et d'analyser les problèmes techniques et économiques qui apparaissent devant les principales phases du cycle de vie du système de production,
- être ouvert et être capable de communiquer avec d'autres applications, de façon à pouvoir s'intégrer au sein d'un système d'aide à la décision,
- être adaptable pour suivre l'évolution des systèmes de production.

Les techniques de simulations discrètes [THI 93] reposent sur les concepts de description d'activités, de partage des ressources, de gestion des files d'attente et de suivi de flux d'entités. La simulation discrète utilise le concept d'objets distincts qui réagissent entre eux à des temps spécifiques et identifiables. L'intervalle s'écoulant entre deux événements correspond à une activité. Le mécanisme de simulation se déroule de façon cyclique à partir de nouveaux éléments dans le modèle, en commençant par la détermination des prochains événements à exécuter, puis par la mise à jour du modèle jusqu'à la sortie des résultats.

Les travaux de Francis Martin ont permis de présenter quelques simulateurs adaptés aux différents domaines d'application sans établir une comparaison entre les simulations [MAR 87]. Un inventaire des méthodes de simulation fourni par Thierry Lenclud [LEN 93], met en évidence les avantages et les inconvénients des différents outils, et pose ainsi les bases d'une méthodologie de modélisation et de simulation des

activités de l'entreprise. Une application de modélisation et de simulation d'un système de production au sein d'une entreprise textile de confection a permis d'examiner la mise en œuvre d'une réflexion stratégique sur le processus de production à partir d'hypothèses d'évolutions incertaines de marchés perturbants [THI 95].

Des études comparatives menées sur les différents langages existants de simulation ont permis de déterminer leurs avantages et inconvénients respectifs [CER 88][EKE 89][COU 90][BAN 91][KEE 91]. Les langages peuvent être regroupés en trois grandes familles d'outils permettant l'élaboration et l'exploitation d'un modèle de simulation :

- les langages généraux tels que FORTRAN, PASCAL, C permettent une modélisation fonctionnelle ou logico-mathématique [CER 88] réalisée à l'aide de primitives de programmation dont les qualités et les défauts sont bien connus [ADI 91],
- les langages de simulation (SLAM II [ALA 89], SIMAN [PED 90], CADENCE, HOCUS [SZY 88]) sont basés sur la simulation à événements discrets. Si ces langages présentent l'avantage de pouvoir s'adapter pratiquement à tous les problèmes de simulation, leur utilisation nécessite une formation avancée. Pour pallier à cet inconvénient, des outils plus spécialisés (logiciels dédiés) ont été élaborés,
- les logiciels dédiés tels que (DOSIMIS, WITNESS, SIMFACTORY) ne requièrent qu'une manipulation simple de la part de l'utilisateur pour l'élaboration du modèle. Leur avantage essentiel est la capacité de modélisation mais l'utilisation de fonctions internes élaborées restreint leurs champs d'application.

D'une façon générale, les modèles existants du système de production peuvent être classés en quatre familles [KIE 86] : les modèles physiques, abstraits, statiques et dynamiques. Ils sont nombreux à avoir été développés pour l'étude des différents systèmes (physique, "informationnel" et "décisionnel") constituant le système de production proprement dit. Le tableau suivant a pour objectif de décrire les complémentarités et les insuffisances des principaux modèles utilisés.

Modèles	Phase d'utilisation				Temps		Domaine d'utilisation			
	Analyse	Pré-étude	étude	réalisation	Statique	Dynamique	décisionnel	informationnel	système physique	
									Partie opérative	Partie Commande
GRAI	•	•			•		•			
MERISE		•	•		•			•		
AXIAL		•	•		•			•		
SSAD		•	•		•			•		
SADT	•	□			•		□	•	•	
Réseaux de files d'attente		•	□			•			•	
Réseaux de Petri	•	•	•			•			•	•
Grafcet	•	•	•	•		•				•
Simulation		•	•			•	□	□	•	•

• : particulièrement bien adapté □ : moins bien adapté

Tableau 1. **Comparatif des différents outils de modélisation**

Ainsi, les réseaux de Petri (RdP) constituent un formalisme simple et de représentation graphique compacte pour les systèmes à événements asynchrones, concurrents et non déterministes [PET 77]. De plus, leur aptitude à la modélisation modulaire et hiérarchisée comme support de représentation et de traitement a été également démontrée [NOY 90]. Les potentialités de simulation des RdP ont été étudiées et mises en œuvre dans le système SAGASSE [BES 91]. Cet outil repose sur l'interaction dynamique entre un "interpréteur" basé sur les RdP colorés qui permet la représentation rigoureuse et précise de l'état courant de l'atelier avec un système expert garant du pilotage et de la gestion.

Suivant le but de modélisation et/ou d'analyse précis, différents formalismes des RdP sont actuellement connus (Rdp colorés, temporisés, à capacité, à objets, ... [DAV 92]) ainsi que leurs exploitations dans différents domaines [DAV 94][PRO 95]. Par exemple, un modèle de RdP, appliquant la méthodologie du "juste à temps" par un système KANBAN a été étudié et permet d'évaluer les performances du système suivant une méthode analytique [MAS 90]. Les RdP de haut niveaux (colorés/prédicats/transition) ont été développés dans un souci d'unification des différents modèles. Les Hypernets [YIM 95][LEF 95] sont définis à partir d'une

sémantique ensembliste conjointement avec une théorie des contraintes permettant une plus grande puissance d'expression, ce qui facilite les phases d'analyse, de pré-étude et d'étude. La théorie additionnelle est celle des fonctions "multivoques" [BER 86], des multi-ensembles, et son application autorise la propagation des contraintes [JAF 94]. Les Hypernets sont en fait une généralisation de nombreux autres RdP et peuvent modéliser graphiquement des programmes logiques contraints [CHA 80] et des problèmes discrets/continus. Un outil de simulation ([YIM 95]) a été développé à cet effet et réalise une implantation et une expérimentation d'un algorithme de semi-décision pour le problème de l'accessibilité d'un marquage pour les Hypernets linéaires. Cette spécificité permet au simulateur de trouver le chemin optimal pour atteindre un but (marquage à atteindre) à partir des faits (places sources). Il fournit une analyse multi-critères où les transitions du réseau comportent les paramètres de décision standards à l'entreprise. La connexion des modèles propres à chaque intervenant permet d'obtenir des renseignements optimisés dans des délais très courts (fonction du temps de simulation des combinaisons).

II.4.2. Représentation graphique

Les places sont représentées graphiquement par des cercles et identifiées par un marquage propre. Les jetons sont de type colorés [JEN 91], c'est à dire qu'ils peuvent prendre un ensemble de caractéristiques écrits sous la forme d'un vecteur v_k : $\text{jeton}_k(x_{k1}, \dots, x_{ki}, \dots, x_{kn})$. La représentation graphique des arcs symbolise le déplacement des jetons entre des places et des transitions. Ils sont nommés par le ou les jetons qu'ils véhiculent. Les transitions représentent des fonctions et des conditions (écrites sous forme de prédicats) qui vont s'appliquer sur les composantes des jetons. Lorsque ces prédicats sont vérifiés, il y a franchissement de la transition et les jetons se positionnent dans les places correspondantes.

Pour les Hypernets, le schéma de représentation d'un maillon se dessine de la manière suivante (Figure 5) :

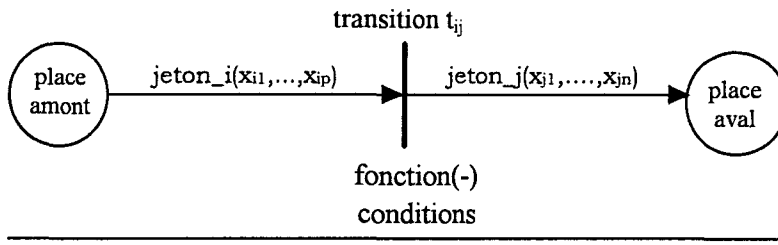


Figure 5. **Ecriture générale d'un maillon hypernet.**

II.4.3. Déclarations

Les marquages initiaux et finaux initialisent la position des jetons aux places voulues. Les valeurs initiales a_{ik} et a_{jk} , affectées lors du passage du jeton dans la place en amont, sont ensuite associées aux variables x_{ik} et x_{jk} , des caractéristiques de la fonction de l'arc correspondant. Les déclarations des marquages Hypernets s'écrivent suivant une écriture propre au langage Prolog :

```
marquage_initial(<<place_amont, <jeton_i(a_i1,...a_ip)>>>).
marquage_final(<<place_aval, <jeton_j(a_j1,...a_jn)>>>).
```

Le réseau permet de faire les liens entre les places et les transitions liées graphiquement, et de représenter le passage des jetons. Les déclarations des noeuds Hypernets sont de la forme suivante :

```
reseau(<transition,
<<place_amont , < jeton_i(x_i1,...x_ip) >>>,
<<place_aval, < jeton_j(x_j1,...x_jn) >>>>
:- fonction(-), {conditions}.
```

Les fonctions et contraintes sont décrites à la fin du réseau. Les fonctions sur les transitions sont déclarées de la manière suivante :

```
fonction(-) {conditions}.
```

Un modèle hiérarchique peut être envisagé où les transitions peuvent être développées en Hypernets pour définir tous les niveaux de l'entreprise : ordonnancement,

atelier de travail, station, machine. Des exemples d'écriture des conditions sur les fonctions sont définies dans les simulations.

II.4.4. Modèle d'une entreprise

Une modélisation de l'entreprise par Hypernets, caractérisant les 3 fonctions principales et les 2 types de flux de la Figure 3, est représentée en Figure 6. Trois circuits principaux du traitement de l'information à partir de la réception de demande dans la place RD peuvent être envisagées :

- Le premier circuit correspond au « déstockage » direct des produits finis. Le jeton commande de la demande aval est réceptionné dans la place RD et un test sur la transition (t1) permet d'évaluer son acceptation par l'entreprise. Par exemple, il faut que:
 - le produit soit réalisable par l'entreprise,
 - la quantité demandée corresponde à des quantités suffisantes,
 - le délai demandé soit acceptable,

Si toutes ces conditions sont réunies alors la demande passe dans la place DPF et une comparaison est effectuée avec le contenu de la place SPFV grâce à la transition (t2). Si la place SPFV satisfait à la demande alors il y a déstockage et le produit demandé est disponible dans la place SAV. Sinon le complément est commandé en amont par la transition (t3).

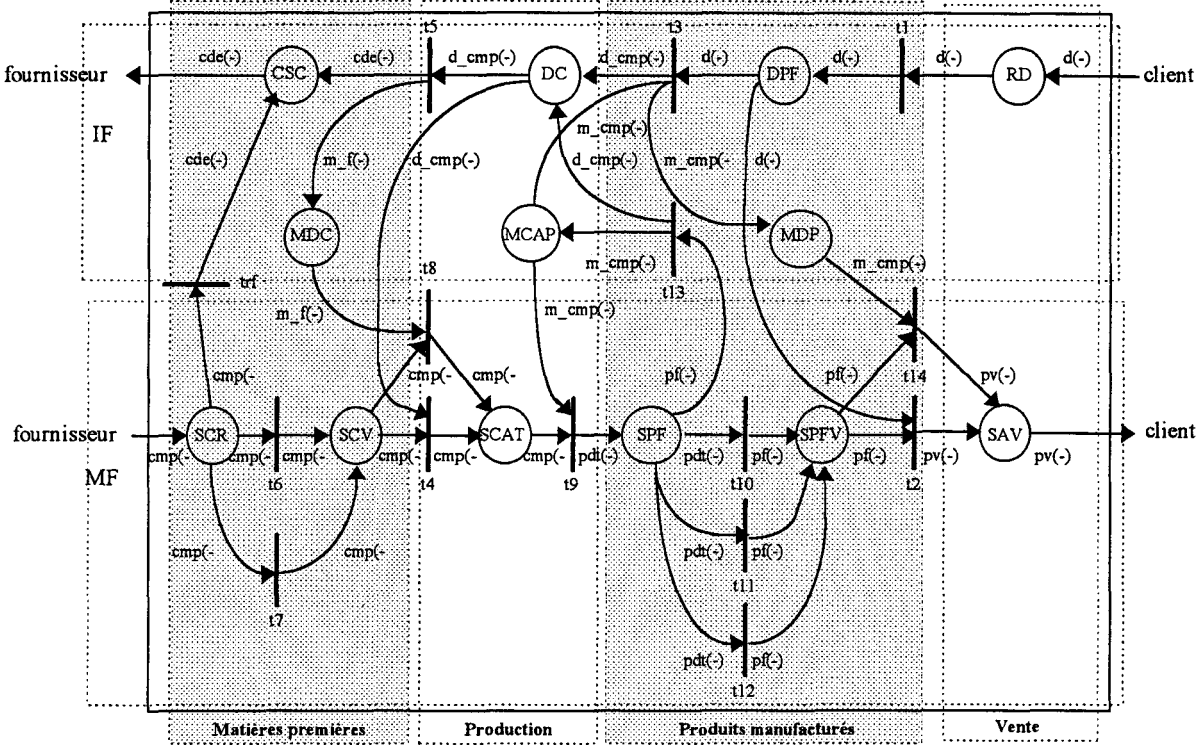


Figure 6. Modélisation de l'entreprise par Hypernets

Places	Description
RD	réception de la demande
DPF	demande en produits finis
DC	demande en composant
MDPF	mémorisation de la demande en produits finis
MCAP	mémorisation de la commande à produire
CSC	commande en composant aux fournisseurs
MDC	mémorisation de la demande en composant
SCR	stock composant en réception
SCV	stock composant validé après les tests de qualité
SCAT	stock composant à transformer
SPF	stock produits finis
SPFV	stock produits finis validés
SAV	stock à vendre

Tableau 2. Description des places du modèle hypernet de l'entreprise

- Le deuxième circuit correspond à la demande en composant et à la production des produits finis. Le complément de la demande est mis en mémoire dans la place MDPF. La position du jeton demande dans les places DC et MCAP traduit la décomposition du produit en ses constituants élémentaires correspondant à la matière première du stock

de l'entreprise. La place MCAP est une mémorisation de cette décomposition et teste avant le lancement de la production, la disponibilité des composants nécessaires à la réalisation des produits finis. Une comparaison du contenu de la place DC avec le stock en composant SCV est effectuée pour évaluer les besoins en stock. En fonction des disponibilités des matières premières, les opérations suivantes sont effectuées dans l'ordre par : la production (t9), le stockage dans la place SPF, la vérification des produits finis (t10), (t11), (t12) et enfin la vente directe grâce à la place mémorisation MDPF et à la transition (t14). En fonction de la qualité des produits finis, une nouvelle commande peut être enclenchée par le franchissement de la transition (t13). Sinon, de la même manière que pour les produits finis précédemment définis, une partie des composants disponibles est déstockée en SCAT et attend que le complément soit commandé en amont vers les fournisseurs par la transition (t5).

- Le troisième circuit définit les commandes aux fournisseurs et la récupération des composants. La demande est enregistrée dans les places CSC et MDC (pour la mémorisation). La réception des composants nécessaires à la fabrication des produits finis s'effectue par la place SCR dont la qualité est vérifiée par la transition (t6). Si la marchandise reçue n'est pas acceptable (ou si elle ne peut subir de remise de prix (exemple (t7)) , alors son retour est assuré par la transition (trf) et engendre ainsi un retour de la commande fournisseur. Le complément des composants est additionné au stock de la place SCAT pour le lancement en fabrication.

Les inscriptions sur les arcs sont des fonctions propres à chaque entreprise. Elles modélisent son comportement interne. Sa validation nécessite une initialisation des paramètres à partir des données réelles et la mise en application du modèle dans différentes entreprises partenaires de la filière.

Les places

Elles correspondent aux différentes étapes de la réception de la demande jusqu'au stock à vendre.

Toutes ces places contiennent un certain nombre de jetons représentés par des vecteurs au niveau des arcs (cf. **Tableau 2**).

Les jetons

L'utilisation des variables associées aux caractéristiques des jetons varient en fonction de leur position dans les différentes places du réseau. Par exemple, les caractéristiques du jeton demande sont définies par les variables n, x, q, p, d . Pour alléger la notation, les composantes du vecteur sont réduites au symbole : (-). Pour notre réseau, les déplacements des jetons sont définis dans les places qui les contiennent.

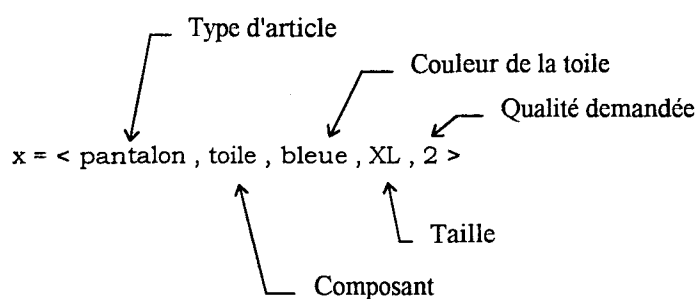
Places	Jetons	Notation simplifiée pour le schema
RD	demande(n, x, q, p, d)	d(-)
DPF	demande(n, x, q, p, d)	d(-)
DC	demande_en_composant(n, c, q, p, d)	d_cmp(-)
MDPF	mémorisation_produit(n, x, q, p, d)	m_p(-)
MCAP	mémorisation_composant(n, c, q, p, d)	m_cmp(-)
CSC	commande(n, c, q, p, d)	cde(-)
MDC	mémorisation_fournisseur(n, c, q, p, d)	m_f(-)
SCR	composant(n, c, q, p, d, s)	cmp(-)
SCV	composant(cq, p) composant(n, c, q, p, d)	cmp(-)
SCAT	composant(n, x, q, p, d)	cmp(-)
SPF	produit(n, x, q, p, d, s)	pdt(-)
SPFV	produit_fini(x, q, p) produit_fini(n, x, q, p, d)	pf(-)
SAV	produit_vendu(n, x, q, p, d)	pv(-)

Tableau 3. Liste des fonctions associées aux jetons du modèle hypernet de l'entreprise

Chacune des variables (n, x, c, q, p, d, s) des fonctions associées aux jetons permettent d'initialiser les valeurs de leurs propres caractéristiques et sont définis telles que :

n corresponde au nom du donneur d'ordres,

x représente la nomenclature propre du produit, par exemple :



c représente la liste des composants nécessaires à la formation du produit fini, par exemple $c = \langle \text{toile, bleue} \rangle$,
q définisse la quantité du produit fini,
p caractérise le prix unitaire du produit fini ou composant,
d indique la date de fabrication ou de livraison, par exemple $d = \langle \text{jour, mois, année} \rangle$,
s représente la cotation de la qualité obtenue après des opérations de transformation sur les produits.

Un exemple d'initialisation des paramètres associés aux jetons est représenté en **Figure 7**. Les caractéristiques du jeton, situé dans la place CD, traduisent une commande, en date du 15 décembre 1997, de 450 pantalons de toile bleue, de taille XL, de qualité 5 et au prix unitaire de 45 FF. Les variables n, x, q, p, d prennent respectivement les valeurs des caractéristiques du jeton telles que :

n=distributeur
 $x = \langle \text{pantalon, toile, bleue, XL, 5} \rangle$
 $q = 450$, $p = 45$ FF
 $d = \langle 15, \text{dec}, 1997 \rangle$.

Le franchissement de la transition t_{ij} ne s'effectue que si les variables de la fonction demande de l'arc répondent aux contraintes définies par :

$q > 10$ et $q < 5000$
 $x = \langle \text{pantalon, -, -, -, -} \rangle$ ou $x = \langle \text{chemise, -, -, -, -} \rangle$
Prix $p > 42$.

Si toutes les conditions sont réunies, les variables associées à la fonction réception de l'arc prennent les valeurs telles que :

$n = n'$
 $x = x'$
 $q = q'$
 $p = p'$
et $d' = d + 1$.

Enfin, les caractéristiques du jeton réception situé dans la place RD, traduisent la réception de la commande initiale.

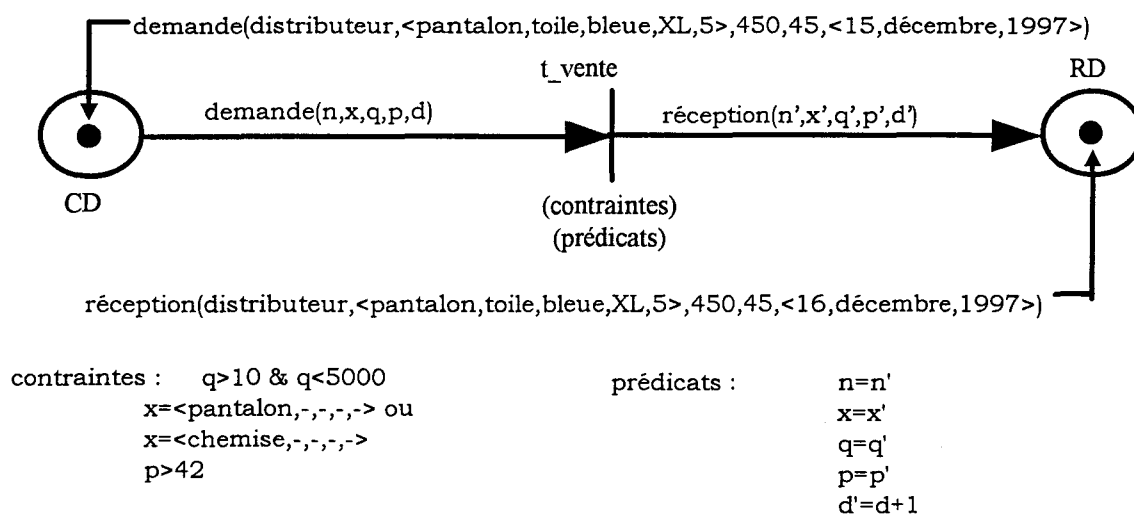


Figure 7. Exemple d'initialisation des paramètres associés aux jetons.

Le nombre et la couleur des jetons correspondants au marquage initial du réseau peuvent être déterminés à partir de la base de données des stocks de l'entreprise.

Les transitions

L'ensemble des contraintes définies pour chaque transition est décrit par le commentaire correspondant.

Transitions	Description	Commentaire
t1	test-acceptabilité	teste si le produit demandé correspond aux possibilités de l'entreprise.. si la quantité demandée est acceptable. si la date de livraison demandée est acceptable
t2	test-spfv	teste si le stock produit fini correspond à la demande
t3	demande-composant	décompose le produit en ses composants propres (suivant la nomenclature du produit)
t4	test-scv	teste si le stock composant correspond à la demande
t5	cde-amont	commande vers les entreprises en amont les composants manquants
t6	test-Qualité-cmp	accepte ou non les composants reçus suivant des critères de qualité
t7	test-rabais	revoit à la baisse le prix des composants de mauvaise qualité
trf	retour-fournisseur	renvoie les malfaçons
t8	destockage-cmp	sort du stock les composants suffisants pour fabriquer les produits finis
t9	production	transforme et assemble les composants suivant la nomenclature des produits finis
t10	test-qualité-pf	teste la qualité des produits finis
t11	test-rabais	revoit à la baisse le prix des produits finis de mauvaise qualité
t12	test-amelioration	teste si le produit, dont les tests qualité et rabais ont été négatifs, nécessite une retouche
t13	nlle-cde	teste une nouvelle commande de produit si l'amélioration est négative
t14	destockage-pf	déstocke les produits finis pour la vente

Tableau 4. Description des transitions du modèle
hypernet.

II.4.5. Modélisation de la filière

(1) Principe

A partir du modèle de l'entreprise (Figure 3), nous pouvons déterminer le modèle de la filière (Figure 8), en utilisant la structure de base de la représentation globale de la filière THD (Figure 4).

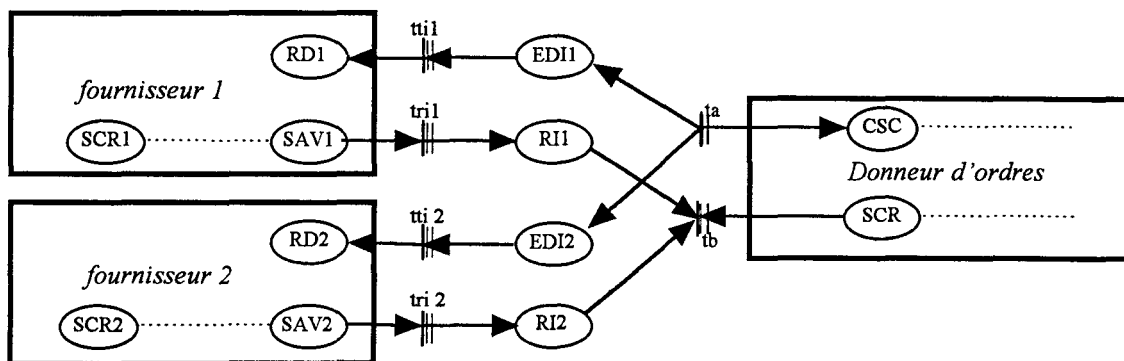


Figure 8. **Modèle de la filière**

Le jeton commande contenu dans la place CSC du donneur d'ordres franchit la transition ta, (contenant les contraintes de temps de départ des commandes) et se décompose en 2 jetons commande identiques dont les informations sont traduites sous format EDI. La réception des demandes dans les places RD1 et RD2 des fournisseurs est établie par le franchissement des transitions tti1 et tti2 (Temps de Transfert d'Information). Les jetons contenus dans les stocks à vendre SAV1 et SAV2 des fournisseurs 1 et 2, répondant à la demande précédente, sont acheminés par l'ensemble des noeuds correspondants aux transitions tri1, tri2 (temps de retour des informations) aux places RI1, RI2 (Réunion Info et/ou flux 1 et 2). Ces jetons contenus dans les places RI1 et RI2 sont sélectionnés ou non par la transition tb (choix du fournisseur) puis sont réceptionnés dans la place du stock composant du donneur d'ordres SCR.

(2) Options de stratégies d'approvisionnement

Un certain nombre de stratégies d'approvisionnement peuvent être envisagées entre un distributeur et deux fournisseurs. La première stratégie correspond au choix du meilleur partenaire fournisseur pour une commande donnée, la deuxième consiste à partager de façon optimale la même commande entre les deux fournisseurs. Une évaluation des différents cas possibles est proposée afin de montrer la faisabilité d'une application de simulation de stratégie entre différents partenaires :

Choix d'un partenaire fournisseur

Un donneur d'ordres formule la même demande en produits vis-à-vis de ses deux fournisseurs. Nous pouvons envisager de laisser le libre choix de la décision au donneur d'ordres ou de lui fournir une solution selon son propre paramétrage initial.

Sans critère de sélection automatique

Le donneur d'ordres récupère les deux jetons issus des fournisseurs 1 et 2 pour effectuer sa propre analyse des informations associées aux jetons. Le décideur dispose de tous les degrés de liberté possibles pour prendre sa décision.

Avec critères de sélection automatique

Le critère de sélection qui comprend un certain nombre de contraintes sélectionnées par le décideur, doit être écrit dans le réseau au niveau de la transition t_b du modèle de la filière THD (Figure 8). Le donneur d'ordres réceptionne un seul des deux jetons issus des fournisseurs 1 et 2 lui permettant ainsi de disposer immédiatement d'une solution possible.

Partage optimal d'une commande

Un donneur d'ordres souhaite s'approvisionner chez son fournisseur principal et compléter la quantité manquante chez un fournisseur secondaire. Par exemple, la demande du donneur d'ordres est de 300 pantalons; elle est décomposée en une quantité X pour le fournisseur 1 et en une quantité Y pour le fournisseur 2 avec la contrainte $X+Y=300$

III. Simulation

III.1. Simulation d'une entreprise³

Trois situations usuelles de commande client sont envisagées dans les exemples suivants :

- le premier exemple recherche la date de livraison "au plus tôt" d'une commande. Le simulateur utilise le principe du chaînage avant pour déterminer les marquages susceptibles d'être atteints à partir du marquage initial,
- le deuxième exemple recherche la date de production "au plus tard" par rapport à une commande fixée. Le simulateur utilise le principe du chaînage arrière pour déterminer les marquages aboutissant au marquage final ainsi déterminé,
- le dernier exemple recherche les dates "au plus tôt" et "au plus tard" de livraison d'une commande. Les marquages initiaux et finaux sont partiellement fixés et les variables non "instanciées" correspondent aux critères à minimiser.

III.1.1. Date de livraison au plus tôt

Commande :

Un distributeur passe une commande à un seul fournisseur en fonction des informations suivantes :

- Date de commande : 2 décembre 1997,
- 110 pantalons de toile bleue de taille XL, de qualité 4, au prix unitaire de 44 FF,
- 510 chemises de soie verte de taille XXL, de qualité 4, au prix unitaire de 47 FF,
- Critère à minimiser : date de livraison au plus tôt.

Nomenclature :

³ La représentation des différents jetons en écriture Prolog apparaît uniquement pour l'application de la première simulation. Par la suite, l'ensemble des résultats de simulation sont regroupés au sein d'un même tableau afin de faciliter les phases d'analyse et d'interprétation.

La nomenclature d'un pantalon de toile bleue de taille XL nécessite 1 unité de toile bleue, 1 unité de boutons et de fil, et 1 unité de ceinture. La nomenclature d'une chemise de soie verte de taille XXL nécessite 1 unité de soie verte et 1 unité de col.

Avant la simulation :

L'écriture du *marquage initial* est traduite par les déclarations suivantes en Prolog :

Commande:

```
<rd,<demande(distributeur,<pantalon,toile,bleue,XL,4>,110,44,<2,dec,97>),
demande(distributeur,<chemise,soie,verte,XXL,4>,510,47,<2,dec,97>)>>
```

Stocks Produit Fini:

```
<spfv,<produit_fini(<pantalon,toile,bleue,XL,4>,100,43.5),
produit_fini(<chemise,soie,verte,XXL,4>,120,46.5)>>
```

Stocks Composants:

```
<scv,<  composant(<toile,bleue>,1500,18),
        composant(<soie,verte>,850,19),
        composant(<bouton_p,fil>,40,1),
        composant(<ceinture>,40,11),
        composant(<col>,50,7)>>,>
```

L'écriture du *marquage final* est traduite par les déclarations suivantes en Prolog :

```
<sav,_>,
<spfv,_>,
<scv,_>,>
```

En début de simulation, les places SAV, SPFV et SCV sont non instanciées comme indiqué par la notation ‘_’.

Après la simulation :

Les valeurs des jetons dans les places SAV, SPFV et SCV nous fournissent la date de livraison au plus tôt et l'état des stocks en produits finis et composants. Les résultats de la simulation en écriture Prolog sont les suivants:

Commande:

```
<sav,<produit_vendu(distributeur,<pantalon,toile,bleue,XL,4>,10,42.87,<9,dec,97>)>+  
<produit_vendu(distributeur,<pantalon,toile,bleue,XL,4>,100,43.5,<4,dec,97>)>+  
<produit_vendu(distributeur,<chemise,soie,verte,XXL,4>,120,46.5,<4,dec,97>)>+  
<produit_vendu(distributeur,<chemise,soie,verte,XXL,4>,390,45.93,<13,dec,97>)>>.
```

La date de livraison au plus tôt des 110 pantalons est le 9 décembre 1997 : 100 pantalons disponibles le 4 décembre 1997 à 43.5 FF et 10 pantalons disponibles le 9 décembre 1997 à 42.87 FF en puisant dans le stock de composants disponibles pour la mise en production.

La date de livraison au plus tôt des 510 chemises est le 13 décembre 1997 : 120 chemises disponibles le 4 décembre 1997 à 46.5 FF et 390 chemises disponibles le 13 décembre 1997 à 45.93 FF en puisant dans le stock de composants disponibles et en commandant les quantités manquantes pour la mise en production.

Stocks Produit Fini:

```
<spfv,<stock(<pantalon,toile,bleue,XL,4>,0)>+<stock(<chemise,soie,verte,XXL,4>,0)>>
```

Les stocks disponibles en produits finis de pantalons et de chemises s'annulent.

Stocks Composants:

```
<scv,<composant(<toile,bleue>,1490,18)>+<composant(<soie,verte>,460,19)>+  
<composant(<bouton_p,fil>,30,1)>+<composant(<ceinture>,30,11)>+<composant(<col>,0)>>.
```

Les stocks de composants disponibles sont diminués des quantités nécessaires pour répondre à la demande du distributeur.

Un récapitulatif des résultats de la simulation figure dans le Tableau 5 et le Tableau 6, permettant de rendre compte des différentes évolutions des paramètres du système logistique de l'entreprise.

Avant la simulation			
commande du distributeur	date commande	02/12/97	
type	Prix unitaire	nombre	critère à respecter
pantalon toile bleue taille XL qualité 4	prix < 44	110	livraison au plus tôt
chemise soie verte taille XXL qualité 4	prix < 47	510	livraison au plus tôt
Après la simulation			
réponse à la commande	date commande	02/12/97	
type	Prix unitaire	nombre	date de livraison
pantalon toile bleue taille XL qualité 4	43.5	100	04/12/97
pantalon toile bleue taille XL qualité 4	42.87	10	09/12/97
chemise soie verte taille XXL qualité 4	46.5	120	04/12/97
chemise soie verte taille XXL qualité 4	45.93	390	13/12/97

**Tableau 5. Résultats de simulation relatifs à la stratégie
de la date de livraison au plus tôt.**

La date de livraison pour la totalité de la commande de 110 pantalons est le 09/12/97. 100 pantalons sont disponibles en stock produits finis et 10 pantalons sont fabriqués pour compléter la commande en utilisant 10 unités de toile bleue, 10 unités de boutons et de fil et 10 unités de ceinture.

La date de livraison pour la totalité de la commande de 510 pantalons est le 13/12/97. 120 chemises sont disponibles en stock produits finis et 390 chemises restent à fabriquer. Pour cela, 390 unités de soie verte prélevées du stock composants et 390 unités de col sont nécessaires. Parmi les 390 unités de col nécessaires, 50 sont disponibles en stock composants et une commande des 340 cols manquants est faite et mémorisée dans la place MDC (Figure 6).

Avant la simulation		
Type	Prix unitaire	Nombre d'unités
Stocks composants		
toile bleue	18	1500
soie verte	19	850
bouton pantalon et fil	1	40
ceinture	11	40
col	7	50
Stock produits finis		
pantalon toile bleue taille XL qualité 4	43.5	100
chemise soie verte taille XXL qualité 4	46.5	120
Après la simulation		
Type	Prix unitaire	Nombre d'unités
Stocks composants		
toile bleue	18	1490
soie verte	19	460
bouton pantalon et fil	1	30
ceinture	11	30
col	7	0
Stock produits finis		
pantalon toile bleue taille XL qualité 4	43.5	0
chemise soie verte taille XXL qualité 4	46.5	0

Tableau 6. **Etat initial et final des stocks du fournisseur**
dans la stratégie de la date de livraison au plus tôt.

La *séquence de tirs* des transitions traduit l'évolution de l'information traitée par la simulation en fonction du modèle Hypernet de l'entreprise (Figure 6 page 33):

test_acceptabilite test_acceptabilite test_spfv test_spfv demande_composant
demande_composant test_scv test_scv test_scv test_scv test_scv cde_amont cde_amont cde_amont
cde_amont retour_fournisseur retour_fournisseur retour_fournisseur retour_fournisseur test_q1
test_q1 test_q1 test_q1 destockage_cmp destockage_cmp destockage_cmp destockage_cmp
production production test_q2 test_q2 destockage_pf destockage_pf.

III.1.2. *Date de production au plus tard*

Un récapitulatif des résultats de la simulation figure dans le Tableau 7 et le Tableau 8, permettant de rendre compte des différentes évolutions des paramétrages du système logistique de l'entreprise. L'objectif de cette simulation est de connaître la date de passage de la commande au plus tard pour être livrée avant la date de livraison fixée.

Commande:

Un distributeur souhaite recevoir 120 pantalons de toile bleue de taille XL, de qualité 4, au prix unitaire de 45 FF et être livré au plus tard le 20 décembre 1997.

Avant la simulation			
commande du distributeur	date commande	02/12/97	
type	Prix unitaire	nombre	critère à respecter
pantalon toile bleue taille XL qualité 4	prix<44	120	livraison < 20/12/97
Après la simulation			
réponse à la commande	date commande	02/12/97	
type	Prix unitaire	nombre	date limite de passage de commande
pantalon toile bleue taille XL qualité 4	43	120	18/12/97

Tableau 7. Résultats de simulation relatifs à la stratégie
de la date de livraison au plus tard.

La date limite de passage de la commande de 120 pantalons pour être livré au plus tard le 20/12/97 est le 18/12/97. A la date de la simulation de la commande du 02/12/97, 100 pantalons sont disponibles en stock produits finis.

Avant la simulation		
Type	Prix unitaire	Nombre d'unités
Stocks composants		
toile bleue	18	100
bouton pantalon et fil	1	1000
ceinture	7	160
Stock produits finis		
pantalon toile bleue taille XL qualité 4	43	200
Après la simulation		
Type	Prix unitaire	Nombre d'unités
Stocks composants		
toile bleue	18	100
bouton pantalon et fil	1	1000
ceinture	7	160
Stock produits finis		
pantalon toile bleue taille XL qualité 4	43	80

Tableau 8. Etat initial et final des stocks du fournisseur
dans la stratégie de la date de livraison au plus tard.

III.1.3. Date au plus tôt et au plus tard de livraison d'une commande

Le simulateur donne la possibilité de propager des intervalles sur les dates, les prix, les quantités. Ainsi, nous définissons un intervalle de temps en utilisant des inégalités sur la variable de temps d. L'écriture de $d \in]20,25[$ pour la date finale est : $\{d > 20, d < 25\}$.

Commande:

Un distributeur souhaite recevoir 120 pantalons de toile bleue de taille XL, de qualité 2, au prix unitaire de 45 FF et être livré entre le 20 et 25 décembre 1997.

Avant la simulation			
commande du distributeur	date commande	02/12/97	
type	Prix unitaire	nombre	critère à respecter
pantalon toile bleue taille XL qualité 2	prix<45	120	livraison entre le 20 et le 25/12/97
Après la simulation			
réponse de la commande	date commande	02/12/97	
type	Prix unitaire	nombre	date limite de passage de la commande
pantalon toile bleue taille XL qualité 2	43	120	du 18 au 23/12/97

Tableau 9. Résultats de simulation relatifs à la stratégie de la date de livraison au plus tôt et au plus tard.

Le simulateur fournit l'intervalle de la date de mise en oeuvre de la commande entre le 18 et 23 décembre 1997 de 120 pantalons de qualité 2 au prix unitaire de 43 FF pour être livré entre le 20 et 25/12/97.

Avant la simulation		
Type	Prix unitaire	Nombre d'unités
Stocks composants		
toile bleue	18	1500
bouton pantalon et fil	1	40
ceinture	11	40
Stock produits finis		
pantalon toile bleue taille XL qualité 2	43	120
Après la simulation		
Type	Prix unitaire	Nombre d'unités
Stocks composants		
toile bleue	18	1500
bouton pantalon et fil	1	40
ceinture	11	40
Stock produits finis		
pantalon toile bleue taille XL qualité 2	43	0

Tableau 10. Etat initial et final des stocks du fournisseur dans la stratégie de la date de livraison au plus tôt et au plus tard.

III.2. Simulation de la filière Textile/Habillement/Distribution

Deux situations usuelles entre les différents partenaires de la filière THD sont envisagées dans les exemples suivants.

Le premier exemple recherche le meilleur fournisseur selon une décision multicritère utilisant le prix et la date de livraison.

Le deuxième exemple détermine la meilleure répartition de la commande entre deux fournisseurs.

III.2.1. Choix d'un partenaire fournisseur avec critère de sélection

Pour obtenir le meilleur des deux fournisseurs (avec une même demande) suivant un critère de sélection, il suffit de modifier l'écriture du réseau en ajoutant une contrainte sur la transition tb. Dans notre exemple, le critère à minimiser combine le prix proposé et la date de livraison possible du fournisseur. En cas d'égalité de la valeur minimale du critère, le choix s'effectue alors en priorité en faveur du fournisseur dont le délai d'approvisionnement est le plus court.

Commande

Deux jetons sont envoyés, datés du 2 décembre 1997, du distributeur au fournisseur 1 et au fournisseur 2 traduisant la commande suivante :

450 pantalons de toile bleue de taille XL, de qualité 4, au prix unitaire maximum de 42 FF en fonction d'une contrainte combinant une date de livraison au plus tôt et un prix.

Avant la simulation				
commande du distributeur	date commande	02/12/97		
type	Prix unitaire	nombre	critère à respecter	
pantalon toile bleue taille XL qualité 2	prix < 42	450	livraison le plus tôt	
Après la simulation				
réponse de la commande	date commande	02/12/97		
type	Prix unitaire	nombre	date de livraison	meilleur fournisseur
pantalon toile bleue taille XL qualité 2	39.06	450	18/12/97	fournisseur 2

Tableau 11. Résultats de simulation relatifs à la stratégie du choix du partenaire.

Les résultats de simulation aboutissent directement au choix du fournisseur 2. La livraison "au plus tôt" des 450 pantalons au prix unitaire de 39.06 FF est le 18 décembre 1997, en puisant 180 pantalons disponibles en stock et en complétant en fabrication les 270 restants. La fabrication des 270 pantalons complémentaires nécessite :

- 270 unités de toile bleue dont 90 sont en stock et compléter par une commande de 180,
- 270 unités de boutons et de fil pris directement dans le stock composants,
- 270 unités de ceinture se décomposant en 50 disponibles en stock et 220 à commander.

Avant la simulation		
Type	Prix unitaire	Nombre d'unités
Stocks composants fournisseur 1		
toile bleue	14	120
bouton pantalon et fil	1	830
ceinture	11	120
Stock produits finis fournisseur 1		
pantalon toile bleue taille XL qualité 4	40	300
Stocks composants fournisseur 2		
toile bleue	17	90
bouton pantalon et fil	1	750
ceinture	13	50
Stock produits finis fournisseur 2		
pantalon toile bleue taille XL qualité 4	41	180
Après la simulation		
Type	Prix unitaire	Nombre d'unités
Stocks composants fournisseur 1		
toile bleue	14	0
bouton pantalon et fil	1	680
ceinture	11	0
Stock produits finis fournisseur 1		
pantalon toile bleue taille XL qualité 4	40	0
Stocks composants fournisseur 2		
toile bleue	17	0
bouton pantalon et fil	1	480
ceinture	13	0
Stock produits finis fournisseur 2		
pantalon toile bleue taille XL qualité 4	41	0

Tableau 12. Etats initiaux et finaux des stocks
fournisseurs pour la stratégie de choix de partenaire.

Les états, l'initial et le final, des stocks du fournisseur 1 sont relatifs à la réponse à la commande du distributeur. Si l'offre du fournisseur 1 est refusée alors les valeurs des quantités en stock retrouvent leurs valeurs initiales.

L'écriture de la transition t_b dans le réseau est traduite par les déclarations suivantes :

Les places en amont de la transition $t_{critère}$ sont identifiées par $ri1$ et $ri2$ (voir Figure 8 page 39) avec l'écriture des jetons associés :

```

reseau      (<t_critere,
<
<ri1, <ri1d(N,X,Q,P,<J,M,A>)>>,
<ri2, <ri2d(N,X,Q,P',<J',M,A>)>>
>,
<
<scr, <meilleur(F,X,Q,P'',<J'',M,A>)>>
>
>
):-

```

A la transition $t_{critère}$ est associée une fonction critère :

```
critere(P,P',J,J',P'',J'',F)
```

Les contraintes sur les informations des jetons sont :

```
{J>=1,J<32,J'>=1,J'<32,J''>=1,J''<32,Q>=0,N=distributeur}.
```

Les différentes écritures de la fonction critere sont définies à la fin de la simulation :

critere(P,P',J,J',P',J')	{{(P×J)=S,(P'×J')=S',S>S',F=fournisseur2}.
critere(P,P',J,J',P,J)	{{(P×J)=S,(P'×J')=S',S<S',F=fournisseur1}.
critere(P,P',J,J',P,J)	{{(P×J)=S,(P'×J')=S',S=S',J<J',F=fournisseur1}.
critere(P,P',J,J',P',J')	{{(P×J')=S,(P'×J')=S',S=S',J>=J',F=fournisseur2}.

Par exemple, la première écriture compare les valeurs s et s' des fournisseurs 1 et 2, telles que $S = P \times J$ avec :

- P le prix proposé par le fournisseur 1,
- J le nombre de jours ouvrables entre la commande et la date de livraison.

Si la valeur de s est supérieure à la valeur de s' , alors le fournisseur 2 est sélectionné.

Le choix des valeurs initiales des paramètres au niveau de la réception $t_{crit\grave{e}re}$ ne laisse passer qu'un seul des deux jetons demande des fournisseurs 1 et 2. Le choix se fait automatiquement.

III.2.2. Complément de commandes

Le distributeur commande 300 pantalons de toile bleue de taille XL et de qualité 2 au prix unitaire de 44 FF et, souhaite s'approvisionner auprès de ses deux fournisseurs. La stratégie consiste à commander la quantité X disponible en stock de produits finis au fournisseur 1 et, compléter la quantité Y manquante au fournisseur 2 avec la contrainte $X+Y=300$. Cette simulation montre la possibilité d'ordonner les traitements des données suivant des priorités choisies.

Commande

Le 2 décembre 1997, le distributeur commande X pantalons de toile bleue de taille XL de qualité 2 et de prix unitaire 44 FF au fournisseur 1 et Y pantalons avec les mêmes caractéristiques au fournisseur 2. Les valeurs X et Y ne sont pas précisées initialement mais sont liées par la contrainte suivante $X+Y=300$.

Avant la simulation				
commande du distributeur	date commande	02/12/97		
type	Prix unitaire	nombre	Critère à respecter	
pantalon toile bleue taille XL qualité 2	prix<44	X	fournisseur1	livraison au plus tôt
pantalon toile bleue taille XL qualité 2	prix<44	Y	fournisseur2	livraison au plus tôt
Après la simulation				
réponse de la commande	date commande	02/12/97		
type	Prix unitaire	nombre	Date livraison	
pantalon toile bleue taille XL qualité 2	43.05	100	fournisseur1	8/12/97
pantalon toile bleue taille XL qualité 2	40.60	200	fournisseur2	17/12/97

Tableau 13. Résultats de simulation relatifs à la stratégie
du partage de la commande.

La date de livraison au plus tôt des 100 pantalons du fournisseur 1, disponibles en stock produits finis, au prix unitaire de 43.5 FF est le 8 décembre 1997. La date de livraison au plus tôt des 200 pantalons du fournisseur 2 au prix unitaire de 40.60 FF est le 17 décembre 1997, dont 180 sont disponibles en stock produits finis et les 20 restants

sont fabriqués en puisant dans le stock des composants. L'état des stocks initiaux et finaux des différents fournisseurs est récapitulé dans le Tableau 14.

Avant la simulation		
Type	Prix unitaire	Nombre d'unités
Stocks composants fournisseur 1		
toile bleue	18	1500
bouton pantalon et fil	1	840
ceinture	11	120
Stock produits finis fournisseur 1		
pantalon toile bleue taille XL qualité 2	43	100
Stocks composants fournisseur 2		
toile bleue	17	90
bouton pantalon et fil	1	750
ceinture	13	50
Stock produits finis fournisseur 2		
pantalon toile bleue taille XL qualité 2	43	180
Après la simulation		
Type	Prix unitaire	Nombre d'unités
Stocks composants fournisseur 1		
toile bleue	18	1500
bouton pantalon et fil	1	840
ceinture	11	120
Stock produits finis fournisseur 1		
pantalon toile bleue taille XL qualité 2	43	0
Stocks composants fournisseur 2		
toile bleue	17	70
bouton pantalon et fil	1	730
ceinture	13	30
Stock produits finis fournisseur 2		
pantalon toile bleue taille XL qualité 2	43	0

Tableau 14. Etats initiaux et finaux des stocks
fournisseurs pour la stratégie du complément de
commande.

III.3. Les capacités du simulateur

III.3.1. Informations spatiales

Une analyse des variations minimales et maximales en temps réel des informations des jetons contenus dans chaque place peut être effectuée. Ils peuvent représenter le stock

de commande de chaque référence (réel, cumulé, physique, en-cours de commande) et les activités des ressources de production (Homme, Machine).

III.3.2. Informations temporelles

Nous pouvons déterminer :

- les délais prévisionnels d'un produit entre la date du passage de sa commande et la date de la livraison.
- le délai minimum d'une commande par la sélection des meilleurs partenaires fournisseurs.
- la productivité maximale de l'entreprise journalière/hebdomadaire/mensuelle en tenant compte de ses contraintes et celles de ses partenaires.

III.3.3. Algorithmes d'optimisation

Nous pouvons calculer :

- les coûts prévisionnels d'un produit entre sa commande et sa livraison,
- la quantité optimale d'un produit par la sélection des meilleurs partenaires fournisseurs,
- les contraintes d'optimisation multicritères fixées par le client qui permettent de signaler les dysfonctionnements de chaque critère par rapport à des seuils préfixés.

IV. Conclusion

Nous avons fourni une représentation des flux logistiques d'une entreprise et de la filière THD dans le cadre de l'application d'une démarche "Quick Response". L'application d'un outil mathématique de modélisation et de simulation : les réseaux de Pétri de haut niveaux de type Hypernets appliqués à la représentation globale de l'entreprise et de sa filière, a permis d'évaluer différentes stratégies logistiques. Les mécanismes "d'instanciation" de variables et de propagation de contraintes des Hypernets fournissent un puissant formalisme pour des cas typiques d'interrogation du simulateur. Les exemples de simulation mettent en évidence la capacité du simulateur à traiter une information afin de fournir une solution envisageable tout en tenant compte de l'ensemble des contraintes initiales. Cependant, la multiplication des variables à "instancier" et des contraintes sur les paramètres ralentissent fortement le temps d'exécution de la simulation et peut, parfois, ne pas fournir de solution. Une réduction des contraintes permet d'augmenter les chances d'obtenir une solution acceptable.

Dans le cadre d'une utilisation industrielle du simulateur, l'exploitation des informations contenues dans les bases de données et leurs traitements en temps réel permettent d'initialiser automatiquement l'ensemble des paramètres de l'entreprise. Cela nécessite quelques extensions sur le formalisme proposé, notamment sur les types d'arcs et le test de validation des transitions. Un rapprochement vers des représentations symboliques peut être abordé par l'utilisation d'une méthode fournissant des requêtes agissant sur les bases de données. Cette nouvelle approche peut être traduite par les réseaux formels [CAD 96].

Ayant à notre disposition cette nouvelle représentation, une évaluation des performances de chacun des partenaires de la filière THD, révélée par des indicateurs standards contenus dans les tableaux de bord [SAU 84], permet d'identifier les dysfonctionnements subsistants dans le cadre de l'application d'une démarche "Quick Response". Le simulateur permet de visualiser les dysfonctionnements sur la base de stratégies prévisionnelles mais la pertinence de l'analyse est liée au modèle de prévision adapté. C'est pourquoi, nous présentons dans le chapitre suivant les modèles les plus couramment utilisées dans l'industrie.

Chapitre 2

Modèles et méthodes de prévision des ventes

I. Introduction

Pour faire face à la demande client, gérer sa production et ses stocks, mais aussi orienter sa politique commerciale (prix, marketing, produits, etc.), une entreprise peut s'intéresser aux prévisions des ventes [MEL 90]. Dans le cadre d'une démarche « Quick Response », le recours aux méthodes de prévision constitue l'une des étapes d'évolution des liens entre partenaires et permet d'ajuster au plus juste les ressources financières et les capacités de production en fonction de la demande du client.

En général dans l'industrie textile, le planning de la production est établi sur un horizon prévisionnel compris entre trois et douze mois. L'affectation des ressources de production s'effectue en deux étapes. La première affectation, dont la valeur est proportionnelle à la quantité totale prévue, permet la mise en place des produits aux différents points de vente. Cette première étape peut être issue d'une démarche prévisionnelle à moyen terme. La deuxième étape consiste à organiser un approvisionnement de la quantité totale restante entre le fournisseur et le distributeur. L'application d'une démarche « Quick Response » repose sur une politique d'approvisionnement en « juste à temps » de la quantité nécessaire et suffisante des produits, suivant une fréquence de livraison définie. Cette politique sollicite fortement les capacités de réactivité du fournisseur. Afin de gérer correctement les ressources de

production, l'utilisation d'un stock tampon peut remédier au manque de réactivité. La taille de celui-ci peut être évaluée par l'utilisation de méthodes prévisionnelles adaptées à l'environnement à court terme

Les contraintes d'utilisation des méthodes de prévision dans l'environnement de vente à court terme des articles textiles sont multiples :

- les nombres élevés de références dus à la nature des articles dont la désignation additionne souvent un modèle, un type, une couleur et une taille,
- les quantités à mettre en production varient de quelques dizaines à plusieurs milliers de produits,
- les durées de vie des articles varient de quelques semaines à plusieurs saisons de vente,
- les influences prononcées des facteurs de différentes natures : économiques, climatiques, liées au calendrier, etc., modifient le comportement de vente des articles textiles,
- les renouvellements partiels ou totaux des articles textiles d'une collection soumise à l'influence du phénomène de mode.

Face à toutes ces contraintes, le choix et l'utilisation du modèle de prévision adéquat et adapté nécessitent une connaissance globale de l'ensemble des méthodes disponibles. L'environnement de vente des articles textiles présente des spécificités propres qui sont souvent à l'origine des difficultés rencontrées lors de l'utilisation des méthodes de prévision sur des données réelles de vente. Une connaissance du contexte économique de l'industrie textile pour évaluer les contraintes potentielles à l'utilisation d'un modèle de prévision est abordée dans ce chapitre.

Ainsi, la première partie de ce chapitre est consacrée à un examen des différentes méthodes usuelles de prévision des ventes et de leurs différentes caractéristiques. La deuxième partie de ce chapitre concerne le problème du choix d'un modèle prévisionnel et les différents critères adéquats à mettre en œuvre pour optimiser la prise de décision [WHE 83] : l'objectif est d'examiner l'ensemble des indicateurs de performance disponibles permettant d'évaluer et de comparer les méthodes de prévision entre elles. La troisième partie présente une analyse des différents travaux effectués sur la comparaison et l'évaluation des modèles de prévision afin d'apporter un ensemble d'indications utiles au choix d'une méthode. La quatrième partie présente différentes applications de modèles de

prévision en fonction d'un objectif précis et apporte un éclairage supplémentaire sur ces contraintes. La majorité des différentes applications est effectuée sur des séries temporelles (ou processus de vente) relatives à des articles textiles. Une série temporelle caractérise une suite d'observations répétées, correspondant à des dates différentes, d'un phénomène de vente. Chacune des applications précédentes est relative à des objectifs de prévision précis et classée en fonction de l'horizon prévisionnel tels que :

- les méthodes adaptées à la vente des articles soumis à l'influence de la mode,
- les méthodes de prévision à court terme adaptées pour les ventes d'articles textiles soumis à de fortes variations,
- les méthodes de prévision ajustées pour les articles saisonniers et permettant le regroupement des coefficients saisonniers,
- les méthodes basées sur l'introduction des variables marketing textiles,
- les méthodes de prévision adaptées à l'estimation des quantités nécessaires à la mise en place de collection,
- les méthodes de prévision à horizon moyen et long terme pour optimiser l'achat de la matière première textile,

Enfin, une application des principales méthodes adaptées à l'environnement à court terme aux données textiles d'une entreprise et une comparaison de ces procédures prévisionnelles entre elles par des mesures de l'erreur différentes fournissent des indications sur leur choix en fonction du contexte économique.

II. Présentation des méthodes de prévision usuelles

L'objectif principal de cette présentation consiste à examiner les différents types de modèle les plus couramment utilisés et d'exposer brièvement leur méthodologie. Chacune de ces méthodes repose sur la définition d'un type de modèle issu de l'analyse des séries temporelles. Différents critères permettent de distinguer les méthodes de prévision et d'établir un classement en considérant soit l'horizon prévisionnel, soit le type de modèle ou soit le domaine d'application [BRO 88][CHA 88][GOU 90][BOW 90]. La présentation des différents types des méthodes de prévision est effectuée suivant un ordre décroissant de leurs utilisations au sein des entreprises textiles. La première catégorie de méthodes les

plus couramment rencontrées sont de nature intuitives et personnalisées à la politique commerciale et logistique de chaque entreprise. Une deuxième catégorie de méthodes, souvent utilisées, se base uniquement sur des processus endogènes qui caractérisent des phénomènes de vente liés aux propres paramètres du produit. Une troisième catégorie de méthodes permet de prendre en compte des processus exogènes qui caractérisent l'influence de phénomènes de vente indépendants de la nature du produit. Enfin, la dernière catégorie englobe l'ensemble des méthodes dont la représentation aboutit à une modélisation non linéaire du processus de vente.

II.1. Notations

Dans un souci d'uniformité de la représentation des différents paramètres utilisés dans ce chapitre, les notations sont les suivantes :

x_t caractérise la valeur à l'instant t du processus de vente (X_t).

$\hat{x}_{t(h)}$ représente la valeur de la prévision calculée à l'instant t pour un horizon de h périodes pour le processus de vente (X_t).

F_t désigne la moyenne lissée (sans résidu) de la série calculée à l'instant t du processus de vente (X_t).

S_t désigne le coefficient saisonnier calculé à l'instant t du processus de vente (X_t).

T_t désigne la tendance calculée à l'instant t du processus de vente (X_t).

Soit l'opérateur de retard L défini tel que : $L^k X_t = X_{t-k}$, $\forall t \in \mathbb{Z}, \forall k \in \mathbb{N}$.

e_t représente l'erreur de prévision calculée à l'instant t et définie telle que :

$$e_t = x_t - \hat{x}_{t-1(1)} .$$

II.2. Méthodes intuitives

Beaucoup plus que sur la statistique, ces techniques sont basées sur l'intuition et sur les éléments que les « managers » ressentent comme devant faire partie du processus de prévision.

Beaucoup d'études ont recensé les utilisations relatives des méthodes subjectives et objectives dans la pratique [CER 75][LAW 83][MEN 84][SPA 84][DAL 87][TAR 89][BAT 90][SAN 94]. L'examen de toutes ces études indique de façon remarquable la forte utilisation, dans 40 à 50 % des entreprises, des techniques de prévision intuitives [DAL 87][TAR 89]. Une étude comparative de Webby et O'Connor [WEB 96] sur les

différents travaux concernant l'importance des facteurs humains d'interactions dans les approches statistiques et subjectives de la prévision des séries temporelles a permis de démontrer le rôle crucial de la connaissance du contexte de la prévision des séries temporelles sur la précision. Webby et O'Connor suggèrent également que la précision globale des méthodes intuitives peut être grossièrement équivalente à celle des méthodes statistiques, mais la contribution majeure des méthodes subjectives réside dans leurs capacités à intégrer ces informations, non relatives aux séries temporelles, dans les prévisions.

Fort de ce constat, une première recommandation dans le choix d'un modèle de prévision consiste à exploiter les ressources existantes de l'entreprise en terme de savoir faire et de connaissance du marché. Cependant, la non formulation mathématique de ce type de modèle empêche tout traitement informatique et automatique des prévisions. Dans le cas d'une gestion de plusieurs milliers d'articles, le recours à une méthode automatique de prévision s'avère indispensable pour traiter une par une chacune des références des produits.

II.3. Méthodes endogènes

Un processus de vente est de nature endogène si son évolution est dépendante de la nature même du produit. L'utilisation d'une nomenclature détaillée et précise des caractéristiques endogènes de chaque produit permet d'améliorer la désignation de l'article et de préciser la nature d'un même comportement de vente.

II.3.1. Principe de décomposition d'une série chronologique.

Les méthodes de prévision peuvent admettre l'existence d'une loi comme base dans l'établissement d'une prévision. L'aptitude d'une technique donnée à fournir une prévision effective dans une situation particulière dépend donc largement de l'adaptation entre la loi qui régit cette situation et la technique susceptible de s'accommoder de cette loi.

La série chronologique devient par supposition le résultat de la juxtaposition de plusieurs composantes dont les 3 formes principales sont :

- la tendance (T_t) qui marque l'allure générale du phénomène de vente et ses variations à moyen et long terme.

- les variations saisonnières (S_t), essentiellement liées au rythme imposé par les saisons météorologiques, que ce soit directement (températures, précipitations, etc.) ou indirectement par les activités économiques et sociales (fêtes, vacances, soldes, mode, etc.).
- les variations accidentelles ou erreurs (e_t) qui résultent de multiples causes et dont l'effet est souvent de faible intensité et de courte durée.

Les notions de tendance et de «saisonnalité» sont d'usage courant dans l'analyse des comportements de vente des articles textiles. Leur utilisation dans une méthode de prévision oblige à connaître précisément leur signification afin d'éviter toute interprétation. Une analyse précise et détaillée des phénomènes à l'origine de leur constitution, ainsi qu'une distinction établie suivant leur horizon d'application, contribuent à une meilleure connaissance des notions de tendance et "saisonnalité".

(1) Notion de tendance dans le comportement de vente des articles textiles.

Deux grands modes de constitution des tendances peuvent être distingués; d'une part, la tendance de mode résultante de la constitution de l'offre et, d'autre part, la tendance du marché résultante de la segmentation de la demande. La tendance de mode résulte de la circulation d'une information de mode dépendante d'un ou plusieurs émetteurs simultanés et d'un certain nombre de relais. L'analyse des tendances au travers des émetteurs et/ou des relais permet de dégager les différents moyens d'action disponibles quant à la constitution de l'offre, mais aussi d'effectuer une segmentation des marchés au travers des comportements vestimentaires.

(2) Notion de "saisonnalité" dans le comportement de vente des articles textiles.

La notion de «saisonnalité» peut être caractérisée par l'influence systématique de l'ensemble des facteurs climatiques, liés au calendrier ou autres, dont les effets sont cycliques. A chaque période de l'année est associé un coefficient propre appelé coefficient saisonnier S_t . Les principaux facteurs à l'origine des oscillations saisonnières des ventes sont :

- les différences climatiques qui sont à l'origine de variations sensibles dans la consommation.

- les fêtes mobiles telles que Pâques ou Pentecôte qui déterminent les ventes de certains produits. Ces fêtes ne sont pas toujours fixes dans le calendrier et entraînent des variations saisonnières "flottantes" d'un mois sur l'autre.
- le calendrier de l'année en cours donne le nombre de jours ouvrables.
- les facteurs psycho-sociologiques qui influencent le comportement des consommateurs.

II.3.2. La méthode de prévision par la moyenne mobile

Les prévisions par la méthode des moyennes mobiles sont utilisées à différentes fins; principalement, pour lisser une série chronologique et extraire le niveau moyen de la série.

(1) La prévision par la moyenne mobile simple

Une formulation de la prévision à l'instant t d'horizon 1 à partir d'une moyenne mobile d'ordre p (caractérisant la taille de la fenêtre d'observation), avec $p \in \mathbb{N}^*$, de la série (X_t) peut s'écrire :

$$\hat{x}_{t(1)} = \frac{1}{p} \sum_{i=0}^{p-1} x_{t-i} \quad (1)$$

Ce schéma n'est valide que si la chronique est stationnaire, c'est à dire dépourvue de tendance.

(2) La prévision par la moyenne mobile double

Dans le cas de la présence d'une tendance marquée, un nouvel opérateur de lissage par moyenne double doit être introduit et défini par $\hat{y}_{t(1)}$ tel que :

$$\hat{y}_{t(1)} = \frac{1}{p} \sum_{i=0}^{p-1} \hat{x}_{t-i(1)} \quad (2)$$

La prévision calculée à l'instant t pour un horizon de h périodes peut s'écrire suivant le modèle linéaire suivant (la démonstration est disponible dans [MAK 83]) :

$$\hat{y}_{t(h)} = \hat{a}_t + \hat{b}_t h \quad (3)$$

$$\text{avec } \hat{a}_t = 2\hat{x}_{t(1)} - \hat{y}_{t(1)} \text{ et } \hat{b}_t = \frac{2}{p-1}(\hat{x}_{t(1)} - \hat{y}_{t(1)})$$

La fonction de prévision possède une forme localement linéaire dont les coefficients $\hat{a}_t$ et $\hat{b}_t$ sont calculées pour minimiser la différence entre la valeur réelle calculée à l'instant $t-h$ et la prévision faite à l'instant t pour l'horizon h par la méthode des moindres carrés. Les formules de $\hat{x}_{t(1)}$ (1) et $\hat{y}_{t(1)}$ (2) sont utilisées pour initialiser les coefficients $\hat{a}_t$ et $\hat{b}_t$ du modèle de prévision $\hat{y}_{t(h)}$ (3).

De façon générale, la méthode de prévision par moyennes mobiles permet de résoudre en première approximation les problèmes posés par l'analyse des séries chronologiques.

II.3.3. Les méthodes de prévision par lissage exponentiel

Les méthodes de prévision par le lissage exponentiel simple et double ont été initialement introduites par Brown [BRO 59].

(1) Classification des méthodes de prévision par lissage exponentiel

Différentes formulations du lissage exponentiel existent dans la littérature [ABR 83][GAR 85]. Par exemple, la classification due à Pegels [PEG 69] comporte 9 cas, selon le type de composante de tendance et de composante saisonnière. Les composantes T_t et S_t sont obtenues par lissage et traduites par les relations suivantes, avec $(\alpha, \nu, \delta) \in [0,1]$:

$$(\text{tendance additive}) \quad T_t = \nu(F_t - F_{t-1}) + (1-\nu)T_{t-1}$$

$$(\text{tendance multiplicative}) \quad T_t = \nu\left(\frac{F_t}{F_{t-1}}\right) + (1-\nu)T_{t-1}$$

$$(\text{saisonnier additif}) \quad S_t = \delta(x_t - F_t) + (1-\delta)S_{t-s}$$

(saisonnier multiplicatif) $S_t = \delta(\frac{x_t}{F_t}) + (1-\delta)S_{t-s}$

La moyenne F_t est toujours calculée par une expression du type :

$$F_t = \alpha P_t + (1-\alpha)Q_t \tag{4}$$

où P_t et Q_t sont définis dans le Tableau 15 ci-dessous.

Composante de tendance		Composante saisonnière		
		1 (aucune)	2 (additive)	3 (multiplicative)
A (aucune)	P_t	x_t	$x_t - S_{t-s}$	$\frac{x_t}{S_{t-s}}$
	Q_t	F_{t-1}	F_{t-1}	F_{t-1}
B (additive)	P_t	x_t	$x_t - S_{t-s}$	$\frac{x_t}{S_{t-s}}$
	Q_t	$F_{t-1} + T_{t-1}$	$F_{t-1} + T_{t-1}$	$F_{t-1} + T_{t-1}$
C (multiplicative)	P_t	x_t	$x_t - S_{t-s}$	$\frac{x_t}{S_{t-s}}$
	Q_t	$F_{t-1} \times T_{t-1}$	$F_{t-1} \times T_{t-1}$	$F_{t-1} \times T_{t-1}$

Tableau 15. Classement des différents types de lissage exponentiel

Le cas A-1 correspond au lissage exponentiel simple, le cas B-1 à la méthode de Holt. Les cas B-2 et B-3 représentent respectivement les méthodes de Winters additives et multiplicatives; mais, le lissage exponentiel double n'apparaît pas. Une autre classification plus vaste, élaborée par Roberts [ROB 82], tient compte uniquement des modèles purement additifs. Enfin, Broze et Mélard [BROz 88] ont effectué un recensement, non exhaustif, des modèles de lissage exponentiel et aboutissent à 27 formulations différentes. Notre étude se limite aux méthodes les plus courantes telles que :

- les méthodes du lissage exponentiel simple et double,
- les méthodes de Holt-Winters additives et multiplicatives.

(2) La prévision par lissage exponentiel simple

Le modèle de prévision du lissage exponentiel simple apparaît comme une moyenne pondérée par le coefficient de lissage α de la dernière réalisation et de la dernière prévision. Soit la relation suivante :

$$\hat{x}_{t(1)} = \alpha x_t + (1-\alpha) \hat{x}_{t-1(1)} \quad (5)$$

avec $F_t = \hat{x}_{t(1)}$ le cas A1 du Tableau 15 est obtenu tel que :

$$F_t = \alpha x_t + (1-\alpha) F_{t-1}$$

(3) La prévision par lissage exponentiel double

La formule précédente permet de calculer une prévision pour des séries chronologiques stationnaires, sans tendance. Dans le cas d'une série avec tendance, un lissage exponentiel double $\hat{y}_{t(1)}$ peut être défini par :

$$\hat{y}_{t(1)} = \alpha \hat{x}_{t(1)} + (1-\alpha) \hat{y}_{t-1(1)} \quad (6)$$

La prévision calculée à l'instant t à l'horizon h peut s'écrire en fonction du modèle linéaire suivant :

$$\hat{y}_{t(h)} = \hat{a}_t + \hat{b}_t h \quad (7)$$

$$\text{avec } \hat{a}_t = 2\hat{x}_{t(1)} - \hat{y}_{t(1)} \text{ et } \hat{b}_t = \frac{\alpha}{1-\alpha} (\hat{x}_{t(1)} - \hat{y}_{t(1)}).$$

La démonstration du calcul des formes des coefficients $\hat{a}_t$ et $\hat{b}_t$ est disponible dans [MEL 90].

(4) Choix du coefficient de lissage α :

Une première méthode consiste à sélectionner des critères d'évaluation d'une méthode de prévision et d'effectuer des simulations sur les différentes valeurs de α . Il suffit alors de choisir la valeur optimale en fonction d'une valeur de α qui améliore le

mieux le critère choisi [MEL 90]; par exemple, la détermination du coefficient de lissage qui minimise la somme des écarts au carré des erreurs de prévision est optimale au sens des moindres carrés. Cependant, la recherche du meilleur coefficient de lissage sur toute la période d'estimation n'implique pas d'être optimum quand la prévision est calculée. Pour supprimer cet inconvénient, des procédures dynamiques de régulation du coefficient de lissage ont été développées [CLA 88][BOU 92].

La valeur optimale instantanée résulte d'un compromis entre : l'inertie liée à l'intégration de données lointaines et la sensibilité aux valeurs récentes. En cas d'erreur de prévision constatée, deux interprétations sont possibles :

- il s'agit d'un accident; le coefficient de lissage doit alors diminuer afin de gommer l'effet de cette valeur anormale;
- il s'agit d'une rupture de tendance durable; le coefficient de lissage doit être augmenté afin d'intégrer plus rapidement cette rupture.

Le rôle de chacun des "signaux d'alertes" est clairement expliqué dans la procédure d'autorégulation du coefficient de lissage proposée par Bourbonnais et Usunier. Cette méthode de régulation permet d'optimiser localement la valeur du coefficient de lissage pour des séries de vente relativement stables [BOU 92] par l'utilisation des deux signaux d'alertes NF_t et AWS_t définis tels que :

$$NF_t = \frac{e_t}{MAD_t} \text{ et } AWS_t = \frac{\sum_{i=1}^t e_i}{MAD_t} \quad (8)$$

$$\text{avec } MAD_t = \frac{1}{t} \sum_{i=1}^t |e_i| \text{ et } MAD_t \neq 0.$$

La valeur MAD_t correspond à la moyenne des erreurs de prévision en valeur absolue calculée jusqu'à l'instant t . Le signal d'alerte NF_t permet de détecter les observations anormales de la série alors que le signal d'alerte AWS_t permet d'identifier un changement de tendance.

La procédure d'autorégulation du coefficient de lissage exponentiel double α peut se résumer ainsi [BOU 92] :

- Une valeur initiale du coefficient de lissage $\alpha=\alpha_0$ est choisie,
- Si la valeur observée du processus de vente est anormale, la valeur absolue du signal d'alerte NF_t augmente et, par conséquent, la valeur de α doit être diminuée d'une valeur de pas choisie $\Delta\alpha$ telle que $\alpha = \alpha - \Delta\alpha$,
- Si un changement de tendance est observé, la valeur absolue du signal AWS_t augmente, alors la valeur de α doit être augmentée de la même valeur de pas $\Delta\alpha$ telle que: $\alpha = \alpha + \Delta\alpha$;
- Si les valeurs absolues de NF_t et AWS_t sont relativement stables à une valeur de seuil près (valeur recommandée par l'auteur [BOU 92] alors α reste inchangée.

(5) Méthode de prévision utilisant les notions de tendance et de «saisonnalité»

Parmi l'ensemble des méthodes permettant d'établir un modèle de prévision sur les notions de tendance et de «saisonnalité», la méthode de Holt-Winters [WIN 60] s'avère être la plus connue et la plus utilisée. Cette méthode permet de combiner linéairement la tendance et la composante saisonnière suivant différents formats :

- par addition, c'est le cas du modèle de Winters additif
- par multiplication, c'est le cas du modèle de Winters multiplicatif

L'application de cette méthode nécessite un historique des coefficients saisonnier de la saison de vente précédente. Soient S la périodicité de la «saisonnalité» des données et α, γ, δ trois valeurs réelles positives comprises entre 0 et 1 caractérisant des coefficients de lissage.

(a) Modèle additif

Les trois composantes de la série temporelle sont calculées à chaque instant t par les relations suivantes.

Lissage de la moyenne.

$$F_t = \alpha(x_t - S_{t-s}) + (1-\alpha)(F_{t-1} + T_{t-1}) \quad (9)$$

Lissage de la tendance.

$$T_t = \gamma(F_t - F_{t-1}) + (1-\gamma)T_{t-1} \quad (10)$$

Lissage de la "saisonnalité".

$$S_t = \delta(x_t - F_t) + (1-\delta)S_{t-s} \quad (11)$$

La valeur de la "saisonnalité" est calculée à la période t de la saison de vente courante afin d'être intégrée aux prévisions futures effectuées à la prochaine saison de vente.

La valeur prévue à l'instant t pour un horizon de h périodes est calculée selon le modèle additif suivant :

$$\hat{x}_{t(h)} = F_t + hT_t + S_{t+h-s} \quad (12)$$

En plus de la valeur prévisionnelle, un intervalle de prévision pour la méthode de Holt-Winters de type additif avec "saisonnalité" peut être calculé en fonction d'un seuil de probabilité fixé par l'industriel [YAR 90].

(b) *Modèle multiplicatif*

Les relations permettant de calculer les trois composantes de la série temporelle à intégrer dans le modèle multiplicatif sont les suivantes :

Lissage de la moyenne.

$$F_t = \alpha \frac{x_t}{S_{t-s}} + (1-\alpha)(F_{t-1} + T_{t-1}) \quad (13)$$

Lissage de la tendance.

$$T_t = \gamma(F_t - F_{t-1}) + (1-\gamma)T_{t-1} \quad (14)$$

Lissage de la "saisonnalité".

$$S_t = \delta \frac{x_t}{F_t} + (1-\delta)S_{t-s} \quad (15)$$

La valeur prévue à l'instant t pour un horizon de h périodes est calculée selon le modèle multiplicatif suivant :

$$\hat{x}_{t(h)} = (F_t + hT_t)S_{t+h-s} \quad (16)$$

De même que pour le modèle additif, un intervalle de prévision pour la méthode de Holt-Winters de type multiplicatif avec «saisonnalité» peut être calculé en fonction d'un seuil de probabilité [CHA 91].

(6) Procédure d'auto-régulation de la méthode de Holt-Winters [VRO 96][VRO 97]

Une procédure d'auto-régulation des paramètres de lissage α et β de la méthode de Holt-Winters a été étudiée par Vroman [VRO 96]. Les performances du modèle de prévision visent deux objectifs :

- obtenir une réaction rapide face à une variation significative de la tendance et/ou de la "saisonnalité"
- reconnaître et "lisser" les événements qui sont purement aléatoires

Ainsi, cette approche consiste à observer la courbe de demandes réelles, et d'adapter la méthode de calcul de la prévision aux résultats de cette observation. Cela implique la définition des paramètres d'observation du contexte de la demande, puis l'initialisation des paramètres de calcul de la prévision par l'utilisation de la théorie floue.

Le modèle de prévision de Holt-Winters et la procédure d'autorégulation FAHWS sont représentés en Figure 9 :

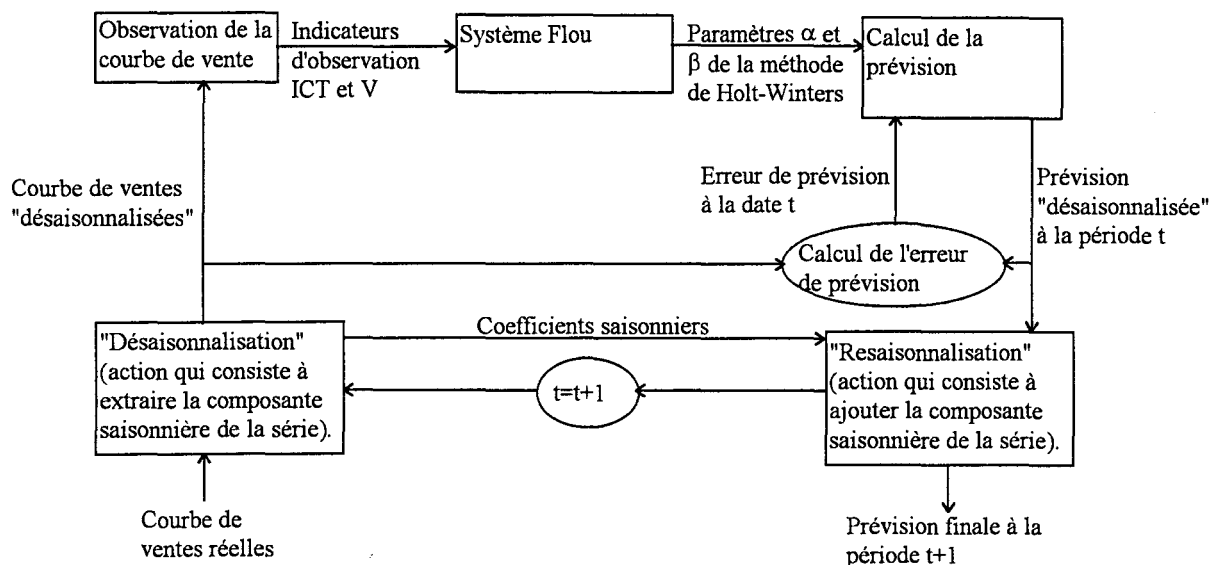


Figure 9. **Présentation de la procédure d'autorégulation des paramètres α et β de la méthode de Holt-Winters.**

A partir des données réelles, une procédure d'analyse du processus de vente permet d'identifier les deux indicateurs d'observation ICT (Indicateur de Changement de Tendence) et V (Variance) à intégrer dans le système flou. Les valeurs des paramètres de lissage α et β obtenues par le système flou permettent de calculer la valeur de la prévision suivant le modèle de Holt-Winters sur des données sans composante saisonnière. Une procédure de "resaisonnalisation" des données fournit la valeur de la prévision calculée à l'instant t pour l'instant $t+1$ (pour un horizon de prévision égal à 1). Les mises à jour des différents paramètres et des coefficients saisonniers sont effectuées à chaque instant t par le calcul de l'erreur de prévision et l'exploitation des données réelles de vente les plus récentes.

II.3.4. Méthode de prévision basée sur un modèle issu d'une analyse de série temporelle.

La méthode de Box et Jenkins [BOX 76] permet, en plusieurs étapes, de trouver un modèle ARMA susceptible de représenter une série chronologique. Des traitements préalables sur les données de vente permettent de rendre stationnaires et non saisonniers des processus de vente afin de pouvoir appliquer la méthode de Box et Jenkins. Cette procédure consiste à formuler des suppositions sous forme d'un modèle ARMA, à les mettre à l'épreuve et à réviser ensuite le modèle en conséquence; ces étapes sont répétées

autant de fois que nécessaire. Une fois le modèle ARMA connu, la détermination des prévisions peut se faire mécaniquement.

Cependant, cette méthode de prévision ne peut s'appliquer qu'à des chroniques relativement stables, peu sujettes aux à-coups de la conjoncture ou aux influences de la politique marketing de la firme. La méthode de Box-Jenkins reste d'une utilisation pratique extrêmement limitée dans le contexte de la prévision à court terme.

II.4. Méthodes exogènes

Une variable exogène est principalement caractérisée par son indépendance aux phénomènes des ventes et plus précisément aux paramètres endogènes des produits. L'introduction de variables exogènes pouvant avoir une influence sur les variables endogènes et, dont les valeurs sont fixées "extérieurement" à ce phénomène, permettent d'introduire un aspect explicatif dans le comportement de vente des articles. Une méthode d'identification des modèles ARMAX (modèle ARMA avec des variables exogènes X) a été proposée par Bierens [BIE 87]. Cependant, les phases d'analyse et d'estimation des paramètres des modèles de prévision basés sur une représentation ARMAX nécessitent une base de connaissance du comportement de vente sur au moins une saison afin d'effectuer une modélisation fiable du point de vue statistique. Or, l'identification d'un modèle de prévision selon une représentation ARMAX en tenant compte d'historiques des ventes relativement courts semble inadapté aux contraintes d'utilisation d'une procédure prévisionnelle dans le contexte économique textile.

L'utilisation de la régression linéaire multiple comme méthode de prévision à caractère purement exogène aux données de vente de l'industrie textile ne semble pas appropriée à une gestion de plusieurs centaines ou plusieurs milliers de produits. L'introduction des variables exogènes dans un contexte de prévision à court terme ne permet pas d'aboutir à des résultats significatifs [JOH 88][TEX 91]. Plusieurs raisons expliquent la non adaptation de ces méthodes à la prévision à court terme :

- Parmi les diverses séries chronologiques possibles, il est difficile d'isoler la (ou les) variable(s) ayant un pouvoir explicatif à court terme.
- La fonction exprimant la relation entre cette (ces) variable(s) et la série à analyser n'est pas connue. Sa définition demande en général une recherche longue, difficile et coûteuse.

- L'introduction d'une variable explicative dans l'analyse prévisionnelle à court terme est d'une telle complexité que le rapport coût/qualité de la prévision devient souvent prohibitif.

II.5. Méthodes non linéaires

La plupart des travaux sur l'analyse des séries temporelles ont été faites avec l'hypothèse que la structure de la série peut être décrite par des modèles linéaires de séries temporelles. Une insuffisance de représentation des modèles linéaires pousse à s'orienter vers les méthodes non linéaires. Des analyses comparatives sur les récents travaux réalisés dans le domaine de la prévision issue de modèle non linéaire, en insistant sur les notions de robustesses, ont été effectuées [TON 90][GOO 92] et aboutissent au même constat : les calculs d'optimisation des différents paramètres des modèles non linéaires reposent sur des estimations statistiques dont la fiabilité nécessite des historiques de vente sur au moins une saison. Une approche de prévision non linéaire par les réseaux de neurones aboutit à des résultats au moins aussi précis que ceux issus de techniques plus classiques pour des séries temporelles stationnaires, sans tendance, sans «saisonnalité», et dont les historiques de vente des séries temporelles testées s'étendent sur plusieurs dizaines d'années [KOU 93]. De même, l'utilisation d'une méthode de filtrage non linéaire permettant d'évaluer la tendance d'une série temporelle fortement "bruitée" aboutit à une valeur de l'erreur quadratique moyenne inférieure à celle d'une moyenne mobile classique [FIZ 86][HUR 89].

III. Critères d'un modèle de prévision

Ayant à sa disposition un aperçu des méthodes de prévision disponibles et adaptées à l'environnement de vente des articles textiles, une préoccupation majeure consiste à choisir la "meilleure" méthode. Tout « manager », qui récemment a dû prévoir avant de décider, est bien au fait de l'importance que revêt le choix de la technique de prévision la mieux adaptée à sa propre situation. Cette prise de décision résulte de la combinaison de plusieurs critères de choix relative aux contraintes d'utilisation des méthodes de prévision dans un environnement de vente d'articles textiles.

III.1. Choix d'une méthode

Les six facteurs dominants dans le choix d'une technique de prévision sont :

- l'horizon temporel : deux aspects de l'horizon temporel sont à considérer dans une méthode de prévision prise individuellement; le premier caractérise la durée de la projection dans le futur et le second concerne le nombre souhaité de périodes pour établir une prévision.
- la loi d'évolution des données : la plupart des méthodes de prévision repose sur une hypothèse sous-jacente de la loi d'évolution des données à prédire.
- le type de modèle : en plus d'une hypothèse de loi de base sous-jacente aux données, un modèle relatif à la situation à prévoir semble plus approprié que d'autres procédures.
- le coût : la mise en œuvre d'une application d'un procédé de prévision engendre des coûts relatifs au développement de la méthode, au stockage des données historiques et prévisionnelles de vente, et à leurs temps d'exploitation.
- la précision : les exigences de précision sont en rapport direct avec le degré de détail recherché.
- la facilité d'application : la compréhension, de la part des décideurs, des opérations relatives au traitement des données prévisionnelles, facilite l'intégration des procédures de prévision comme outils efficaces d'aide à la décision.

Le fait que certaines techniques sont beaucoup plus appropriées aux prévisions à court terme, tandis que d'autres soient capables de traiter efficacement des prévisions à long terme, est dû à la caractéristique d'horizon temporel identifiable dans toute méthode de prévision. Cette caractéristique est en relation étroite avec les moyens qu'utilise chaque technique pour établir sa prévision et avec la quantité de données dont elle a besoin. Principalement, un historique de vente représente les éléments de base d'une analyse méthodique permettant d'établir un modèle de prévision et oblige à définir certaines propriétés relatives à son utilisation :

- il ne se compose que de valeurs connues et calculées, qui sont effectivement réalisées.
- il est homogène dans le temps. Pour respecter cette condition d'homogénéité, une solution consiste à suivre à la fois les deux séries chronologiques, la nouvelle et l'ancienne, afin que, lorsque le nouvel historique est suffisant, il puisse être utilisé à son

tour. L'historique sur lequel porte la prévision est alors l'agrégat des deux séries. Cette solution est préconisée dans le cas de sociétés modifiant fréquemment les codes d'articles lors de la création de nouveaux produits venant remplacer des références similaires.

- il comprend un certain nombre de valeurs. Il semble logique d'écrire que plus l'historique est long, meilleure est la qualité de l'analyse et par voie de conséquence, la prévision.

Le critère le plus important pour la sélection d'un modèle de prévision est la précision d'une procédure de prédiction [CAR 82][MAH 86][COL 92a][WIT 92][YOK 95]. Cependant, d'autres critères comme étant comparables en importance tels que : l'interprétation d'un critère, la flexibilité d'utilisation, la disponibilité des variables et l'exploitation d'une méthode de prévision sont souhaitables pour la sélection et l'évaluation des techniques de prévision. Les comparaisons entre les techniques de prévision doivent tenir compte d'un ensemble de critères de sélection permettant d'être combinés et pondérés selon l'objectif recherché dans le choix multicritères d'un modèle. De façon générale, les praticiens ont classé les critères relatifs à l'application et à l'implantation d'un système de prévision comme la flexibilité, les facilités d'utilisation et d'implantation, de façon plus performante que ne l'ont fait les experts en prévision.

III.2. Evaluation d'une méthode

Les mesures "objectives" de la qualité d'une prévision dépendent du but recherché en terme de contrainte d'utilisation d'un modèle dans un environnement particulier. Les différentes formulations des critères d'évaluation permettent une analyse détaillée des méthodes de prévision en terme de précision, de stabilité aux données aberrantes, de réaction aux changements rapides, etc..

Soient les « n » erreurs de prévision $e_1, \dots, e_n$ de la méthode de prévision à comparer et les " n " erreurs de prévision $e_1^{RW}, \dots, e_n^{RW}$ pour la méthode de prévision de type "marche aléatoire" (Random Walk).

Les mesures suivantes de l'erreur sont généralement utilisées pour évaluer les résultats des méthodes de prévision.

- l'erreur quadratique moyenne (Root Mean Square Error : RMSE)

$$\text{RMSE}(e) = \sqrt{\frac{1}{n} \sum_{i=1}^n e_i^2} \quad (17)$$

Contrairement au critère MSE (Mean Square Error : $\text{MSE} = \text{RMSE}^2$), le critère RMSE est exprimé dans les mêmes unités que les données.

- l'erreur absolue moyenne en pourcentage (Mean Absolute Percentage Error : MAPE)

$$\text{MAPE}(e) = \frac{1}{n} \sum_{i=1}^n \frac{|e_i|}{x_i} \quad (18)$$

Ce critère correspond à la fonction de coût en valeur absolue et en pourcentage; c'est un nombre sans dimension et il faut que $x_i > 0$.

- l'erreur absolue moyenne médiane en pourcentage (Median Absolute Percentage Error : MdAPE)

$$\begin{aligned} \text{si } n=2q+1; \text{ MdAPE}(e) &= \frac{|e_{\frac{n+1}{2}}|}{x_{\frac{n+1}{2}}} & \text{si } n=2q; \text{ MdAPE}(e) &= \frac{|e_{\frac{n}{2}+1}|}{x_{\frac{n}{2}+1}} \end{aligned} \quad (19)$$

- l'erreur absolue relative moyenne harmonique (Geometric Mean of the Relative Absolute Error : GMRAE)

$$\text{GMRAE}(e, e^{\text{RW}}) = \left(\prod_{i=1}^n \frac{|e_i|}{|e_i^{\text{RW}}|} \right)^{\frac{1}{n}} \quad (20)$$

- l'erreur absolue relative moyenne médiane (Median Relative Absolute Error : MdRAE)

$$\begin{aligned} \text{si } n=2q+1; \text{ MdRAE}(e, e^{\text{RW}}) &= \frac{|e_{\frac{n+1}{2}}|}{|e_{\frac{n+1}{2}}^{\text{RW}}|} & \text{si } n=2q; \text{ MdRAE}(e, e^{\text{RW}}) &= \frac{|e_{\frac{n}{2}+1}|}{|e_{\frac{n}{2}+1}^{\text{RW}}|} \end{aligned} \quad (21)$$

Dans un souci d'uniformiser les comparaisons entre les méthodes de prévision, une méthode statistique a été développée sur les bases du critère MSE : le logarithme moyen du rapport de l'erreur au carré (LMR). La statistique (LMR) nécessite une quantité de mesures significatives et ne semble pas appropriée à un contexte d'évaluation à court terme [THO 90].

Pour fournir une aide au choix des mesures de l'erreur, un classement de chacune de ces mesures a été effectué suivant différents critères d'évaluation [ARM 92]. Le critère de fiabilité mesure la capacité d'une méthode de prévision à fournir les mêmes résultats en considérant les mêmes données. Le critère de validité de la modélisation permet d'évaluer la procédure d'estimation de la précision des méthodes d'extrapolation. Les résultats relatifs à ces deux critères sont issus d'une procédure empirique de comparaison basée sur la corrélation moyenne entre les différents résultats d'évaluation de l'erreur de prévision obtenue sur chacune des séries temporelles. Deux autres critères subjectifs sont également utilisés : la sensibilité de la précision d'une méthode de prévision aux changements de l'un des paramètres du modèle et la facilité d'interprétation d'une mesure de l'erreur de prévision.

De façon générale, les mesures de l'erreur RMSE (17), MAPE (18) et GMRAE (20) permettent d'évaluer correctement les changements obtenus sur la précision lors d'une modification d'un des paramètres du modèle de prévision. La mesure de l'erreur GMRAE est préférable car la mesure de l'erreur RMSE ne s'avère pas fiable et la formulation de la mesure de l'erreur MAPE introduit une erreur systématique pour les faibles valeurs de la prévision [ARM 92]. Pour sélectionner la méthode de prévision la plus précise parmi l'ensemble des procédures évaluées, le critère de l'erreur médiane (MdRAE (21)) est recommandé quand peu de séries sont disponibles. Quand beaucoup de séries temporelles sont disponibles pour évaluer le modèle de prévision, le critère de fiabilité devient de moindre importance et la mesure de l'erreur MdAPE (19) semble appropriée pour sélectionner les méthodes les plus précises [ARM 92].

De façon globale, afin d'évaluer la validité d'un modèle de prévision, il est souhaitable de considérer un ensemble de critères de mesure, aussi bien objectif que subjectif [CHA 92].

IV. Comparaison entre les différentes méthodes de prévision et validation des modèles pour un horizon de prévision h fixé

Le classement des différentes méthodes existantes de prévision peut s'opérer selon plusieurs critères. Beaucoup de ces classifications sont axées sur la technique : par exemple, la distinction entre méthodes statistiques et non statistiques ou entre méthodes endogènes et méthodes exogènes. Chambers [CHA 71] propose un cadre quelque peu différent, fondé davantage sur l'utilisation fonctionnelle de la prévision que sur les caractéristiques mathématiques des techniques. Le critère de classification se base sur le concept marketing du cycle de vie du produit pour identifier les caractéristiques importantes de situations de prévision à différentes phases de ce cycle. Ces caractéristiques sont assorties aux différentes méthodologies pour choisir les plus appropriées à chaque phase. Finalement, l'ensemble des critères de classification permet d'acquérir une connaissance approfondie et détaillée de l'application des différentes méthodes de prévision en fonction de divers domaines d'utilisation. Ainsi, les résultats d'études comparatives sur les modèles de prévision en fonction du domaine d'application et mesurés suivant les critères d'évaluation en adéquation avec l'objectif d'utilisation du modèle prévisionnel permettent de sélectionner la procédure la plus adéquate parmi l'ensemble des méthodes testées.

Dans différentes études [KIR 66][LEV 67][LEN 95][KRA 72], trois techniques de méthodes de prévision sont comparées : la procédure par moyenne mobile, par lissage exponentiel et par la régression linéaire. En terme d'exactitude évaluée mois par mois, la méthode de prévision par la technique du lissage exponentiel donne les meilleurs résultats. Lorsque l'horizon est repoussé à six mois ou plus, le lissage exponentiel et la moyenne mobile donnent des résultats similaires. Le modèle de régression linéaire utilisé dans cette étude l'emporte, enfin, pour des prévisions à plus long terme (un an ou plus). Pour l'essentiel, ces recherches montrent que les modèles de lissage exponentiel sont en général supérieurs en précision pour la prévision à court terme.

D'après Lewandowski [LEW 74], la procédure d'analyse de Box-Jenkins peut être considérée, du point de vue méthodologique, comme une démarche statistiquement fiable pour l'analyse et la prévision de séries chronologiques mais son utilisation pour la

prévision à court terme des séries chronologiques reste limitée [WHE 74]. Newbold et Granger [NEW 74] ont constaté de faibles améliorations en terme de précision quand ils ont comparé les modèles ARIMA avec la méthode du lissage exponentiel de Holt-Winters. Toujours d'après Lewandowski [LEW 74], l'introduction d'une série exogène dans le cadre d'une analyse à court terme des séries chronologiques ne permet pas d'aboutir à des prévisions plus précises sauf si la série explicative est assez bien connue (tant en ce qui concerne son évolution passée que celle de son futur proche); la prise en compte d'une variable explicative peut présenter alors des avantages substantiels.

Makridakis et Hibon [MAK 79], Makridakis et al. [MAK 82][MAK 93], soutiennent fortement l'idée que les méthodes simples, sous entendu des méthodes d'exécution mécanique de la prévision (tel que le lissage exponentiel), fournissent de meilleurs résultats, en moyenne, que les spécifications plus complexes des modèles ARIMA. D'autres études empiriques aboutissent au même résultat appuyant le fait que les modèles simples des séries temporelles font au moins aussi bien que les méthodes sophistiquées [FIL 83][HUS 85][GEU 86][SCH 86][WAT 87][KOE 88][COL 92b]. La plupart de ces études reposent sur des comparaisons effectuées à partir d'ensembles de données différentes.

Bartolomei et Sweet [BAR 89] ont comparé les 2 modèles saisonniers additifs et multiplicatifs de la méthode de Holt-Winters en utilisant 47 séries temporelles issues de la base de données de la compétition M2 [MAK 84]. Bien que les modèles de Holt-Winters fournissent toujours de meilleurs résultats en terme d'ajustement de l'erreur au carré, ces modèles n'ont donné une meilleure prévision que dans 55% des cas des séries temporelles en comparaison avec les modèles généraux de lissage exponentiel. Miller et Liberatore [MIL 93] ont examiné les résultats de simulation sur les 3 paramètres de lissage exponentiel du modèle de Holt-Winters à partir d'une famille de 28 séries temporelles; en particulier, ils ont étudié si l'utilisation de paramètres permettant d'estimer la tendance améliore la précision. Ils concluent que le modèle de Holt-Winters à 4 paramètres (c'est à dire, le modèle saisonnier augmenté d'un quatrième paramètre d'estimation de la tendance) fournit les prévisions les plus précises pour les modèles évalués (principalement des méthodes de lissage exponentiel).

Dans une étude récente de Fildes et Makridakis [FIL 95], l'introduction d'une flexibilité additionnelle associée à la méthode de Box-Jenkins peut conduire, en moyenne, à la réduction de l'erreur de prévision.

De façon globale, selon le domaine d'application et l'objectif d'utilisation de la procédure de prévision, chaque méthode présente des avantages et des inconvénients :

- les méthodes de décomposition saisonnière sont simples dans leur principe et dans leur application, nécessitent que peu de paramètres et sont relativement résistantes vis à vis des données aberrantes. Leurs utilisations restent limitées quand les caractéristiques des composantes ne sont pas fortement représentées par les séries chronologiques et ne permettent pas d'inclure de l'information extérieure.
- les méthodes de lissage exponentiel sont très simples à comprendre et à mettre en œuvre mais le choix optimal de la méthode particulière n'est pas toujours évident. Il n'est pas non plus possible d'inclure de l'information extérieure et reste sensible, en terme de qualité de la précision, aux données aberrantes.
- les méthodes de prévisions basées sur la technique de la régression linéaire multiple, permettent d'incorporer des variables explicatives mais nécessitent l'application d'une procédure de choix sur les paramètres exogènes et une démarche de spécification du modèle en plusieurs étapes. Les calculs sont plus intensifs que dans les méthodes précédentes et leur utilisation à court terme semble inutile.
- les méthodes de prévision basées sur les modèles de séries chronologiques de type ARMA nécessitent une expertise dans la mise au point du modèle, qui une fois élaboré, permet de déterminer des prévisions de façon automatique et les intervalles de prévision correspondants. La quantité de calcul est importante et l'introduction d'information extérieure ne semble pas possible.

Or, il est évident que la prévision de plusieurs milliers de séries de ventes dans un contexte de gestion des stocks ou la prévision d'une variable d'importance pour l'entreprise ne sont pas des problèmes identiques. Lorsque beaucoup de séries chronologiques sont concernées, seule une méthode très rapide, quasi-automatique, peut être envisagée, par exemple le lissage exponentiel simple. Pour une prévision d'une variable d'intérêt stratégique, seule une méthode raisonnée peut être prise en considération. Même dans le cas où la prévision concerne de nombreuses séries,

l'application de méthodes raisonnées permet de choisir de manière appropriée la méthode automatique la plus avantageuse. L'emploi de techniques plus raffinées d'estimation des paramètres de la méthode de prévision rend possible l'amélioration de la qualité des modèles mais n'engendre pas obligatoirement une meilleure précision [BOU 92].

A la lecture de ce paragraphe, la méthode de prévision "idéale" semble ne pas exister. Partant de ce constat, Granger [GRA 80] propose de ne pas utiliser une seule technique de prévision mais une combinaison de méthodes et de comparer les résultats respectifs.

Nous proposons d'aborder dans le paragraphe suivant, une description du contexte de la prévision dans l'industrie textile qui permet d'identifier les contraintes spécifiques d'utilisation des modèles de prévisions et améliorer leurs précisions.

V. Méthodes de prévision spécifiques appliquées à l'industrie textile.

Face à toutes ces contraintes d'utilisation des méthodes de prévision dans l'environnement économique des ventes d'articles textiles, différentes méthodologies ont été développées dans la littérature. Dans la plupart des cas, ces méthodes permettent de répondre à une ou plusieurs contraintes par rapport à un contexte économique et un objectif de prévision précis.

V.1. Contexte de la prévision

Les techniques de prévision doivent rapidement fournir les informations nécessaires pour connaître le style et le besoin en vêtements recherchés par les consommateurs [CAN 80]. Les différentes contraintes en terme de logistique globale (production et distribution des produits) des acteurs de la filière de l'industrie du vêtement peuvent être traduites sous plusieurs formes :

- les confectionneurs de l'industrie textile disposent actuellement de 4 à 8 collections par an. Chaque collection implique un renouvellement total ou partiel de la gamme de produits. Le temps disponible pour la création, la planification, la production et la distribution devient de plus en plus court,

- les conditions économiques très fluctuantes de l'industrie du vêtement ont considérablement augmenté le risque d'atteindre des niveaux de stock élevés,
- les variations de la demande peuvent engendrer des niveaux importants de stocks d'articles finis,
- même si un produit a du succès, sa durée de vie peut être courte et oblige à une gestion précise des quantités à fabriquer,
- des sous-estimations des potentiels du marché conduisent à des ventes manquées qui ne peuvent pas être rattrapées dans un marché économique très rapide,
- les articles arrivés en fin de collection et non vendus ne peuvent être présentés pour la prochaine saison de vente,
- la concurrence peut très rapidement s'emparer d'un modèle à la mode ayant du succès et le reproduire très rapidement,
- les rabais, opérés sur les prix des articles pour réduire le stock, réduisent les profits et les intérêts financiers,
- l'absence d'historique de vente réduit fortement le choix des modèles de prévision,
- la prévision de commandes pour la totalité de la saison doit être calculée 6 mois à l'avance afin de lancer en fabrication les quantités optimales : des prévisions sous estimées entraînent un « sur-stock » qu'il faut solder en fin de saison,
- en cours de saison, la prévision de réassortiment ("le réassort") doit s'ajuster le plus finement à la demande,
- enfin, la multiplicité des références : article $\times$ coloris $\times$ taille (prévision des préférences par des couleurs et des achats par taille) sollicite fortement les capacités de traitement de l'information des entreprises.

L'environnement économique de l'industrie de l'habillement est particulièrement versatile. Les incertitudes de la fluctuation de la demande gouvernent le rythme de l'activité économique et fournissent des opportunités de vente conduisant soit à des succès soit à des échecs. L'intuition et la connaissance du marché ont toujours été un des ingrédients nécessaires au succès de la prévision des ventes [GAR 81]. Quand les transactions commerciales deviennent de plus en plus importantes en valeurs financières et en volume, les enjeux sont plus grands et les risques aussi. Dans ce cas, l'expérience

doit être structurée et mise au service de la gestion de l'entreprise pour atteindre l'objectif souhaité.

Dans l'environnement économique de la mode, l'une des principales tâches de la prévision est d'aider le décideur à l'éclairer sur les produits à fort potentiel de succès économique dans le cadre de la création de lignes de produits. De bonnes prévisions peuvent permettre de prendre des décisions le plus tôt possible dans la vie du produit et ainsi réduire les coûts des engagements financiers tels qu'ils peuvent être impliqués dans les réductions des stocks. Avec aucune connaissance spécifique du comportement du marché, une société peut planifier une première production représentant 30% de la valeur totale prévue de ses détaillants. De façon évidente, les articles avec un fort ou un faible succès peuvent être décelés mais leur identification n'est pas suffisante pour effectuer un renouvellement en conséquence. Dans le cas des produits textiles, les critères, liés au produit, susceptibles d'expliquer les ventes (c'est à dire finalement les réactions positives des consommateurs à l'offre) sont nombreux. Il s'agit notamment du coloris, de la texture, du prix, de la forme, etc. [BOU 92].

Une comparaison sur la forme des allures de vente des articles sur toute la saison précédente peut apporter quelques indications. Les allures peuvent être comparées avec des méthodes statistiques appliquées afin d'établir des regroupements entre les allures communes. La détermination des tendances des gammes de produits des indices saisonniers et cycliques permet de décomposer les courbes de vente des produits et de comparer les facteurs spécifiques à la demande. De telles données n'assurent pas d'effectuer une prévision du marché correcte mais les chances de s'en approcher sont grandes grâce à cette connaissance très fine du marché.

V.2. Présentation des différentes méthodes textiles existantes

Différentes applications de méthodes de prévision dans l'industrie textile sont présentées en fonction de leur horizon prévisionnel.

V.2.1. Horizon à court terme

(1) Influence de la mode

(a) Analyse de Garcia-Cuevas et al. [GAR 81]

Cette étude considère la nature, les conditions et les problèmes d'une prévision de la demande dans le domaine des articles de mode. Des critères sont établis pour effectuer la sélection des techniques de prévision. Un ensemble de résultats d'applications est donné à partir des valeurs de ventes d'articles de mode. Les résultats suggèrent que les méthodes de décomposition classique sont plus adaptées et relativement peu coûteuses en terme d'aide à la prévision pour la gestion des stocks dans l'industrie de l'habillement.

(b) Méthode MODEFA

Les séries chronologiques de produits soumises à l'influence de la mode sont souvent de très courtes durées. Leur traitement est délicat et volumineux. Ces séries sont de nature très irrégulières et très fortement saisonnières, ce qui complique encore l'emploi des méthodes conventionnelles de prévision [LEW 74]. Pour tenter de résoudre l'ensemble de ces problèmes, une analyse peut être effectuée sur la courbe des ventes cumulées, ce qui permet d'éliminer en partie les variations aléatoires de la série.

Une méthodologie de prévision des ventes, tenant compte de l'ensemble des contraintes précédemment citées, est décrite par Lewandowski [LEW 74] et se décompose en trois parties :

- le fichier spécifique (S.I.V. - Système d'Information des Ventes), contenant :
 - la matrice des analogies permettant la recherche des caractéristiques pour des articles similaires,
 - les valeurs des ventes annuelles et les profils saisonniers pour chacun des articles (actuels et ceux des trois dernières années),
 - les valeurs de l'analyse des ventes réelles et des "actions spéciales" des articles.
- le système d'information permet, soit par voie de recherche analogique (pour des produits similaires), soit par voie directe, de mettre en place les caractéristiques initiales au début de la saison.

- le système prévisionnel contrôle automatiquement les procédures « d'auto-adaptation ». A chaque période, des prévisions de ventes sont calculées à partir des valeurs actualisées et des nouvelles informations ("actions spéciales" ou autres).

L'avantage d'un tel système réside dans une détection très rapide des changements inattendus des ventes d'articles soumis à la mode, et permet un changement rapide de la politique d'approvisionnement.

Cependant, cette méthode n'est accompagnée d'aucunes précisions supplémentaires permettant de comprendre les formulations employées. Les seules applications réalisées et non comparées à d'autres méthodes reposent sur l'utilisation des données de vente des produits de l'industrie automobile, dont les contraintes de gestion des pièces détachées sont fortement similaires aux traitements de références des articles textiles.

(2) Ventes d'articles textiles soumis à de fortes variations

L'article de Staci Bonner [BON 96] relate brièvement les performances d'un logiciel basé sur l'exploitation des données des points de vente par liaison EDI permettant d'effectuer des prévisions pour les articles de faible durée de vie et d'historique de vente instable. La méthodologie de prévision repose sur l'exploitation des données et une méthode de reconnaissance de formes basée sur les allures de vente permettant le regroupement des articles dont les ventes sont similaires. Une interface graphique conviviale représente les allures de vente des produits, en tenant compte de différents facteurs tels que : la répartition géographique des points de vente, les événements marketing et la météorologie. Cependant, aucune explication sur les formulations utilisées dans la procédure de vente et aucun résultat d'application sur des données réelles ne sont fournis.

V.2.2. Horizon à moyen terme

(1) Regroupement des coefficients saisonniers

Les méthodes de prévision de la demande des produits fortement saisonniers, nécessitent une estimation des composantes saisonnières basées sur les données historiques par article. L'étude de Withycombe [WIT 89] suggère que l'estimation saisonnière peut être plus précise si les demandes pour des produits similaires sont d'abord associées dans une même gamme de produits. Les indices saisonniers calculés sur

les produits d'une même gamme sont utilisés comme une estimation des composantes saisonnières pour chaque article. Une application de cette technique, sur 6 gammes de produits différents et en combinant les indices saisonniers, aboutit à une réduction de l'erreur moyenne de prévision de 20%. La technique de la combinaison d'indices saisonniers présente 2 avantages en plus de l'amélioration de la prévision :

- la réduction des effort et des coûts de traitement des indices saisonniers,
- la capacité d'ajuster une composante saisonnière à de nouveaux produits.

En combinant chaque produit dans une gamme appropriée, les calculs des indices saisonniers sont réalisés seulement une fois pour la gamme de produit au lieu d'être répétés pour chaque produit.

Une autre approche de regroupement d'indices saisonniers en vue d'améliorer les méthodes de prévision par décomposition est formulée par Bunn et Vassilopoulos [BUN 93] et aboutit aux mêmes résultats. Les travaux précédents ont montré l'intérêt de l'approche initiée par Withycombe [WIT 89] pour l'utilisation d'un regroupement de type moyenne mobile pondérée des indices saisonniers d'une unité de stock tampon. Des indications laissent présager que la simple moyenne des indices saisonniers est plus robuste quand elle est appliquée à des historiques de données courts et à de faibles valeurs de ventes [DAL 74]. La formation de groupes distincts de produits et d'indices saisonniers relativement stables est réalisée par une procédure de classification. Cependant, aucune précision n'est fournie sur la méthodologie de classification utilisée et la validité des regroupements ainsi obtenus.

(2) Introduction des variables marketing textile

Le désir de promouvoir les ventes conduit les entreprises à une activité commerciale extrêmement variée. Cette activité se concrétise, dans la politique commerciale à court terme, par un grand nombre de types d'actions. Fortement présentes dans la majorité des processus de vente des articles textiles; les variables marketings sont caractérisées par :

- la réduction temporaire des prix de vente,
- les campagnes publicitaires,
- les actions promotionnelles sur les lieux de vente.

Une procédure de prévision à deux étapes proposée par Smith, McIntyre et Achabal [SMI 94] tente d'aborder le problème lié aux critères d'évaluation pour l'analyse des promotions. Abraham et Lodish [ABR 87] suggèrent qu'une procédure d'estimation adéquate doit faire face à des contraintes réelles pour satisfaire les notions suivantes :

- être insensible aux irrégularités des données telles que les erreurs de programmation de promotion au niveau du calendrier,
- être capable de travailler à grande échelle et avoir une vue globale de la situation,
- être capable de travailler avec les données disponibles,
- et avoir une démarche pragmatique dans l'utilisation des critères adéquats.

La première étape de la procédure est appliquée aux séries temporelles des articles similaires regroupés dans une seule classe de produit. Cette étape repose sur une estimation, par la méthode des moindres carrés ordinaires, des données regroupées de la saison précédente pour déterminer les effets saisonniers et les valeurs initiales des variables liées à l'environnement et aux actions marketings. Le regroupement des données s'avère nécessaire afin de constituer des observations suffisamment longues.

Dans les exemples traités, les données sont agrégées suivant une répartition par magasins, couleurs, tailles, et proviennent d'une sous classe de produit d'un distributeur de vêtements. De telles agrégations et regroupements aboutissent sûrement à des coefficients biaisés. Une étude sur les biais résultant d'agrégation est proposée par Blattberg et Neslin [BLA 90], et les problèmes d'agrégation d'une analyse en régression linéaire sont abordés dans les études de Bass et Wittink [BAS 75][BAS 78].

La seconde étape de la procédure consiste à appliquer une méthode des moindres carrés pour ré-estimer les paramètres clés de chaque "programme promotionnel", définis tels que chaque sous-classe de produit possède des campagnes de promotions identiques. L'objectif de cette étape consiste à réduire le biais résultant de l'agrégation des paramètres de regroupement estimés dans la première étape de la procédure et de corriger les paramètres ayant subi un fort changement dans l'environnement du marché à partir de la saison précédente. Le manque de données suffisantes à chaque période ne permet pas de rectifier les variables observées de fréquences irrégulières telles que les effets saisonniers, bien que la seconde étape puisse correctement mettre à jour les paramètres clés des programmes promotionnels, tels que les rapports des ventes et les prix de base. Dans

cette étape, tous les paramètres non mis à jour conservent leurs valeurs calculées à la première étape. L'hypothèse faite ici laisse entendre que les coefficients saisonniers sont plus stables que les paramètres clés de chaque programme promotionnel. Par exemple, en raison des changements de tendances ou de la présence de la concurrence, les rapports des ventes ou la réponse aux déclassements pour beaucoup de sous-gammes d'articles peuvent changer d'une saison à l'autre, tandis que les variations relatives à l'effet saisonnier peuvent apparaître plus stables.

Toutes ces actions diverses appelées "actions spéciales" ont, dans la plupart des cas, une influence très forte sur les ventes "normales". Elles perturbent tellement l'évolution des autres caractéristiques, comme la moyenne ou la «saisonnalité», qu'une interprétation par des méthodes classiques s'avère alors difficile, voire impossible.

Face aux bons résultats obtenus en terme de prévision des quantités à fournir pour répondre aux "actions spéciales" promotionnelles, les auteurs de cette étude tiennent à signaler l'importance de la prise en compte de l'effet synthétique de telles actions dans le processus prévisionnel. Il s'agit, non seulement d'une condition indispensable à l'analyse et donc à la prévision, mais encore d'une base fondamentale pour le contrôle des effets des actions commerciales destinées à promouvoir les ventes.

V.2.3. Horizon à long terme

(1) Estimation des quantités nécessaires à la mise en place de collection textile

Une autre application d'une méthode de prévision utilisée dans l'industrie de l'habillement repose sur la projection des commandes clients [CAN 80]. A partir de ces ventes reçues dans les premières semaines de la saison, une première estimation des ventes totales par article, couleur et style est effectuée. Le concept est identique à une analyse de profil de vente dans le cadre d'une élection. A partir d'un mode de scrutin simplifié, des notations, traduisant des intuitions de comportements de vente issus d'un personnel restreint et qualifié, permettent d'estimer une valeur finale de la quantité prévue à la vente avec un certain degré de précision. Cependant, dans une élection, aucune pénalité n'est attribuée à de mauvaises prévisions. Dans l'industrie textile, les erreurs ont des répercussions financières immédiates sur les profits des acteurs de la filière. Les résultats obtenus par cette première méthode ont permis de mieux gérer les ressources de

production mais les prévisions sont relativement imprécises en raison du peu d'informations disponibles.

Deux autres approches de prévision permettant d'estimer les quantités nécessaires pour la mise en place d'une collection peuvent être présentées [BOU 92] : l'une est fondée sur des mesures "perceptuelles" (analyse des mesures conjointes, sur des données émanant des consommateurs potentiels, au départ subjectives) et l'autre sur des séries de données de ventes et de prises de commandes (données objectives au départ). Elles se distinguent néanmoins fortement par leur cadre d'application.

- la première approche s'applique plutôt à la définition du nouveau produit optimal, en terme de combinaison d'attributs, par rapport aux préférences des consommateurs, avant le lancement. Elle peut servir aussi en termes prévisionnels. Compte tenu de son coût, sa mise en œuvre ne peut être effectuée que pour des produits représentant un chiffre d'affaire conséquent.
- la deuxième approche vise à extrapoler, à partir d'une base limitée de prises de commandes lors des premières semaines, les ventes d'articles de collection, dont le lancement a déjà été décidé. Le but est de corriger la production le plus tôt possible de façon à obtenir un stock minimal d'invendus en fin de période de vente de la collection. Elle peut finalement être automatisée et s'applique alors aisément à des références fines.

La première approche est plutôt qualitative, la deuxième est quantitative. Souvent opposées, ces deux approches sont en fait complémentaires. Malheureusement, aucune application de ces méthodes ne sont disponibles et ne permettent pas d'illustrer leurs différentes potentialités.

Une méthodologie différente d'un système prévisionnel basé sur l'utilisation des informations contenues dans les commandes d'articles textiles est décrite par Bug et Fischer [BUG 93]. Afin d'appuyer la prévision, un système d'analyse expérimentale de processus aléatoires a été développé et testé, permettant de pronostiquer et de contrôler les entrées de commandes d'articles textiles dits de "fantaisie". L'instrument d'analyse et de formation de classe choisi repose sur une méthode d'analyse appelée "Cluster" et son outil (CLAN), et permet d'effectuer une classification des clients en fonction du temps, dans le sens prévisionnel. Aucune indication sur la méthodologie utilisée n'est fournie dans cette

étude. En parallèle, deux autres concepts ont été développés et mis en pratique : ARIMAT et PROCLAN. Le modèle ARIMAT est basé sur la prévision des séries temporelles en fonction du modèle ARIMA. La prise en compte d'un minimum de données provenant des trois saisons précédentes pour prévoir les séries temporelles conditionnées par les saisons s'avère nécessaire. A cette fin, ARIMAT crée des séries temporelles à partir des données des saisons précédentes, par l'aspect statistique, permettant la formulation de "passés virtuels" de séries temporelles individuellement courtes, avec comme base de modèle des séries : la modélisation ARIMA. Ce type de prévision s'est traduit par de bonnes expériences : toutefois, l'utilisateur est très sollicité sous l'angle de ses connaissances. Les utilisateurs de systèmes statistiques prévisionnels peuvent introduire dorénavant davantage de composants basés sur les connaissances, ce qui permet de renforcer la précision de la prévision.

Afin d'obtenir la prévision la plus rapide possible avec un bon coefficient de confiance, un système prévisionnel a été établi, incluant en ARIMAT les classes de séries de temps structurelles. Ce procédé suppose la mise en service de 10 à 20 % des entrées de commandes de la saison en cours. Il nécessite une modification à brève échéance (analyse des points aberrants), des conditions de marchés, des articles et des clients. Ce système, désigné PROCLAN, désormais mis en place, fournit aux responsables une meilleure base de décision que l'estimation habituelle avec les corrections du service des ventes.

(2) Optimisation de l'achat de la matière première textile

L'application d'une méthode de prévision, pour estimer les prix du coton, permet d'optimiser les ressources financières des entreprises attribuées à l'achat des matières premières. Trois séries de prix ont été analysées [ZHA 91] (de fréquence hebdomadaire, mensuelle et trimestrielle) et caractérisées par des modèles ARIMA. Les performances de prévision des modèles ARIMA ont été comparées à des méthodes plus simples telles que : "la marche aléatoire", le lissage exponentiel simple et la moyenne mobile simple. La méthode de prévision de type "marche aléatoire" consiste à initialiser la valeur de la prévision à l'instant t pour un horizon de 1 période par la valeur de la vente réelle obtenue à l'instant t . Les résultats ont montré que les prévisions des modèles ARIMA sont plus précises, par rapport à la mesure de l'erreur (RMSE), que toutes les autres méthodes pour les données à fréquence mensuelle et trimestrielle. Les données à la semaine présentent

une structure stochastique instable. Si une ré-estimation constante du modèle n'est pas possible, le modèle de prévision de type "marche aléatoire" est la méthode la plus adaptée.

V.3. Analyse comparative des méthodes de prévision textiles pour un horizon de prévision $h=1$.

Cette étude, que nous avons réalisé sur des données de vente d'articles textiles issues d'une entreprise française spécialisée dans les vêtements de sport pour l'année 1995, repose sur l'application et l'évaluation des méthodes de prévision adaptées à un horizon à court terme. L'objectif de prévision consiste à évaluer de façon hebdomadaire les demandes en produits finis afin de gérer les stocks tampons et les ressources de production au plus juste. Cette estimation s'effectue dans le cadre d'une démarche « Quick Response » dont la priorité est de fournir la quantité demandée dans un délai de livraison très court (de 1 à 3 jours).

Les méthodes de prévision envisagées sont dans l'ordre :

- la moyenne mobile double d'ordre $p=5$ (Moy mob),
- le lissage exponentiel double avec le coefficient $\alpha=0.3$ (lis expo1),
- le lissage exponentiel double avec une procédure d'autorégulation [BOU 92] du paramètre α (lis expo2) dont la valeur initiale est 0.3,
- la méthode de Holt-Winters selon le schéma additif (HW add) avec $\alpha=0.3$ et $\beta=0.7$,
- la méthode de Holt-Winters selon le schéma multiplicatif (HW mul) avec $\alpha=0.3$ et $\beta=0.7$,
- la méthode de Holt-Winters selon le schéma multiplicatif avec une procédure d'autorégulation [VRO 97] des paramètres de lissage α et β (Fuzzy HW mul) dont les valeurs initiales sont $\alpha=0.3$ et $\beta=0.7$.

Chacune de ces méthodes est appliquée à un échantillon de 59 séries de vente des articles extraites de la base de données par rapport à un seul client. Une méthode fournit des valeurs prévisionnelles à l'horizon 1. Les coefficients saisonniers sont :

- soit estimés par l'industriel en fonction de sa connaissance des ventes des produits,
- soit calculés à partir des allures des centres de classe représentatifs des ventes similaires des articles des saisons précédentes.

Les résultats (en Annexe 1) concerne 59 articles soumis à 6 méthodes de prévision et évaluées par 5 mesures de l'erreur.

Une première représentation est illustrée en Figure 10 par le taux de représentation des articles dont la valeur minimale du critère est obtenue par la méthode utilisée. Par exemple, 42.37% des 59 articles ont une valeur minimale du critère RMSE par la méthode (Fuzzy HW mul), 28.81% des articles par la méthode (lis expo2), 15.25% des articles par la méthode (lis expo1), 6.78% des articles par la méthode (HW add), 5.08% des articles par la méthode (Moy Mob) et 1.71% des articles par la méthode (HW mul).

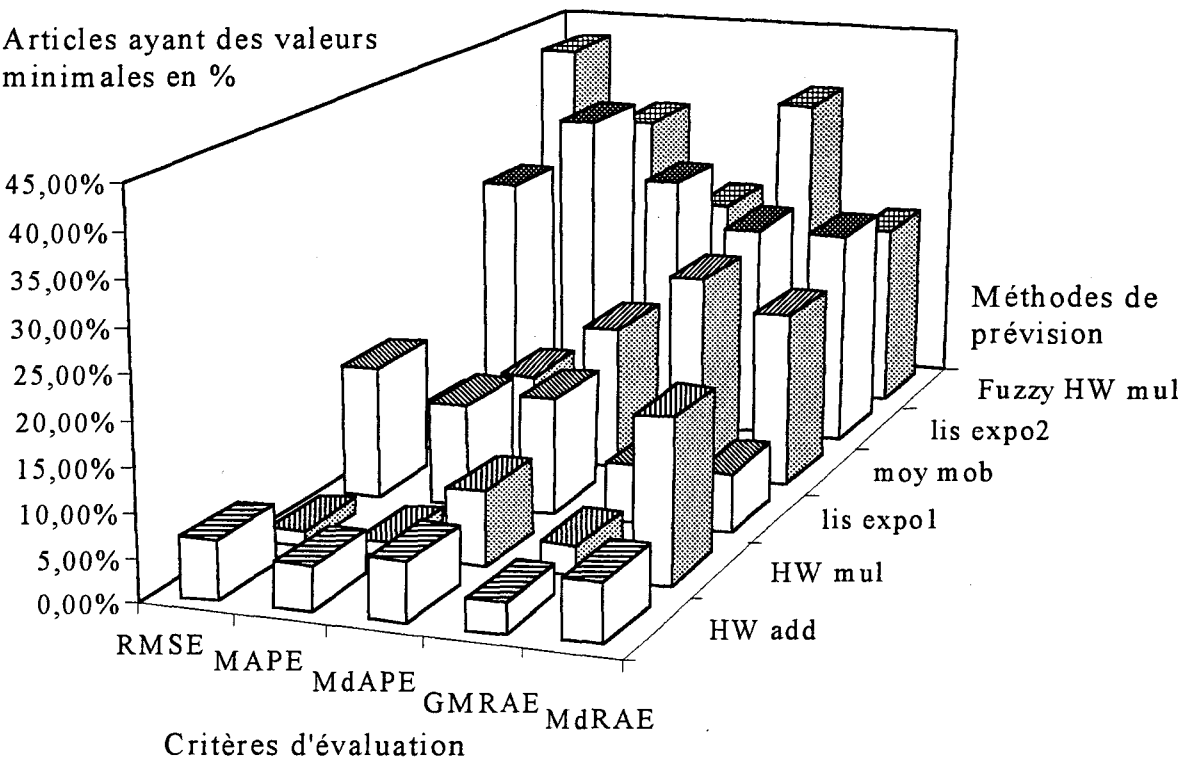


Figure 10. Evaluation des méthodes de prévision en fonction des valeurs minimales des mesures de l'erreur.

Les valeurs minimales des mesures de l'erreur (RMSE) et (GMRAE) sont le plus souvent représentées par la méthode (Fuzzy HW mul) et traduisent une meilleure précision de la valeur prévisionnelle en fonction de la modification des paramètres de lissage. La valeur minimale de la mesure (MAPE) est représentée par la méthode (lisexpo 2), mais ce critère introduit un biais systématique pour la plupart des valeurs prévisionnelles à chaque période en raison des faibles valeurs réelles de vente proches de

l'unité. La valeur de la mesure de l'erreur (MdAPE) ne semble pas adaptée pour l'évaluation des méthodes entre elles en raison du nombre peu important de séries temporelles examinées. La valeur de la mesure de l'erreur (MdRAE) minimale la plus représentée correspond à la méthode de prévision (lisexpo2) et indique que cette procédure fournit des valeurs prévisionnelles plus précises. Ce résultat peut s'expliquer par la procédure d'autorégulation des coefficients de lissage de la méthode (Fuzzy HW mul) qui est basée uniquement sur la minimisation de la mesure de l'erreur (RMSE) et non pas sur la mesure (MdRAE). De ce fait, les paramètres de lissage optimaux correspondants de cette méthode fournissent une valeur prévisionnelle précise en terme d'adaptation de la méthode de prévision aux séries de données mais moins précise en valeur globale.

Une seconde représentation (Figure 11) des valeurs maximales des mesures de l'erreur en fonction de la méthode de prévision indiquent que les méthodes de Holt-Winters, selon le schéma additif et multiplicatif, fournissent les prévisions les moins adaptées et les moins précises des 6 méthodes évaluées. La non régulation des paramètres de lissage en fonction de l'évolution des ventes introduit une erreur systématique dans la valeur prévisionnelle de vente.

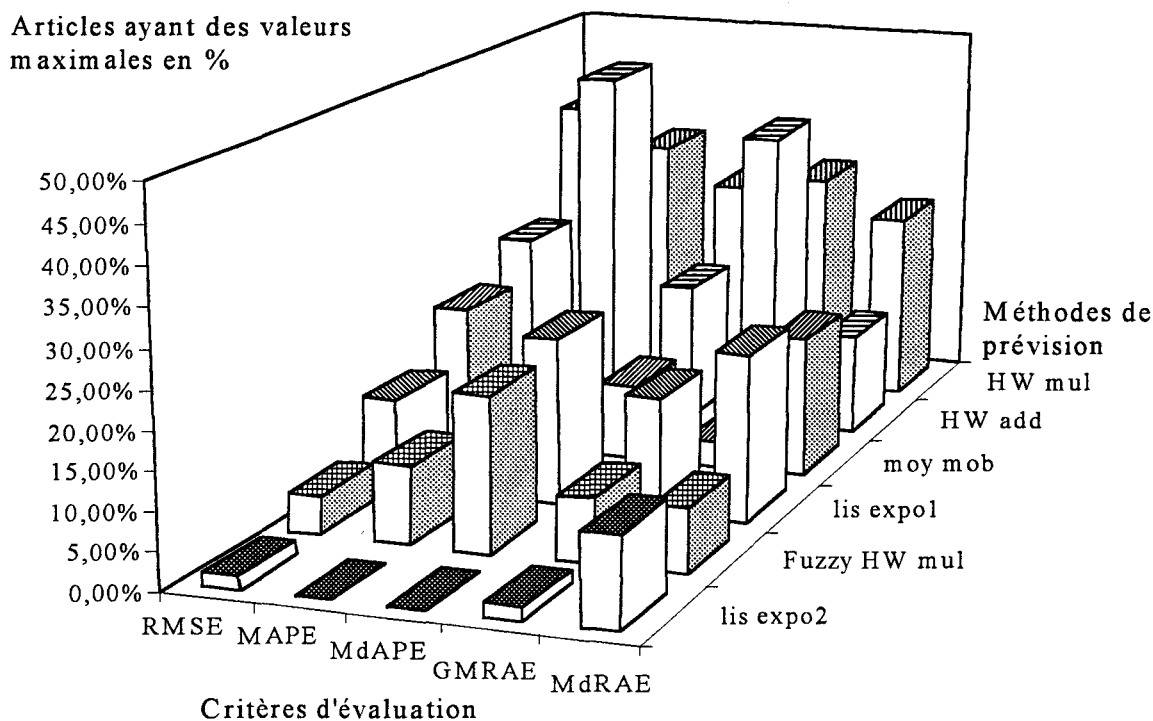


Figure 11. Evaluation des méthodes de prévision en fonction des valeurs maximales des mesures de l'erreur.

Un test statistique peut être effectué pour déterminer si les valeurs d'une même erreur de mesure issues de deux méthodes de prévision peuvent avoir une même valeur moyenne sous l'hypothèse que les écarts types sont égaux. La variable statistique t construite à partir des valeurs moyennes de deux échantillons x et y issue d'une loi de distribution normale s'écrit :

$$t = \frac{\bar{x} - \bar{y}}{s} \text{ avec } s \text{ l'écart type supposé identique pour les deux échantillons } x \text{ et } y.$$

Si la valeur du test est égale à 0 alors les valeurs moyennes de x et y sont statistiquement identiques suivant une valeur de seuil fixé ε . Les valeurs limites de l'intervalle de confiance correspondant à la valeur de seuil ε sont fournies ainsi que la probabilité associée à la valeur significative du test. Par exemple, les moyennes de l'erreur de mesure RMSE sont testées pour chacune des méthodes de prévision prises deux à deux suivant une valeur de seuil $\varepsilon=0.01$ (cf. Tableau 16).

Méthode 1	Méthode 2	test	probabilité	limite inférieure	limite supérieure
Moymob	Lisexpo1	0	0.9267	-13.22	12.32
Moymob	Lisexpo2	0	0.4175	-8.171	15.54
Moymob	Hwadd	0	0.5604	-16.49	10.48
Moymob	Hwmul	0	0.2292	-23.56	8.679
Moymob	HWFuzzy	0	0.6501	-11.08	15.73
Lisexpo1	Lisexpo2	0	0.3791	-8.123	16.39
Lisexpo1	Hwadd	0	0.6293	-16.4	11.28
Lisexpo1	Hwmul	0	0.2671	-23.41	9.426
Lisexpo1	HWFuzzy	0	0.5982	-10.98	16.54
Lisexpo2	Hwadd	0	0.1804	-19.69	6.311
Lisexpo2	Hwmul	0	0.06634	-26.84	4.592
Lisexpo2	HWFuzzy	0	0.7839	-14.27	11.56
Hwadd	Hwmul	0	0.4956	-21.41	12.55
Hwadd	HWFuzzy	0	0.3349	-9.094	19.76
Hwmul	HWFuzzy	0	0.1332	-7.147	26.68

Tableau 16. Résultats des tests statistiques d'égalités
des valeurs moyennes de l'erreur de mesure RMSE
pour l'ensemble des méthodes de prévision.

Les valeurs moyennes de l'erreur de mesure RMSE pour l'ensemble des méthodes de prévisions sont identiques, suivant une valeur de seuil $\varepsilon=0.01$ (99%), et ne permettent pas d'aboutir à une comparaison significative entre les méthodes de prévision. Une représentation de l'erreur médiane du critère RMSE pour chaque méthode, de ses quartiles inférieurs (25%) et supérieurs (75%), de ses valeurs minimales et maximales correspondant à 1.5 fois la valeur de l'écart inter-quartile, permet de visualiser la validité du test sur l'égalité des moyennes obtenues (cf. Figure 12).

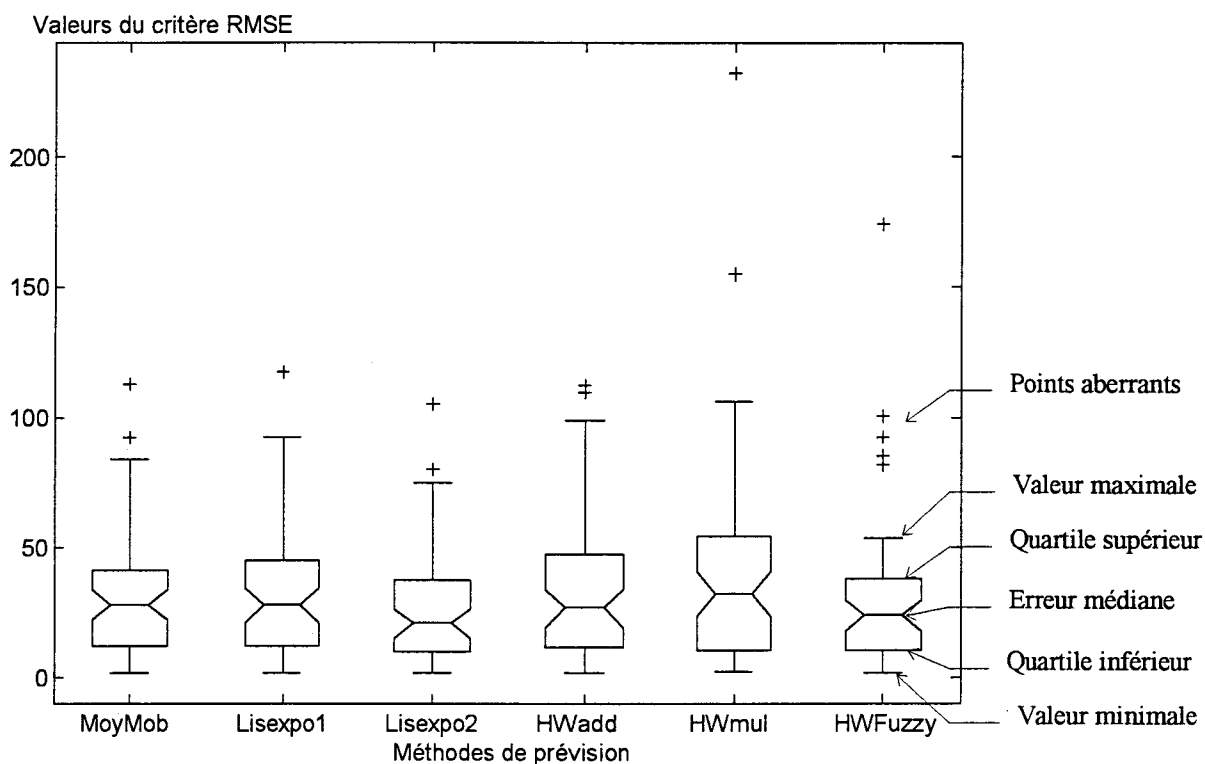


Figure 12. Représentation des erreur médianes, des valeurs minimales et maximales, des quartiles inférieurs et supérieurs de la valeur de la mesure de l'erreur RMSE pour chaque méthode de prévision.

Une représentation des mesures de l'erreur GMRAE (cf. Figure 13) de chacune des procédures prévisionnelles permet de distinguer les méthodes de prévision de Holt-Winter additif et multiplicatif des méthodes de Holt-Winter avec procédure de régulation floue et le lissage exponentiel dynamique. La distinction entre ces deux dernières méthodes ne peut être faite ni sur ce critère et ni sur les autres. Les valeurs moyennes obtenues pour chaque mesure de l'erreur de ces deux méthodes sont statistiquement identiques à un seuil près $\varepsilon=0.01$.

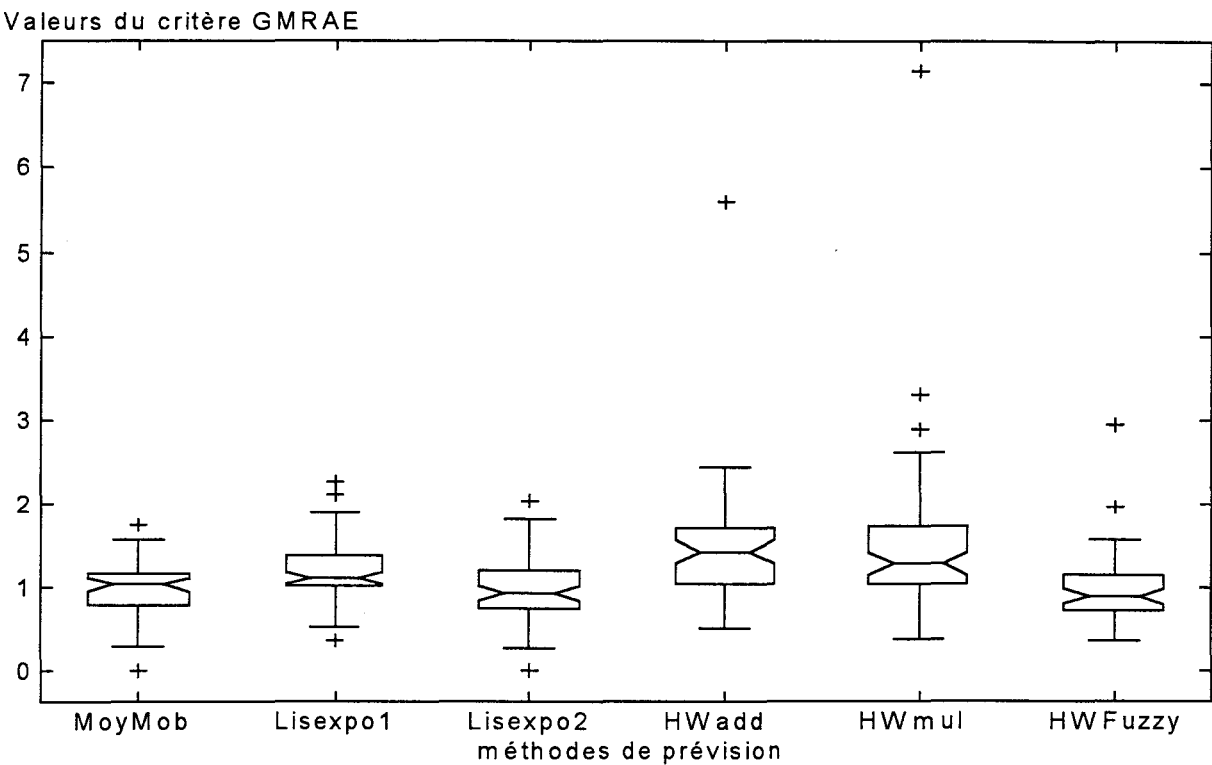


Figure 13. Représentation des erreur médianes, des valeurs minimales et maximales, des quartiles inférieurs et supérieurs de la valeur de la mesure de l'erreur GMRAE pour chaque méthode de prévision.

Cette étude a permis d'évaluer les 6 méthodes de prévision en fonction de différentes mesures de l'erreur. En fonction des résultats obtenus, les méthodes de prévision les plus adaptées à l'environnement de vente de ces 59 articles textiles sont celles qui possèdent une procédure d'autorégulation des paramètres en fonction de l'évolution des ventes telles que : la méthode du lissage exponentiel double avec procédure d'autorégulation et la méthode de Holt-Winters multiplicatif avec procédure d'autorégulation. Un compromis doit être trouvé entre la simplicité d'expression de la méthode et la précision de la prévision souhaitée [MAK 83].

VI. Conclusion

Ce chapitre, consacré aux modèles et méthodes de prévision des ventes, montre qu'il s'avère difficile de fournir des "recettes" prêtes à l'emploi ou des méthodes de prévision "presse bouton" en raison de la multitude de contraintes et critères à prendre en compte en fonction d'une situation donnée. Cependant, certaines indications peuvent apparaître et orienter le choix d'un type de modèle de prévision par rapport à son environnement économique. Dans le contexte de l'horizon de prévision à court terme, et pour fournir les outils d'aide à la décision nécessaires dans le cadre d'un concept "Quick Response", les méthodes simples sont fortement conseillées. Elles permettent d'effectuer une décomposition des différentes caractéristiques principales d'une série temporelle, introduites au sein d'une procédure itérative rapide, telle que le lissage exponentiel. L'utilisation de plusieurs mesures de l'erreur de prévision permet d'évaluer les différentes caractéristiques des méthodes d'extrapolation et fournit une aide au choix du modèle de prévision adapté. Différents paramètres exogènes, sélectionnés en fonction des intuitions des décideurs et traduisant leur connaissance intrinsèque de l'environnement de vente, peuvent être ajoutés indirectement au modèle de prévision dans une procédure d'optimisation d'un coefficient de lissage en fonction d'un critère de validité choisi.

Un éclairage supplémentaire sur les contraintes d'utilisation des méthodes de prévision dans le contexte de la vente des articles textiles permet de mieux définir les capacités attendues des procédures d'extrapolation. Les différentes méthodologies utilisées mettent en évidence les diverses formes de l'utilisation des méthodes de prévision en fonction d'objectifs différents. L'application des 6 méthodes de prévision, et leurs évaluations par des mesures différentes de l'erreur sur des données de vente réelles entre un confectionneur et un distributeur dans le cadre d'une démarche « Quick Response », a permis de mettre en valeur les capacités d'adaptation et de précision des méthodes de lissage utilisant une procédure d'autorégulation de leurs propres paramètres.

L'identification d'un modèle de prévision sur une saison de vente et sa validation sur la saison suivante n'est pas envisagée dans notre approche. Les incertitudes issues de différentes contraintes liées à l'environnement économique de vente des articles textiles

viennent s'ajouter à l'incertitude propre à la prévision. Ces contraintes sont principalement dues à :

- la durée de vie très courte de l'article fortement liée au renouvellement total ou partiel de la collection en raison de l'effet de mode. L'ensemble des données de vente ne constitue pas un ensemble continu et stable d'une saison de vente à l'autre,
- le manque de connaissance actuelle sur la transposition d'un modèle de prévision adapté à une forme d'allure caractéristique pour un article précis à un article dont les modalités des caractéristiques qualitatives sont similaires,
- la nature même de l'influence dynamique et aléatoire des phénomènes exogènes sur les données de vente.

Enfin, un des points importants soulignés dans l'ensemble des études précédentes, et plus particulièrement celles de Withycombe [WIT 89] et Bunn et Vassilopoulos [BUN 93], concerne la dimension du problème de prévision dû au nombre important de références à considérer. Leurs approches préconisent une analyse des séries temporelles par des méthodes de classification ou de reconnaissance de formes. Cette approche a déjà été abordée par la classification des modèles de prévision de type ARMA appliqués à des séries temporelles comportant au minimum 200 périodes d'observations [PIC 90][TON 90][SHA 92]. Nous suggérons que l'information issue de ces analyses permette d'obtenir des informations plus globales et robustes sur le processus de vente des articles tout en permettant de réduire la taille du problème sans nuire à la précision de la prévision.

Nous nous proposons ainsi d'introduire dans le chapitre suivant les outils et les méthodes de classification permettant d'analyser l'ensemble des séries temporelles des ventes d'articles textiles.

Chapitre 3

L'apport de la classification pour la prévision.

I. Introduction

La connaissance intrinsèque des comportements de vente des articles repose souvent sur les intuitions des responsables marketing de l'entreprise dont la fiabilité peut s'avérer satisfaisante en fonction de l'expérience acquise. L'examen minutieux de l'ensemble des ventes, en tenant compte de la nomenclature détaillée des articles textiles, peut s'avérer très long. Un traitement automatique de ces données de vente permet, d'une part, de réduire le temps consacré à cette étude, et, d'autre part, de confronter les intuitions des responsables marketing avec les résultats issus d'une classification des articles selon un critère d'agrégation défini. L'écriture de ce critère peut reposer sur des notions quantitatives et/ou qualitatives et peut inclure, en plus, des traductions intuitives d'association de paramètres qualitatifs.

Les objectifs de classification en terme de constitution de classes reposent sur la formulation du critère de regroupement entre les articles. Elle est fondée sur la notion de mesure de ressemblance entre articles et classes d'articles. La représentation d'un article sous la forme d'un objet symbolique permet d'intégrer les différents caractères quantitatifs et qualitatifs des données dans le but de conserver le maximum d'informations au

moment de l'agrégation des données [DID 96]. Une formulation du critère de classification sur la base de données quantitatives, que sont les allures de vente des articles textiles, permet d'identifier les différentes formes de vente des articles représentées par les centres respectifs de classes. Une formulation du critère de classification sur des paramètres qualitatifs peut aboutir à une classification commerciale permettant d'élaborer le classement des informations au sein d'une base de données. Le croisement des résultats d'une classification à caractère quantitatif à ceux d'une classification à caractère qualitatif permet d'associer des informations communes aux centres des classes et enrichir la connaissance du comportement de vente des articles en fonction de leurs paramètres endogènes. Les centres de classes ainsi caractérisés peuvent être analysés en vue de leur attribuer un modèle de prévision respectif. Le choix du modèle de prévision adéquat est d'autant plus précis que la caractérisation des centres de classe représentatifs d'un ensemble d'articles est complète.

Les différents types de méthodes de classification existantes sont caractérisés par le type d'information recherchée (une hiérarchie, un arbre, une partition, une typologie). Par exemple, un classement suivant un critère peut être réalisé pour distinguer les méthodes de classification qui aboutissent à l'obtention d'une hiérarchie de l'ensemble des différentes partitions, et de celles qui ne l'autorisent pas. Une hiérarchie permet de connaître l'historique de la construction des classes à laquelle peut être associée un "coût" relatif à chacune de ces agrégations. Ce coût caractérise généralement la distance minimale entre les deux partitions successives; la hiérarchie ainsi obtenue est appelée une hiérarchie indicée par la valeur de ce coût.

L'information contenue dans une hiérarchie indicée nous permet d'identifier les partitions dont les classes sont significatives par rapport à une contrainte de précision souhaitée. Les centres de classes correspondants à la partition "optimale" et calculés à partir de chacun des éléments contenus dans leurs classes respectives apportent une information significative et réduite de l'information totale de l'ensemble des articles. L'utilisation d'une méthode de classification appropriée qui permet d'édifier une hiérarchie indicée par sa représentation sous la forme d'un dendogramme, fait l'objet d'une première application.

L'information contenue dans la représentation des centres de classes en fonction des variables qualitatives et quantitatives constitue la source essentielle d'informations



nécessaires pour l'analyse des séries temporelles. De ce fait, l'utilisation d'une méthode de classification permettant de positionner les centres de classes par la minimisation d'une fonction objective aboutit à une caractérisation précise de cette information et fait l'objet d'une seconde application.

Cependant, le choix de la « meilleure » méthode de classification est étroitement lié au problème de la définition d'un critère permettant d'évaluer la qualité de représentation des partitions. La validation de la classification fait référence à des procédures qui évaluent les résultats de l'analyse de classe selon une approche quantitative et objective. Cette évaluation peut être complétée par l'utilisation de contraintes de "compacité" sur les classes de la partition permettant d'ajouter une valeur de précision autour des allures des centres de classes.

Enfin, pour compléter la représentation symbolique des centres de classes par rapport à des paramètres exogènes (phénomènes indépendants de la nature du produit), une évaluation de l'influence de ces phénomènes extérieurs sur le comportement de vente des produits, par l'intermédiaire d'un coefficient de corrélation, permet de quantifier la valeur de la relation de dépendance linéaire.

La description de notre approche de classification des articles textiles peut se représenter par la Figure 14 suivante :

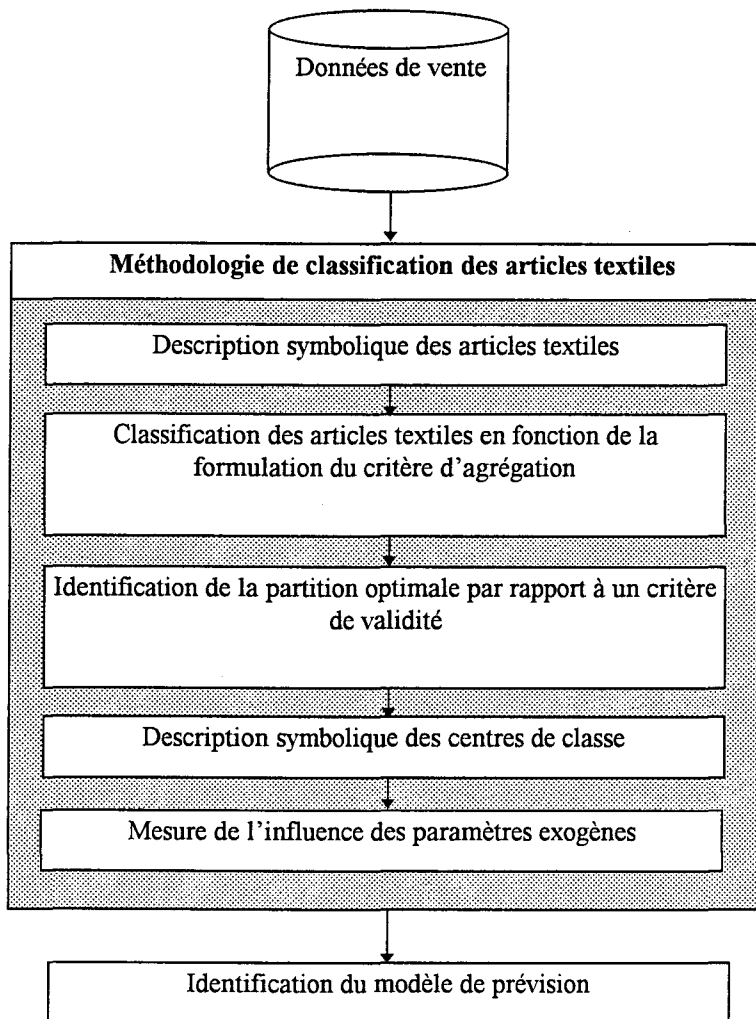


Figure 14. Proposition d'une approche de classification en vue de réduire la complexité du problème de l'identification des méthodes de prévision.

Dans la première partie, une évaluation est effectuée :

- sur la base des travaux réalisés sur les méthodes de classification et leurs capacités propres, pour justifier le choix de la méthode de classification par hiérarchie et de celle par partition,
- sur les critères d'agrégation et les mesures de ressemblance entre objets permettent de justifier le choix des différents paramètres.

Ensuite, une description des algorithmes de classification sélectionnés est proposée, suivie d'une évaluation des travaux réalisés sur les critères de validité des partitions permettant de justifier le choix du paramètre adéquat. Une description symbolique des articles textiles permet d'initialiser l'ensemble des paramètres nécessaires à l'application des méthodes de classification.

Dans une deuxième partie, des applications de classification sur les caractéristiques symboliques des articles textiles sont réalisées et analysées.

Enfin, dans une troisième partie, une méthode est développée afin d'identifier les influences des phénomènes extérieurs aux données sur le comportement des ventes.

II. Proposition d'une approche de classification.

En fonction de la nature des données et de l'objectif à atteindre, différents types de classification peuvent être utilisés. Parmi les méthodes de classification automatiques par partition, la fonction objective de l'algorithme des " c " classes moyennes tend à minimiser la variance intra-classe [DUD 73]. La convergence de la fonction objective vers une valeur finie, et ceci quelque soit la distance utilisée, permet d'obtenir une représentation localement optimale de la partition [SEL 84]. Le temps de calcul de cet algorithme est de l'ordre de $nTcS$, où n est le nombre d'objets, T la dimension de chaque objet, c le nombre désiré de classes, et S le nombre d'itérations nécessaires pour arriver à un seuil de convergence souhaité. La valeur de S est dépendante de la position initiale des centres de classe, des différents modes de distribution des données, et de la taille du problème de classification.

M. Delattre [DEL 79] s'est intéressé à la recherche de la partition qui maximise l'écart entre les classes et minimise le diamètre intra-classe. L'écart entre deux classes est mesuré par l'indice d'agrégation du lien minimum. Le diamètre d'une classe correspond à

la dissimilarité de ses deux individus les plus éloignés. Cette approche a été améliorée par Zeng [ZEN 91] en définissant le critère d'agrégation comme une combinaison linéaire des deux écarts précédents tout en identifiant les frontières des intervalles pour lesquelles la valeur du critère d'agrégation des objets est minimum. De ce fait, les partitions optimales sont obtenues pour les valeurs minimales du critère d'agrégation [ZEN 93]. Cependant, les hypothèses d'applications limitent l'utilisation d'un tel algorithme car elles reposent sur une répartition « multi-modale » des points autour des centres de classes. Différentes applications sur des données textiles [HAP 96a] aboutissent à des classes de formes "allongées", principalement dues à l'effet de taille du critère du saut minimum utilisé dans la définition de la distance entre deux classes. De plus, le temps de calcul est de l'ordre de n^2c dû à la taille de la matrice des distances entre les " n " individus contenant $\frac{n(n-1)}{2}$ éléments distincts, multipliée par le nombre de classes c variant de 2 à $n-1$.

II.1. Proposition des méthodes de classification

II.1.1. Classification ascendante hiérarchique

Selon Lebart et al. [LEB 95], les méthodes de classification peuvent être regroupées en trois catégories : les méthodes par partitions, les méthodes par hiérarchie et par arbre, les méthodes mixtes (couplage d'une analyse factorielle et d'une classification). Les méthodes de classification hiérarchique permettent d'édifier une hiérarchie de la construction des classes dont l'ordre peut être ascendant ou descendant. Parmi l'ensemble des méthodes de classification ascendante hiérarchique, l'utilisation de l'algorithme des voisins réciproques [RHA 80][JUA 82a] permet d'optimiser les ressources mémoires et les temps de calcul [JUA 82b] en agrégeant, de façon non complètement itérative, l'ensemble des paires de données les plus proches à chaque étape de calcul de la matrice de proximité. De bons résultats sont obtenus si les classes sont compactes et bien séparées entre elles, de telle sorte qu'il n'y ait, ni ambiguïté, ni incertitude, associée à l'affectation des objets dans leurs classes respectives [KAM 85].

II.1.2. Classification par partition floue

La théorie des ensembles flous développée par Zadeh [ZAD 65] permet à un objet d'appartenir avec un certain degré à une classe. Ce degré dit "degré d'appartenance" prend

sa valeur entre 0 et 1. L'étape fondamentale dans l'algorithme de classification floue réside dans la définition de la fonction d'appartenance. Les algorithmes de classification floue génèrent des partitions qui minimisent une "fuzzification" induite en suivant la même démarche que les algorithmes de classification non floue basés sur le critère de l'erreur au carré. Le résultat d'un algorithme de classification floue n'inclut pas seulement une partition mais aussi de l'information additionnelle contenue dans la forme des valeurs des degrés d'appartenance. Dans une étude sur les méthodes de classification floue [YAN 93], une présentation chronologique des premières applications met en évidence l'évolution de la formulation des fonctions objectives utilisées. La première formulation de la fonction objective [RUS 69] utilisée au sein des algorithmes de classification floue permet d'introduire la notion de classe floue. Cependant, la convergence de cette fonction objective vers une valeur minimale locale n'est pas prouvée. L'algorithme des " c " classes moyennes floues basé sur l'optimisation d'une fonction objective J_c définie par Dunn [DUN 74] et généralisée par Bezdek [BEZ 81], par l'introduction d'une valeur quelconque du degré de flou, permet de converger vers une valeur minimale locale unique [BEL 96].

II.2. Critères d'agrégations

Un critère d'agrégation caractérise la connaissance de la mesure de ressemblance entre des groupes d'individus [DID 82]. Celui utilisé pour notre application est le critère de Ward [WAR 63]. Plusieurs travaux comparatifs, réunis dans l'étude de Milligan [MIL 79], concluent que le critère de Ward, appelé le critère de la variance minimale, fournit de meilleurs résultats de classification par rapport aux autres critères d'agrégation tels que : le critère du lien minimum [SOK 58], le critère du lien maximum [SOR 48], le critère de la distance moyenne [SOK 58], le critère de la distance moyenne pondérée [JAM 78] et le critère des centres de gravité [SOK 58]. Le critère de Ward, appliqué dans une méthode de classification, est basé sur des notions bien connues, dans les procédures de l'analyse de la variance et d'autres procédures statistiques, de minimisation de l'erreur au carré. Son principe consiste à regrouper les éléments qui minimisent la perte d'inertie inter-classes due au passage de la partition en " c " classes à la partition en " $c-1$ " classes.

La combinaison d'une analyse en composantes principales [GRO 72][FOU 85][SAN 89], qui permet de maximiser la représentation de l'inertie totale d'un nuage de points, avec une méthode de classification utilisant le critère de Ward comme indice

d'agrégation, aide à la représentation optimale de l'information contenue dans chaque classe [ROU 85]. Le critère d'agrégation de Ward entre deux classes (C_p) et (C_q) , contenant respectivement n_p et n_q éléments est défini tel que :

$$D[(C_p), (C_q)] = \frac{n_p n_q}{(n_p + n_q)} d(a_p, a_q) \quad (22)$$

avec :

- a_p et a_q les centres des classes (C_p) et (C_q) ,
- $d(a_p, a_q)$ la distance entre les centres des classes a_p et a_q .

La distance entre deux articles peut être définie différemment suivant le type de données utilisées.

Le critère d'agrégation utilisé ne doit pas créer une inversion en détruisant la relation de voisin le plus proche qui existait au préalable entre deux éléments. Cette condition de non inversion (appelée aussi "critère de la médiane") [DID 82] est vérifiée également par le critère de Ward.

II.3. Mesure de ressemblance entre les articles.

Les définitions suivantes [DID 82] permettent d'introduire différents exemples de mesures de ressemblances :

Indice de "dissimilarité"

Un indice de "dissimilarité" sur un ensemble Ω est une application s qui vérifie les trois propriétés suivantes :

- s une application de $\Omega \times \Omega$ dans $\mathbb{R}^+$.
- s est symétrique : $s(x, x') = s(x', x)$, $\forall (x, x') \in \Omega \times \Omega$.
- $s(x, x) = 0$, $\forall x \in \Omega$.

Distance

Une distance est un indice de "dissimilarité" qui vérifie en plus les deux propriétés suivantes :

- $s(x, x') = 0 \Leftrightarrow x = x'$
- $s(x, x') \leq s(x, x'') + s(x'', x')$ (l'inégalité triangulaire), $\forall x, x', x'' \in \Omega$.

II.3.1. distance entre des données de type quantitatif

La distance de Mahalanobis [MAH 36] est utilisée comme une distance en classification [EVE 74] pour les données de type quantitatif. La distance de Mahalanobis permet d'incorporer la corrélation entre les caractéristiques des données et de rendre standard chaque caractéristique à une moyenne nulle et une variance unitaire.

Une approche de pondération locale des variables, suivant la classe d'individu considérée, peut être envisagée par l'utilisation de techniques de classification automatique dites de « distance et de codage adaptatif » [GOV 75][DID 80].

II.3.2. distance entre des données de type qualitatif

De nombreux codages permettent de traiter les variables qualitatives (codage à modalités avec ou sans ordre, codage disjonctif complet [DID 82]) afin d'obtenir une matrice de variables quantitatives nominales, binaires ou rationnelles. De même, de nombreuses mesures de ressemblance [SOK 58][SOK 63] existent pour identifier les objets similaires en terme d'attributs qualitatifs communs significativement codés. Deux types de traitement des caractéristiques qualitatives des objets symboliques sont possibles :

- soit par reconnaissance syntaxique des variables qualitatives dont la représentation est significative [GOW 95],
- soit par traduction des similarités subjectives entre les objets, suggérées par l'utilisateur, au sein d'une matrice à variables quantitatives.

II.4. Présentation des algorithmes retenus

Parmi les deux grandes familles de méthode de classification par hiérarchie ou par partition, nous avons choisi respectivement l'algorithme des voisins réciproques et l'algorithme des "c" classes moyennes floues.

II.4.1. Classification ascendante hiérarchique par la méthode des voisins réciproques

Une façon naturelle de définir les classes est d'utiliser les propriétés des plus proches voisins ; un objet doit être affecté à la même classe que celle de son plus proche voisin.

Comme indiqué précédemment, les résultats de classification obtenus par les algorithmes accélérés tel que l'algorithme des voisins réciproques sont équivalents à ceux

des algorithmes de base en classification ascendante hiérarchique [BEN 85]. Différentes formulations, optimisant le critère du temps d'exécution, de l'algorithme des voisins réciproques aboutissent à des résultats à peu près identiques avec une préférence pour l'architecture du programme HIVOR. Celui-ci semble particulièrement adapté au traitement rapide des fichiers de données de grandes tailles [JUA 82b].

L'algorithme des voisins réciproques reposent sur deux grandes étapes [BOC 74][BRI 77][BRU 78]. La première étape consiste à trouver tous les couples d'objets tels que chaque élément soit associé avec son plus proche voisin ; ces couples sont appelés les voisins réciproques. La seconde étape consiste à fusionner tous les couples de voisins réciproques et de créer ainsi une nouvelle classe. Ces deux étapes sont exécutées alternativement jusqu'à ce que tous les objets soient contenus dans la même classe.

La définition du concept de voisin réciproque est la suivante :

Soit $X = \{x_1, x_2, \dots, x_n\}$ un ensemble de n articles à classer.

Soit P une partition de l'ensemble $P(X)$ des partitions sur X avec

$$P = \{(C_1), (C_2), \dots, (C_m)\} \text{ et } n = \sum_{i=1}^m \text{card}(C_i) \text{ (l'abréviation card désigne le cardinal de la classe),}$$

Chaque classe (C_k) est représentée par son centre de classe respectif a_k .

Par définition, (a_k) et (a_l) sont deux voisins réciproques si $v[(C_l)] = (C_k)$ et $v[(C_k)] = (C_l)$.

Soient $(C_k), (C_l)$ deux classes différentes de P :

$$v[(C_k)] = \{ \forall (C_l) \subset P, (C_l) \neq (C_k), d[(C_k), (C_l)] = d^*[(C_k)] \} \quad (23)$$

$$d^*[(C_k)] = \inf \{ d[(a_k), (a_l)], \forall (a_l), \text{ avec } (a_k) \neq (a_l) \} \quad (24)$$

Une description complète de l'algorithme des voisins réciproques est disponible dans [BEN 82]; les principales étapes peuvent se résumer ainsi :

Soit P la partition initiale telle que : $P = \{(C_1), (C_2), \dots, (C_n)\}$ et $(C_1) = \{x_1\}$, $(C_2) = \{x_2\}, \dots, (C_n) = \{x_n\}$.

1. Calculer $d(a_k, a_l)$ entre tous les centres des classes a_k et a_l ; puis rechercher les voisins $v[(C_k)]$ pour tout a_k centre de classe de (C_k) .

2. Trouver l'ensemble M des couples de voisins réciproques tel que :

- $M = \{(a_k, a_l) / a_k \neq a_l, v[(C_l)] = (C_k) \text{ et } v[(C_k)] = (C_l)\}$
- Initialiser la partition $C = \emptyset$

3. Faire pour chaque couple $(a_k, a_l) \in M$

- Fusionner dans la nouvelle classe $(c) = (C_k) \cup (C_l)$
- Enlever les éléments (a_k) et (a_l) de la partition P tels que : $P = P - \{a_k\} - \{a_l\}$
- Ajouter à la partition C la nouvelle classe (c) telle que : $C = C + (c)$

4. Arrêter la procédure si $\text{card}(P) + \text{card}(C) = 1$.

5. Calculer des valeurs des distances entre :

- les classes existantes si $(C_k) \in P$ et les nouvelles classes $(C_l) \in C$
- les nouvelles classes $(C_l) \in C$

6. Fusionner les deux partitions P et C telles que $P = P \cup C$ et retourner à l'étape n°2.

II.4.2. Algorithme de classification des "c" classes moyennes floues

L'algorithme des "c" classes moyennes floues est basé sur la notion d'une distance entre les objets en correspondance et l'optimisation d'une fonction objective. Les positions de tous les articles ainsi que leurs degrés d'appartenances sont utilisés pour calculer les valeurs des centres de classes. De même, la convergence de l'algorithme de classification floue ISODATA, qui est un cas particulier de l'algorithme des "c" classes moyennes floues avec un degré de flou m égal à 2, est démontrée [BEZ 76]. Elle permet d'aboutir à une partition floue de l'ensemble des données localement optimales. L'algorithme des "c" classes moyennes floues est également bien adapté pour former des classes de données de formes hyper-sphérique. Plusieurs généralisations de cet algorithme permettent de détecter des formes de classes dont les structures sont hyper-ellipsoïdales ou non [BEZ 81a][BEZ 81b][GUF 78][GUN 83]. Mais, le choix de paramètres de seuil déterminant la forme théorique des différentes classes possibles et constituant une base d'apprentissage n'est pas toujours évident pour ces généralisations. Les paramètres de l'algorithme des "c" classes moyennes floues, sont simplement limités par la recherche des valeurs initiales des degrés d'appartenance des objets aux différents centres de classes ou aux valeurs initiales des positions des centres de classes.

Une présentation simplifiée de l'algorithme des " c " classes moyennes floues se compose des principales étapes suivantes :

Soit un nombre de classes c donné, le critère à minimiser est :

$$J_c = \sum_{i=1}^n \sum_{k=1}^c \mu_{ik}^m d^2(x_i, a_k) \quad (25)$$

avec :

- a_k le centre de la classe (C_k)
- m le degré de flou tel que $m \in [1, +\infty[$
- $d^2(x_i, a_k)$ la distance au carré entre l'article x_i et le centre a_k ,
- μ_{ik} le degré d'appartenance floue de l'article x_i au centre a_k .

L'algorithme des « c » classes moyennes floues peut se résumer ainsi :

1. Initialiser le numéro de l'itération t à 1 et les degrés d'appartenances μ_{ik} pour chaque x_i appartenant à la classe (C_k) tel que :

$$\sum_{k=1}^c \mu_{ik} = 1 \quad (26)$$

2. Calculer les valeurs des positions des centres a_k , pour k variant de 1 à c , en utilisant

$$a_k(t) = \frac{\sum_{i=1}^n \mu_{ik}^m x_i(t)}{\sum_{i=1}^n \mu_{ik}^m}, \forall t \in \{1, \dots, T\}, \quad (27)$$

3. Mettre à jour les degrés d'appartenance flous μ_{ik} en utilisant

$$\mu_{ik} = \frac{\left(\frac{1}{d(x_p, a_k)} \right)^{\frac{2}{m-1}}}{\sum_{j=1}^c \left(\frac{1}{d(x_p, a_j)} \right)^{\frac{2}{m-1}}} = \left[\sum_{j=1}^c \left(\frac{d(x_p, a_k)}{d(x_p, a_j)} \right)^{\frac{2}{m-1}} \right]^{-1} \quad (28)$$

4. Calculer la fonction objective J_e à l'itération e et ajouter à e une unité
5. Répéter les points 2) 3) et 4) jusqu'à ce que la différence entre la valeur de J_e à l'itération e et à l'itération $e-1$ soit inférieure à une valeur de seuil ε initialement fixée :

$$J_e(n) - J_{e-1}(n) \leq \varepsilon \quad (29)$$

II.5. Critère de validité utilisé

Chacun des deux algorithmes retenus aboutit à une partition floue ou non floue en fonction d'un nombre de classes fixées. Plus le nombre de classes s'approche du nombre d'individus à classer, plus précise est la classification. Un compromis doit être fait pour aboutir à une représentation de l'information globale suffisamment précise par rapport à un critère tout en considérant un nombre de classes le plus faible possible. L'utilisation d'un critère de validité sur les partitions obtenues permet d'atteindre cet objectif.

Les deux principales propriétés d'une classe sont la compacité et la séparabilité. La compacité mesure la cohésion interne entre les objets d'une classe tandis que la séparabilité mesure les séparations entre les différentes classes. Une classe est valable si elle est suffisamment compacte et séparée des autres. A chaque catégorie de la méthode de classification est associée un type de critère de validité. La distinction peut être faite entre les critères de validité pour les méthodes par hiérarchie, les méthodes par partitions et les méthodes mixtes, sur la formulation de la mesure utilisée. Tous les indices de validité des classes pris individuellement mesurent la compacité et la séparabilité.

En fonction des nombreux indices de validité existants, le critère de validité D_1 , défini par Dunn [DUN 74], permet d'identifier les classes « compactes et séparées », pour lequel une partition optimale est obtenue si $D_1 > 1$.

La formulation du critère D_1 est la suivante en considérant les classes (C_k) avec k variant de 1 à c et les articles x_i avec i variant de 1 à n :

$$D_1 = \frac{\min_{1 \leq i \leq c} \left(\min_{i+1 \leq j \leq c} (\text{dis}(c_i, c_j)) \right)}{\max_{1 \leq k \leq c} (\text{dia}(c_k))}, \quad (30)$$

avec :

- $\text{dia}(c_k) = \max_{x_i, x_j \in c_k} (d(x_i, x_j))$ (sensiblement proche de la définition d'une compacité),
- $\text{dis}(c_i, c_j) = \min_{x_i \in c_i, x_j \in c_j} (d(x_i, x_j))$ (équivalent à la séparabilité).

Le principal inconvénient de l'utilisation directe de cet indice de validité réside dans le temps de calcul qui peut devenir extrêmement long, quand le nombre de classes et d'articles à mettre sous forme de partition augmente [DAV 79][JAI 88].

Le critère de validité des partitions floues J , défini par Xie et al. [XIE 95], dépend de l'ensemble des données, de la nature géométrique de la mesure de distance et de la distance entre les centres de classes. Ce critère est indépendant du type d'algorithme de classification floue utilisé et conserve les mêmes propriétés d'unicité mathématique que le critère D_1 . La validité de la fonction objective contenue dans l'expression de ce critère repose sur l'identification de la partition, parmi toutes les autres et sans aucune hypothèse quant à la distribution et à la répartition des objets, pour lesquelles les classes sont les plus séparées et les plus compactes. Une évaluation de la performance de ce critère, comparé aux autres critères, conduit à de meilleurs résultats [XIE 95].

II.5.1. Pour les partitions floues

La dispersion de l'ensemble des données en fonction de la partition en " c " classes floues est définie, pour un degré de flou m égal à 2 par :

$$\sigma = \sum_{k=1}^c \sum_{i=1}^n \mu_{ik}^2 d^2(x_i, a_k) \quad (31)$$

La compacité de la partition en " c " classes floues est définie par :

$$\pi = \frac{\sigma}{n} \quad (32)$$

La séparabilité de la partition en "c" classes floues est définie par :

$$s = d_{min}^2 \quad (33)$$

$\forall p, q \in [1, c]$, d_{min} représente la distance minimale entre les centres de classes (a_p, a_q) :

$$d_{min} = \min_{p \neq q} d(a_p, a_q) \quad (34)$$

A partir des définitions suivantes, le critère de validité S est défini par la relation suivante [XIE 95] :

$$S = \frac{\pi}{s} \quad (35)$$

II.5.2. Pour les partitions non floues

L'équivalent du critère S pour des partitions non floues repose sur les mêmes définitions de compacité et de séparabilité hormis que le degré d'appartenance de chaque article par rapport à chaque centre de classe est défini par la relation suivante :

$$\text{si } x_i \in (C_m) \text{ alors } \mu_{im} = 1 \text{ sinon } \forall k = \{1, \dots, c\} \text{ avec } k \neq m, \mu_{ik} = 0. \quad (36)$$

II.5.3. Formulation d'une contrainte de compacité

Pour chaque type d'algorithme, le nombre optimal de classes est déterminé par les différentes valeurs minimales du critère de validité tout en respectant un critère de précision, de nature globale ou locale, sur la compacité des classes.

La compacité ainsi définie caractérise la dispersion des articles autour d'un centre de classe. Un point éloigné du centre de classe se traduit pour des données quantitatives par une allure de vente de l'article textile différente de la moyenne des allures de la classe. De ce fait, l'introduction de valeurs de contrainte sur le critère de la compacité permet de

limiter la dispersion des allures de ventes des différents articles autour du centre de classe. Cette limitation peut être locale, à chaque période de l'allure de vente, et peut se formuler de façon additive ou multiplicative à partir de la valeur du centre de classe à l'instant t . Cette limitation peut être aussi globale, sur toute la période de vente et peut aussi se formuler de façon additive ou multiplicative sur la valeur de vente du centre de classe. Les différentes formulations d'une valeur de contrainte de compacité reflète le choix de la représentation de la valeur limite moyenne de vente des articles textiles par l'industriel. Ainsi, une valeur minimale locale du critère de validité S peut ne pas respecter la contrainte de compacité fixée par l'industriel et oblige à considérer une autre valeur minimale du critère correspondant à un nombre de classes plus grand. Une application sur des données quantitatives des articles textiles [HAP 96b][RAB 97] permet de mettre en évidence le rôle spécifique de la contrainte de compacité sur la moyenne des allures de vente des centres de classes et la précision de représentation correspondante. Les différentes expressions des contraintes sont décrites en fonction des relations suivantes :

(1) locale additive

Soit la classe C_k contenant k articles.

A chaque instant t , l'allure de vente $r_i(t)$ de l'article x_i peut être comprise entre deux limites z_t et y_t autour de la valeur $r_{ak}(t)$ de l'allure du centre de classes a_k telle que :

$$(r_{ak}(t) - y_t) \leq r_i(t) \leq (r_{ak}(t) + z_t).$$

Deux cas peuvent être envisagés :

cas 1. Pour tout $b \in \mathbb{R}$, $\forall t, y_t = z_t = b$

cas 2. A chaque instant t , y_t et z_t peuvent être définies telles que :

$$y_t = \max_{x_i \in C_k} (r_i(t)) \text{ et } z_t = \min_{x_i \in C_k} (r_i(t))$$

(2) globale additive

La contrainte de compacité globale additive peut être calculée à partir de la contrainte locale additive pour les 2 cas en fonction de la formulation de la compacité π .

Deux cas peuvent être envisagés :

cas 1. $\forall t$, pour tout article $x_i \in C_k$ et $x_j \in C_k$, $|r_i(t) - r_j(t)| \leq 2b$

d'où $d(x_i, a_k) \leq 2bT$

et donc $\pi \leq \frac{4b^2T^2}{n}$

cas 2. $\forall t$, pour tout article $x_i \in C_k$ et $x_j \in C_k$, $|r_i(t) - r_j(t)| \leq y_{\max} - z_{\min}$

avec $y_{\max} = \max_{t \in [1, T]} (y_t)$ et $z_{\min} = \min_{t \in [1, T]} (z_t)$

d'où $d(x_i, a_k) \leq T(y_{\max} - z_{\min})$

et donc $\pi \leq \frac{(y_{\max} - z_{\min})^2 T^2}{n}$

(3) locale multiplicative

A chaque instant t , l'allure de vente $r_i(t)$ de l'article x_i peut être comprise entre deux valeurs f_t et e_t autour de la valeur $r_{ak}(t)$ de l'allure du centre de classes a_k telle que :

$$(1-f_t) * r_{ak}(t) \leq r_i(t) \leq (1+e_t) * r_{ak}(t).$$

Deux cas peuvent être envisagés :

cas 1. Pour tout $c \in \mathbb{R}$, $\forall t$, $f_t = e_t = c$

cas 2. A chaque instant t , f_t et e_t peuvent être définies telles que :

$$f_t = \max_{x_i \in C_k} \left(\frac{r_i(t)}{r_{ak}(t)} - 1 \right) \text{ avec } r_{ak}(t) \neq 0,$$

$$e_t = \min_{x_i \in C_k} \left(1 - \frac{r_i(t)}{r_{ak}(t)} \right) \text{ avec } r_{ak}(t) \neq 0.$$

(4) globale multiplicative

La contrainte de compacité globale multiplicative peut être calculée à partir de la contrainte locale multiplicative pour les 2 cas en fonction de la formulation de la compacité π .

Deux cas peuvent être envisagés :

cas 1. $\forall t$, pour tout article $x_i \in C_k$ et $x_j \in C_k$, $|r_i(t) - r_j(t)| \leq 2c * r_{ak}(t)$

$$\text{d'où } d(x_i, a_k) \leq 2c \sum_{t=1}^T r_{ak}(t)$$

$$\text{et donc } \pi \leq \frac{4b^2 \left(\sum_{t=1}^T r_{ak}(t) \right)^2}{n}$$

cas 2. $\forall t$, pour tout article $x_i \in C_k$ et $x_j \in C_k$, $|r_i(t) - r_j(t)| \leq (f_{\max} - \ell_{\min}) * r_{ak}(t)$

avec $f_{\max} = \max_{t \in [1, T]} (f_t)$ et $\ell_{\min} = \min_{t \in [1, T]} (\ell_t)$

$$\text{d'où } d(x_i, a_k) \leq T(f_{\max} - \ell_{\min}) * \sum_{t=1}^T r_{ak}(t)$$

$$\text{et donc } \pi \leq \frac{(f_{\max} - \ell_{\min})^2 T^2 \left(\sum_{t=1}^T r_{ak}(t) \right)^2}{n}$$

II.6. Description symbolique des articles textiles

L'objectif de la procédure de classification d'objets symboliques, combinant à la fois des caractéristiques endogènes de nature qualitatives et quantitatives, est d'élaborer des classes de façon significative tout en minimisant la perte d'informations. Le regroupement des articles, en fonction du critère d'agrégation considéré, ne doit pas, pour autant, chercher à expliquer une variable par rapport aux variables qualitatives (principe de l'analyse discriminante sur variables qualitatives [DID 82]). Ainsi, le critère d'agrégation permet d'introduire la traduction de données intuitives sur le comportement général de vente des articles textiles par rapport à leur environnement économique.

II.6.1. Définitions des objets symboliques

Les objets symboliques sont définis par une conjonction logique d'événements liant les valeurs et les variables. Les variables peuvent prendre différentes valeurs. Les principaux types des objets symboliques sont dans l'ordre croissant de complexité : les objets événement, les objets assertion et les objets horde [DID 88].

Soient p variables $(y_1, y_2, \dots, y_p)$ définies sur un ensemble donné Ω d'objets élémentaires où chaque variable y_i est une application de $\Omega \rightarrow O_i$ avec O_i une observation de l'ensemble des y_i . Ici, y_i est de type quantitatif si O_i est un point ou un intervalle dans

$\mathfrak{R}$, de type qualitatif si O_i est fini, nominal si O_i est non ordonné, et ordinal si O_i est ordonné.

Un événement. Un événement est une paire contenant une valeur et une variable y_i définie par $e_i = [y_i = V_i]$ où la variable y_i prend ses valeurs dans V_i , avec $V_i \subset O_i$. De façon formelle, un événement $e_i = [y_i = V_i]$ est une expression symbolique qui représente une application $e_{y_i V_i} : \{\text{vrai, faux}\}$ de telle sorte que $e_{y_i V_i}(\omega) = \text{vraie}$, avec $\omega \in \Omega$, si et seulement si $y_i(\omega) \in V_i$.

Par exemple, l'événement e_1 indique que la variable *Type* est soit égale à la valeur *Piscine* ou soit à la valeur *Gym* : $e_1 = [Type = \{Piscine, Gym\}]$.

L'événement e_2 indique que la variable *Genre* est soit égale à *Homme* ou soit à *Femme* tel que : $e_2 = [Genre = \{Homme, Femme\}]$

L'événement e_3 indique que la variable *Prix* varie entre les valeurs 100 et 150 tel que : $e_3 = [Prix = [100-150]]$.

Un objet de type Assertion. Un objet assertion est une conjonction d'événements. Il est défini par :

$$a = [y'_1 = V_1] \wedge [y'_2 = V_2] \wedge \dots \wedge [y'_k = V_k]$$

où $y'_i \in Y = \{y_1, y_2, \dots, y_p\}$ et $V = (V_1, \dots, V_k) \subset (O_i)^k$. L'ensemble de tous les objets élémentaires de Ω conformes à la définition de l'objet assertion a est défini par $|a|_\Omega$ où

$|a|_\Omega = \{\omega \in \Omega \mid y'_i(\omega) \in V_i \text{ pour } i = 1, 2, \dots, k\}$. De façon formelle, a qui définit une expression symbolique $\Lambda_i = [y'_i = V_i]$ représente une application $a_{YV} : \Omega \rightarrow \{\text{vrai, faux}\}$ telle que $a_{YV}(\omega) = \text{vraie}$ si et seulement si $\forall i y'_i(\omega) \in V_i$.

Par exemple, l'objet assertion $a1$ peut être défini comme une conjonction logique des événements e_1 , e_2 et e_3 précédents, tel que :

$$a1 = [Type = \{Piscine\}] \wedge [Genre = \{Homme\}] \wedge [Prix = [100-150]].$$

L'objet assertion $a2$ peut être défini comme une conjonction logique des événements e_1 , e_2 et e_3 précédents, tel que :

$$a2 = [Type = \{Gym\}] \wedge [Genre = \{Femme\}] \wedge [Prix = 110].$$

Un objet de type Horde. Un objet de type horde h est une conjonction, d'événements ou d'objets assertions, définie par :

$$h = [y'_1(u_1) = V_1] \wedge [y'_2(u_2) = V_2] \wedge \dots \wedge [y'_k(u_k) = V_k]$$

où $u = \{u_1, u_2, \dots, u_k\}$ sont les objets élémentaires, $y'_i \in \{y_1, y_2, \dots, y_p\}$, et $V_i \subset O_i$.
 De façon formelle, b est une expression symbolique qui représente une application $b_{YV} : \Omega^k \rightarrow \{\text{vrai, faux}\}$ telle que $b(w_1, w_2, \dots, w_k)$ est vraie si et seulement si $\forall i = 1, 2, \dots, k$ et $\forall j = 1, 2, \dots, q, y'_i(w_j) \in V_i$.

Par exemple, l'objet horde b peut être défini par une conjonction logique des objets assertion a_1 et a_2 tel que :

$$\begin{aligned} b = & [Type(article1) = \{Piscine\}] \wedge [Genre(article1) = \{Homme\}] \wedge \\ & [Prix(article1) = [100-150]] \wedge [Type(article2) = \{Gym\}] \wedge \\ & [Genre(article2) = \{Femme\}] \wedge [Prix(article2) = 110] \end{aligned}$$

II.6.2. Notations utilisées

Chaque article x_i , avec $i \in \{1, \dots, n\}$, est représenté de façon symbolique par un objet de type assertion, résultant d'une conjonction logique d'objets événement définie à partir d'un ensemble de différentes caractéristiques de nature quantitatives et qualitatives. De façon générale, x_i appartient à l'ensemble Ω qui est un produit cartésien $\times$ de l'ensemble des caractéristiques A_k tel que :

- A_k représente l'ensemble des $k^{\text{ièmes}}$ caractéristiques,
- $\Omega = A_1 \times \dots \times A_k \times \dots \times A_r$ où r est le nombre maximum des caractéristiques,
- $x_i \in \Omega, x_i = [A_{i1} \times A_{i2} \times \dots \times A_{ir}]$ avec $A_{ik} \in A_k, \forall k \in \{1, \dots, r\}$.

Dans toutes nos applications, l'ensemble Ω des caractéristiques des articles textiles issu d'une base de données d'une société spécialisée dans les articles de sport, est composé de :

$$\Omega = Q \times Y \times G \times D \times M \times U \times C \quad (A_1 = Q, A_2 = Y, A_3 = G, A_4 = D, A_5 = U, A_6 = C)$$

tels que :

- Q représente l'ensemble des quantités vendues de tous les articles,
- Y regroupe l'ensemble des types,
- G désigne l'ensemble des genres,
- D englobe l'ensemble des désignations,
- M réunit l'ensemble des modèles,

- U représente l'ensemble des tissus,
- C regroupe l'ensemble des coloris.

Exemple :

- $Y = \{\text{bonnet, Enfant, Femme, Homme}\}$, (la variable bonnet est traité comme un cas particulier parmi les genres)
- $G = \{\text{Piscine, Gym}\}$,
- $D = \{\text{Bonnet de bain, Maillot de bain, Short, Juste au corps, Collant, Cycliste}\}$,
- M et U codé sur 4 chiffres,
- $C = \{\text{bianco, marino, nero, denim bq bleu, v2 bleu noir vert,}\}$.

A chaque article x_i est associé un vecteur des valeurs de vente $Q_i \in \mathbb{N}^T$, tel que :

$$\forall t \in [1, T], Q_i = (q_i(1); q_i(2); ; q_i(T))^t \text{ (où } t \text{ désigne la transposée du vecteur),}$$

Afin d'éviter l'effet de taille entre les différentes valeurs de vente des articles textiles, un vecteur R_i de dimension T , représentant son allure de vente, est défini tel que :

$$R_i = (r_i(1); r_i(2); ; r_i(T))^t \tag{37}$$

$$\text{avec } r_i(t) = \frac{q_i(t)}{\sum_{i=1}^T q_i(t)}$$

Ainsi, l'article n°2 de la classification commerciale du Tableau 17 est représenté par un objet de type assertion :

article2= [$Q_2 \wedge$ Piscine $\wedge$ Enfant $\wedge$ Maillot de bain $\wedge$ 5016 $\wedge$ 2406 $\wedge$ v2 bleu noir vert]

avec $Q_2 \in \mathbb{N}^{52}$:

$$Q_2 = [0;0;0;0;0;0;0;0;0;0;0;0;0;11;11;41;59;71;71;92;62;62;61;62;77;202;271;271;264;259;195;141;112;112;114;123;123;102;100;100;55;52;29;29;47;47;42;35;0;0;0]^t$$

Les données considérées représentent un échantillon de 59 articles de la première saison de vente, relatif à un seul client, dont les allures sont regroupées sur 52 périodes (une période correspond à une semaine). La classification commerciale de l'entreprise des articles en fonction des paramètres endogènes (genre, type et désignation) est représentée

dans le Tableau 17 en 11 classes distinctes. Le critère d'agrégation de cette classification repose sur ces variables qualitatives et constitue un cas particulier de représentation des données.

Classe	Genre	Type	Désignation	Numéro de code d'article														
1	Piscine	bonnet	Bonnet de bain	1														
2	Piscine	Enfant	Maillot de bain	2	3	4	5	6	7	8	9	10	11	12	13	14	15	
3	Piscine	Femme	Maillot de bain	16	17	18	19	20										
4	Piscine	Homme	Maillot de bain	21	22	23	24	25										
5	Piscine	Homme	Short	26	27	28	29	30	31									
6	Gym	Enfant	Juste au corps	32	33	34	35											
7	Gym	Enfant	Collant	36	37	38	39											
8	Gym	Enfant	Cycliste	40	41	42	43											
9	Gym	Femme	Juste au corps	44	45	46	47	48	49	50	51							
10	Gym	Femme	Collant	52	53	54												
11	Gym	Femme	Cycliste	55	56	57	58	59										

Tableau 17. **Classification commerciale des 59 articles.**

II.6.3. *Caractérisation des centres de classe*

La principale difficulté dans la classification consiste à définir les caractéristiques du centre de classe en fonction des paramètres des articles tout en minimisant la perte d'information.

(1) pour la classification ascendante hiérarchique

La caractérisation des centres de classe a_p repose sur les variables endogènes quantitatives et qualitatives de l'ensemble des articles de la classe (C_p) .

Pour les données de type quantitatif, l'allure du centre de classe est calculée à partir de la moyenne des valeurs des articles appartenant à la classe (C_p) .

Pour les variables endogènes qualitatives, la définition d'une moyenne sur des paramètres qualitatifs n'existe pas, seul un pourcentage de représentation de ces variables peut être calculé.

Soit $\mathcal{A}_k = \{A_{1k}, ...A_{jk}, ...,A_{hk}\}$, l'ensemble des "h" valeurs possibles de la caractéristique endogène k .

Le centre de classe a_p est représenté par les "h" valeurs de l'ensemble des caractéristiques qualitatives $\mathcal{A}_{ik}$ auxquelles sont associés "h" coefficients de représentation respectifs $(\alpha a_{p,k(1)}, ..., \alpha a_{p,k(i)},, \alpha a_{p,k(h)})$.

Ainsi, pour tout article x_i appartenant à la classe (C_p) :

si la valeur de la caractéristique A_{ik} de l'article x_i correspond à la $j^{\text{ème}}$ valeur de la caractéristique A_k alors $\alpha_{a_p,k(j)} = \alpha_{a_p,k(j)} + 1$ (où $\alpha_{a_p,k(j)}$ représente le coefficient associé à la $j^{\text{ème}}$ valeur de la caractéristique qualitative A_k du centre de classe a_p).

Les coefficients du centre de classe ainsi obtenus sont normalisés entre 0 et 1 en divisant chaque valeur par la somme des coefficients et mis sous forme d'un vecteur $\alpha_{a_p,k}$ tels que :

$$\alpha_{a_p,k} = \frac{1}{\sum_{j=1}^h \alpha_{a_p,k(j)}} (\alpha_{a_p,k(1)}, ..., \alpha_{a_p,k(j)}, ..., \alpha_{a_p,k(h)})^t \quad (38)$$

Par exemple, soit la classe C_1 contenant les articles n°1, 2 et 3 de la classification commerciale du Tableau 17 avec type(1) = bonnet, type(2) = Enfant et type(3) = Enfant. La variable qualitative Y du centre de classe C_1 est représentée par : $Y = \{\text{bonnet, Enfant, Femme, Homme}\}$. Les coefficients de représentation de la variable Y associés au centre de la classe C_1 sont regroupés sous le vecteur $\alpha_{a_1,Y}$ tel que :

$\alpha_{a_1,Y} = (\frac{1}{3}, \frac{2}{3}, 0, 0)$ où la variable qualitative Y du centre de la classe C_1 est représentée par : $\frac{1}{3} \{\text{bonnet}\}$ et $\frac{2}{3} \{\text{Enfant}\}$.

(2) pour la classification par partition floue

Une description des caractéristiques qualitatives du centre de classe floue peut également être obtenue en considérant les valeurs des différents degrés d'appartenances associées aux paramètres endogènes des articles. La $k^{\text{ème}}$ caractéristique qualitative $A_{a_p,k}$ du centre de classe a_p définie suivant " h " valeurs différentes est calculée à partir des degrés d'appartenance μ_{ip} de tous les articles au centre de classe a_p .

Le centre de classe a_p est représenté par les " h " valeurs de l'ensemble des caractéristiques qualitatives A_{ik} auxquelles sont associés les " h " coefficients de représentation respectifs

$(\beta_{a_p,k(1)}, ..., \beta_{a_p,k(j)}, ..., \beta_{a_p,k(h)})$ tels que :

Pour tout article x_i appartenant à la classe (C_k) :

si la valeur de la caractéristique $\mathcal{A}_{ik}$ de l'article x_i correspond à la valeur de la $j^{\text{ème}}$ caractéristique $\mathcal{A}_k$ alors $\beta_{a_p, k(j)} = \beta_{a_p, k(l)} + \mu_{ip}$

Les coefficients du centre de classe ainsi obtenus sont normalisés entre 0 et 1 en divisant chaque valeur par la somme des coefficients et mis sous la forme d'un vecteur $\beta_{a_p, k}$ tels que :

$$\beta_{a_p, k} = \frac{1}{\left(\sum_{i=1}^n \mu_{ip} \right)} (\beta_{a_p, k(l)}, \dots, \beta_{a_p, k(l)}, \dots, \beta_{a_p, k(h)})^t \quad (39)$$

II.6.4. Mesure de ressemblance entre les articles

La traduction de l'objectif de la classification réside dans la formulation de la mesure de ressemblance entre deux articles.

(1) Pour les données de type quantitatif

L'expression de la distance au carré de Mahalanobis pour les articles x_i et x_k :

$$d_{M_{ik}}^2(x_i, x_k) = (R_i - R_k)^t M_{ik}^{-1} (R_i - R_k) \quad (40)$$

avec M_{ik} la matrice des variances-covariances (symétrique et définie positive et M_{ik}^{-1} désigne l'inverse de la matrice M_{ik}).

Si M_{ik} est égale à la matrice identité alors la distance au carré de Mahalanobis est la même que la distance Euclidienne au carré, soit :

$$d^2(x_i, x_k) = (R_i - R_k)^t (R_i - R_k) \quad (41)$$

(2) Pour les données de type qualitatif

La distance $d(a_p, a_q)$ entre deux centres de classes a_p et a_q repose sur la définition de la similitude $s_{-k}(a_p, a_q)$ de la caractéristique endogène k des centres de classes et de son coefficient de pondération correspondant δ_k .

$$d(a_p, a_q) = \frac{1}{r} \sum_{k=1}^r \delta_k * s_{-k}(a_p, a_q) \quad (42)$$

$$\text{avec } \sum_{k=1}^r \delta_k = 1$$

(a) pour la classification ascendante hiérarchique

La similitude de la caractéristique qualitative k entre les centres de classes a_p et a_q est calculée à partir des paramètres endogènes A_{ik} et A_{jk} des articles x_i et x_j appartenant respectivement aux classes (C_p) et (C_q) . Soit :

$$s_{-k}(a_p, a_q) = \frac{1}{n_p * n_q} \sum_{i=1}^{n_p} \sum_{j=1}^{n_q} s(A_{ik}, A_{jk}) \quad (43)$$

La similitude $s(A_{ik}, A_{jk})$ des caractéristiques qualitatives A_{ik} et A_{jk} des articles x_i et x_j peut être définie de deux façons :

- Soit par une procédure de reconnaissance de chaîne de caractères. Par exemple, si la caractéristique A_{ik} est similaire à la caractéristique A_{jk} alors $s(A_{ik}, A_{jk})=0$ sinon $s(A_{ik}, A_{jk})=1$.
- Soit par une matrice de similitude entre les valeurs de la caractéristique qualitative k . Cette matrice de similitude est équivalente à une matrice de codage des valeurs de la variable qualitative en valeurs quantitatives traduisant les connaissances intuitives du responsable de l'entreprise par rapport à la nomenclature des produits de l'entreprise. En règle générale, plus les caractéristiques qualitatives sont similaires, plus la valeur quantitative de la similitude est proche de 0 et, inversement, proche de 1. Par exemple, les chaînes de caractères des coloris "9040noir" et "noir9425" sont différentes mais représentent une même couleur noire suivant différentes tonalités. La similitude $s("9040noir", "noir9425")$ entre ces deux coloris possède une valeur proche de 0. De ce fait, la similitude entre les caractéristiques qualitatives A_{ik} et A_{jk} des articles x_i et x_j peut être définie par $s(A_{ik}, A_{jk})$.

Dans notre application [HAP 97], la formulation de $s(a_p, a_q)$ s'écrit de la façon suivante :

$$s(a_p, a_q) = (\delta_1 * s_Y(a_p, a_q) + \delta_2 * s_G(a_p, a_q) + \delta_3 * s_M(a_p, a_q) + \delta_4 * s_U(a_p, a_q) + \delta_5 * s_C(a_p, a_q) + \delta_6 * s_D(a_p, a_q)) \quad (44)$$

Les fonctions de similitudes : s_Y , s_G , s_M et s_U , sont basées sur une procédure de reconnaissance de chaînes de caractères. Les fonctions de similitudes : s_C et s_D , reposent sur une matrice de codage des différentes valeurs des caractéristiques qualitatives.

(b) pour la classification par partition floue

La similitude de la caractéristique qualitative k entre les centres de classes a_p et a_q est calculée par :

$$s_k(a_p, a_q) = d_2(Aa_{p,k}, Aa_{q,k}) \quad (45)$$

où d_2 représente la distance Euclidienne entre les vecteurs $Aa_{p,k}$ et $Aa_{q,k}$ des coefficients quantitatifs associés aux valeurs des caractéristiques.

III. Applications textiles

Les données utilisées pour nos applications permettent d'illustrer notre démarche et sont représentées dans le Tableau 18 suivant :

Num	Quantitatif		Qualitatif				
	Quantités vendues	type	genre	désignation	modèle	tissu	coloris
1	(cf. annexe 3)	PISCINE	BONNET L	Bonnet	null	null	assortis
2		PISCINE	ENFANT	MB 1 P Fille	5016	2406	v1 bleu/noir/vert
3		PISCINE	ENFANT	MB 1 P Fille Basic	5373	2519	carminio io rouge
4		PISCINE	ENFANT	MB 1 P Fille Basic	5373	2519	nero ik
5		PISCINE	ENFANT	MB 1 P Fille Basic	5373	2519	royal q1 bleu
6		PISCINE	ENFANT	MB 1 P Fille Imprimé	5395	2827	03 bleu/noir/denim
7		PISCINE	ENFANT	MB 1 P Fille	5429	2625	denim bq bleu
8		PISCINE	ENFANT	MB 1 P Fille	5429	2625	geranio 71
9		PISCINE	ENFANT	MB 1 P Fille	5429	2625	marino ro
10		PISCINE	ENFANT	MB Slip Garçon	7130	2406	v1 bleu/noir/vert
11		PISCINE	ENFANT	MB Slip Garçon	7130	2406	v2 noir/rouge/gris
12		PISCINE	ENFANT	MB Slip Garçon	7142	2827	03 bleu/noir
13		PISCINE	ENFANT	MB Slip Garçon	7146	2519	carminio io rouge
14		PISCINE	ENFANT	MB Slip Garçon	7146	2519	nero ik
15		PISCINE	ENFANT	MB Slip Garçon	7146	2519	royal q1 bleu
16		PISCINE	FEMME	MB Basic 1 P	1622	2500	canard ru
17		PISCINE	FEMME	MB Basic 1 P	1622	2500	indigo mq bleu
18		PISCINE	FEMME	MB Basic 1 P	1622	2500	nero ik
19		PISCINE	FEMME	MB Basic 1 P	1622	5500	noir nero ik
20		PISCINE	FEMME	MB 1 P Femme forte	136B	2500	noir nero ik
21		PISCINE	HOMME	MB Slip	7646	2500	noir nero ik
22		PISCINE	HOMME	MB Slip	7646	2500	roy royal q1
23		PISCINE	HOMME	MB Slip	7646	2500	vert bottiglia vw
24		PISCINE	HOMME	MB Slip	7725	5904	v1 gris rouge noir
25		PISCINE	HOMME	MB Slip	7725	5904	v2 bleu vert noir
26		PISCINE	HOMME	Short	7920	3107	v1 noir/bleu
27		PISCINE	HOMME	Short	7920	3107	v2 noir/rouge
28		PISCINE	HOMME	Short	7925	3107	v1 noir/bleu
29		PISCINE	HOMME	Short	7925	3107	v2 noir/rouge
30		PISCINE	HOMME	Short	7945	3233	v1 noir/bleu
31		PISCINE	HOMME	Short	7945	3233	v2 noir/rouge
32		GYM	ENFANT	Juste au corps 1 Pièce Fille	5071	2519	bianco ao
33		GYM	ENFANT	Juste au corps 1 Pièce Fille	5071	2519	ciliega 96 saumon
34		GYM	ENFANT	Juste au corps 1 Pièce Fille	5071	2519	denim bq bleu
35		GYM	ENFANT	Juste au corps 1 Pièce Fille	5071	2519	nero ik
36		GYM	ENFANT	Collant fille	5991	2519	bianco ao
37		GYM	ENFANT	Collant fille	5991	2519	ciliega 96 saumon
38		GYM	ENFANT	Collant fille	5991	2519	denim bq bleu
39		GYM	ENFANT	Collant fille	5991	2519	nero ik
40		GYM	ENFANT	Cycliste Garçon	7233	2519	bianco ao
41		GYM	ENFANT	Cycliste Garçon	7233	2519	ciliega 96 saumon
42		GYM	ENFANT	Cycliste Garçon	7233	2519	denim bq bleu
43		GYM	ENFANT	Cycliste Garçon	7233	2519	nero ik
44		GYM	FEMME	Juste au corps 1 Pièce	1420	2073	0743 rosso
45		GYM	FEMME	Juste au corps 1 Pièce	1420	2073	gris standard
46		GYM	FEMME	Juste au corps 1 Pièce	1492	2529	indigo mq
47		GYM	FEMME	Juste au corps 1 Pièce	1492	2529	nero ik
48		GYM	FEMME	Juste au corps 1 Pièce	1625	2097	9040 noir
49		GYM	FEMME	Juste au corps 1 Pièce	1625	2666	66015 gomet.jn/oran
50		GYM	FEMME	Juste au corps 1 Pièce	1625	2685	25 bleu tie die
51		GYM	FEMME	Juste au corps 1 Pièce	1629	2841	01 gris noir blc/noi
52		GYM	FEMME	Collant Femme	1881	2022	noir 9425
53		GYM	FEMME	Collant Femme	1881	2073	gris standard
54		GYM	FEMME	Collant Femme	1881	2519	noir nero ik
55		GYM	FEMME	Cycliste Femme	3811	2022	noir 9425
56		GYM	FEMME	Cycliste Femme	3811	2073	gris standard
57		GYM	FEMME	Cycliste Femme	3811	2097	9040 nero
58		GYM	FEMME	Cycliste Femme	3811	2519	indigo mq
59		GYM	FEMME	Cycliste Femme	3811	2519	noir nero ik

Tableau 18. Représentation des variables quantitatives
et qualitatives des 59 données utilisées.

III.1. Application de la méthode de classification ascendante hiérarchique

Une classification suivant les caractéristiques quantitatives des données, en utilisant l'algorithme des voisins réciproques, permet de regrouper les articles dont les allures de vente sont similaires au sein des différentes classes.

III.1.1. Application aux allures de vente

L'application de l'algorithme des voisins réciproques suivie de la procédure d'évaluation des partitions par le critère de validité S , est représentée Figure 15 en faisant varier le nombre de classes c de 2 à 20. La valeur maximale des classes est de 59 (59 articles de la classification commerciale), mais la valeur de c_{max} peut être choisie en fonction d'une connaissance a priori du nombre de classes à atteindre ou en fonction du nombre d'individus n tel que $c_{max} = \frac{n}{3}$; en effet, l'essentiel de l'information souhaitée se situe généralement entre c et c_{max} [XIE 95]. La "meilleure" partition correspond à la valeur minimale des minimums locaux du critère de validité S .

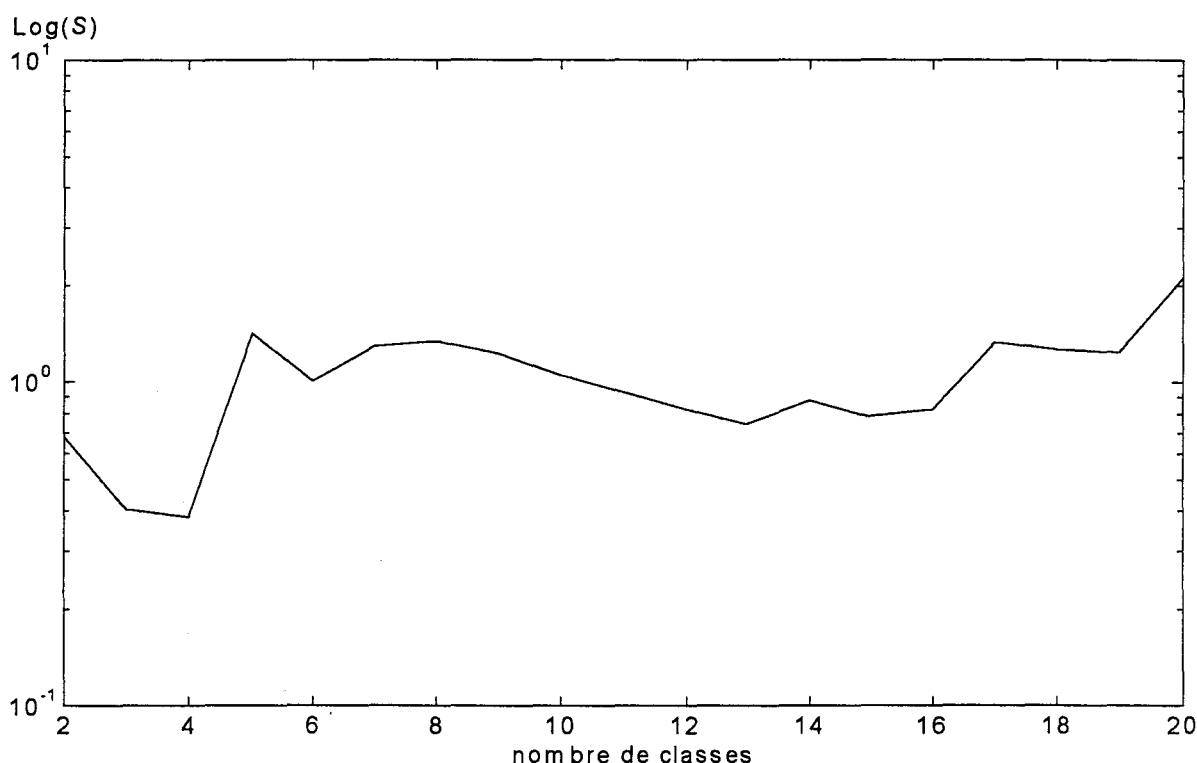


Figure 15. Evaluation des partitions de c classes des données quantitatives.

La première valeur minimale locale du critère J correspond à une partition de 4 classes. Ainsi, la partition des 4 classes est obtenue "en coupant avec une ligne horizontale" le dendogramme en Figure 16, représentant la hiérarchie indiquée, entre les noeuds 115 et 114. Le numéro du noeud terminal 117 ($117=2n-1$ avec $n=59$ articles) correspond à une partition constituée de deux classes. Chaque noeud du dendogramme caractérise une partition différente dont le contenu des classes est directement lié aux contenus des noeuds de numéro inférieur.

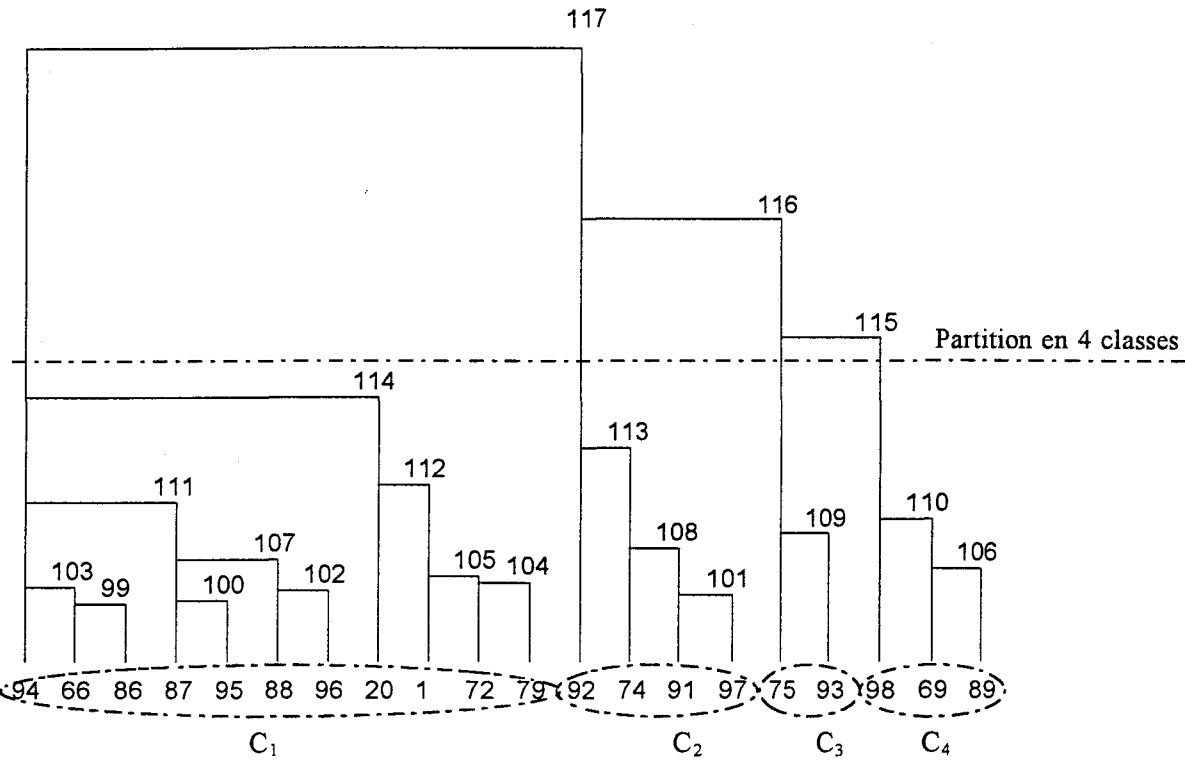


Figure 16. **Hiérarchie restreinte en 20 noeuds terminaux de la classification des données quantitatives.**

La représentation (Figure 17) suivant les deux premiers axes factoriels principaux du nuage de points, dont l'appartenance à chaque classe est identifiée par un symbole respectif, permet de visualiser les résultats de la classification.

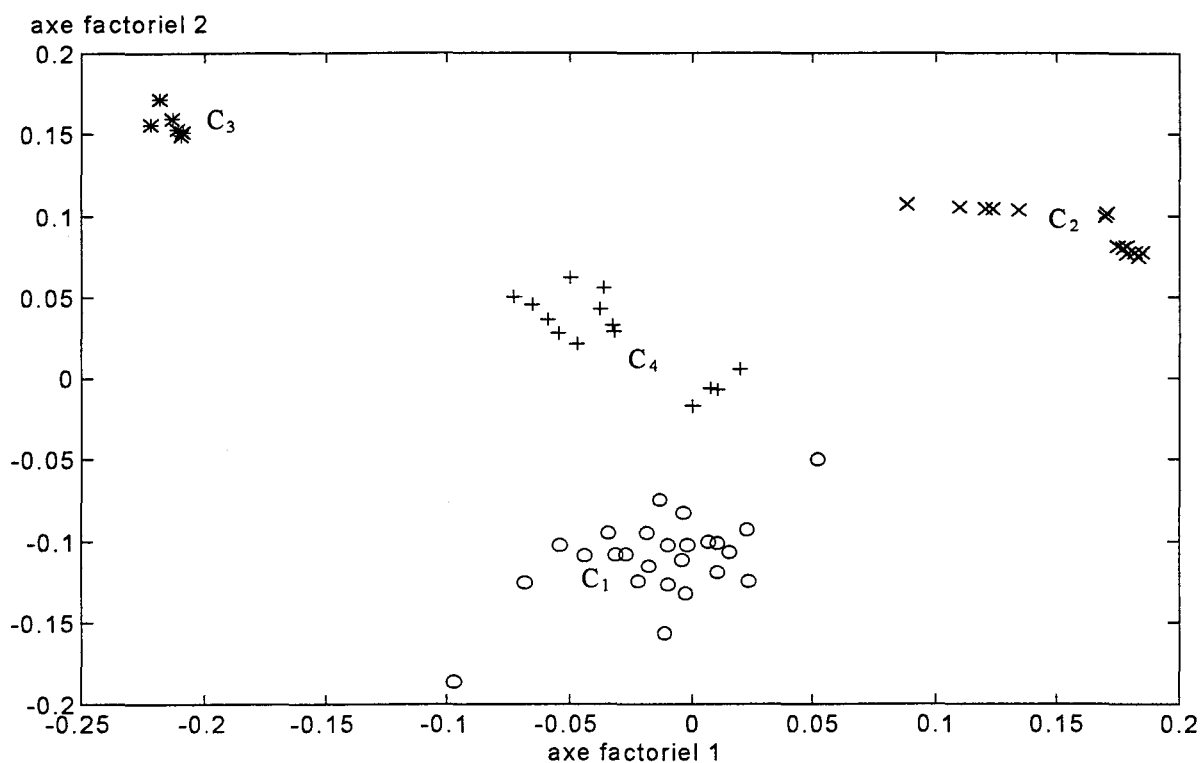


Figure 17. Visualisation de la répartition du nuage de points et des 4 classes en fonction des deux premiers axes factoriels (taux de représentation : 66.83%).

L'histogramme des valeurs croissantes, des taux de représentation du nuage de points, associées à chaque axe factoriel (Figure 18) permet d'évaluer la validité obtenue par la combinaison des axes factoriels correspondants. Ainsi, les 5 premiers axes factoriels aboutissent à une représentation de 90.54% de l'inertie totale du nuage de points.

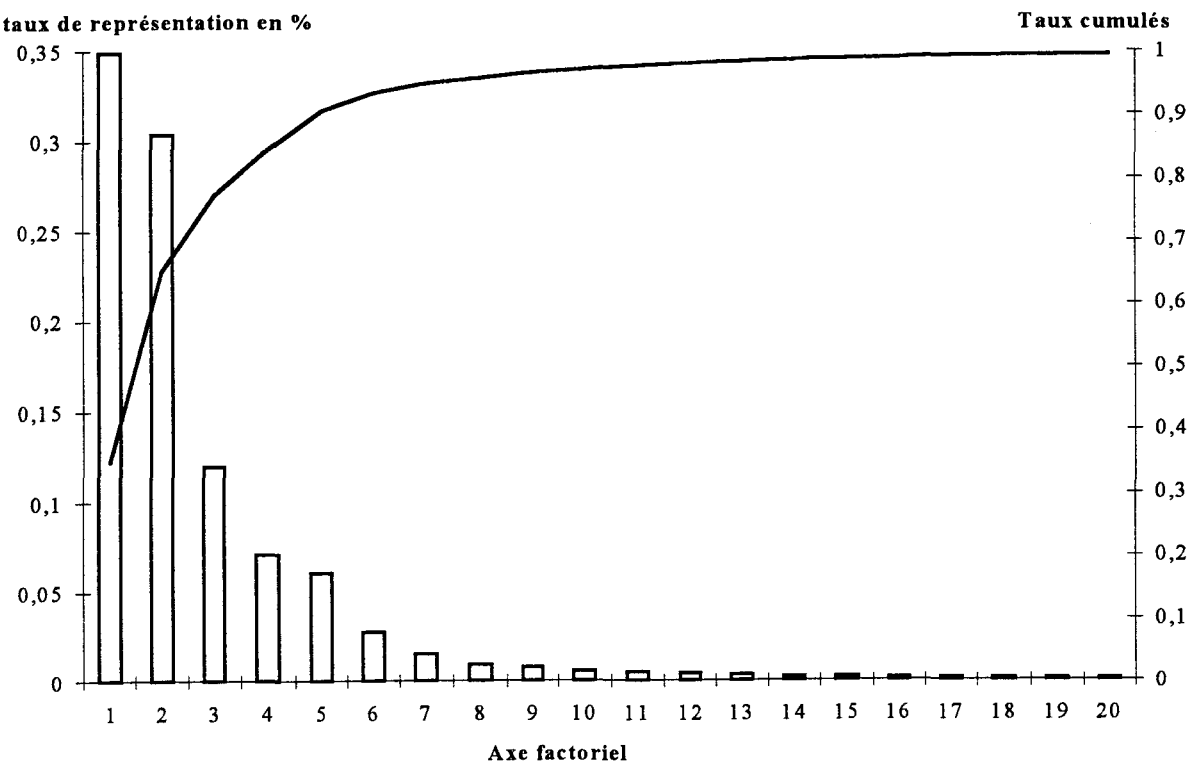


Figure 18. Représentation des valeurs des taux associées
à chaque axe factoriel.

Les allures de vente des 4 centres de classe de la partition sont représentées en
Figure 19 :

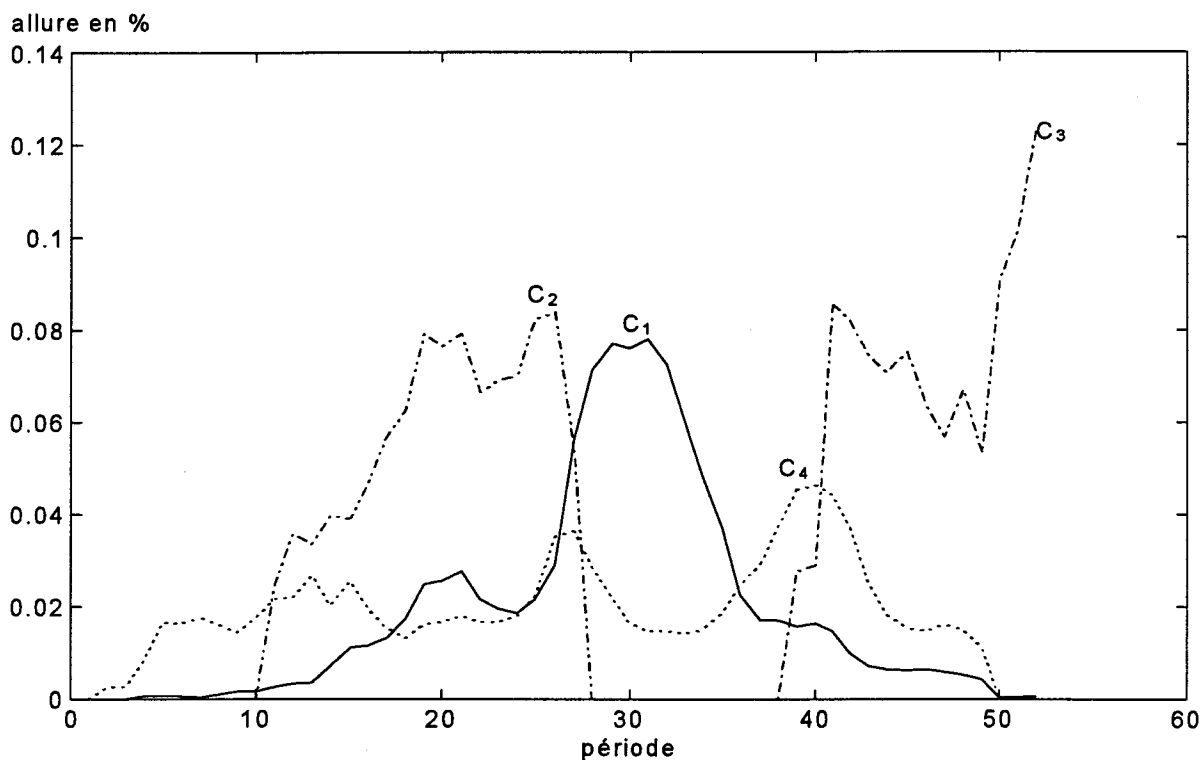


Figure 19. Représentation des allures des 4 centres de classe de la classification des données quantitatives.

L'objectif de classification des articles suivant les valeurs de vente est atteint en représentant les allures des centres de classes.

Les 4 classes représentées par leurs centres contiennent les numéros d'articles suivants:

$$C_1 = [1;2;3;4;5;6;7;8;9;10;11;12;13;14;15;16;17;18;19;20;48;51;55;57;59]$$

$$C_2 = [21;22;23;24;25;44;45;46;47;49;50;53;56;58]$$

$$C_3 = [26;27;28;29;30;31]$$

$$C_4 = [32;33;34;35;36;37;38;39;40;41;42;43;52;54]$$

Les compacités π des classes sont respectivement :

$$\pi(C_1) = 3.6233$$

$$\pi(C_2) = 2.8049$$

$$\pi(C_3) = 0.3817$$

$$\pi(C_4) = 0.4202$$

Une première contrainte sur la taille de la classe peut être effectuée en choisissant une valeur de compacité minimale π par classe. Un test d'acceptation des classes constituées peut être ajouté au cours de l'évaluation des partitions par le critère de validité S afin de déterminer le nombre de classes optimal satisfaisant à la contrainte de compacité. Cependant, le choix d'une valeur minimale de compacité peut sembler trop abstrait pour l'industriel. L'utilisation d'une contrainte locale sur l'allure du centre de classe permet d'obtenir une représentation plus pertinente de la dispersion des allures de vente des articles autour de la valeur moyenne. Les limites supérieures et inférieures par rapport au centre de classe à chaque période peuvent être déterminées en considérant respectivement les valeurs maximales et minimales des allures des articles à chaque période. L'allure du centre de classe C_4 et ses limites supérieures et inférieures calculées à chaque période, sont représentées en Figure 20 ainsi que les valeurs du quartile supérieur (regroupant 75% des valeurs des allures des articles du centre de classe C_4) et celles du quartile inférieur (25% des valeurs des allures des articles du centre de classe C_4).

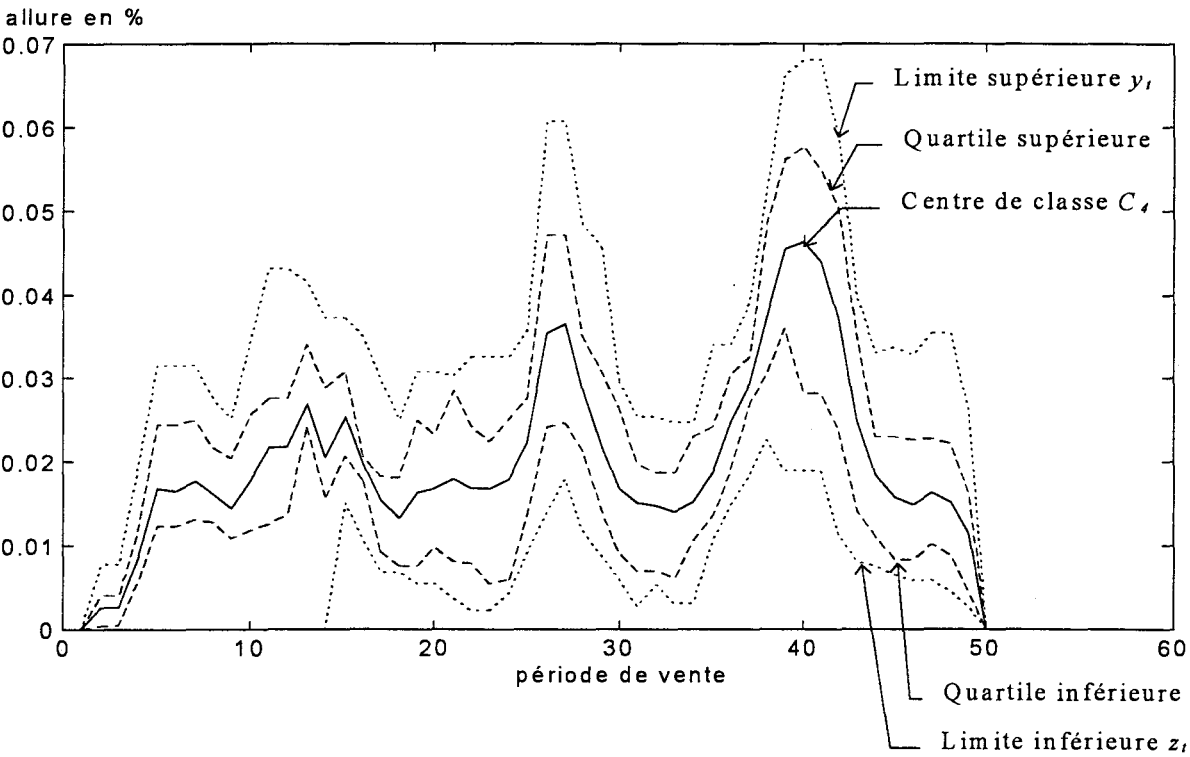


Figure 20. Représentation de l'allure du centre de classe C_4 de ses limites maximales, de ses quartiles inférieures et supérieures.

A chaque période, l'écart entre les valeurs des limites inférieures et supérieures avec l'allure du centre de classe peut s'écrire de façon additive ou multiplicative. Par exemple, les écarts supérieurs et inférieurs additifs sont calculés à chaque période à partir des valeurs limites supérieures et inférieures de la classe et des valeurs de l'allure du centre de classe (Figure 21). Les valeurs de l'écart supérieur multiplicatif sont calculées, à chaque période, à partir de la division de la valeur de la limite supérieure par la valeur de l'allure du centre de classe à laquelle est retranchée une unité. Les valeurs de l'écart inférieur multiplicatifs sont obtenues par la différence entre une unité et la division de la valeur de la limite inférieure par la valeur de l'allure du centre de classe (Figure 22).

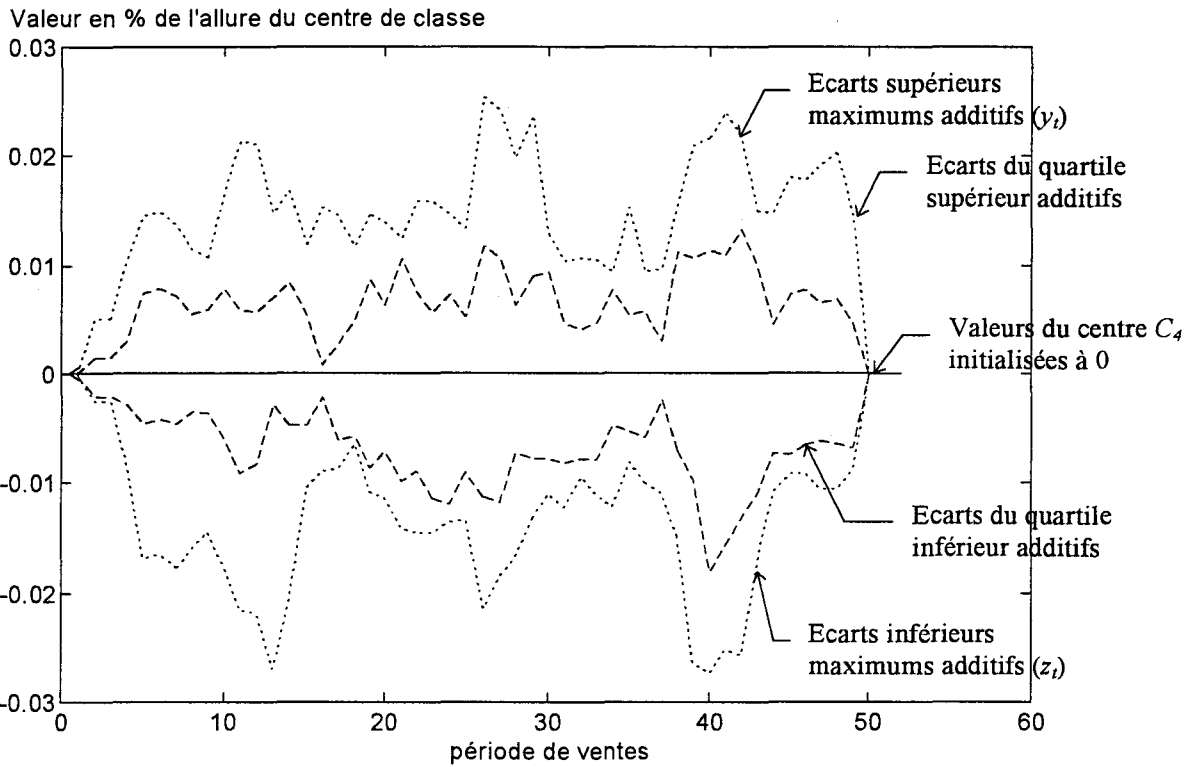


Figure 21. Valeurs des écarts additifs locaux des maximums et quartiles supérieurs et inférieurs du centre de classe C_4 .

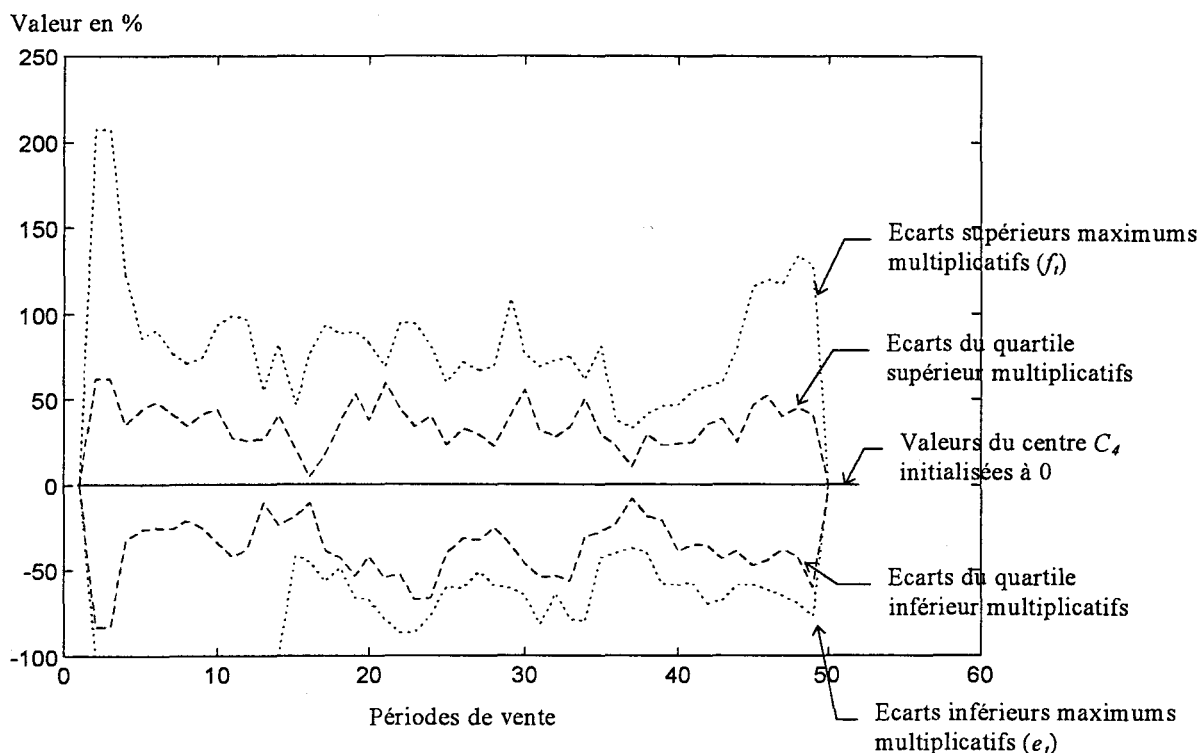


Figure 22. Valeurs des écarts multiplicatifs locaux des maximums et des quartiles supérieurs et inférieurs du centre de classe C_4 .

Le choix d'une contrainte locale ou globale sur les allures des articles d'une classe ou sur la compacité de la partition obtenue permet d'introduire une notion de "précision" des résultats de la classification. Cela permet également de choisir un nombre de classes suffisamment représentatif du degré de détail souhaité au niveau de l'information obtenue par la classification. La constitution des classes est effectuée en tenant uniquement compte des allures de vente des articles. La représentation (cf. Tableau 19) des partitions issues de la classification à caractère quantitatif, à laquelle est associée les taux de représentation des variables qualitatives des centres de classes, permet d'établir les relations entre les valeurs de vente et les paramètres des articles. Les centres de classes possèdent ainsi une représentation symbolique complète par rapport au taux maximum de représentation des variables qualitatives en %.

	Centre C ₁		Centre C ₂		Centre C ₃		Centre C ₄	
type	Piscine	80%	Piscine	64%	Piscine	100%	Gym	100%
genre	Enfant	56%	Femme	64%	Homme	100%	Enfant	85%
désignation	Maillot de bain	76%	Juste au corps 1 pièce	42%	Short	100%	Juste au corps 1 pièce fille	28.57%
couleur	Noire	64%	Gris standard	21.4%	V1 noir bleu	80%	Nero IK	35.7%
modèle	1622	20%	7646	21.4%	7925	66%	5071	28.6%
tissu	2519	25%	2073	28.5%	3107	66%	2519	92%

Tableau 19. Croisement des partitions des classifications
sur les données qualitatives et quantitatives.

Par exemple, les articles représentés par le centre de classe C_1 sont, au maximum, à 80% de type Piscine, à 56% de genre Enfant, à 76% relatif à la désignation Maillot de bain, à 64% de couleur Noire, à 20% de modèle 1622 et à 25% relatif au tissu 2519. Certaines valeurs des paramètres endogènes des articles sont représentées par une allure unique de vente telle que, par exemple, la classe C_3 des articles "Piscine Homme short". L'association de la caractéristique endogène tissu "2519" avec les articles de type "gym" et de genre "enfant" permet d'expliquer en grande partie le même comportement de vente des articles de la classe C_4 (classification quantitative). Les valeurs du taux de représentation de la variable qualitative inférieures à 50% ne permettent pas d'associer une explication du comportement de vente à la variable considérée.

Cette description permet d'attribuer une allure de vente, voire un modèle de prévision adapté, aux nouveaux articles de la collection dont les caractéristiques endogènes sont similaires à celles d'un centre de classe existant.

III.1.2. Application aux variables qualitatives

Les résultats obtenus en utilisant les fonctions de similitude sur les caractéristiques qualitatives des 59 articles ne s'avèrent pas significatifs. La mesure de ressemblance entre les variables qualitatives est définie par la formule (44). En effet, la procédure de construction d'un centre de classe sur des données qualitatives repose sur la différence existante entre les poids de représentation des variables qualitatives de chaque article. Mais, la combinaison de deux variables qualitatives n'aboutit pas à la description d'une nouvelle variable qualitative. Cet inconvénient majeur limite la description du centre de classe aux valeurs des variables qualitatives existantes.

Si la fonction de similitude globale concerne une seule variable qualitative alors le résultat de classification correspond à un simple tri effectué selon les différentes valeurs de cette variable.

Si la fonction de similitude globale combine deux ou plusieurs variables qualitatives alors le résultat de classification correspond à la détermination de l'ensemble des combinaisons existantes des différentes valeurs de ces variables.

III.2. Application de la méthode de classification par partition floue

L'introduction d'une méthode de classification floue permet de décrire de façon non catégorique les différentes influences des variables qualitatives de tous les articles sur les différents centres de classes.

III.2.1. Application aux allures de ventes

La simulation est réalisée sur les mêmes données quantitatives issues du Tableau 17 de l'algorithme précédent en faisant varier le nombre de classes c de 2 à 20, et pour chaque valeur de c , 10 valeurs initiales différentes des valeurs des degrés d'appartenance sont effectuées avec une valeur de degré de flou m égale à 2 (valeur souvent préconisée dans les publications et pour laquelle la convergence vers une valeur unique est prouvée [YAN 93]). Chacune des partitions floues de « c » classes est évaluée par le critère de validité J dont les valeurs sont représentées en Figure 23.

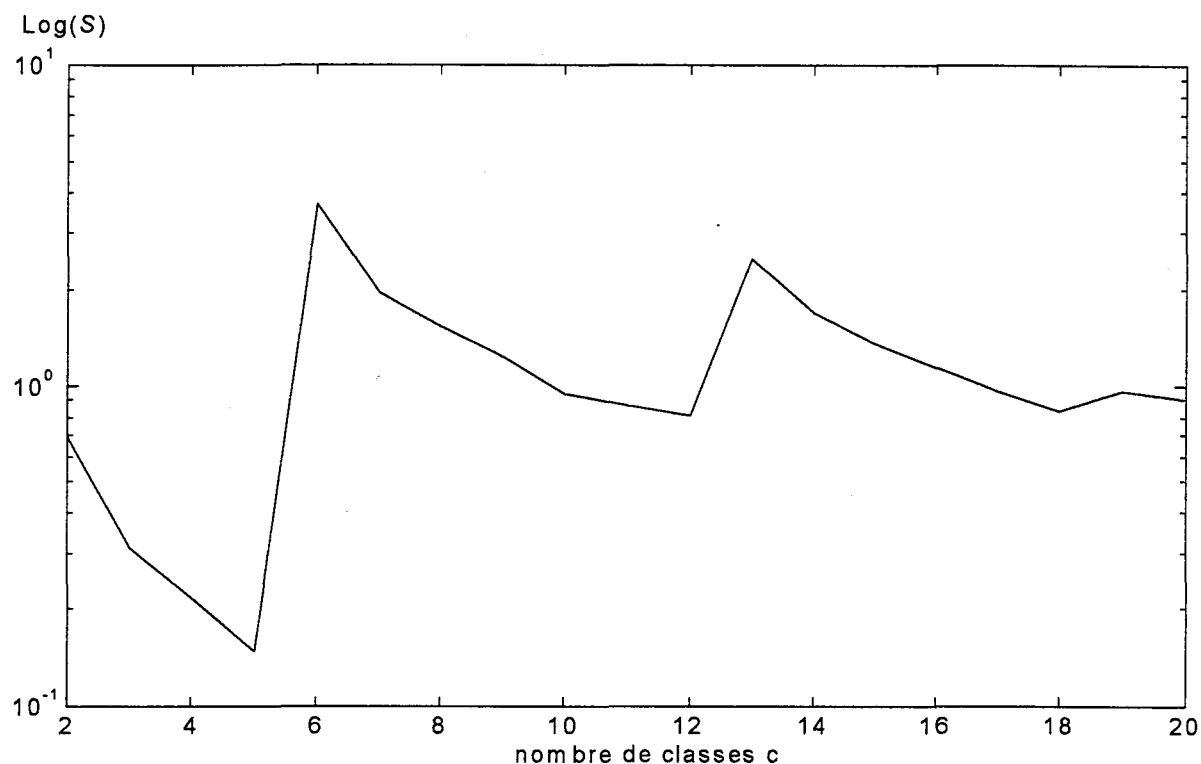


Figure 23. Evaluation des " c " partitions floues des données quantitatives.

La partition floue peut également se représenter en Figure 24 par des valeurs croissantes des degrés d'appartenance de chaque élément à classer en fonction des centres de classes [RAB 97]. Cette représentation permet de visualiser le degré de flou de la représentation du nuage de points. La majorité des articles possède une valeur du degré d'appartenance à une classe supérieure à 0.5, ce qui laisse supposer que les points sont "bien" regroupés autour de leurs centres respectifs.

degrés d'appartenance des codes
d'articles aux centres de classe

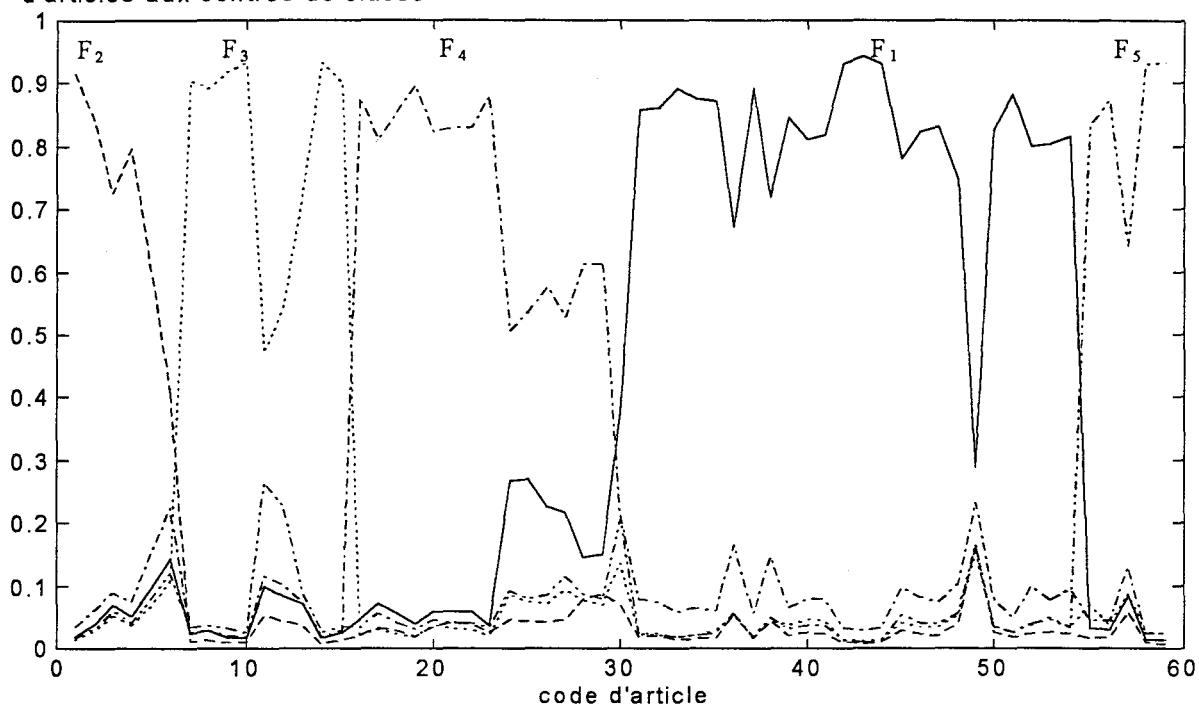


Figure 24. Visualisation des degrés d'appartenance des
codes d'articles⁴ par rapport aux centres des classes.

L'information complémentaire de la représentation de la position des centres de classes (symbolisé par des o) suivant les deux premiers axes factoriels permet de visualiser en Figure 25 les résultats de la classification par minimisation d'une fonction objective.

⁴ Afin d'obtenir une représentation des degrés d'appartenance des centres de classes dans un ordre croissant des valeurs maximales, les codes d'articles utilisés ne correspondent pas aux numéros des articles.

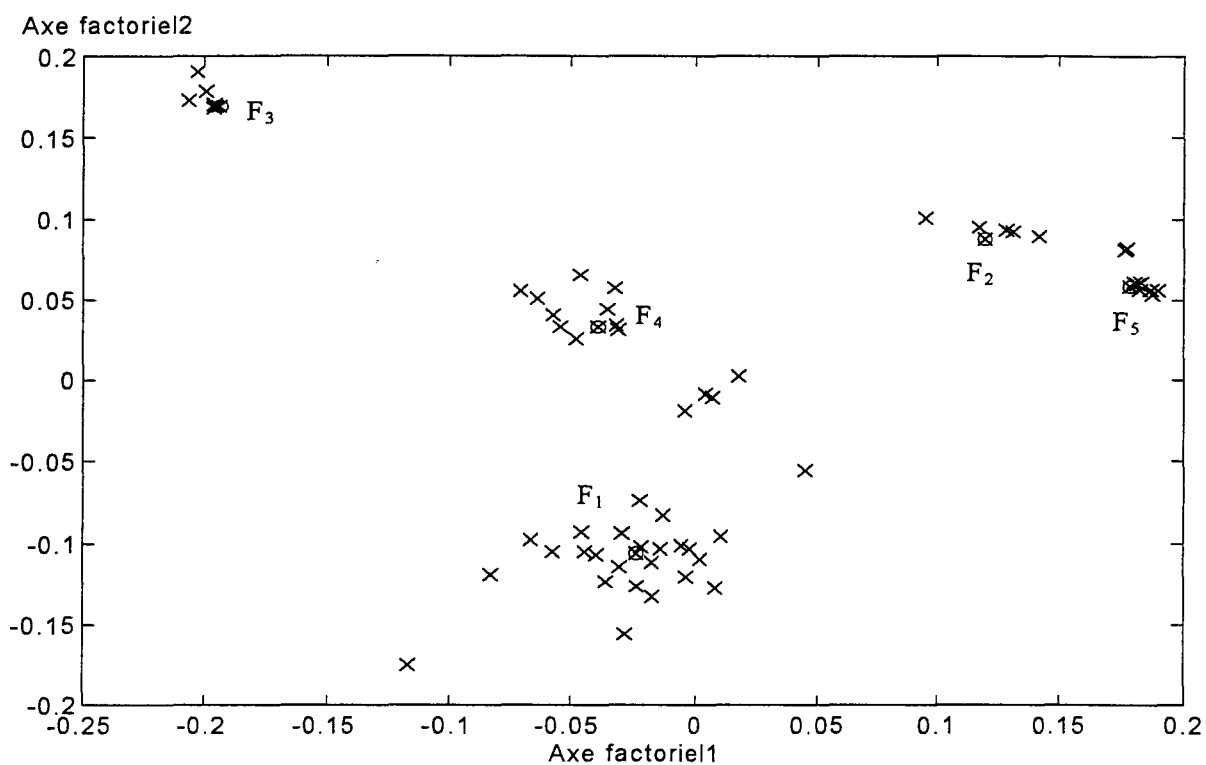


Figure 25. Visualisation de la répartition du nuage de points et des 5 centres de classes en fonction des deux premiers axes factoriels (taux de représentation : 66.83%).

Les allures (cf. Figure 26) de chaque centre de classe sont calculées à partir de la valeur de chaque degré d'appartenance de tous les éléments à classer afin d'optimiser la fonction objective J_0 .

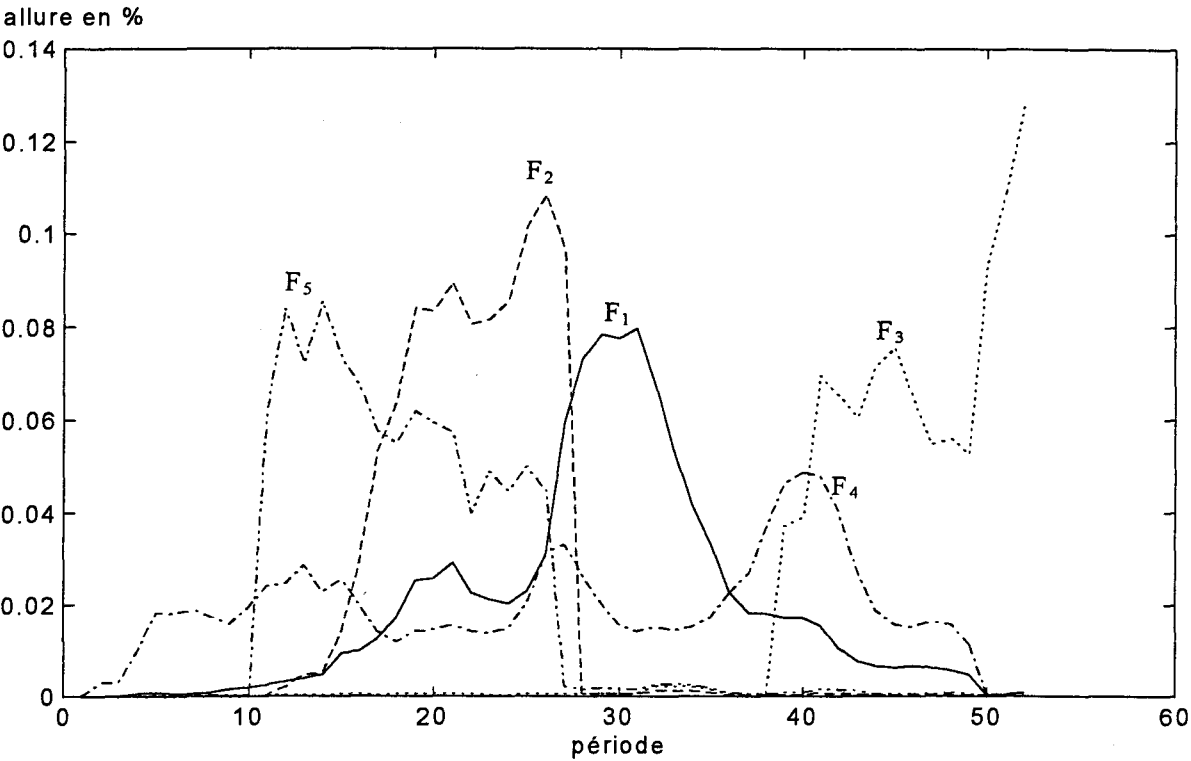


Figure 26. Représentation des allures des 5 centres de classes floues.

L'objectif de classification des articles suivant les caractéristiques de vente est atteint par la représentation différente des allures des centres de classes. De plus, la description des caractéristiques qualitatives des centres de classes est représentée en fonction de l'ensemble des paramètres endogènes des articles à classer, pondérés par leurs degrés d'appartenance respectifs. Dans notre exemple, les centres de classes sont décrits par les taux de représentations les plus forts des caractéristiques majoritaires en fonction de leurs descriptions symboliques globales.

	Centre F ₁		Centre F ₂		Centre F ₃		Centre F ₄		Centre F ₅	
type	Piscine	72.47%	Gym	85.29%	Piscine	84.81%	Gym	80.18%	Piscine	75.24%
genre	Enfant	59.05%	Femme	82.98%	Homme	74.05%	Enfant	72.66%	Homme	62.42%
désignation	Maillot de bain slip garçon	23.40%	Juste au corps 1 pièce	49.82%	Short	72.22%	Juste au corps 1 pièce fille	25.08%	Maillot de bain	57.64%
couleur	Nero IK	13.19%	Gris standard	26.94%	V1 noir bleu	37.77%	Nero IK	18.10%	Nero	16.00%
modèle	1622	14.11%	3811	20.60%	7920	29.68%	5071	25.08%	7646	32.15%
tissu	2519	34.71%	2073	36.49%	3107	55.30%	2519	72.71%	2500	36.23%

Tableau 20. Représentation des caractéristiques qualitatives majoritaires des différents centres de classes floues.

Les comportements de vente, représentés par les allures de centres de classes, peuvent être expliqués par la nature différente de certaines caractéristiques qualitatives. Mais, la réciproque n'est pas forcément vraie. Cependant, la différence observée entre les natures qualitatives peut avoir une influence indirecte sur le comportement de vente en fonction de phénomènes extérieurs caractérisés par des variables exogènes. Par exemple, le comportement de vente des articles de type Piscine Enfant n'est pas forcément dû à la nature des produits mais peut être fortement influencé par un paramètre exogène lié au calendrier tel que la rentrée des classes. Ainsi, la connaissance du degré et de l'intensité de l'influence des paramètres exogènes sur les allures des centres de classes permet de fournir une information supplémentaire sur le choix adéquat du modèle de prévision.

III.3. Comparaison des résultats obtenus

Les mêmes caractéristiques qualitatives et quantitatives sont obtenues pour les centres de classes C_1 et F_1 , C_3 et F_3 , C_4 et F_4 . Les fortes valeurs des taux de représentation des coefficients associés aux variables qualitatives permettent d'effectuer une distinction nette de ces articles entre eux. Le centre de classe C_2 restant possède les taux de représentation des caractéristiques qualitatives les plus forts des centres F_2 et F_5 . La différence de représentation peut s'expliquer par la formulation de la fonction objective de l'algorithme des "c" classes moyennes flous qui positionne les centres de classes par la minimisation d'un critère de compacité. Ce critère tend à agréger les points du nuage le plus près des centres de classe. De ce fait, les centres F_2 et F_5 possèdent des taux de représentation des valeurs des variables qualitatives fortement différents mais dont les allures de vente des articles se situent aux mêmes instants et possèdent une même durée de vie.

De ce fait, l'utilisation d'un algorithme basé sur la minimisation d'une fonction objective permet d'aboutir à un positionnement et une description plus fine des variables qualitatives des centres de classes. Mais, l'information obtenue se limite à cette description de la partition optimale. Si une description de l'historique de la constitution des classes est souhaitée, le recours à un algorithme de classification hiérarchique s'avère alors nécessaire.

IV. Mesure de l'influence des paramètres exogènes sur les données de vente

De façon générale, un paramètre exogène peut s'interpréter comme un phénomène extérieur dont la source n'est pas associée aux caractéristiques propres du produit. Les paramètres exogènes reflètent plus un environnement économique, ou encore des habitudes d'achat, pouvant conduire à des résultats non prévisibles parmi les variables endogènes. Le choix des variables exogènes pouvant influencer les ventes d'articles textiles, parmi l'ensemble des phénomènes, reste à l'appréciation de l'industriel possédant à lui seul une connaissance intuitive de son environnement économique. Nous pouvons citer un certain nombre de variables exogènes d'usage courant et facilement disponibles.

Les données, relatives aux phénomènes de consommation, issues de base de données de l'INSEE (Institut Nationale Supérieure des Etudes Economiques) peuvent être obtenues à faible coût, ou par l'intermédiaire de sociétés spécialisées telles que : SECODIP ou CTCOE, pour un coût d'acquisition plus élevé.

Parmi les données relatives à des phénomènes météorologiques, issues de la base de données de Météo France, pouvant influencer les ventes d'articles textiles, la variable température est sélectionnée (en degrés Celsius). Celle-ci, de fréquence hebdomadaire, regroupe les valeurs de température des 22 stations météorologiques réparties dans toute la France (Abbeville; Ajaccio; Beauvais; Besançon; Bordeaux; Caen; Chateauroux; Clermont-Ferrand; Dijon; Lille; Lyon-Bron; Marseille; Montpellier; Nancy-Essey; Nantes; Orléans; Paris-Montsouris; Poitiers; Reims; Rennes; Strasbourg; Toulouse-Blagnac).

Les données relatives à des phénomènes de type "actions spéciales" telles que : l'effet de la publicité, l'effet de mode et les opérations promotionnelles, sont des données caractéristiques à une situation économique précise et ne sont pas exposées dans ce chapitre.

Dans notre application, l'ensemble des transformations appliquées sur la variable exogène permet de rendre linéaire une variable, supposée au préalable non linéaire. Le calcul de la corrélation est effectué à partir de cet ensemble de fonctions et des allures de vente des centres de classe suivant une fenêtre d'observation fixe. La translation suivant l'axe des temps de la fenêtre d'observation d'un paramètre par rapport à un autre permet

d'introduire ensuite un décalage temporel entre les séries de vente. La valeur maximale de la valeur de corrélation entre l'allure de vente du centre de classe et la fonction appliquée à la variable exogène permet de déterminer le type de fonction correspondant et la valeur maximale du décalage. Ce calcul de corrélation envisage ainsi les deux procédures suivantes :

IV.1. Procédure de transformation des variables exogènes.

L'application de l'ensemble Φ des « k » fonctions, avec $\Phi = \{f_1, \dots, f_k\}$, à chaque série temporelle S_p de la variable exogène p , avec $S_p = (x_1 \dots x_{T1})^t$, aboutit au calcul des valeurs initiales d'une matrice $X(p)$ des variables exogènes transformées telle que chaque ligne caractérise le résultat de l'application de la fonction f_k à la variable exogène. La fonction f_j s'applique de la façon suivante :

$$\begin{aligned} \text{pour tout } x_t \in S_p \text{ avec } t \in [1, T1], \quad \mathbb{R} &\rightarrow \mathbb{R} \\ x_t &\rightarrow f_j(x_t) \end{aligned}$$

Afin d'illustrer notre méthode, l'ensemble des fonctions choisies est volontairement restreint aux trois procédures f_1, f_2 et f_3 telles que :

$$f_1(x_t) = x_t \tag{46}$$

$$f_2(x_t) = \sqrt{x_t} \tag{47}$$

$$f_3(x_t) = \nabla x_t = x_t - x_{t-1} \text{ pour } t \in [2, T1] \tag{48}$$

IV.2. Procédure de corrélation entre variables exogènes et allures de vente.

Soit le vecteur $f_j(S_p)$ issu de la matrice $X(p)$ exogène p avec $S_p = (s_1 \dots s_{T1})^t$,
et $X = \{x_1, x_2, \dots, x_n\}$ un ensemble de « n » articles chacun représenté par un vecteur des allures de vente R_i de dimension T avec $R_i = (r_i(t1); r_i(t2); \dots ; r_i(T))^t$.

Dans nos applications, la fenêtre de corrélation entre les allures de vente et les paramètres exogènes est égale à la durée de vie de chaque article, avec un décalage " l " de la fenêtre d'observation de taille fixe. Dans le cadre de l'analyse des allures de vente dans le

domaine du court terme, l'influence des paramètres exogènes sur les allures de vente des articles textiles n'est calculée que pour un retard maximum l de 30 semaines. De ce fait, la fenêtre de corrélation fixe oblige à respecter la condition suivante pour tous les articles :

$$Tl > T+l$$

avec T représentant le nombre total de périodes de ventes de chaque article
et Tl le nombre de périodes total de chaque paramètre exogène.

Un vecteur de corrélation $corr(R_i, S_p)$ est défini entre l'allure de vente de l'article i et le paramètre exogène p tel que :

$$corr(R_i, S_p) = (corcoef(R_i(t), S_p(t)) \dots \dots corcoef(R_i(t-k), S_p(t)) \dots \dots corcoef(R_i(t-l), S_p(t)))^t.$$

avec $k \in [0, l]$ et

$$corcoef(R_i(t-k), S_p(t)) = \frac{cov(R_i(t-k), S_p(t))}{var(R_i(t-k)) var(S_p(t))} \quad (49)$$

les opérateurs cov et var désignent respectivement la covariance et la variance.

La valeur de corrélation ainsi calculée, et évaluée par un test statistique, permet de conclure à la présence ou non d'une dépendance linéaire entre les deux variables.

IV.3. Application au paramètre exogène température

Dans le cadre d'une identification à court terme des influences des variables exogènes sur les allures de vente des articles textiles, les paramètres sélectionnés pour nos applications doivent avoir une fréquence hebdomadaire. Les données météorologiques température, précipitation et facteur d'insolation répondent à ce critère de sélection. Seule la variable température est exposée dans ce chapitre afin d'illustrer notre démarche. Les allures de ventes correspondantes sont les allures des 5 centres de classes identifiées par la méthode de classification floue sur les données à caractère quantitatifs en Figure 26.

L'évolution de la température moyenne hebdomadaire pour l'ensemble des stations météorologiques est représentée sur la Figure 27 avec comme date de début le lundi 03/01/94 et comme date de fin le dimanche 31/12/95, soit 2 ans (120 périodes d'une semaine).

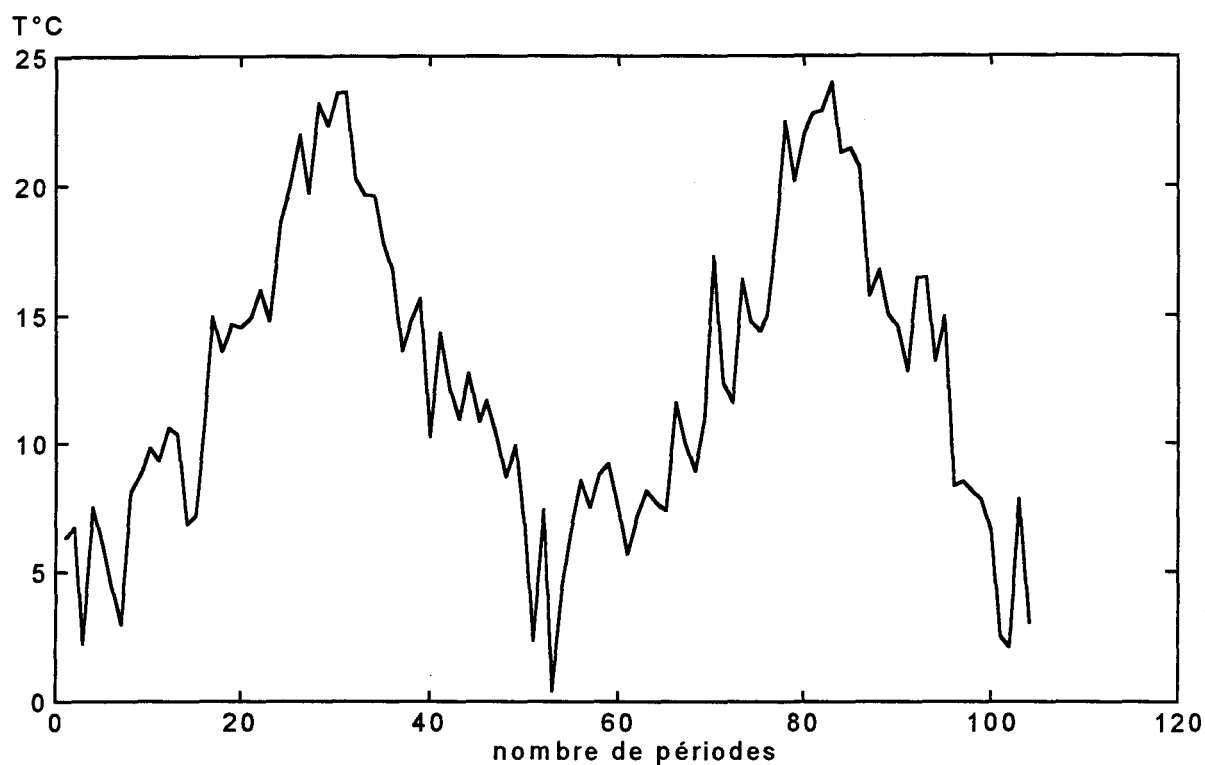


Figure 27. Evolution de la température moyenne nationale T en degré Celsius en fonction du temps

L'application de la procédure de transformation de la variable exogène température en fonction de l'ensemble des fonctions $\Phi = \{f_1, f_2, f_3\}$ précédemment choisi aboutit à une matrice de différentes allures de la variable exogène, dont les fonctions f_1 (identité), f_2 (racine carrée) et f_3 (différentiation première) sont représentées en Figure 28.

Valeurs des fonctions sur la température $T^{\circ}\text{C}$

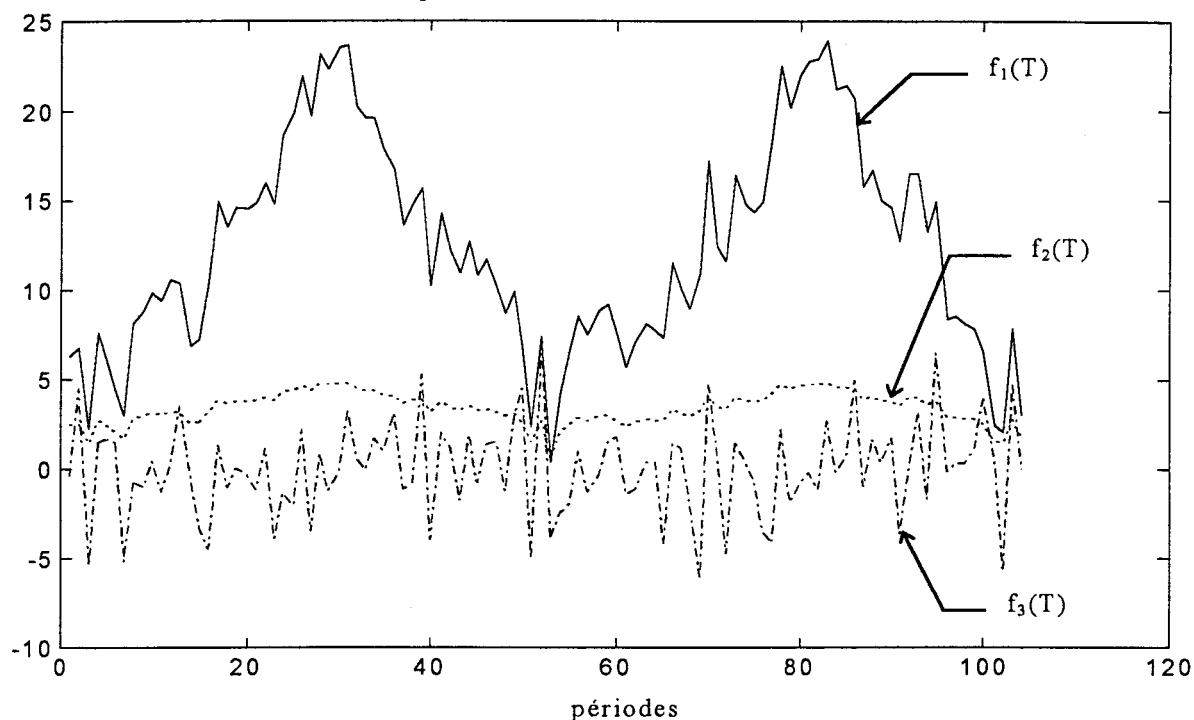


Figure 28. Représentation de l'application des fonctions identité, racine carrée et différentiation première sur le vecteur exogène température.

L'ensemble des fonctions Φ est également appliqué sur les allures des centres de classes flous (F_1, F_2, F_3, F_4, F_5). Les corrélations sont calculées entre les fonctions appliquées aux allures des centres de classe et au paramètre exogène température. Les valeurs maximales des coefficients de corrélation obtenues pour chaque centre de classes sont représentées en Figure 29. Par exemple, la corrélation maximale de +0.92, entre la fonction f_2 (racine carré) appliquée à l'allure de vente du centre F_2 et la fonction f_1 (identité) appliquée au vecteur exogène température T , est obtenue pour un décalage de 2 périodes entre ces deux séries. Cela traduit une évolution dans le même sens du comportement de vente des articles, représenté par le centre de classe F_2 , et du phénomène exogène température, décalé de 2 périodes.

Valeurs de corrélation entre les centres de classe et le paramètre température

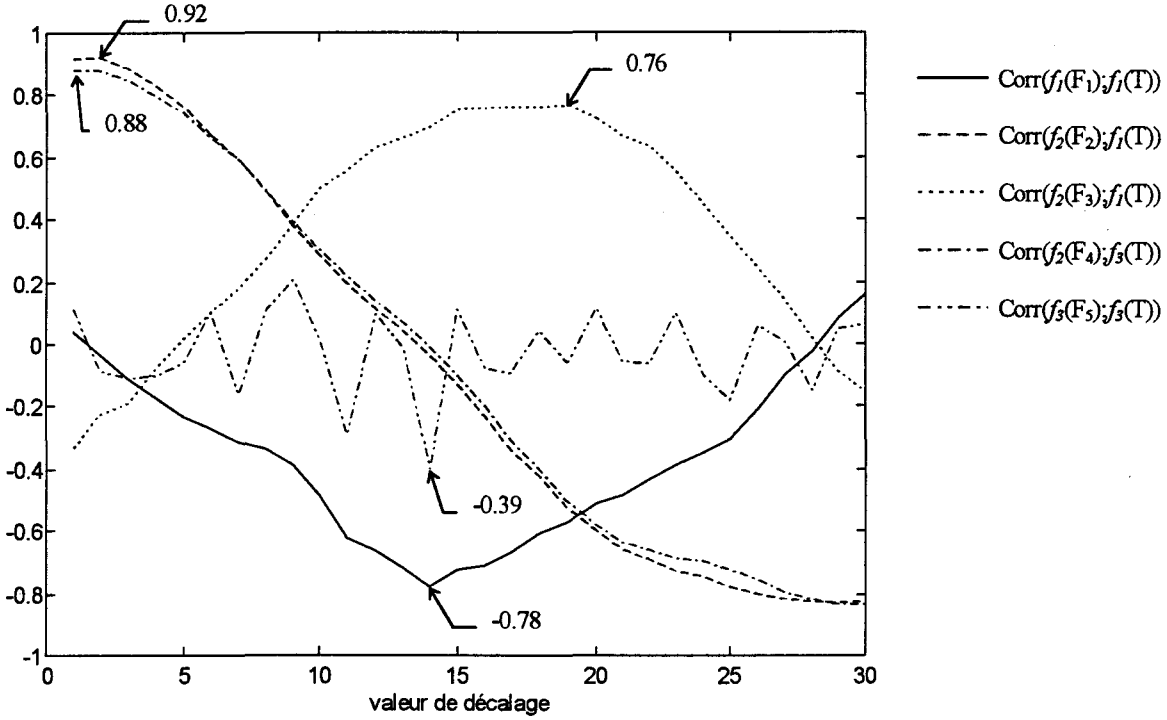


Figure 29. Représentation des valeurs de corrélation entre les allures des centres de classes et les fonctions appliquées à la variable exogène température aboutissant à la valeur maximale de corrélation.

Un test relatif à l'existence d'une corrélation linéaire est effectué sur chaque valeur maximale des coefficients de corrélation en supposant que la population est caractérisée par une loi de Student à $n-2$ degrés de liberté, dont les éléments sont de la forme

$$\frac{\text{corr}(x,y) \cdot \sqrt{n-2}}{\sqrt{1-\text{corr}^2(x,y)}} \text{ [LAV 81].}$$

Si le coefficient de corrélation $\text{corr}(x,y)$ donne à l'expression $\frac{\text{corr}(x,y) \cdot \sqrt{n-2}}{\sqrt{1-\text{corr}^2(x,y)}}$ une valeur supérieure à $+t_{n-2,\frac{\alpha}{2}}$ ou inférieure à $-t_{n-2,\frac{\alpha}{2}}$, alors la corrélation linéaire est significative au niveau de la précision α imposé.

	F ₁	F ₂	F ₃	F ₄	F ₅
<i>corr(x,y)</i>	-0.78	0.92	0.76	0.88	-0.39
$\frac{corr(x,y)*\sqrt{n-2}}{\sqrt{1-corr^2(x,y)}}$	-6.59	12.69	6.30	9.89	-2.30
$ t_{n-2;\frac{\alpha}{2}} =2.76$ ($\alpha=1\%$)	significatif	significatif	significatif	significatif	non significatif

Tableau 21. **Résultats des tests statistiques sur les valeurs de corrélation.**

Les seuls résultats directement exploitables, dans le cadre d'une utilisation de la variable exogène température dans le contexte de la prévision à court terme, correspondent à la corrélation entre les allures de vente des centres de classe F₄ et F₂ et la variable température décalée respectivement de une et deux périodes. Au delà de 10 périodes, l'intérêt de l'utilisation du paramètre exogène température semble difficilement interprétable pour expliquer le comportement de vente des articles.

IV.4. Avantages et inconvénients de la méthode.

Cette méthode permet de fournir un ensemble d'informations sur les influences directes ou indirectes des phénomènes extérieurs sur les comportements de vente des produits. La connaissance plus précise des comportements de vente des articles permet d'élaborer des modèles de prévision adaptés à l'environnement économique. Cependant, le choix arbitraire des fonctions *f_i* et des paramètres exogènes parmi l'ensemble des indices relatifs à des phénomènes extérieurs influant sur les ventes peut ne pas fournir d'explication adéquate. Enfin, l'existence de relation de dépendance non linéaire entre les variables exogènes n'est pas entièrement envisagée dans cette approche.

V. Conclusion

En fonction de la nature et du type des données à analyser, les procédures de traitement sont différentes. Une procédure de traitement par une analyse en composantes principales sur des données normalisées permet de visualiser l'information contenue dans le nuage de points à observer ; mais, l'appréciation de la répartition des points ainsi obtenue dans des classes respectives repose uniquement sur une estimation visuelle et personnelle de l'utilisateur. L'obtention d'une hiérarchie par un algorithme de classification ascendant (ou descendant) hiérarchique permet de connaître l'historique de la constitution des classes au fur et à mesure de la procédure d'agrégation. Un algorithme de recherche accélérée tel que l'algorithme des voisins réciproques nous a permis d'aboutir à l'élaboration d'une hiérarchie indicée et des partitions correspondantes. De plus, la visualisation de cette hiérarchie est rendue possible par un dendogramme. Cependant, la recherche de la partition optimale, celle qui maximise l'information avec un minimum de données, repose également sur une appréciation visuelle de l'aspect du dendogramme par l'utilisateur. L'obtention d'une partition, par rapport à un critère d'optimisation fondé sur une fonction objective tout en n'affectant pas directement chaque individu à une seule classe, est rendue possible par l'utilisation d'un algorithme de classification floue, par exemple l'algorithme des " c " classes moyennes floues. Cependant, le choix du nombre de classes et donc de la taille de la partition optimale reste à l'appréciation de l'utilisateur. L'utilisation d'un critère de validité permet de résoudre partiellement le choix de la partition qui rend optimale la compacité et la séparabilité des différentes classes, et ce pour n'importe quel type d'algorithme. Ceci n'assure pas pour autant d'atteindre la partition optimale globale d'un problème de classification de données en catégories. L'introduction d'une contrainte de compacité globale, traduisible localement sur les allures de ventes, permet de mieux définir la précision souhaitée au niveau des résultats de classification.

La caractérisation, la plus précise, de l'essentiel de l'information contenue dans un nuage de points, par une représentation symbolique complète des centres de classes, peut aboutir à une meilleure connaissance du comportement de ventes de l'ensemble des articles au sein d'une collection. Cette description peut être approfondie par l'examen de

influence des phénomènes extérieurs (caractérisés par des paramètres exogènes) sur les ventes des articles. Les méthodologies de classification proposées dans ce chapitre constituent des aides à une analyse des données de vente des articles textiles et fournissent l'essentiel de l'information pour l'identification d'un modèle de prévision adapté.

La validation de notre approche de classification en vue d'améliorer la prévision des ventes n'est pas actuellement envisagée pour la raison principale suivante : le manque de connaissance actuelle sur l'évolution dynamique des partitions optimales identifiées d'une saison à l'autre. Une identification correcte des règles de passage d'une partition évaluée d'une saison précise à l'autre doit être envisagée. Des tentatives d'évaluation de ces règles sur notre exemple traité ne permettent d'aboutir à une analyse satisfaisante. Ceci est principalement dû à la difficulté d'une part de sélectionner les bonnes variables exogènes pour les mêmes articles et d'autre part de prendre en compte les modalités des caractéristiques qualitatives des nouveaux articles.

La partition identifiée par le critère de validité peut s'avérer être "optimale" localement. Afin d'explorer l'ensemble des solutions optimales de répartition des données au sein des différentes classes, l'utilisation d'un algorithme évolutif, fondé sur des opérations de transformations génétiques des partitions significativement codées, permet de se rapprocher plus rapidement de la répartition optimale globale. Cette amélioration de notre approche de classification constitue l'objet du prochain chapitre. La procédure de classification par les algorithmes génétiques va autoriser une exploration aléatoire de l'espace de recherche tout en la "guidant" et permettre d'identifier la solution quasi-optimale parmi l'ensemble des minimums locaux.

Chapitre 4

Amélioration de l'approche de classification par des procédures génétiques.

I. Principe de la méthodologie proposée

Les algorithmes de type évolutifs peuvent se fonder sur les idées et les techniques issues de la théorie génétique. Un processus "évolutionnaire" est simulé avec " n " individus qui représentent les points situés dans l'espace de recherche. A chaque pas de l'algorithme génétique, appelée une génération, chaque individu est évalué en fonction de l'ensemble des individus regroupés au sein d'une population. La population évolue plus ou moins rapidement ; à chaque génération, les solutions relativement « bonnes » sont reproduites, tandis que les solutions relativement « mauvaises » sont éliminées. La distinction entre les différentes solutions réside dans l'utilisation d'une fonction objective qui joue le rôle de l'environnement dans lequel la population doit « génétiquement » évoluer pour survivre [MIC 92]. Des opérateurs génétiques de mutation et de recombinaison ("crossover" en anglais) assurent les modifications dans l'espace de recherche. L'opérateur de recombinaison combine des portions de gènes issues de la coupure de deux chromosomes, sélectionnés en fonction d'une valeur de probabilité P_c , pour former deux

nouveaux individus. Une opération de mutation doit remplacer les chromosomes d'un individu sélectionnés en fonction d'une valeur de probabilité P_m et aboutir à un nouvel individu. L'opérateur de mutation permet une recherche proche d'un point de l'espace par petits "sauts" autour de l'optimum alors que l'opérateur de recombinaison occasionne des "sauts" plus importants d'un point de l'espace de recherche vers un autre. Toute la difficulté d'un algorithme génétique réside dans la définition des opérateurs génétiques adéquats qui permettent d'explorer l'espace de recherche de façon non totalement aléatoire et sans perte de l'information acquise au cours des générations. De façon générale, les algorithmes génétiques permettent d'explorer et d'exploiter l'espace de recherche de toutes les solutions possibles en se rapprochant le plus de l'optimum global. Tout en conservant une population de solutions potentielles, ils encouragent la formation et l'échange d'informations.

Seules les méthodes de classification basées sur l'optimisation d'une fonction objective peuvent envisager l'emploi d'une procédure évolutive définie à partir des opérateurs génétiques. Différents types de codage de partition existent et reposent uniquement sur des affectations non floues des éléments aux classes. La fonction objective utilisée dans notre méthodologie est celle de l'algorithme des " c " classes moyennes.

Dans la première partie, les éléments constituant un algorithme génétique de classification sont présentés tels que :

- les différents types de codage,
- les types d'opérateurs de recombinaison et de mutation, et leurs heuristiques associées,
- les valeurs des probabilités de mutation et de recombinaison.

Dans une deuxième partie, des applications sur cet algorithme sont présentées.

1.1. Principaux codages utilisés

1.1.1. Par affectation [JON 91]

Le premier type de codage d'une partition est défini en fonction des " n " éléments et des " c " classes à construire. Chaque valeur (chromosome) de l'individu, prise dans l'ordre croissant, correspond à l'affectation de l'élément à un numéro de classe k avec $k \in \{1, \dots, c\}$.

Par exemple, les 10 éléments répartis dans les 3 classes constituent une partition P telle que :

$$P = \{(1\ 2\ 3\ 6) (5\ 8) (4\ 7\ 9\ 10)\}$$

L'individu i correspondant à la partition P par le codage par affectation est représenté par:

$$i = (1\ 1\ 1\ 3\ 2\ 1\ 3\ 2\ 3\ 3)$$

I.1.2. Par permutation [JON 91]

Le deuxième type de codage d'une partition par permutation repose sur la définition de séparateurs. Les " n " éléments sont codés par leurs numéros de classement dans la représentation d'une partition et les " c " classes sont séparées par " $c-1$ " séparateurs dont les valeurs se situent entre " $n+1$ " et " $n+c$ ".

Par exemple, l'individu i correspondant à la partition $P = \{(1\ 2\ 3\ 6) (5\ 8) (4\ 7\ 9\ 10)\}$ est représenté par le codage par permutation :

$$i = (1\ 2\ 3\ 6\ \underline{11}\ 5\ 8\ \underline{12}\ 4\ 7\ 9\ 10)$$

où les numéros 11 et 12 représentent les séparateurs entre les 3 classes de la partition P .

I.1.3. Par permutation avec heuristique de décodage [JON 91]

Le troisième type de codage repose sur un codage par permutation suivi d'une heuristique de décodage des articles appartenant aux classes. L'interprétation de cette représentation fait appel à une heuristique dont l'utilisation nécessite d'importantes ressources de traitement des données. Les " k " premiers éléments de chaque classe contenus dans la représentation de l'individu i sont utilisés pour initialiser les " k " classes. Les objets restants sont ajoutés selon une règle : premier entré - premier sorti (FIFO) suivant l'ordre dans lequel ils apparaissent dans la chaîne de caractères et, sont placés dans la classe qui conduit à la meilleure valeur de la fonction objective.

Par exemple, à partir de l'individu $i = (1\ 2\ 3\ 6\ \underline{11}\ 5\ 8\ \underline{12}\ 4\ 7\ 9\ 10)$ correspondant à la partition P , les 3 premiers éléments sélectionnés de chaque classe sont : 1, 5 et 4 ainsi que les séparateurs 11 et 12.

Le nouvel individu correspondant est $i^* = (1\ * \ * \ * \ \underline{11}\ 5\ * \ \underline{12}\ 4\ * \ * \ *)$ où les $*$ représentent les futures valeurs possibles.

Le premier élément restant de l'individu i est le numéro 2 pour lequel son affectation à la classe 2 aboutit à une valeur minimale de la fonction objective.

L'individu i^* correspondant devient : $i^*=(1 \text{ * * * } \underline{11} \text{ } 5 \text{ } 2 \text{ } \underline{12} \text{ } 4 \text{ * * *})$

L'élément suivant restant de l'individu i est le numéro 3 pour lequel son affectation à la classe 1 aboutit à une valeur minimale de la fonction objective.

L'individu i^* correspondant devient : $i^*=(1 \text{ } 3 \text{ * * } \underline{11} \text{ } 5 \text{ } 2 \text{ } \underline{12} \text{ } 4 \text{ * * *})$

Au final, les classes peuvent contenir un nombre d'éléments différents de l'individu initial.

Par exemple, l'individu i^* correspondant s'écrit : $i^*=(1 \text{ } 3 \text{ } \underline{11} \text{ } 5 \text{ } 2 \text{ } 6 \text{ } 8 \text{ } 10 \text{ } \underline{12} \text{ } 4 \text{ } 7 \text{ } 9)$

1.1.4. Comparaison

Les résultats obtenus sur deux applications différentes [JON 91] des différents codages à partir d'un même algorithme génétique de classification ne permettent pas de mettre en valeur un type de codage par rapport à un autre. L'approche de codage par permutation utilisant une heuristique de décodage donne de meilleurs résultats que le codage par permutation simple dans les deux applications mais nécessite des temps de calcul beaucoup plus importants. Pour l'une des applications donnée, les résultats sont meilleurs par le codage par permutation avec heuristique que le codage par affectation ; pour l'autre application, les résultats sont inversés. Les auteurs de cette étude suggèrent d'utiliser le codage par affectation pour minimiser les temps de calcul [JON 91]. Ce codage par affectation permet aussi d'appliquer des opérateurs de mutation standard.

1.2. Opérateurs de recombinaison

Chaque chromosome i (avec $i = \{1, \dots, n\}$) représente une partition de " c " classes et se trouve constitué de " m " gènes (" m " objets à classer).

Chaque gène j (avec $j = \{1, \dots, m\}$) du chromosome i noté x_{ij} représente l'affectation de l'objet j au numéro de classe k (avec $k = \{1, \dots, c\}$).

Soit la population initiale $\Pi_n^{(0)}$ de " n " chromosomes à la génération ($g=0$).

De façon générale, la procédure de recombinaison s'applique à chaque chromosome i de la population $\Pi_n^{(g)}$ de la façon suivante [MIC 92] :

1. Attribuer à chaque chromosome i une variable aléatoire $p_i(i)$ dont la valeur se situe entre 0 et 1.

2. Constituer une sous-population $\Pi_b^{(g)}$ de " b " chromosomes par la sélection des individus de la population initiale $\Pi_n^{(0)}$ de " n " chromosomes à la génération g dont la valeur $p_c(i)$ est inférieure à une probabilité "crossover" définie P_c .
3. Si la taille de la sous-population $\Pi_b^{(g)}$ est impaire; alors il suffit de sélectionner dans la sous-population restante un chromosome de façon aléatoire.
4. Sélectionner deux individus de la sous-population $\Pi_b^{(g)}$
Déterminer la variable de coupure b des deux chromosomes par le tirage au hasard d'une valeur entre 2 et $m-1$.
Appliquer l'opérateur de recombinaison *cross* de telle sorte que les deux chromosomes $chr(p)$ et $chr(q)$ se combinent pour former deux nouveaux chromosomes $chr(p^*)$ et $chr(q^*)$ tel que :

$$\forall (p, q) \in \Pi_b^{(g)},$$

$$chr(p) = (x_{p1} \dots x_{pb} x_{pb+1} \dots x_{pm}) \text{ et } chr(q) = (x_{q1} \dots x_{qb} x_{qb+1} \dots x_{qm}),$$

$$chr(p^*) = (x_{p1} \dots x_{pb} x_{qb+1} \dots x_{qm}) \text{ et } chr(q^*) = (x_{q1} \dots x_{qb} x_{pb+1} \dots x_{pm})$$
5. Affecter les nouveaux chromosomes $chr(p^*)$ et $chr(q^*)$ à la sous-population $\Pi_b^{(g+1)}$ des " b " chromosomes pour la génération suivante $g+1$.
6. Répéter à partir de l'étape 5 jusqu'à ce que la taille de la sous-population $\Pi_b^{(g)}$ soit égale à 0.

La nouvelle population $\Pi_n^{(g+1)}$ de " n " chromosomes pour la génération $g+1$ est obtenue en combinant la sous-population $\Pi_b^{(g+1)}$ des " b " chromosomes pour la génération $g+1$ obtenue ci-dessus et la sous-population $\Pi_{n-b}^{(g)}$ des " $n-b$ " chromosomes restants de la génération précédente g .

Pour le codage par affectation choisi, l'étude de Jones et al. [JON 91] présente une comparaison statistique des différents opérateurs de recombinaison sur la valeur moyenne de la fonction d'évaluation obtenue par les individus optimaux.

Les différents opérateurs sont les suivants :

Cas 1. Opérateur de recombinaison avec rejet.

La méthode du rejet complète la procédure de recombinaison précédemment définie par la condition suivante :

Si deux individus recombinaisonnés représentent une partition de moins de " c " classes, les deux chromosomes sont rejetés de la population finale et remplacés par les chromosomes initiaux.

Cas 2. Opérateur de recombinaison avec rejet et renumérotation.

La renumérotation des gènes des individus dans l'ordre croissant permet d'éviter la représentation d'une partition identique par deux codages différents.

Cas 3. Opérateur de recombinaison uniforme avec rejet.

La procédure de recombinaison uniforme [SYS 89] "construit" des chromosomes pour la nouvelle population $\Pi_n^{(g+1)}$ à partir des individus de la population précédente par la recombinaison de l'ensemble des 2^n solutions possibles. Un critère d'arrêt relatif au nombre de solutions permet de limiter la recherche des combinaisons possibles et réduit ainsi les temps de calcul.

Cas 4. Opérateur de recombinaison uniforme avec rejet et renumérotation.

Cet opérateur regroupe l'ensemble des trois méthodes précédentes.

Cas 5. Opérateur de recombinaison par couples de classes.

La procédure de recombinaison par classe repose sur une heuristique [WHI 89] qui permet de construire deux chromosomes, à partir des deux individus de la population précédente, en fonction des $\frac{k(k-1)}{2}$ couples d'éléments possibles de chaque classe de " k " éléments. Les partitions ainsi obtenues sont ensuite validées par l'examen de l'ensemble des combinaisons de classe existantes.

Comparaisons

Les meilleurs résultats sont obtenus pour le cas n°5 de l'opérateur de recombinaison à partir d'une même application. Le seul inconvénient de cet opérateur réside dans son temps de calcul qui peut être de l'ordre de c^4 [JON 91] avec " c " classes et devient donc inutilisable pour des grandes valeurs du nombre de classes. L'opérateur de recombinaison uniforme peut aboutir à des temps de calcul assez longs. La méthode de renumérotation ne permet pas d'aboutir à des résultats significatifs en terme de précision et de temps de calcul à cause du phénomène de redondance lors de l'utilisation de l'opérateur de mutation. Enfin, l'opérateur de recombinaison avec rejet (cas n°1), commandé par les auteurs [JON 91], fournit des résultats suffisamment précis et plus

rapides que tout autre opérateur. Dans nos applications, cet opérateur de recombinaison (cas n°1) est utilisé.

1.3. Opérateurs de mutation

La procédure généralement utilisée pour une opération de mutation peut être définie par les étapes suivantes [MIC 92] :

Soit la population de " n " chromosomes $\Pi_n^{(g)}$ disponibles à la génération g , chacun des chromosomes caractérise une partition de " c " classes contenant " m " gènes dont la valeur x_{ij} représente la valeur de l'affectation du point j à un centre de classe a_k avec $k \in \{1, \dots, c\}$.

La procédure de mutation pour l'ensemble des individus (ou chromosomes) de la population $\Pi_n^{(g)}$ s'écrit :

1. Attribuer à chaque chromosome i de l'ensemble des individus une variable aléatoire $p_m(i)$ dont la valeur se situe entre 0 et 1.

2. Pour chaque gène j allant de 1 à m

si $p_m(j) < P_m$ (probabilité de mutation) alors l'opérateur de mutation mut est appliqué sur le gène j de l'ensemble des individus de la population $\Pi_n^{(g)}$ tel que :

initialement $x_{ij} = k_1$ avec $k_1 \in \{1, \dots, c\}$; $mut(x_{ij}) = k_2$ avec $k_2 \in \{1, \dots, c\}$ et $k_1 \neq k_2$

sinon la valeur d'affectation du j^{me} gène reste inchangée

3. fin de boucle pour la valeur j .

Pour le codage par affectation choisi, la sélection des gènes à muter s'effectue de façon complètement aléatoire et reste la même pour tous les individus de la population, ce qui peut engendrer un nombre important de générations supplémentaires sans l'emploi de quelques heuristiques. Ainsi, la première heuristique concernant l'opérateur de mutation est définie de telle sorte que le gène le plus éloigné de son centre de classe respectif possède une probabilité de mutation non nulle. Cette heuristique « oriente » la sélection du gène, pour l'opération de mutation, permettant de se rapprocher le plus rapidement de l'optimum global. Cette heuristique permet à un seul gène d'être muté ou pas. La deuxième heuristique affecte une probabilité de mutation par gène et par individu pondérée par la distance de chaque gène par rapport à son centre d'affectation (notion de

compacité). Avec cette heuristique, l'ensemble des gènes d'un chromosome possède une probabilité de mutation. Le gène le plus éloigné de sa classe a une plus grande chance d'être muté. Enfin, la troisième heuristique, repose sur le même principe que la deuxième heuristique sauf que la valeur du coefficient de pondération tient compte de l'ensemble des compacités des classes. L'ensemble des gènes d'un chromosome possède une probabilité de mutation dont le plus éloigné de son centre, parmi l'ensemble des classes, a plus de chance d'être muté.

Heuristique 1

Cette procédure de mutation pour l'ensemble des individus (ou chromosomes) de la population $\Pi_n^{(g)}$ s'écrit :

1. Attribuer à chaque chromosome i de l'ensemble des individus une variable aléatoire $p_m(i)$ dont la valeur se situe entre 0 et 1.
2. Pour chaque chromosome i allant de 1 à n

Pour chaque gène j allant de 1 à m

Calculer pour chaque gène j du chromosome i la distance euclidienne $d_2(y_j, a_k)$ entre le point y_j et le centre de classe a_k correspondant à la valeur d'affectation x_{ij} .

fin de boucle pour j

Identifier le gène j^* tel que : $d_2(y_{j^*}, a_k) = \max_{j \in \{1, \dots, m\}} (d_2(y_j, a_k))$

Affecter la probabilité P_m au gène optimal j^* correspondant

si $p_m(i) < P_m$ (probabilité de mutation) alors l'opérateur de mutation mut est appliqué sur le gène j^* du chromosome i de la population $\Pi_n^{(g)}$ tel que :

initialement $x_{ij} = k_1$; $mut(x_{ij}) = k_2$ avec $k_2 \in \{1, \dots, c\}$ et $k_1 \neq k_2$

sinon la valeur d'affectation du gène j^* reste inchangée

3. fin de boucle pour la valeur i .

Heuristique 2

Cette procédure de mutation pour l'ensemble des individus (ou chromosomes) de la population $\Pi_n^{(g)}$ s'écrit :

1. Attribuer à chaque chromosome i et à chaque gène j une variable aléatoire $p_m(ij)$ dont la valeur se situe entre 0 et 1.

2. Pour chaque chromosome i allant de 1 à n

Pour chaque gène j allant de 1 à m

Calculer pour chaque gène j du chromosome i la distance euclidienne $d_2(y_{ji}, a_k)$ entre le point y_{ji} et le centre de classe a_k correspondant à la valeur d'affectation x_{ji} .

fin de boucle pour j

Calculer un coefficient de pondération c_j tel que :

Pour tous les gènes $j \in \{1, \dots, m\}$, $c_j = \frac{d_2(y_{ji}, a_k)}{\sum_{x_{ji} \in C_k} d_2(y_{ji}, a_k)}$ avec C_k désignant la classe de

centre a_k

Calculer la moyenne $\bar{c}_j$ des valeurs de c_j pour l'individu i

Calculer les nouvelles valeurs de c_j par translation d'une valeur $|\bar{c}_j - 1|$ pour obtenir une valeur moyenne $\bar{c}_j$ égale à 1.

Identifier les gènes à muter tels que :

pour tout $j \in [1, m]$,

si $p_m(j) < c_j P_m$ alors l'opérateur de mutation mut est appliqué sur le gène j du chromosome i de la population $\Pi_n^{(g)}$ tel que :

initialement $x_{ji} = k_1$; $mut(x_{ji}) = k_2$ avec $k_1 \in \{1, \dots, c\}$ et $k_1 \neq k_2$

sinon la valeur d'affectation du j^{me} gène reste inchangée

3. Fin de boucle pour la valeur i .

Heuristique 3

Cette procédure de mutation pour l'ensemble des individus (ou chromosomes) de la population $\Pi_n^{(g)}$ repose sur les mêmes étapes de l'heuristique 2. La différence réside dans la formulation du coefficient de pondération c_{ij} qui permet de tenir compte de l'ensemble des compacités des classes.

Pour tous les gènes $j \in \{1, \dots, m\}$, $c_j = \frac{d_2(y_{ji}, a_k)}{\sum_{k=1}^c \sum_{x_{ji} \in C_k} d_2(y_{ji}, a_k)}$ avec C_k désignant la classe de

centre a_k

I.4. Algorithme génétique de classification

La structure de base de l'algorithme génétique de classification pour un codage par affectation peut être représentée de la façon suivante :

Soit une population initiale (génération : $g=0$) $\Pi_n^{(0)}$ de " n " chromosomes.

1. Initialiser g à 1.
2. Calculer la fonction objective $E_v(i)$ pour chaque chromosome i issu de la population $\Pi_n^{(g-1)}$
3. Calculer la fonction de mérite $f(i)$ pour chaque chromosome i telle que :

$$f(i) = E_{max} - E_v(i) \text{ avec } E_{max} = \max_{i \in [1,n]} (E_v(i)) \quad (50)$$

4. Calculer la fonction de probabilité de sélection $p(i)$ et la fonction de probabilité cumulée $Cp(i)$ pour chaque chromosome i telle que :

$$p(i) = \frac{f(i) - \min_{j \in [1,n]} (f(j))}{\sum_{i=1}^n (f(i) - \min_{j \in [1,n]} (f(j)))} \text{ et } Cp(i) = Cp(i) + p(i) \quad (51)$$

5. Sélectionner les " n " nouveaux individus de la population $\Pi_n^{(g)}$ à partir de la fonction de probabilité cumulée et des individus de la population $\Pi_n^{(g-1)}$ par la stratégie de la « roulette russe » [GOL 89].
6. Appliquer l'opérateur de recombinaison sur les " m " individus sélectionnés dans la population $\Pi_n^{(g)}$ en fonction d'une probabilité de "crossover" P_c .
7. Appliquer l'opérateur de mutation sur les gènes sélectionnés de chaque chromosome de la population $\Pi_n^{(g)}$ en fonction d'une probabilité de mutation P_m .
8. Ajouter 1 à la variable g et retourner à l'étape 1 jusqu'à ce que g soit égale à max_gen (maximum de générations).

1.5. Choix des valeurs du couple (P_c et P_m)

De façon générale, la valeur de la probabilité P_m est faible pour éviter de modifier complètement l'individu et permettre de s'approcher le plus possible de l'optimum global. La valeur de la probabilité P_c doit avoir une valeur suffisante pour espérer une exploration totale de l'espace de recherche.

Le couple (P_c , P_m) des valeurs recommandées par Schaffer et al. [SCH 89] correspond à :

cas 1. $P_c=0.6$ et $P_m = 0.0091$

Le couple (P_c , P_m) des valeurs recommandées par Hesser et Männer [HES 90] correspond à :

cas 2. $P_c=0.6$ et $P_m = 0.011$

Une approche probabiliste [GRE 95], appuyée par des résultats statistiques issus d'une analyse de la variance, suggère d'utiliser une valeur de P_m égale à 0.011 et une valeur de P_c égale à 0.6.

Le couple (P_c , P_m) des valeurs recommandées par Bastian [BAS 97] correspond à :

cas 3. $P_c=0.95$ et $P_m = \frac{1}{n}$ avec " n " caractérisant la longueur de l'individu, c'est

à dire le nombre d'articles à classer.

L'ensemble des valeurs recommandées pour P_m sont proches de 0.01. Cependant, dans une étude de Von Laszewski [VON 91], une valeur de P_m égale à 0.2 est utilisée pour trouver la solution optimale.

Chacune de ces études est relative à un algorithme génétique de classification couplé avec des opérateurs de recombinaison et de mutation adaptés. Différentes valeurs des probabilités associées aux opérateurs sont possibles. Dans nos applications, les cas n°1, 2 et 3 des valeurs initiales des probabilités de mutation et de recombinaison sont utilisés pour permettre une analyse comparative des résultats de classification.

II. Applications des stratégies adoptées

Initialement, l'idée contenue dans l'étude de Phanendra Babu et al. [PHA 93] consiste à évaluer un ensemble de partition de " n " éléments répartis dans " c " classes pour aboutir à une répartition optimale. L'individu optimal obtenu sert de partition initiale pour l'algorithme des " c " classes moyennes (ou k-means). Les temps de convergence de l'algorithme de classification sont réduits en utilisant cette répartition initiale. Notre approche de classification repose sur le même type de programme évolutif, fondé sur un codage par affectation, et sur l'optimisation d'une fonction objective dont le but est d'aboutir à la répartition optimale globale. La fonction objective E_v utilisée est celle de l'algorithme des " c " classes moyennes qui minimise, pour l'ensemble des classes, la distance d entre l'ensemble des points $x_j(i)$, correspondant à une partition i , à leurs centres de classes respectifs $a_k(i)$, telle que :

$$E_v(i) = \sum_{k=1}^c \sum_{j=1}^{m_k(i)} d^2(x_j(i), a_k(i)) \quad (52)$$

avec :

- c le nombre total de classes,
- $m_k(i)$ le nombre de points contenus dans la classe $C_k(i)$,
- $C_k(i)$ représentant la classe de centre $a_k(i)$ de la partition i .

Les trois heuristiques précédentes ont été appliquées sur l'opérateur de mutation P_m afin d'orienter l'exploration de l'ensemble des partitions optimales et réduire ainsi le nombre de générations nécessaires.

Pour l'ensemble des simulations de l'algorithme génétique de classification, une répartition de 100 articles dans 4 classes (avec 25 articles par classe) est obtenue en distribuant uniformément les points de la classe autour de son centre, avec un diamètre maximum de 0.3, dont la représentation est donnée en Figure 30.

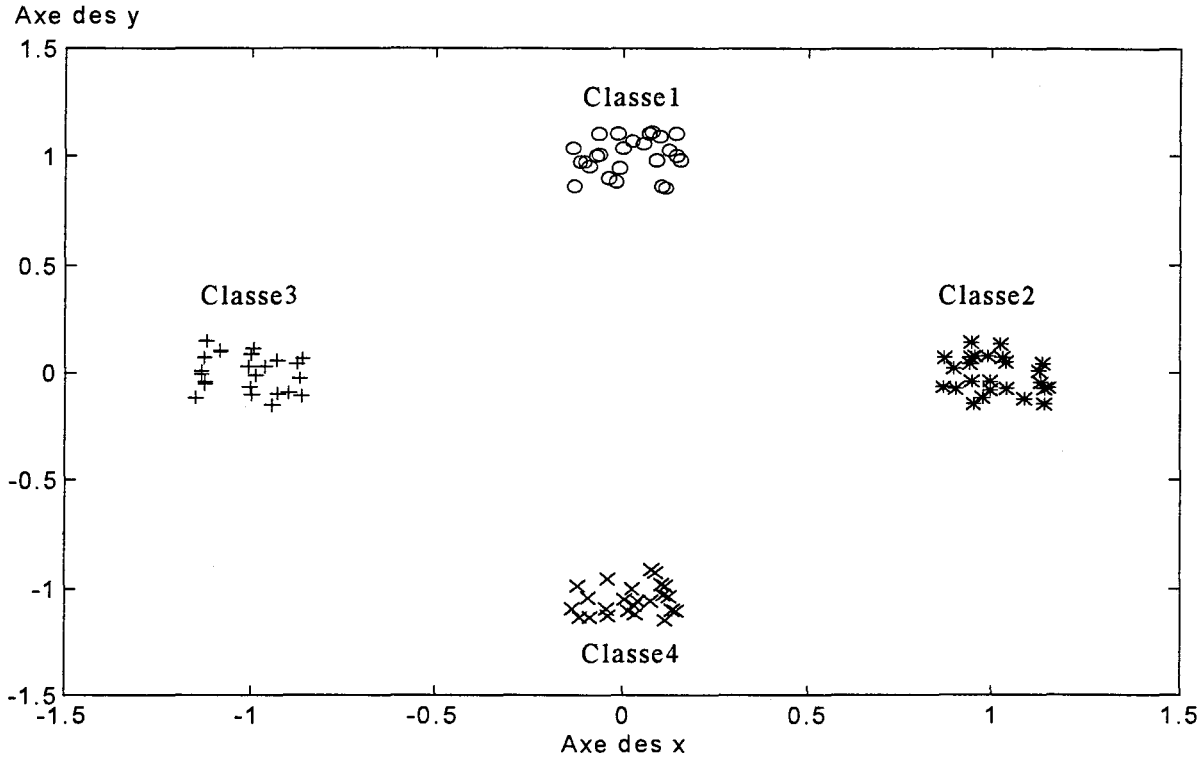
$$\text{Classe4} = (76, 77, \dots, 100)$$


Figure 30. **Représentation du nuage de points répartis en 4 classes.**

L'erreur de classification en pourcentage est calculée à partir des partitions obtenues par l'algorithme génétique de classification et la répartition initiale. Par exemple, la répartition initiale P_{init} est codée par :

$$P_{init}=(\underbrace{1111111111111111}_{N=16}, \underbrace{2222222222222222}_{N=16}, \underbrace{3333333333333333}_{N=16}, \underbrace{4444444444444444}_{N=16})$$

L'un des chromosomes de la population aboutit à la partition optimale P_{opt} codée telle que :

[illegible]

Les 4 gènes différents, de la partition optimale P_{opt} , parmi les 100, engendrent une erreur de classification de 4%.

Deux partitions codées différemment et représentant une même partition aboutissent à une même valeur d'erreur de classification lors de la procédure d'évaluation

des répartitions de classe par individu de la population. Le nombre d'individus, considéré pour chacune des simulations, est de 100. De ce fait, la probabilité de mutation P_m recommandée par Bastian [BAS 97] est de 0.01.

Pour chaque valeur de P_c et P_m choisi, 10 essais par simulation de l'algorithme génétique de classification sont effectués. L'ensemble des partitions de la population est évaluée par rapport à la partition optimale à atteindre. L'individu optimal correspond à la valeur minimale des erreurs de classification.

II.1. Procédures sur l'opérateur de mutation

II.1.1. sans heuristique

cas 1. $P_c=0.6$ et $P_m = 0.0091$

La valeur de l'erreur moyenne de classification est égale à 45.4% et son écart type correspondant est de 2.07% pour 10 essais réalisés jusqu'à 500 générations.

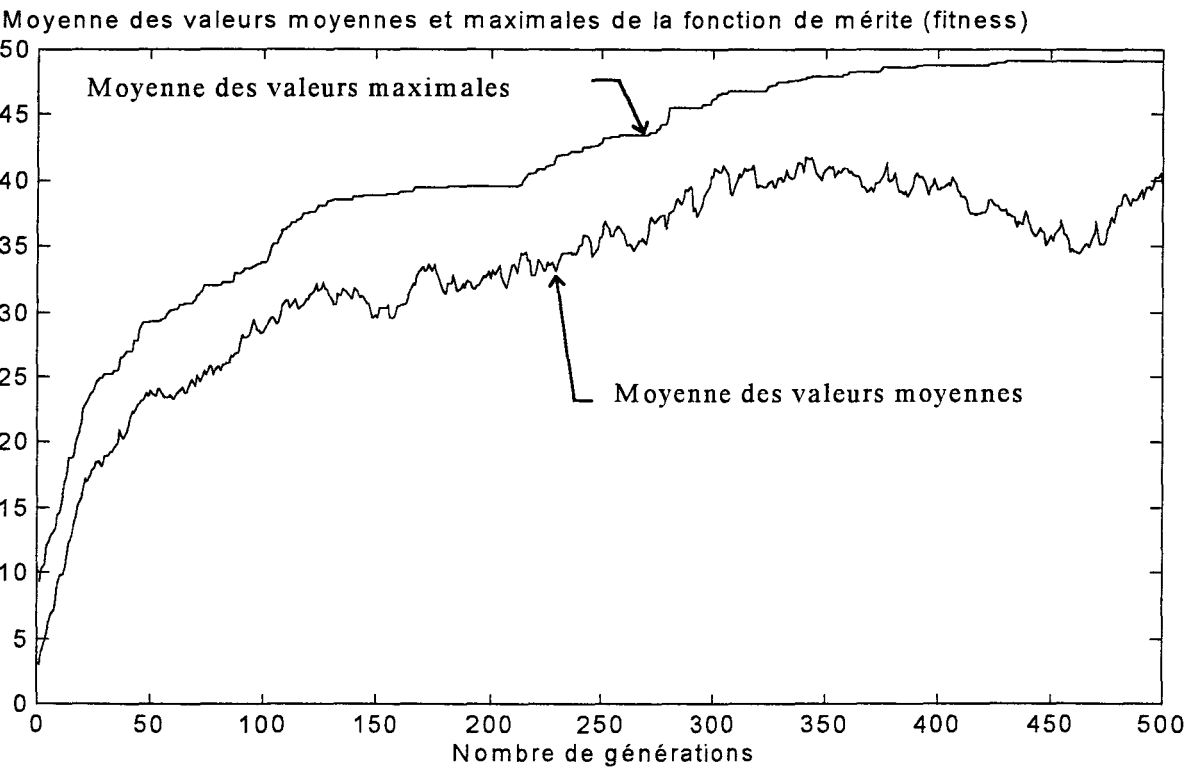


Figure 31. Evolution de la fonction de d'évaluation (fitness) pour $P_c=0.6$, $P_m=0.0091$ et sans heuristique sur l'opérateur de mutation

L'évolution croissante de la moyenne des valeurs maximales de la fonction de mérite (« fitness ») sur plusieurs centaines de générations traduit une exploration lente de l'ensemble des solutions dans l'espace de recherche. Plus la fonction de mérite augmente, plus l'ensemble des individus de la population « s'approche » de l'optimum global. L'examen de la moyenne des valeurs moyennes de la fonction de mérite au cours des générations permet de visualiser l'évolution de la population. Une tendance croissante de cette moyenne traduit un rapprochement de l'ensemble des individus vers une répartition optimale ; au contraire d'une tendance décroissante qui traduit un éloignement.

cas 2. $P_c=0.6$ et $P_m = 0.011$

L'erreur moyenne de classification est de 43.2% et son écart type de 6.45%.

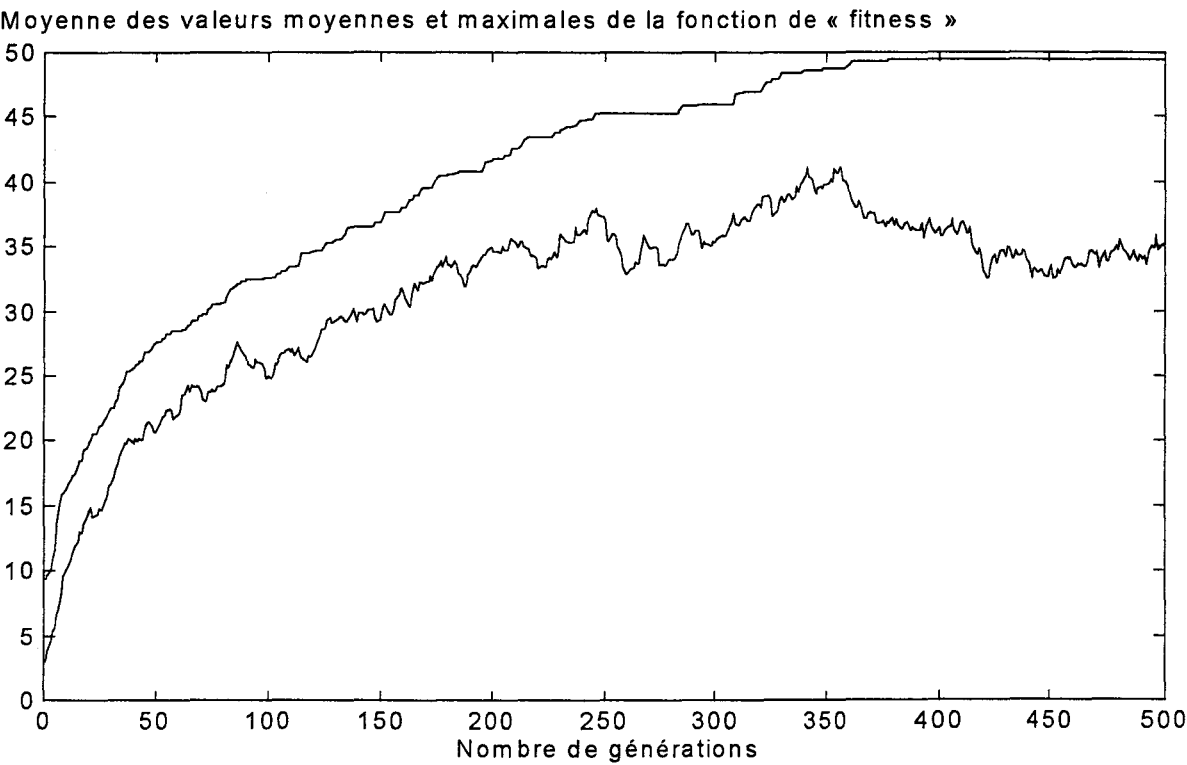


Figure 32. Evolution de la fonction de mérite
(« fitness ») pour $P_c=0.6$, $P_m=0.011$ et sans heuristique
sur l'opérateur de mutation

cas 3. $P_c=0.95$ et $P_m = 0.01$

L'erreur moyenne de classification est de 48.6% et son écart type de 1.5%.

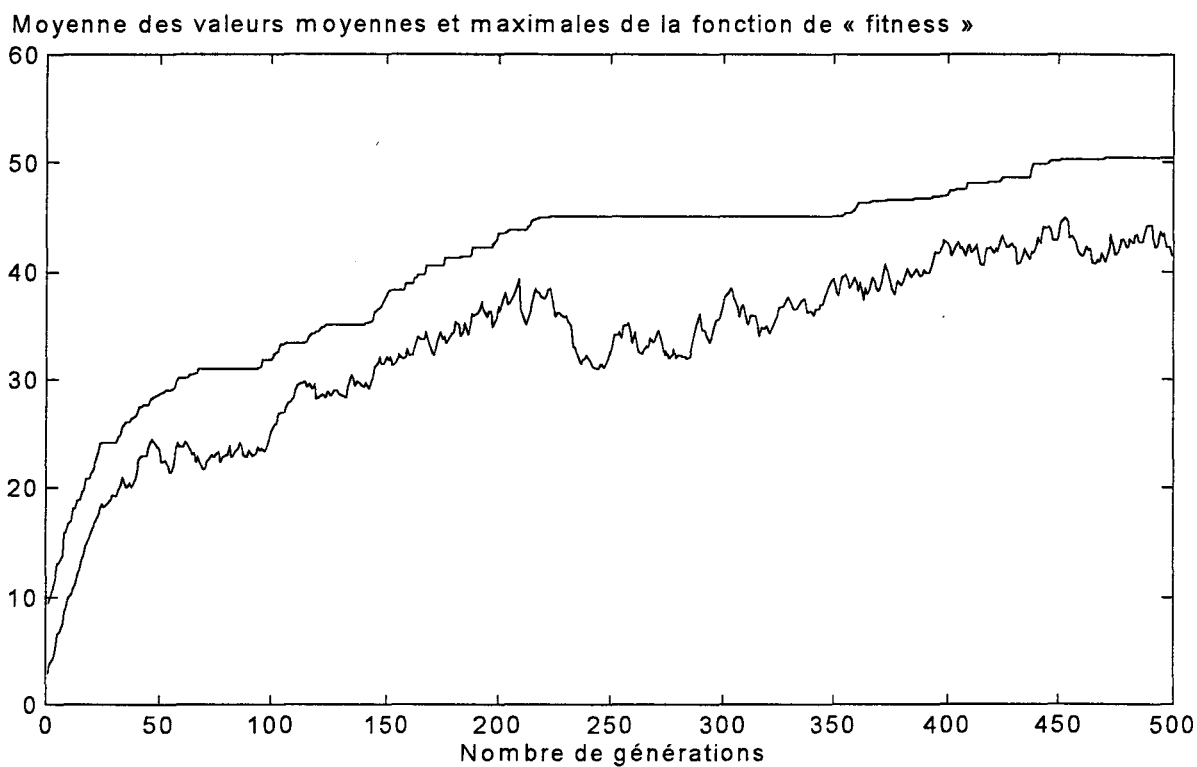


Figure 33. Evolution de la fonction de mérite
(« fitness ») pour $P_c=0.95$, $P_m=0.01$ et sans heuristique
sur l'opérateur de mutation

II.1.2. avec heuristiques

(1) heuristique 1

cas 1. $P_c=0.6$ et $P_m = 0.0091$.

L'erreur moyenne de classification correspond à 54.66% et l'écart type correspondant est de = 2.06%.

Moyenne des valeurs maximales et moyennes de la fonction de « fitness »

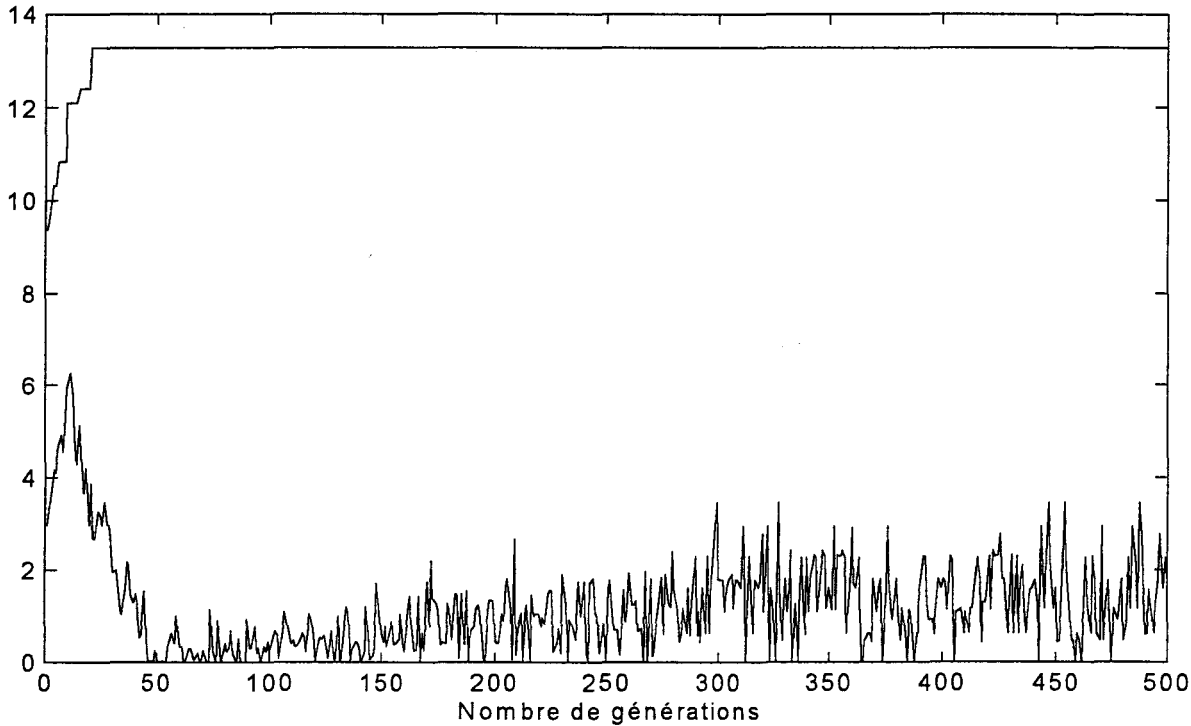


Figure 34. Evolution de la fonction de mérite
(« fitness ») pour $P_c=0.6$, $P_m=0.0091$ avec l'heuristique 1
sur l'opérateur de mutation

Au bout d'une dizaine de générations, les individus de la population ne subissent plus aucune évolution de la part des opérateurs génétiques. L'heuristique 1 sur l'opérateur de mutation semble inadapté à la procédure de recherche de la solution dans l'espace tout entier. Ce cas de figure se répète pour les *cas 2* et *3*.

(2) heuristique 2

cas 1. $P_c=0.6$ et $P_m = 0.0091$

L'erreur moyenne de classification est de 19.6% et son écart type de 13.57%.

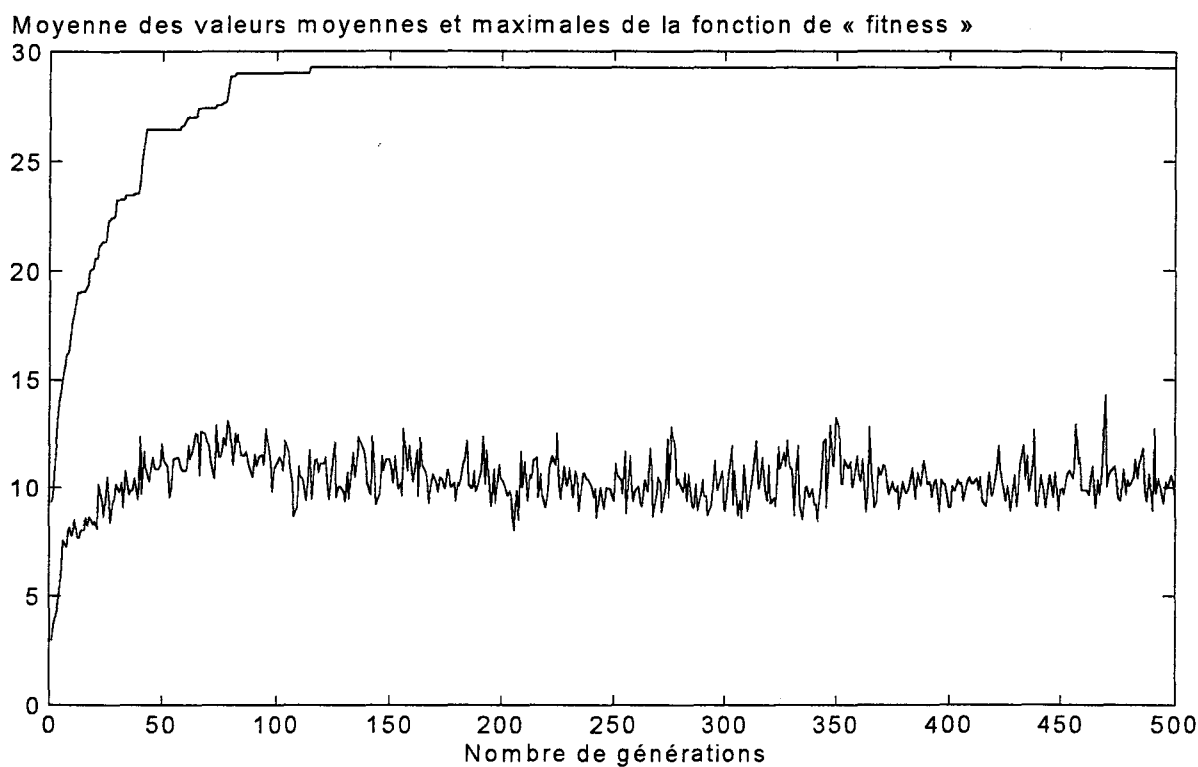


Figure 35. Evolution de la fonction de mérite
(« fitness ») pour $P_c=0.6$, $P_m=0.0091$ avec l'heuristique 2
sur l'opérateur de mutation

cas 2. $P_c=0.6$ et $P_m = 0.011$

l'erreur moyenne de classification est de 21.8% et l'écart type de 22.26%.

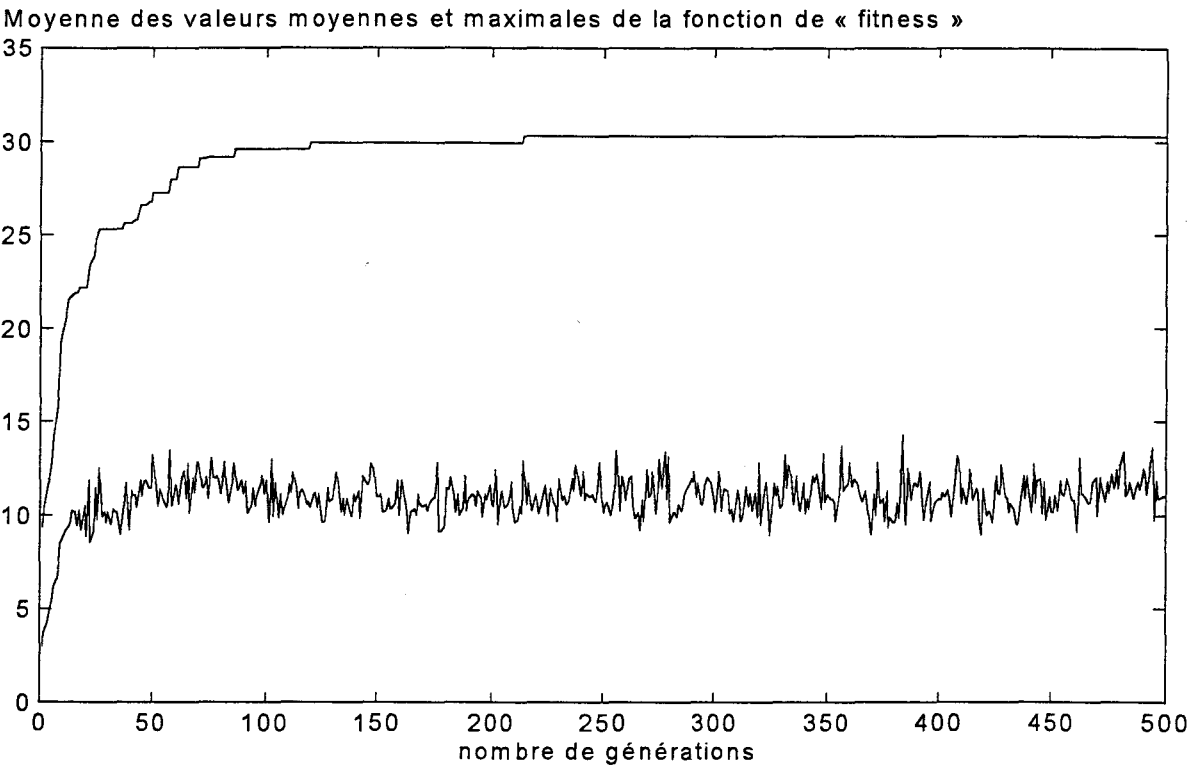


Figure 36. Evolution de la fonction de d'évaluation
(« fitness ») pour $P_c=0.6$, $P_m=0.011$ avec l'heuristique 2
sur l'opérateur de mutation

cas 3. $P_c=0.95$ et $P_m = 0.01$

l'erreur moyenne de classification est de 10.6% et l'écart type est de 11.43%.

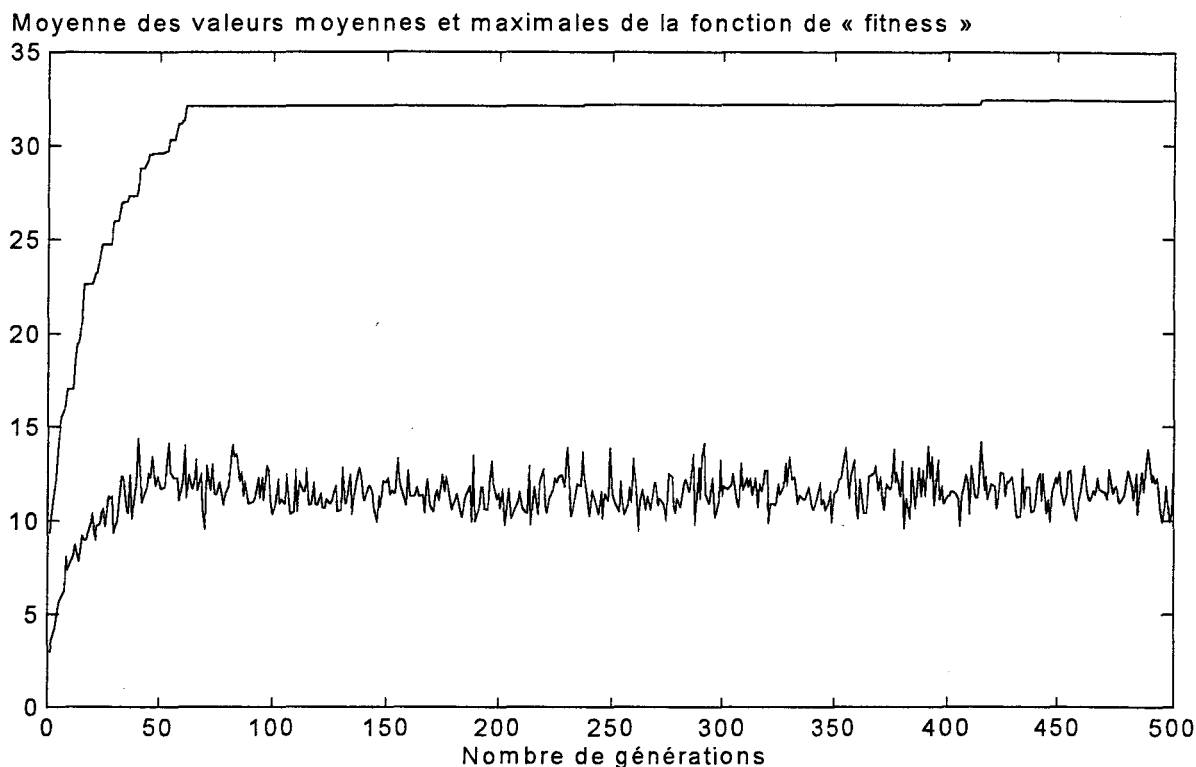


Figure 37. Evolution de la fonction de d'évaluation
(« fitness ») pour $P_c=0.95$, $P_m=0.01$ avec l'heuristique 2
sur l'opérateur de mutation

Pour l'ensemble des valeurs des probabilités utilisées, l'heuristique 2 sur l'opérateur de mutation permet de trouver très rapidement (au bout de 100 générations) une valeur minimale locale mais ne parvient pas à explorer les autres solutions de l'espace de recherche.

(3) heuristique 3

cas 1. $P_c=0.6$ et $P_m = 0.0091$

L'erreur moyenne de classification est de 1.4% et son écart type de 0.89%.

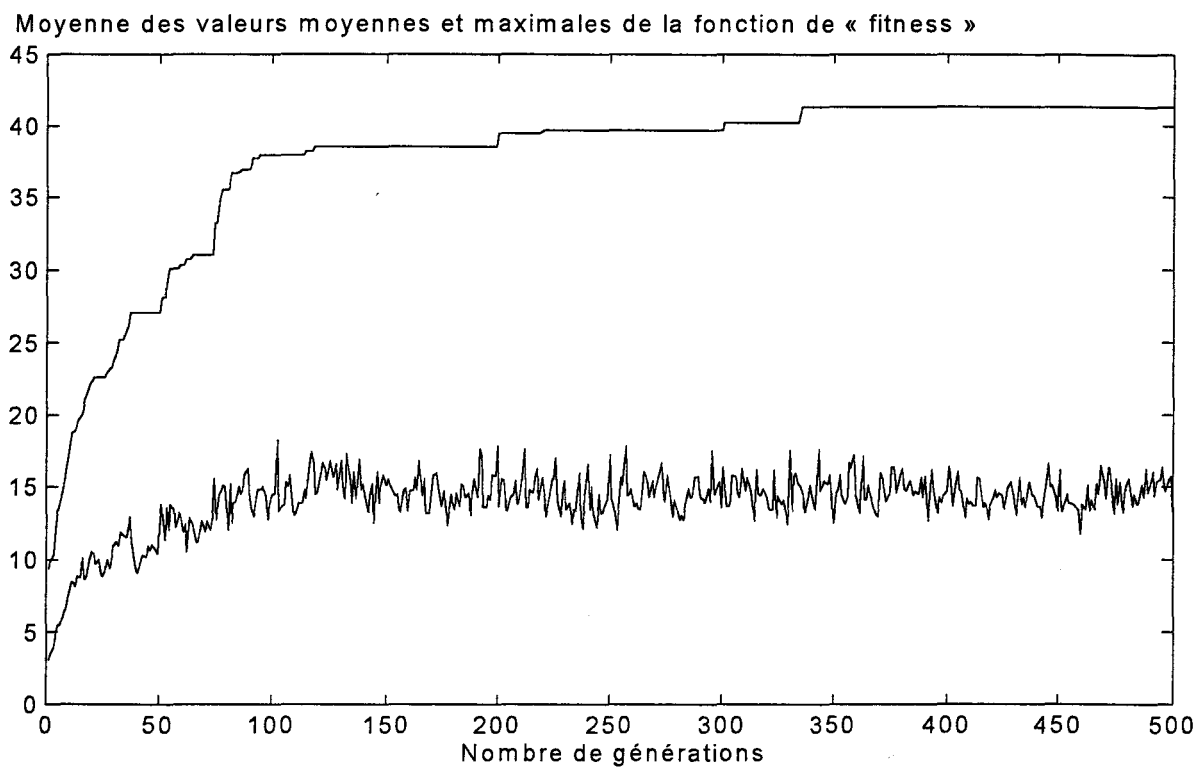


Figure 38. Evolution de la fonction de d'évaluation
(« fitness ») pour $P_c=0.6$, $P_m=0.0091$ avec l'heuristique 3
sur l'opérateur de mutation

cas 2. $P_c=0.6$ et $P_m = 0.011$

L'erreur moyenne de classification est égale à 2.6% et son écart type à 1.14%.

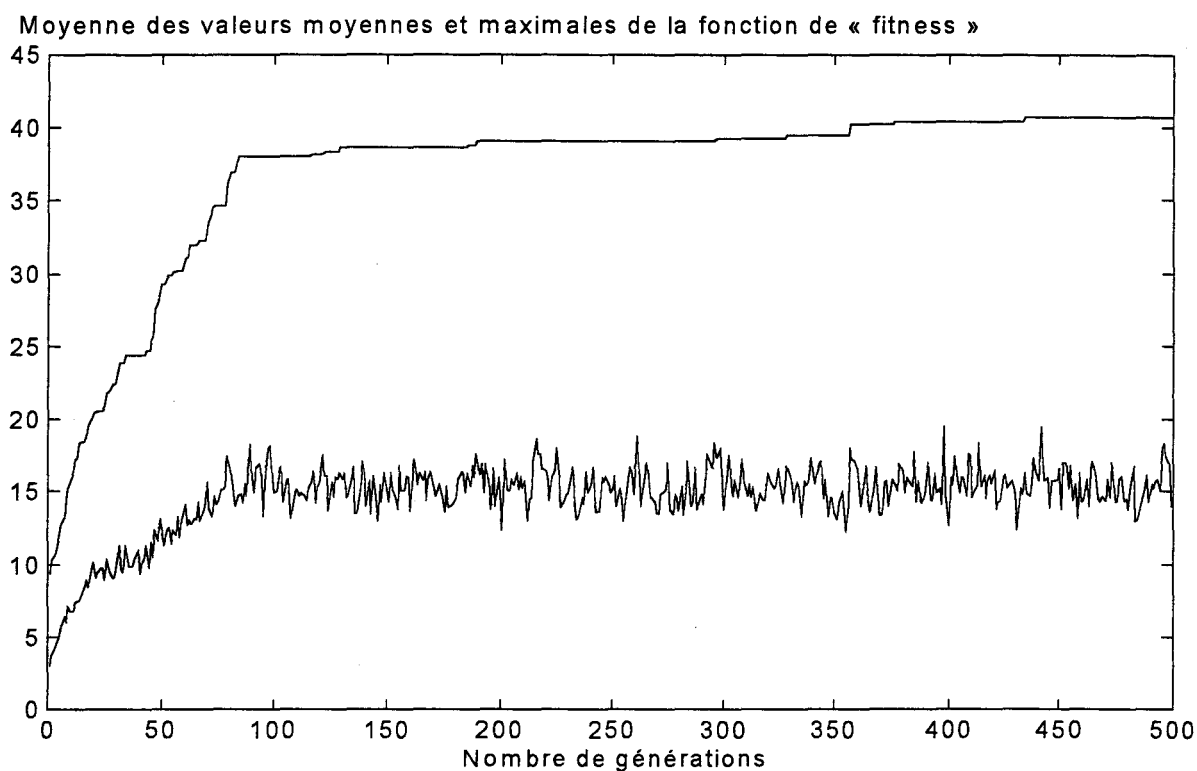


Figure 39. Evolution de la fonction de d'évaluation
(« fitness ») pour $P_c=0.6$, $P_m=0.011$ avec l'heuristique 3
sur l'opérateur de mutation

cas 3. $P_c=0.95$ et $P_m = 0.01$

L'erreur moyenne de classification est de 1.8% et l'écart type correspondant de 1.44%.

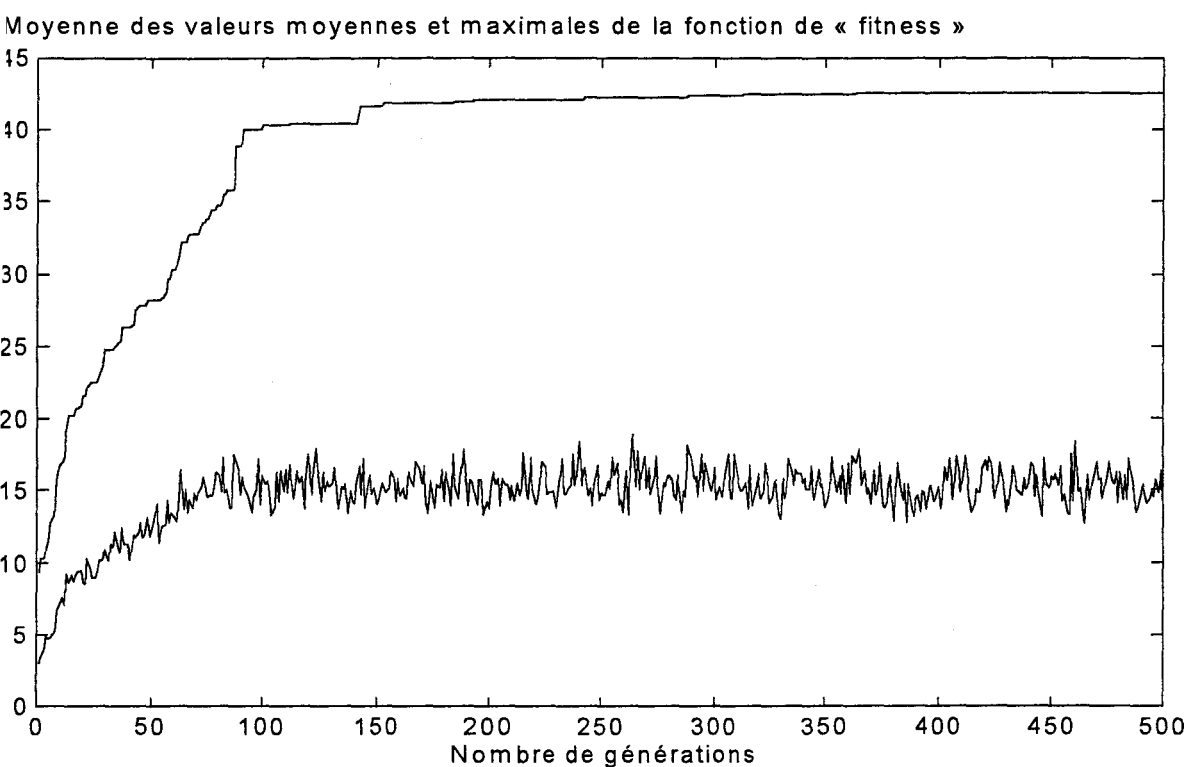


Figure 40. Evolution de la fonction de d'évaluation
(« fitness ») pour $P_c=0.95$, $P_m=0.01$ avec l'heuristique 3
sur l'opérateur de mutation

L'heuristique3 sur l'opérateur de mutation, permet d'approcher la répartition optimale globale au bout de 100 à 150 générations, puis affine la recherche de l'optimum pour les générations restantes.

II.1.3. Récapitulatif des simulations effectuées sur l'algorithme de classification

L'heuristique 3 aboutit à la meilleure valeur de l'erreur de classification avec une faible dispersion autour de cette moyenne pour le *cas 1* (cf. Tableau 22).

Opérateurs de mutation	Erreur moyenne de classification	Ecart type
sans heuristique	45.4%	2.07%
avec l'heuristique 1	54.6%	2.06%
avec l'heuristique 2	19.6%	13.57%
avec l'heuristique 3	1.4%	0.89%

Tableau 22. Récapitulatif des résultats par rapport au
cas 1 : $P_c=0.6$ et $P_m=0.0091$.

Les résultats sont à peu près identiques aux précédents pour le *cas 2* (cf. Tableau 23).

Opérateurs de mutation	Erreur moyenne de classification	Ecart type
sans heuristique	43.2%	6.45%
avec l'heuristique 1	57.8%	2.69%
avec l'heuristique 2	21.8%	22.26%
avec l'heuristique 3	2.6%	1.14%

Tableau 23. Récapitulatif des résultats par rapport au
cas 2 : $P_c=0.6$ et $P_m=0.011$.

Les résultats sont à peu près identiques aux précédents pour le *cas 3* (cf. Tableau 24).

Opérateurs de mutation	Erreur moyenne de classification	Ecart type
sans heuristique	48.6%	1.5%
avec l'heuristique 1	56.2%	3.22%
avec l'heuristique 2	10.6%	11.43%
avec l'heuristique 3	1.8%	0.44%

Tableau 24. Récapitulatif des résultats par rapport au
cas 3 : $P_c=0.95$ et $P_m=0.01$.

La différence obtenue entre les heuristiques 2 et 3 est principalement due à la formulation du coefficient multiplicateur c_{ij} sur les probabilités de mutation. Le coefficient c_{ij} de l'heuristique 2 est relatif au calcul de la distance au sein d'une même classe alors que le coefficient c_{ij} de l'heuristique 3 englobe l'ensemble des distances entre les classes.

Afin de montrer l'intérêt de notre méthodologie de classification par une procédure génétique, une comparaison a été effectuée avec l'algorithme des « c »-classes

moyennes (méthode des "k-means") sur les données représentées en Figure 30. L'évaluation de 1000 tirages aléatoires de la partition initiale en 4 classes par cet algorithme aboutit dans 68% des cas à la partition optimale. Une démarche de recherche aléatoire pure de la partition optimale peut nécessiter un nombre important de tirages aléatoires de la partition initiale d'un point de vue statistique. Le recours à des heuristiques permettant de "guider" le choix des partitions initiales issues d'une démarche aléatoire peut permettre de réduire le nombre de tirages nécessaires.

II.2. Application aux données réelles de vente textile

Une application à partir des 59 données de vente réelle sur 52 périodes issues de la même entreprise textile a été réalisée. Dans une première application, les valeurs initiales des probabilités de mutation et de recombinaison sont celles du *cas 1* : $P_c=0.6$ et $P_m = 0.0091$. Les chromosomes sont codés par affectation et l'heuristique de mutation choisie correspond à l'heuristique 3.

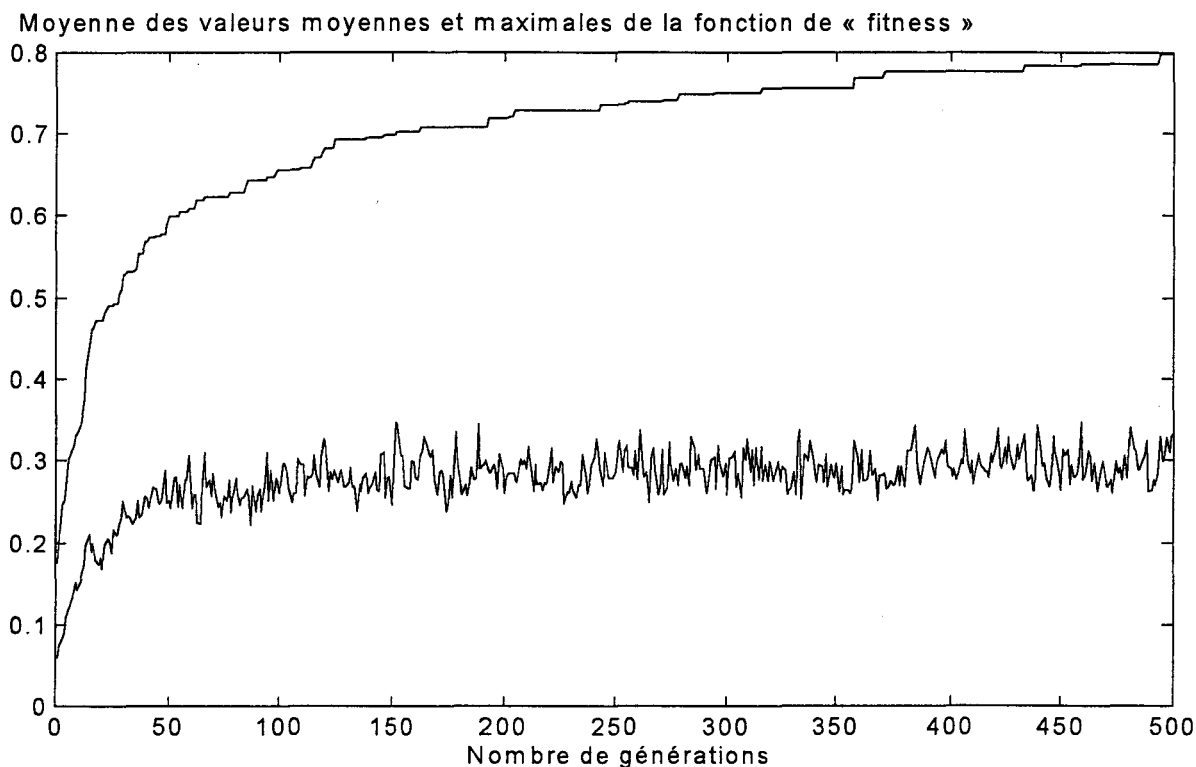


Figure 41. Evolution de la fonction de mérite (« fitness ») pour $P_c=0.6$, $P_m=0.0091$ avec l'heuristique 3 sur l'opérateur de mutation pour les données réelles.

Les résultats obtenus par cette méthode de classification par l'algorithme génétique sont identiques (à 3 points près) à ceux obtenus par la méthode de classification par partition non floue des voisins réciproques (cf. Figure 17 page 127).

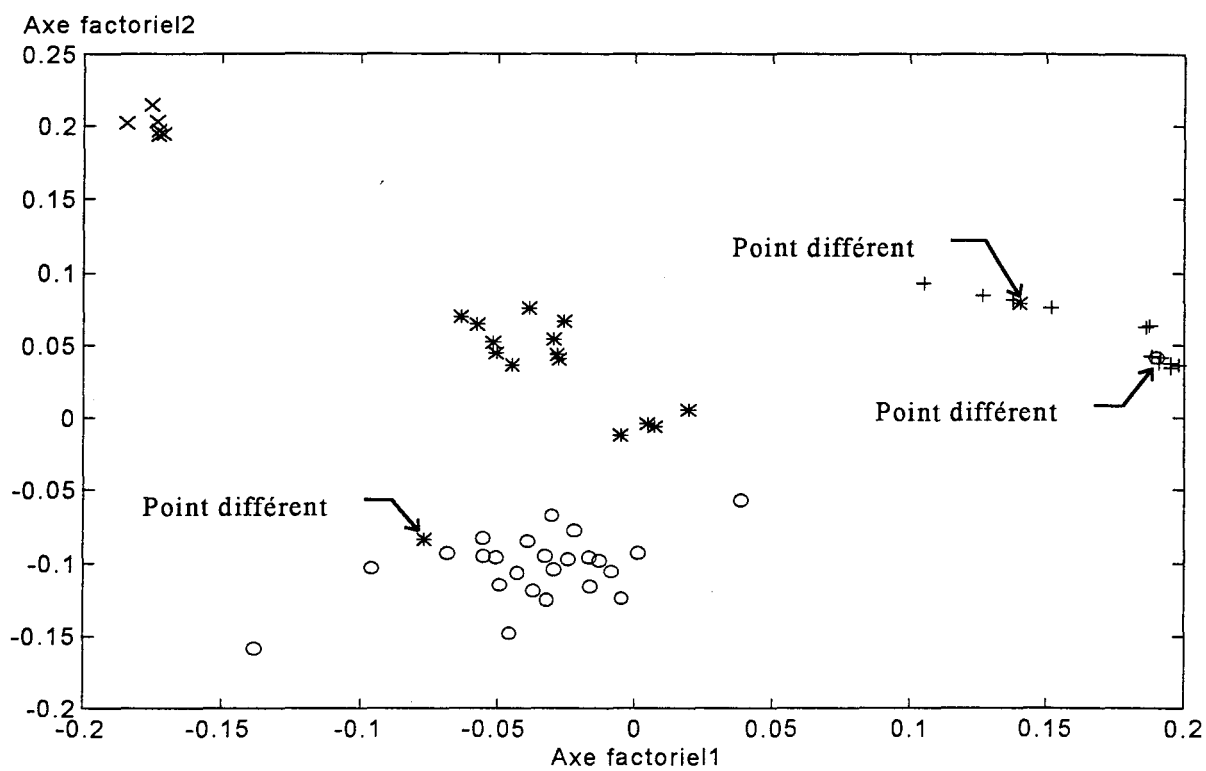


Figure 42. Représentation de la partition "optimale" pour 4 classes suivant les 2 premiers axes factoriels.

III. Conclusion

L'intérêt de l'utilisation d'un algorithme génétique de classification réside dans la capacité à explorer l'ensemble des solutions pour atteindre la répartition optimale globale. Cependant, cette approche dépend de nombreux paramètres tels que :

- le type de codage de la partition en " c " classes,
- les formulations des opérateurs de recombinaison et de mutation,
- les valeurs initiales des probabilités de mutation et de recombinaison.

Ainsi, le choix des paramètres et des fonctions adaptés permet d'atteindre des partitions optimales pour un nombre moyen de 300 générations. Le caractère aléatoire relatif à la sélection des individus et, aux procédures de mutation et de recombinaison, empêche la détermination d'un nombre fixe de générations pour une erreur de classification souhaitée.

A l'issue de cette étude, l'application de l'algorithme génétique de classification sur les données a permis d'identifier la partition optimale globale (ou de s'en approcher). L'individu optimal correspondant peut être utilisé comme partition initiale d'un algorithme de classification basé sur l'optimisation d'une fonction objective. Cette initialisation peut permettre d'atteindre de façon quasi-permanente et plus rapidement la solution quasi-optimale de la classification [PHA 93].

Conclusion générale

La mise en place d'une démarche logistique globale, par exemple de type "Quick Response", entre les différents partenaires de la filière Textile/Habillement/Distribution permet d'ajuster les ressources de production et de distribution de chacun des acteurs en fonction du besoin du client final. Le succès de cette application repose essentiellement sur l'échange d'informations nécessaires et suffisantes entre tous les partenaires. Un premier modèle d'organisation des flux d'informations et de matières entre les acteurs textiles permet d'établir les bases de cet échange. Différentes simulations de stratégies d'approvisionnement ont permis d'ajuster la planification des ressources de production et de distribution de chacun des acteurs de la filière. Une extension du formalisme des « Hypernets » vers les réseaux formels (basés sur une représentation symbolique) autorise le traitement des informations dans le cadre d'une utilisation industrielle couplée avec des bases de données réelles.

L'élaboration d'une stratégie d'approvisionnement repose généralement sur une connaissance prévisionnelle du marché. A cette fin, une première présentation des méthodes de prévision adaptées au contexte économique de l'industrie du textile a été réalisée. A partir de cet inventaire, une évaluation des différentes méthodes de prévision a permis d'identifier les potentialités de chacune d'entre elles. Quelques recommandations sont proposées afin d'aider au choix du modèle de prévision adapté à un environnement économique. L'examen d'un ensemble de méthodes de prévision utilisées dans l'industrie textile suivant différents objectifs et contextes économiques apporte un éclairage supplémentaire sur les contraintes et les spécificités d'application de ces procédures. Dans le cadre d'une démarche « Quick Response », les méthodes de prévision basées sur des procédures de lissage avec « auto-adaptation » des coefficients semblent particulièrement adaptées et fournissent des résultats plus précis que d'autres méthodes, comme l'a montré

l'analyse comparative des méthodes de prévision. L'évaluation des méthodes de prévision par plusieurs mesures de l'erreur permet d'augmenter la qualité de l'analyse relative au choix du modèle adapté.

Le traitement des informations issues des acteurs de la filière, pour l'identification des modèles de prévision, sollicite les ressources de gestion des données de chaque entreprise. Une première approche d'analyse de données, sous l'aspect de la classification, a permis de réduire la taille du problème de traitement de l'information tout en minimisant la perte. Une évaluation des partitions par un critère de validité mesurant les notions de compacité et séparabilité entre classes, chacune représentative d'une information spécifique, aboutit à cette réduction de la complexité. La mesure de la précision obtenue, au sein de chacune des classes par une contrainte de compacité à caractère globale ou locale, permet à l'industriel de valider cette réduction d'informations en fonction de ses données de vente. Une expression de l'information réduite sous forme qualitative et quantitative, exprimée sous forme symbolique, est ainsi obtenue. Une méthode d'identification de l'influence de phénomènes exogènes sur le comportement de vente des articles permet également d'enrichir la connaissance du phénomène de vente d'un contexte économique spécifique.

Enfin, le recours à un algorithme évolutif basé sur des opérateurs génétiques, a permis de s'assurer du caractère global de la partition optimale, tout en minimisant le nombre de générations nécessaires. L'utilisation de l'heuristique 3 sur l'opérateur de mutation contribue fortement à la minimisation du temps de recherche de la partition optimale globale.

Une extension de nos travaux peut être envisagée par :

- l'ébauche d'indicateurs de performance standards permettant d'évaluer les performances de chacun des partenaires quel que soit son type d'activité,
- l'introduction des paramètres exogènes dans les procédures d'autorégulation des paramètres de lissage des méthodes de prévision correspondantes,
- la création de mesures de ressemblance basées sur la définition des différentes formes des objets symboliques. Une approche de regroupement
- l'amélioration des modèles et des méthodes de prévision adaptés aux ventes des articles textiles,

- l'optimisation de la recherche du couple des valeurs de probabilité de mutation et de recombinaison en fonction du choix de codage de la partition.

Dans une perspective à moyen terme, la méthodologie de réduction de la complexité en vue d'améliorer la prévision des ventes sera validée sur des ensembles de données de ventes issus de différents partenaires de la filière THD. Cet ensemble hétérogène de données permettra d'adapter notre méthodologie aux contraintes spécifiques de chaque partenaire. Cette validation ne peut être effectuée que si la connaissance sur l'évolution dynamique des modèles de prévision et des partitions, relatives à un type d'allure de vente, est acquise. L'identification d'un modèle de prévision adapté à l'environnement de vente incertain des ventes d'articles textiles est l'objet d'une première thèse de doctorat actuellement en cours de réalisation. L'analyse de l'évolution dynamique des partitions optimales d'une saison de vente à l'autre constitue l'un des axes majeurs du thème de recherche de l'équipe Organisation et Gestion de la filière THD du laboratoire GEMTEX⁵.

Dans le cadre du projet AIDE⁶, nous espérons développer l'ensemble de ces perspectives, au sein de l'équipe afin de répondre rapidement à la demande croissante et complexe des différents segments de son marché, en proposant un logiciel performant, véritable outil de traitement de l'information nécessaire à de meilleures prévisions des ventes et simulations des stratégies prévisionnelles.

⁵ Génie des Matériaux Textiles de l'ENSAIT (Ecole Nationale Supérieure des Arts et Industries Textiles).

⁶ Relatif à la conception d'un logiciel d'Analyse et d'Intégration des Données EDI pour une meilleure prévision technique et commerciale soutenu par le ministère de l'industrie et en partenariat avec des entreprises textiles.

Annexe 1

Evaluation de la prévision, d’horizon 1, par la moyenne mobile double d’ordre 3 :

article	RMSE	MAPE	MdAPE	GMRAE	MdRAE
1	92,07485	0,532997	0,944444	0	0
2	38,93665	0,358412	0,40057	1,527328	0,96701
3	35,94896	0,340052	0,616959	1,388039	0,1162
4	52,43393	0,335863	0,045906	0	1,49573
5	36,06867	0,319975	0,988701	1,008732	6,03448
6	24,00988	0,425531	0,183215	1,032502	1,26667
7	10,60412	0,575818	0,455026	0,693263	0,04444
8	12,10408	0,399687	0,477954	0,776067	0,11111
9	5,66549	0,642646	0,577778	1,148404	1,33333
10	84,0403	0,474762	0,104211	1,189764	0,52482
11	75,80227	0,516154	0,208138	1,203969	0,1133
12	61,37842	0,665353	0,131369	1,39765	13,679
13	76,01623	0,441683	0,234379	1,556985	1,94082
14	112,9144	0,32524	0,110036	1,171286	1,0652
15	81,65478	0,278494	0,333333	1,218876	0,87093
16	29,0095	0,82282	0,336806	1,350409	7,71605
17	28,03292	0,529063	0,503843	1,059294	0,51471
18	74,39666	0,517946	0,330854	1,270347	0,89622
19	41,85624	0,558851	0,057692	1,091197	2,28283
20	53,37325	2,547169	0,606061	1,167936	1,02564
21	38,61957	0,312207	0,229905	0,970955	0,51491
22	19,99886	0,364081	0,192593	1,129558	0,38272
23	18,29388	0,570845	0,873457	1,11122	0,47737
24	31,37811	0,328519	0,196825	0,984128	0,1358
25	32,12467	0,287932	0,199495	1,078549	2,47222
26	1,997715	0,267802	0,061728	0	0,66667
27	2,242188	0,6814	0,492063	0,608129	3,88889
28	3,118879	0,291731	0,007407	0,369446	3,11111
29	1,968079	0,318945	0,015873	0,300109	0,22222
30	1,700835	0,623071	0,055556	0,942161	3,66667
31	2,066846	1,469907	1,444444	0,886335	0,52778
32	17,29312	0,515358	0,420054	0,391694	0,73611
33	30,31913	0,457535	0,449864	1,014566	3,12222
34	25,33475	0,514502	0,479532	1,115037	3,12963
35	30,62983	0,317726	0,1	1,088345	1,80556
36	18,32881	0,578248	0,375228	0,797162	0,3125
37	15,53699	0,654023	0,725926	1,241742	0,48889
38	17,88036	0,65628	0,082789	1,167142	1,47778
39	40,66401	0,446855	0,040936	0,961773	0,4
40	15,06435	0,300821	0,408846	0	2,95238
41	13,00299	0,414698	0,502058	0,936387	2,55556
42	19,47539	0,340883	0,504425	0,936341	0,37317
43	62,74886	0,283846	0,280822	0,932551	1,0721
44	5,430413	0,21592	0,253968	0,385866	1,33333
45	7,409137	0,365781	0,102564	1,037498	1,51852
46	8,847654	0,33613	0,235294	1,059509	0,71605
47	19,05725	0,312626	0,106061	0,687043	10,6111
48	45,97398	0,483696	0,05078	1,083877	1,72917
49	1,755134	0,26421	0,206349	0,565718	0,11111
50	15,23767	0,297511	0,083333	1,133997	0,58333
51	34,72701	0,559916	0,281959	1,740962	1,29218
52	35,71126	0,274309	0,195122	1,101581	25,7222
53	7,996913	0,448012	0,239316	0,940594	0,47222
54	30,40343	0,222021	0,1	0,955408	1,64444
55	50,80614	0,41676	0,078777	0	2,46784
56	14,73521	0,319429	0,132937	1,343269	1,13131
57	43,10287	0,270117	0,269972	0,810888	0,33075
58	10,36988	0,231753	0,206349	0,791202	1,42222
59	38,28782	0,310624	0,094601	1,360973	0,74262

RMSE : Root Mean Square Error (formule n°17)

MAPE : Mean Absolute Percentage Error (formule n°18)

MdAPE : Median Absolute Percentage Error (formule n°19)

GMRAE : Geometric Mean of the Relatice Absolute Error (formule n°20)

MdRAE : Median Relatice Absolute Error (formule n°21)

Evaluation de la prévision, d'horizon 1, par la méthode du lissage exponentiel double avec $\alpha=0.3$ (Lisexpo1) .

article	RMSE	MAPE	MdAPE	GMRAE	MdRAE
1	81,13855	0,561315	0,468701	0,72125	0,54461
2	39,13578	0,347061	0,543411	1,403953	1,85466
3	37,03871	0,304281	0,610215	1,112686	0,5015
4	50,3394	0,349305	0,100069	1,371991	1,96793
5	37,14203	0,373756	0,832073	1,462061	5,07852
6	23,15159	0,40824	0,54287	1,411138	1,35043
7	9,879084	0,617463	0,534061	1,210554	0,17311
8	12,63709	0,424131	0,539202	1,048486	1,20867
9	5,494959	0,524728	1,060904	1,02493	0,57851
10	87,72145	0,568947	0,286638	2,107877	1,26413
11	80,38141	0,64067	0,464248	1,581422	1,08645
12	64,21362	0,551782	0,294633	1,217204	2,07934
13	82,41546	0,510007	0,378947	1,621953	2,91947
14	117,2735	0,322886	0,23582	1,108257	1,67463
15	92,35496	0,397767	0,428073	1,887609	1,22234
16	33,73117	1,105042	0,461184	2,116617	0,1479
17	32,63412	0,759499	0,572812	1,636205	1,06873
18	85,3447	0,884596	0,418146	1,862512	1,13267
19	50,52409	0,450363	0,054234	1,324189	2,0016
20	45,40961	2,197459	0,818129	1,175758	1,38453
21	33,00381	0,258717	0,079029	0,7002	0,0848
22	19,27334	0,36533	0,203755	1,042889	1,43651
23	17,62219	0,506014	0,129008	0,535805	0,95322
24	28,94885	0,306559	0,09351	1,159203	1,93203
25	29,78159	0,283703	0,203004	1,025721	3,29138
26	2,055758	0,344058	0,375643	1,378255	3,13786
27	1,67336	0,535891	0,328863	0,687413	2,89897
28	3,023032	0,319888	0,038049	0,707523	1,28614
29	1,929959	0,26349	0,088107	0,361156	0,16207
30	1,452653	0,528555	0,48066	1,015867	3,10064
31	1,552896	1,036757	0,36532	0,850355	0,71763
32	15,45133	0,436978	0,456937	1,06564	0,66516
33	27,7568	0,451245	0,638409	1,07591	1,23412
34	23,3788	0,448356	0,595931	0,952762	1,71129
35	31,14108	0,330806	0,232248	1,326816	2,52707
36	16,85018	0,524346	0,433715	0,796148	0,53916
37	15,40963	0,538717	0,848887	1,020788	0,9872
38	16,83638	0,573087	0,31651	0,843174	1,09842
39	43,3089	0,458329	0,255268	1,353816	2,41894
40	12,50342	0,233813	0,424845	0,776009	2,05135
41	11,77122	0,342594	0,517143	0,698923	1,14106
42	15,27648	0,27339	0,417175	1,057958	0,34269
43	55,21556	0,260287	0,267165	1,399475	0,79535
44	5,67268	0,291841	0,014916	0,72594	0,07831
45	6,873285	0,38286	0,268354	1,212267	0,92656
46	7,633699	0,333999	0,035039	1,079973	1,0233
47	15,78127	0,337375	0,105872	1,330205	5,55334
48	60,01527	0,588381	0,162258	1,776173	1,07911
49	2,112224	0,381674	0,294753	1,040072	1,01839
50	13,23852	0,292878	0,222184	1,026593	1,88959
51	41,89131	0,713518	0,108051	2,256313	2,89221
52	33,83386	0,2606	0,247258	1,256796	36,3108
53	6,083351	0,465789	0,474836	1,121758	1,79105
54	30,0335	0,222737	0,113127	1,232993	1,83398
55	51,89584	0,398466	0,012254	1,174781	2,29859
56	13,69307	0,296107	0,156429	1,052276	0,18406
57	52,484	0,400577	0,113343	1,037112	0,64632
58	10,38683	0,278721	0,000697	0,823717	0,59437
59	37,66629	0,31529	0,170258	1,701908	1,21886

Evaluation de la prévision, d'horizon 1, par la méthode du lissage exponentiel double avec procédure d'autorégulation de α (Lisexpo2).

article	RMSE	MAPE	MdAPE	GMRAE	MdRAE
1	74,66446	0,52315	0,671774	0,757853	0,80837
2	39,49316	0,347666	0,179308	1,384133	1,38486
3	29,31977	0,303024	0,460114	1,362167	0,77571
4	46,85255	0,263482	0,0384	1,168913	1,55619
5	36,98244	0,302456	0,849963	1,063335	5,18771
6	20,75106	0,371016	0,233499	1,156006	0,83168
7	9,634465	0,942103	0,503488	1,263715	0,58948
8	9,634003	0,388443	0,144223	0,943233	0,89683
9	5,689049	0,549616	0,740354	1,254099	2,0543
10	73,60044	0,465063	0,127744	1,687785	0,94516
11	58,5456	0,431492	0,209428	1,308375	0,52045
12	68,849	0,606559	0,034156	0	0,86949
13	59,65219	0,343201	0,201053	0,894762	0,87304
14	79,85404	0,257889	0,17311	0,795609	0,03117
15	104,9856	0,33374	0,056407	1,821187	1,56106
16	19,99462	0,463784	0,059319	1,054862	0,08426
17	25,72448	0,35377	0,423928	0,912508	0,93005
18	73,73869	0,545645	0,027819	1,349109	0,07536
19	31,16745	0,305296	0,209705	0	2,00369
20	38,99349	1,571827	0,204573	0,882833	0,3462
21	36,21367	0,272636	0,231313	0,938967	0,79029
22	17,56464	0,312023	0,013655	0,934499	0,88199
23	18,51199	0,461469	0,344091	0,59505	1,23844
24	33,82539	0,400424	0,170727	1,470709	1,7084
25	29,51875	0,252794	0,018075	0,797125	0,4432
26	1,712359	0,225838	0,019491	0,285221	0,02667
27	1,864702	0,511385	0,40347	0,727116	1,26333
28	2,9742	0,329662	0,087519	0,531053	2,56948
29	2,055558	0,296717	0,030245	0,43738	1,25771
30	1,615159	0,453371	0,400782	0,922112	2,20701
31	1,869927	1,351016	0,053564	1,078648	0,83638
32	15,33146	0,444818	0,288994	1,072122	1,13165
33	22,06553	0,338278	0,041955	0,615722	0,21203
34	16,74238	0,298894	0,183362	0,75558	0,10872
35	27,23038	0,260003	0,176334	1,204953	4,18203
36	14,49793	0,432005	0,065549	0,750247	1,06563
37	12,47529	0,464095	0,643572	0,83342	0,56421
38	14,87651	0,529654	0,28102	1,058054	0,25118
39	35,73171	0,316153	0,084774	0,879804	0,80161
40	11,83525	0,218037	0,416531	0,704138	0,49563
41	11,24269	0,313083	0,515198	0,707737	0,42241
42	16,53018	0,280782	0,48626	0,831072	0,79432
43	56,38041	0,258585	0,235784	1,270808	0,60653
44	5,685651	0,317612	0,000176	0,506107	0,00092
45	6,963511	0,350395	0,194856	1,08379	1,28391
46	6,943772	0,308023	0,329302	0,77625	0,05944
47	15,58575	0,362889	0,343152	1,139277	3,41221
48	43,86739	0,345741	0,069062	0,90729	2,10994
49	2,002503	0,33493	0,110063	0,906906	2,42051
50	13,70205	0,301763	0,169679	1,108121	0,3564
51	37,41118	0,573511	0,099752	2,039568	2,66911
52	29,68826	0,21673	0,209155	0,979261	27,3977
53	6,292023	0,486505	0,481867	1,25014	0,08381
54	21,9866	0,167154	0,230158	0,7023	0,59994
55	51,79031	0,412819	0,079678	1,609347	5,08645
56	12,22521	0,269242	0,018149	0,96206	2,25271
57	31,28123	0,26866	0,096381	0,725911	0,59309
58	9,227275	0,236272	0,062235	0,911024	0,71052
59	33,23477	0,228916	0,164006	1,363908	0,25429

Evaluation de la prévision, d'horizon 1, par la méthode de Holt-Winter additif avec

$\alpha=0.3$ et $\beta=0.7$

article	RMSE	MAPE	MdAPE	GMRAE	MdRAE
1	112,1139	0,710911	1,294421	0,678112	0,55253
2	41,77054	0,514712	0,450299	2,08779	1,39225
3	28,30983	0,294268	0,42684	1,428033	0,57637
4	46,30191	0,456268	0,092356	1,67107	0,89133
5	32,04339	0,407827	0,489957	1,112488	2,99043
6	24,65912	0,557766	0,649692	1,63808	0,6276
7	11,40179	1,008565	0,681199	1,937689	2,35146
8	10,76466	0,510539	0,477596	1,103374	0,7176
9	5,738333	0,737606	0,966606	1,487306	0,2186
10	88,52383	0,635203	0,263466	1,959921	0,92552
11	86,56132	0,848153	0,635697	1,871905	0,94018
12	63,56477	0,858523	0,21629	1,740725	10,5888
13	77,76455	0,461176	0,333221	1,523375	2,23801
14	109,6523	0,37884	0,099812	1,246309	0,75766
15	89,03989	0,392498	0,252804	1,645058	0,73194
16	28,20824	1,379199	0,267561	2,041008	1,15828
17	32,2587	1,06364	0,295221	1,988718	1,09143
18	68,88623	1,142426	0,207113	1,675211	0,56103
19	55,67027	0,930317	0,230211	2,061338	4,62744
20	79,55679	9,12852	2,596391	5,592433	4,39389
21	30,75516	0,303146	0,287789	1,098966	0,34877
22	14,17852	0,265558	0,063355	0,834151	0,74987
23	17,81518	0,622422	0,209032	0,907075	1,4522
24	26,53747	0,241199	0,21541	0,769701	0,59388
25	27,65313	0,224626	0,06159	1,025796	1,894
26	2,092135	0,316749	0,477247	1,008907	1,63597
27	1,812496	0,620134	0,053927	0,951632	2,04348
28	3,113614	0,818077	0,349928	0,812291	4,70636
29	2,024605	0,405477	0,31892	1,157109	4,09292
30	1,76427	0,684687	0,792798	1,182397	4,50596
31	1,83578	1,159965	1,00635	0,511304	1,4498
32	14,85329	0,408122	0,211192	0,883373	0,23579
33	30,56213	0,618671	0,170976	1,576312	1,87176
34	23,46315	0,436272	0,333242	0,971684	2,61529
35	32,9626	0,330799	0,046348	1,389812	9,57688
36	18,15019	0,720695	0,500695	1,24405	0,57558
37	16,56344	0,675347	0,725556	1,483924	1,13508
38	19,39641	0,6894	0,406122	1,684041	1,77766
39	45,56143	0,537731	0,224593	1,646443	2,28915
40	17,89468	0,532441	0,326528	1,714931	1,79835
41	14,02635	0,588588	0,432261	1,408351	0,46219
42	22,14633	0,57097	0,187231	1,2634	0,27915
43	98,67253	0,567835	0,128167	2,273093	0,71841
44	5,613937	0,263446	0,103073	0,654743	0,54113
45	5,91007	0,315637	0,434992	0,592705	0,01511
46	6,635332	0,352401	0,079108	1,007584	1,50428
47	11,4592	0,355945	0,210565	1,090695	1,75721
48	75,95164	1,978424	0,399572	1,576839	1,1447
49	3,380464	0,320331	0,221399	0,841907	1,41133
50	17,33637	0,575466	0,011262	1,746352	1,12921
51	38,07837	0,945782	0,004374	2,42864	1,17858
52	45,94345	0,374693	0,403821	1,791069	23,642
53	6,073448	0,49297	0,506827	1,458571	1,63902
54	43,0816	0,425212	0,898915	1,911139	0,23594
55	68,36627	0,651075	0,216152	1,143702	2,84017
56	12,36958	0,336041	0,183099	1,049569	0,56776
57	58,93248	0,378931	0,022633	1,353647	0,04506
58	9,699657	0,305317	0,251447	1,566496	2,47321
59	47,93647	0,40771	0,474622	1,510045	0,69894

Evaluation de la prévision, d'horizon 1, par la méthode de Holt-Winter multiplicatif avec $\alpha=0.3$ et $\beta=0.7$

article	RMSE	MAPE	MdAPE	GMRAE	MdRAE
1	232,23026	1,989851	3,429671	2,399357	1,31615
2	44,380414	0,456632	0,514464	2,021897	1,71436
3	35,498555	0,298182	0,587656	1,348722	0,83382
4	40,724175	0,362408	0,157313	1,431997	0,97755
5	36,766093	0,370804	0,606047	1,085203	3,69898
6	26,637358	0,42559	0,729791	1,213233	0,34019
7	12,309713	0,953725	0,753323	1,804614	1,14949
8	9,456626	0,41713	0,403906	0,847517	0,57293
9	7,040052	0,839154	1,226635	1,420104	0,52168
10	75,731997	0,468768	0,095512	1,870307	0,81048
11	68,673118	0,515415	0,503895	1,49236	0,77852
12	72,931165	0,676342	0,183056	1,945731	2,96034
13	81,456707	0,467802	0,370045	1,217961	2,26631
14	106,22057	0,354387	0,058156	1,278845	0,72783
15	86,098266	0,331596	0,281938	1,439073	0,65607
16	61,536454	1,162718	0,288329	2,230061	3,46073
17	54,093054	0,936193	0,038102	2,315714	0,7898
18	154,6398	0,813688	0,358599	2,613167	0,97137
19	72,941711	1,715327	1,862815	3,300114	2,58578
20	55,383289	2,598056	1,26967	1,663622	2,14867
21	35,778203	0,33126	0,09204	0,962698	0,04761
22	26,385454	0,453355	0,255985	1,240903	2,04065
23	25,109117	0,747234	0,426352	1,021294	1,80912
24	38,089574	0,421922	0,071236	1,61625	2,02461
25	39,694905	0,442298	0,256739	1,509608	4,90577
26	3,58073	0,586413	0,988894	2,115324	3,67022
27	2,19115	0,75947	0,269219	1,155592	2,42195
28	2,403671	0,393105	0,029298	0,392303	0,00398
29	2,327245	0,553359	0,101167	1,153029	2,13347
30	2,170834	0,818891	1,129285	0,95916	4,55535
31	2,224717	1,393426	1,548534	0,738297	1,86
32	15,683354	0,401036	0,187509	0,939367	0,29379
33	29,125865	0,530189	0,445437	1,192657	0,42901
34	23,657137	0,409611	0,611467	0,850308	0,00003
35	37,354439	0,340489	0,046585	1,148044	11,5441
36	16,733649	0,600234	0,483232	0,884641	0,15949
37	18,559801	0,688551	0,883732	1,281654	1,34432
38	20,825287	0,664173	0,585628	1,574715	1,74195
39	54,33453	0,6137	0,210566	1,755752	3,48156
40	13,526318	0,381402	0,188506	1,110833	1,36924
41	10,019369	0,382349	0,352994	0,907846	0,17527
42	18,793672	0,426921	0,027421	1,287495	0,19743
43	88,162744	0,430883	0,055167	1,33088	0,31869
44	6,198536	0,434824	0,552311	0,944095	2,89963
45	5,842241	0,462207	0,208236	1,028692	0,89073
46	6,956354	0,385665	0,271623	0,942558	0,67912
47	11,109151	0,306957	0,112662	0,515391	3,93233
48	53,930271	0,65268	0,388574	1,296801	0,91584
49	5,914273	0,917241	0,639511	2,882152	6,61744
50	40,440564	1,176615	1,281769	7,15531	23,7807
51	32,03158	0,585024	0,015682	1,726617	0,60807
52	41,917479	0,357063	0,269202	1,712527	24,4501
53	6,895847	0,816693	1,183316	1,605372	3,22446
54	35,538803	0,368445	0,694056	1,22817	0,44886
55	66,187148	0,626541	0,062334	1,891514	1,13228
56	16,382593	0,542897	0,775544	1,407753	0,79716
57	49,807153	0,424832	0,309111	1,218157	0,03341
58	8,613055	0,326612	0,090673	0,772083	1,2962
59	45,664106	0,353354	0,27799	1,212623	0,67221

Evaluation de la prévision, d'horizon 1, par la méthode de Holt-Winter multiplicatif avec procédure de régulation des coefficients par système flou.

article	RMSE	MAPE	MdAPE	GMRAE	MdRAE
1	173,99095	1,3258	0,884575	1,212778	0,485365
2	24,324052	0,275468	0,199968	0,896486	0,569908
3	26,020353	0,23602	0,269298	0,723772	0,828834
4	37,656393	0,300184	0,233797	1,149258	0,696185
5	23,8524	0,251432	0,183034	0,829844	1,117137
6	15,008789	0,260636	0,11366	0,512506	0,881774
7	11,466286	0,722493	0,19559	1,127168	0,488782
8	7,806821	0,306058	0,388314	0,528062	0,004543
9	5,643707	0,727905	0,88562	1,147262	0,192162
10	50,487436	0,310887	0,16665	0,890064	0,60438
11	44,539493	0,345903	0,199567	0,722741	0,152113
12	49,747609	0,463639	0,172361	1,099568	4,652235
13	51,502219	0,311661	0,095335	0,780699	0,953666
14	100,22841	0,251579	0,331578	0,665741	0,219681
15	85,104179	0,277907	0,098269	0,975706	1,218067
16	41,428532	0,784597	0,148692	1,451248	0,848747
17	28,64771	0,549554	0,359541	1,265824	1,340323
18	92,249989	0,514145	0,036742	1,473535	0,099526
19	38,371462	1,155591	0,463684	1,567815	2,912969
20	46,803706	1,28182	1,317376	0,729739	2,229405
21	28,207502	0,260095	0,100249	0,817233	0,328637
22	20,490315	0,311178	0,114551	0,487998	0,360861
23	18,294288	0,356273	0,27253	0,369381	0,519147
24	28,368325	0,237926	0,001052	0,437039	0,114879
25	28,598773	0,233897	0,144005	0,555331	0,56757
26	1,923339	0,269757	0,086074	0,811311	1,130362
27	2,07255	0,659105	0,076258	1,114705	1,732253
28	1,785619	0,349988	0,151312	0,850383	3,091899
29	2,230781	0,509781	0,654988	1,520067	2,214907
30	2,371956	0,893184	0,048531	1,139672	2,0644
31	2,319252	1,814227	1,047458	1,142439	1,442402
32	13,929984	0,325628	0,136828	0,694062	1,62458
33	25,487218	0,451014	0,011158	0,95936	1,449602
34	19,782035	0,330168	0,251217	0,803608	1,679358
35	26,441014	0,241625	0,305526	0,755113	10,63435
36	13,181692	0,423736	0,401753	0,655707	1,256148
37	12,874272	0,51255	0,760432	0,884322	0,984132
38	12,652825	0,457257	0,44305	0,903069	0,621573
39	34,826252	0,414043	0,632823	1,036759	0,57781
40	11,990052	0,299	0,17496	0,686822	0,167539
41	10,037122	0,319966	0,019747	0,633567	0,848682
42	18,236213	0,344549	0,108983	1,007008	0,426709
43	81,744395	0,359325	0,498228	1,431757	0,020351
44	4,263293	0,417049	0,319703	1,168201	1,678442
45	3,523728	0,384945	0,222314	0,811087	0,266648
46	4,230838	0,435633	0,331518	0,656143	0,604482
47	10,728841	0,456183	0,578059	1,185629	4,855955
48	42,471532	0,547717	0,319602	1,230634	0,841695
49	4,941864	0,860183	0,666929	1,964796	7,221969
50	33,007348	1,107259	0,451076	2,944696	5,936093
51	20,951923	0,39004	0,027328	1,151416	0,926352
52	34,409064	0,263309	0,07396	1,112712	9,947682
53	5,740786	0,654565	0,389229	0,761433	1,540139
54	32,772741	0,329259	0,315473	1,101206	0,250184
55	53,954425	0,448297	0,250339	1,507958	4,220071
56	16,776232	0,633956	0,78987	1,431269	1,899521
57	37,533673	0,273988	0,00653	0,85085	1,208134
58	5,125767	0,345975	0,270698	0,816327	1,107812
59	35,485248	0,252435	0,377595	0,713212	0,130524

Annexe 2

Récapitulatif des communications au cours des travaux de recherches

- 2 revues internationales en 1997.
- 4 congrès internationaux en 1997.
- 3 congrès internationaux en 1996.
- 2 congrès internationaux et 1 congrès national en 1995.

Liste des publications envoyées dans des revues internationales ou des conférences.

(avec comité de lecture)

Année 1997

- Revue JESA (Journal Européen des Systèmes Automatisés).
Boussu F., Rabenasolo B., Happiette M., Vasseur C., *Simulation de la filière Textile*, RAIRO-APII-JESA, vol 31, n°4, pp 741-768, juin 1997.
- Revue Journal of the Textile Institute.
Boussu F., Lefort A., Happiette M., Rabenasolo B., Yim P., *Modelling and Simulation of the Textile Channel by Hypernets*, Soumis et accepté dans la revue Journal of the Textile Institute, sept 97.
- Résumé étendu soumis pour le congrès CESA 98 IMACS/IEEE à Tunis du 1-4 avril 1998.
Boussu F., Rabenasolo B., Happiette M., Vasseur C., *Determination of seasonality by a window clustering method*, soumis et accepté.
- Congrès Textile en Roumanie du 23 au 24 octobre 1997.
Happiette M., Boussu F., Vroman P., Rabenasolo B., Vasseur C., *Optimization of informational treatment . A response to economical strategy of the textile channel*, 11th Romanian Conference of Textiles and Leathership, Iasi, Romania, vol 2, pp 23-28, October 23-24 1997.
- Congrès IFAC - IMS à Seoul (Corée) du 21-23 juillet 1997.
Happiette M., Boussu F., Rabenasolo B., Zeng X., *Clustering of symbolic objects : Decision aid for production planning*, 4th IFAC Workshop on Intelligent Manufacturing Systems, Seoul, Korea, pp 341-346, July 21-23, 1997.

- Congrès ITAA à Lyon du 10-12 juillet 1997.

Boussu F., Vroman P., Happiette M., Rabenasolo B., Vasseur C., *Amélioration des performances de la filière Textile par la simulation et la prévision des ventes*, Congrès ITAA (International Textile and Apparel Association) , Lyon, France, 10-12 juillet 1997.

- Congrès IFAC - CIS à Belfort du 20-22 mai 1997.

Rabenasolo B., Happiette M., Boussu F., *Sales Forecasting under uncertain environment. Fuzzy classification in Textile Distribution*, IFAC-CIS 97 Workshop, Belfort, France, vol 3, pp 478-483, May 20-22 1997.

Année 1996

- Congrès IEEE - SMC à Beijing (Chine) du 14-17 octobre 1996.

Boussu F., Happiette M., Rabenasolo B., *Sales partition for forecasting into textile distribution network*, IEEE/SMC International Conference on Systems, Man and Cybernetics, Beijing, China, vol 4, pp 2868-2873, October 14-17 ,1996.

- Congrès FUCAM à Mons (Belgique) du 9-11 septembre 1996.

Boussu F., Happiette M., Rabenasolo B., *Textile stock optimization by sales forecasting*. Workshop on production, planning and control, at Mons, Belgium, 9-11 september 1996.

- Congrès CIMAT à Grenoble du 29-31 Mai 1996.

Boussu F., Happiette M., Lefort A., Rabenasolo B., Yim P., *Modélisation et simulation de la filière textile par hypernets*, 5th Int. Conf. on CIMAT'96 , Computer Integrated Manufacturing and Automation Technology, France, Grenoble, pp 375-380, May 29-31 1996.

Happiette M., Boussu F., Rabenasolo B., Zeng X., *L'approche de la classification dans les problèmes de prévision de ventes . Application à la filière Textile/Habillement/Distribution*, 5th Int.Conf on CIMAT'96 , Computer Integrated Manufacturing and Automation Technology, France, Grenoble, pp 402-407, May 29-31 1996.

Année 1995

- Congrès en Tunisie

Boussu F., Happiette M., Rabenasolo B., Godart F., Vasseur C., *Modélisation et simulation de la filière Textile/Habillement/Distribution*, 2^{ème} congrès maghrébin de génie électrique, CMGE'95, Tunisie, Tunis, 16-17 septembre 1995.

- Congrès System Sciences à Wroclaw (Pologne) du 12 au 15 septembre 1995.

Happiette M., Boussu F., Rabenasolo B., Godart F., Vasseur C., *Conception of integrated systems logistic into textile distribution network*, 12th Int. Conf. on Systems Science. Wroclaw, Poland, September 12-15 1995.

(sans comité de lecture)

- Journée logistique à l'ENSAIT (8 juin 1995).

Happiette M., Boussu F., Rabenasolo B., Vasseur C., *Tableau de bord et indicateurs de performances logistiques, la réactivité dans la filière Textile/ Habillement /Distribution*, Journée technique et scientifique de l'ENSAIT du 8 juin 1995.

Annexe 3

Num	périodes de ventes																									
	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	19	20	21	22	23	24	25	26
1	0	0	0	0	0	0	0	0	0	0	0	0	0	350	350	350	200	200	200	200	150	150	50	50	100	100
2	0	0	0	0	0	0	0	0	0	0	0	0	0	0	11	11	41	59	71	71	92	62	62	61	62	77
3	0	0	0	0	0	0	0	0	0	0	0	0	0	20	30	21	30	38	51	48	77	48	53	49	53	64
4	0	0	0	0	0	0	0	0	0	0	0	0	0	14	35	35	41	98	111	117	137	137	135	124	124	150
5	0	0	0	0	0	0	0	0	0	0	0	0	0	10	28	28	36	50	56	70	77	70	58	58	58	69
6	0	0	0	0	0	0	0	0	0	0	0	0	0	22	22	22	32	49	49	58	38	38	34	38	48	
7	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	6	7	7	7	7	4	4	11
8	0	0	0	0	0	0	0	0	0	0	0	0	0	0	3	3	9	11	17	27	30	27	27	14	27	37
9	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	2	9	10	9	3	2	3	4	6	8
10	0	0	0	0	0	0	0	0	0	0	0	0	0	0	18	18	60	127	249	249	293	194	177	177	181	297
11	0	0	0	0	0	0	0	0	0	0	0	0	0	16	16	16	42	76	161	161	240	158	134	134	145	226
12	0	0	0	0	0	0	0	0	0	0	0	0	0	0	120	120	120	134	165	165	176	176	176	176	176	246
13	0	0	0	0	0	0	0	0	0	0	0	0	0	22	22	63	93	152	152	211	135	107	107	132	224	
14	0	0	0	0	0	0	0	0	0	0	0	0	0	55	55	111	207	292	292	292	274	274	274	305	527	
15	0	0	0	0	0	0	0	0	0	0	0	0	0	34	34	80	149	193	193	229	178	184	184	195	405	
16	0	0	0	0	0	0	7	27	29	30	62	62	62	62	53	53	53	56	62	82	82	68	45	45	55	57
17	0	0	0	0	0	0	9	9	31	31	41	45	45	38	33	35	35	43	77	77	84	56	40	40	36	41
18	0	0	0	0	0	0	5	31	83	103	122	122	118	107	105	105	105	115	243	243	230	133	167	149	167	177
19	0	0	0	60	60	38	15	15	10	10	13	13	24	26	26	29	33	40	52	52	49	49	49	60	82	88
20	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
21	0	0	0	0	0	0	0	0	0	0	99	148	143	148	165	165	147	147	188	188	158	107	115	115	137	137
22	0	0	0	0	0	0	0	0	0	0	69	119	99	119	99	93	75	75	75	71	59	47	31	31	37	37
23	0	0	0	0	0	0	0	0	0	0	64	103	71	103	71	57	30	30	36	36	19	19	14	14	15	15
24	0	0	0	0	0	0	0	0	0	0	129	138	129	138	120	106	105	96	105	95	129	65	119	100	100	75
25	0	0	0	0	0	0	0	0	0	0	125	147	125	147	119	107	99	84	88	84	88	66	109	95	109	95
26	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
27	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
28	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
29	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
30	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
31	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0
32	0	12	12	18	49	49	49	31	31	40	43	43	65	45	48	45	21	15	15	24	24	24	22	22	26	41
33	0	12	12	34	90	74	74	74	62	81	127	127	120	110	110	56	20	20	29	29	24	20	16	16	29	41
34	0	12	12	48	77	77	77	44	44	87	89	89	86	86	86	55	23	18	16	16	20	20	11	11	30	38
35	0	24	24	50	91	91	93	103	94	103	143	143	116	108	78	72	68	35	57	43	53	53	55	56	86	90
36	0	4	4	14	17	17	21	27	21	26	26	26	41	25	29	24	19	15	19	19	28	19	20	20	45	61
37	0	5	5	20	26	26	27	27	27	30	33	30	44	27	27	15	10	10	10	14	5	3	3	8	12	45
38	0	0	0	8	30	30	30	28	18	28	35	35	44	20	19	19	12	9	7	7	7	7	7	7	17	51
39	0	6	6	48	59	59	91	102	102	108	130	130	149	140	140	95	76	60	36	36	52	56	65	91	118	133
40	0	2	2	13	26	26	26	16	16	20	23	23	36	36	36	30	30	26	31	35	51	43	51	54	57	103
41	0	1	1	1	12	12	18	18	15	12	12	17	40	17	27	15	25	25	32	32	31	31	31	38	38	81
42	0	1	1	17	33	33	37	37	31	30	30	35	58	44	58	44	38	38	74	74	72	72	72	60	60	113
43	0	15	15	50	108	108	108	111	115	116	208	208	213	181	213	181	169	169	250	250	250	286	286	286	313	438
44	0	0	0	0	0	0	0	0	0	0	0	0	2	2	6	7	15	17	21	21	27	31	31	31	29	29
45	0	0	0	0	0	0	0	0	0	0	0	0	0	0	1	5	11	15	27	27	39	27	28	28	33	33
46	0	0	0	0	0	0	0	0	0	0	0	2	2	2	6	16	22	30	34	34	25	25	25	29	37	40
47	0	0	0	0	0	0	0	0	0	0	0	3	3	3	15	16	33	38	66	66	64	63	49	49	69	82
48	0	0	0	0	0	0	0	0	0	0	0	1	11	11	34	52	56	71	112	112	106	99	99	99	113	121
49	0	0	0	0	0	0	0	0	0	0	0	1	3	3	5	7	7	10	14	12	16	12	11	9	9	7
50	0	0	0	0	0	0	0	0	0	0	0	1	21	21	21	42	44	44	64	64	53	49	53	66	72	66
51	0	0	0	0	0	0	0	0	0	0	0	0	4	4	9	31	54	58	78	78	83	76	76	76	97	106
52	0	0	0	0	0	0	0	0	0	0	0	0	0	0	125	125	107	89	89	108	108	86	80	86	80	105
53	0	0	0	0	0	0	0	0	0	0	0	0	0	0	2	16	16	19	13	17	13	10	10	26	31	
54	0	0	0	0	0	0	0	0	0	0	0	0	0	0	77	77	77	73	95	78	82	78	82	94	94	76
55	0	0	0	0	0	0	0	0	0	0	0	74	74	74	101	101	101	106	166	169	216	169	134	117	117	168
56	0	0	0	0	0	0	0	0	0	0	0	0	0	0	3	33	47	51	56	56	56	45	61	61	64	73
57	0	0	0	0	0	0	0	0	0	0	0	0	0	18	18	60	70	96	103	138	138	171	117	117	121	157
58	0	0	0	0	0	0	0	0	0	0	0	0	0	0	6	6	18	27	33	42	42	42	47	55	53	55
59	0	0	0	0	0	0	0	0	0	0	0	0	0	0	77	77	77	88	137	137	137	97	97	97	127	132

	périodes de ventes																																																			
Num	27	28	29	30	31	32	33	34	35	36	37	38	39	40	41	42	43	44	45	46	47	48	49	50	51	52																										
1	250	250	300	300	250	150	150	150	100	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0																											
2	202	271	271	264	259	195	141	112	112	114	123	123	102	100	100	55	52	29	29	47	47	42	35	0	0	0																										
3	120	251	264	264	271	190	184	138	135	130	90	90	90	90	85	62	37	37	37	37	37	36	21	0	0	0																										
4	325	330	354	365	380	380	276	156	156	154	186	153	169	153	98	71	43	38	43	40	44	40	40	0	0	0																										
5	187	189	273	273	302	177	177	162	131	125	109	102	102	102	75	43	28	28	27	21	24	24	24	0	0	0																										
6	150	151	151	121	121	94	66	51	46	26	26	27	29	27	29	27	15	9	9	9	11	11	5	0	0	0																										
7	17	60	60	49	49	25	21	21	21	16	16	18	20	21	21	21	10	9	6	5	1	1	1	0	0	0																										
8	48	89	95	89	80	63	63	35	30	30	26	34	29	33	29	17	10	12	7	9	6	4	4	0	0	0																										
9	9	24	35	24	35	22	15	13	13	8	8	15	15	12	9	6	6	6	6	6	6	5	5	0	0	0																										
10	562	562	562	557	506	467	279	174	98	95	98	88	88	88	79	41	40	25	25	33	33	24	16	0	0	0																										
11	489	527	527	527	390	355	203	151	95	76	76	76	67	67	63	33	16	14	14	20	21	21	19	0	0	0																										
12	424	434	434	414	414	362	228	92	72	63	61	63	69	76	76	20	13	13	22	22	31	19	14	0	0	0																										
13	523	530	550	550	549	457	303	261	224	135	122	122	122	109	79	45	42	35	32	32	32	11	11	0	0	0																										
14	864	864	891	861	924	827	585	384	384	384	299	284	232	230	230	134	93	93	97	97	86	77	44	0	0	0																										
15	667	672	718	718	786	602	564	465	453	358	257	198	171	168	148	95	88	60	50	60	61	61	54	0	0	0																										
16	138	192	201	202	204	204	148	117	57	8	8	8	12	13	12	7	5	5	9	9	8	8	4	0	0	0																										
17	91	159	198	198	198	193	174	99	64	54	16	16	10	10	11	11	8	8	11	11	10	8	8	0	0	0																										
18	359	493	505	493	625	421	389	349	120	99	41	41	41	46	18	16	16	16	17	17	17	15	15	0	0	0																										
19	171	217	306	306	339	282	226	204	132	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0																										
20	0	0	0	0	0	204	226	204	132	44	15	9	9	22	22	9	2	2	6	6	6	4	2	0	0	0																										
21	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0																										
22	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0																										
23	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0																										
24	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0																										
25	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0																										
26	0	0	0	0	0	0	0	0	0	0	0	0	0	8	9	12	12	11	11	10	9	8	7	7	11	16																										
27	0	0	0	0	0	0	0	0	0	0	0	0	5	5	5	4	4	9	9	7	3	3	3	9	9	10																										
28	0	0	0	0	0	0	0	0	0	0	0	0	1	1	8	8	8	12	14	15	15	14	14	22	38	38																										
29	0	0	0	0	0	0	0	0	0	0	0	0	6	6	6	5	5	5	8	7	8	8	8	13	13	24																										
30	0	0	0	0	0	0	0	0	0	0	0	0	0	0	7	7	6	6	6	4	4	4	5	8	5	2																										
31	0	0	0	0	0	0	0	0	0	0	0	0	0	0	7	7	6	2	2	2	2	5	1	1	1	4																										
32	49	49	31	23	23	22	23	23	23	23	29	36	74	74	71	64	34	19	10	10	15	15	7	0	0	0																										
33	73	73	63	58	58	51	51	50	50	60	85	110	152	152	101	79	73	66	40	35	30	26	8	0	0	0																										
34	54	54	48	43	37	37	39	40	40	67	82	107	121	121	121	104	71	37	27	27	27	28	24	0	0	0																										
35	93	93	87	61	52	52	42	42	46	63	76	134	184	184	160	124	74	57	47	37	38	36	28	0	0	0																										
36	61	28	17	9	9	9	4	4	14	24	28	46	47	89	89	77	48	30	15	15	16	21	21	0	0	0																										
37	45	29	14	12	6	7	7	14	19	25	42	65	88	88	87	68	53	35	32	30	30	30	11	0	0	0																										
38	51	33	18	8	8	7	7	11	15	35	37	67	83	86	86	66	33	25	25	29	29	24	14	0	0	0																										
39	133	56	41	28	13	28	29	33	64	116	133	176	234	234	252	234	164	108	108	104	137	104	100	0	0	0																										
40	103	82	56	48	43	43	37	41	41	46	46	43	37	33	33	31	24	19	17	16	16	15	15	0	0	0																										
41	81	67	63	41	33	33	32	32	47	47	47	45	52	39	39	33	19	11	9	8	8	8	4	0	0	0																										
42	113	107	74	63	60	59	59	59	60	73	73	73	73	62	50	48	28	18	18	18	26	11	11	0	0	0																										
43	438	307	285	235	157	157	161	201	215	285	285	198	166	166	166	101	70	70	72	73	94	107	107	0	0	0																										
44	29	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0																										
45	33	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0																										
46	37	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0																										
47	82	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0																										
48	237	283	305	413	477	413	350	327	303	119	119	63	28	28	25	25	33	60	62	60	45	45	45	0	0	0																										
49	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0																										
50	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0																										
51	186	278	279	278	279	252	223	140	60	29	25	21	17	21	28	28	21	15	12	13	13	14	14	0	0	0																										
52	87	87	61	53	53	41	41	41	74	76	138	181	200	181	175	153	120	118	120	95	101	95	87	0	0	0																										
53	26	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0																										
54	67	63	63	52	72	70	70	70	90	118	146	181	215	215	205	188	138	122	115	122	132	132	97	0	0	0																										
55	271	309	330	359	378	359	321	321	321	202	74	71	53	71	78	78	78	58	44	30	30	30	18	0	0	0																										
56	64	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0																										
57	243	267	267	364	396	364	264	220	178	70	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0																										
58	53	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0	0																										
59	266	266	257	255	257	268	268	249	170	125	90	90	70	91	70	55	36	35	35	40	40	37	37	37	39	72																										

Bibliographie

- [AAC 89] Arthur Andersen Consulting, *Quick Response : A study of Costs and Benefits to Retailers of Implementing Quick Response and Supporting Technologies*, Technical Report, New York, NY, 1989.
- [ABR 83] Abraham B., Ledolter J., *Statistical Methods for Forecasting*, Wiley, New York, 1983.
- [ABR 87] Abraham M.M., Lodish L.M., *Promoter : an automated promotion evaluation system*, Marketing Science, n°6, Spring, pp 101-123, 1987.
- [ADI 91] Adiga S., Glassey R., *Object oriented simulation to support research in manufacturing systems*, International Journal of Production Research, vol 29, n°12, pp 2529-2542, 1991.
- [ALA 89] Alan A., Pritsker B., Jack R. D., *SLAM II, network models for decision support*, Prentice Hall, 1989.
- [ARM 92] Armstrong J. S., Collopy F., *Error measures for generalizing about forecasting methods : Empirical comparisons*, International Journal of Forecasting, n°8, pp 69-80, 1992.
- [BAN 91] Banks J., *The simulator : New member of the simulation family*, Interfaces 21, pp 76-86, March-April 1991.
- [BAR 89] Bartolomei S. M., Sweet A. L., *A note on a comparison of exponential smoothing methods for forecasting seasonal series*, International Journal of Forecasting, n°5, pp 111-116, 1989.
- [BAS 75] Bass F.M., Wittink D.R., *Pooling issues and Methods in Regression Analysis with examples in Marketing Research*, Journal of Marketing Research, n°12, pp 414-425, November 1975.
- [BAS 78] Bass F.M., Wittink D.R., *Pooling issues and Methods in Regression Analysis : Some further reflections*, Journal of Marketing Research, n°15, pp 277-279, May 1978.
- [BAS 97] Bastian A., *Evolutionary Computation and Its Industrial Application*, Tutorial Presentation, EUFIT'97, 5th European Congress on Intelligent Techniques and Soft Computing, Aachen, Germany, September 8, 1997.
- [BAT 90] Batchelor R., Dua P., *Forecaster ideology, forecasting technique, and the accuracy of economic forecasts*, International Journal of Forecasting, n°6, pp 3-10, 1990.
- [BEL 96] Bellman R., Kalaba R., Zadeh L. A., *Abstraction and pattern classification*, J. Math. Anal. And Appl.2, pp 581-586, 1996.
- [BEN 82] Benzecri J.P., *Construction d'une classification ascendante hiérarchique par la recherche en chaîne des voisins réciproques*, Les Cahiers de l'analyse des données, vol 7, n°2, pp 209-218, 1982.
- [BEN 85] Benzecri F., Benzecri J.P., *Démonstration de l'équivalence des résultats des algorithmes accélérés à ceux de l'algorithme de base en CAH*, Les cahiers de l'analyse des données, vol X, n°3, pp 257-271, 1985.
- [BER 86] Berge C., *Espaces topologiques-fonctions multivoques*, Dunod, Paris, 1986.

- [BES 91] Besombes B., Ladet P., *La formalisation de la connaissance dans un système d'aide au pilotage d'ateliers flexibles*, actes du 3^{ème} congrès de Génie Industriel, Tome 1, Tours, 20-22 mars 1991.
- [BEZ 76] Bezdek J. C., *A physical interpretation of fuzzy ISODATA*, IEEE Transactions On Systems, Man and Cybernetics SMC6, pp 387-389, 1976.
- [BEZ 81] Bezdek J.C., *Pattern recognition with fuzzy objective function algorithm*, Plenum Press, New York, 1981.
- [BEZ 81a] Bezdek J. C., Coray C., Gunderson R., Watson J., *Detection and characterization of cluster substructure. I. Linear Structure : Fuzzy c-lines*, SIAM J. Appl. Math. 40 (2), pp 339-357, 1981.
- [BEZ 81b] Bezdek J. C., Coray C., Gunderson R., Watson J., *Detection and characterization of cluster substructure. II. Fuzzy c-varieties and convex combinations thereof*, SIAM J. Appl. Math. 40 (2), pp 358-372, 1981.
- [BIE 87] Bierens H.J., *Armax model specification testing, with an application to unemployment in the Netherlands*, Journal of Econometrics, n°35, pp 161-190, 1987.
- [BLA 90] Blattberg R.C., Neslin S.A., *Sales promotion : Concepts, Methods and Strategies*, Englewood Cliffs, Prentice Hall, 1990.
- [BOC 74] Bock H. H., *Automatische Klassifikation*, Vandenhoeck & Ropprecht, Göttingen, 1974.
- [BOU 92] Bourbonnais R., Usunier J.C., *Pratique de la prévision des ventes*, Economica, 1992.
- [BOU 94] Bounine J., *Le textile français a tout faux*, Journal du Textile, n°1393, pp 2-3, 21 novembre 1994.
- [BOU 97a] Boussu F., Rabenasolo B., Happiette M., Vasseur C., *Simulation de la filière Textile*, RAIRO-APII-JESA, vol 31, n°4, pp 741-768, juin 1997.
- [BOU 97b] Boussu F., Lefort A., Happiette M., Rabenasolo B., Yim P., *Modelling and Simulation of the Textile Channel by Hypernets*, Soumis et accepté dans la revue Journal of the Textile Institute, sept 97.
- [BOW 90] Bowerman B.L., O'Connell R.T., *Forecasting and Time Series : An applied approach*, Duxbury Press, 1990.
- [BOX 76] Box G.E.P., Jenkins G.M., *Time series analysis - forecasting and control*, Prentice Hall, 1976.
- [BRI 77] Brisse H., Grandjouan G., *Un procédé de classification par agrégation d'un effectif nombreux*, Bulletin de l'A.S.U., 3, Paris, 1977.
- [BRO 59] Brown R. G., *Statistical forecasting for inventory control*, New York, MacGraw Hill, 1959.
- [BRO 88] Brockwell P.J., Davis R.A., *Time Series : Theory and Methods*, Springer Verlag, New York, 1988.

- [BROz 88] Broze L., Melard G., *Maximum likelihood estimation for generalized exponential smoothing*, Eighth International Symposium on Forecasting, Amsterdam (Pays-Bas), 13-15 June 1988.
- [BRU 78] Bruynooghe M., *Méthodes nouvelles en classification de données taxinomiques nombreuses. Statistique et analyse des données*, Les Cahiers de l'analyse des données, vol 3, n°1, pp 7, 1978.
- [BRU 94] Brunet H., Le Denn Y., *La démarche logistique*, AFNOR gestion, 1994.
- [BUG 93] Bug P., Fischer T., *Information Systems in Textile Marketing and Sales*, ITS Textile Leader, pp 66, december 1993.
- [BUN 93] Bunn D. W., Vassilopoulos A. I., *Using group seasonal indices in multi-item short-term forecasting*, International Journal of Forecasting, n°9, pp 517-526, 1993.
- [CAD 96] Cadivel C., Yim P., Lefort A., *Spécifications des traitements avec les réseaux formels*, rapport interne LAIL, EC Lille - URA CNRS D1440, juin 1996.
- [CAN 80] Cantor J., *The Forecasting Model - A Study in Market Simulation*, J.S.N. International, december 1980.
- [CAR 82] Carbone R., Armstrong J.S., *Evaluation of extrapolative forecasting methods : Results of a survey of academicians and practionners*, Journal of Forecasting, n°1, pp 215-217, 1982.
- [CER 75] Cerullo M.J., Avila A., *Sales forecasting practices : a survey*, Managerial planning, pp 33-39, september/october 1975.
- [CER 88] Cernault A., *La simulation des systèmes de production*, Ed. Cepadues, 1988.
- [CHA 71] Chambers J.C., *How to choose the right forecasting technique*, Harvard Business Review, july-august 1971.
- [CHA 80] Charniack E., Riesbeck C., McDermott D., *Artificial Intelligence Programming*, Laurence Erlbaum Associates, Publishers, Hillsdale, New Jersey, 1980.
- [CHA 88] Chatfield C., *The analysis of Time Series - An introduction*, Chapman and Hall, 1988.
- [CHA 91] Chatfield C., Yar M., *Prediction intervals for multiplicative Holt-Winters*, International Journal of Forecasting, n°7, pp 31-37, 1991.
- [CHA 92] Chatfield C., *A commentary on error measures*, International Journal of Forecasting, n°8, pp 99-111, 1992.
- [CLA 88] McClain J.O., *Dominant tracking signals*, International Journal of Forecasting, n°4, pp 563-572, 1988.
- [COL 92a] Collopy F., Armstrong J.S., *Expert opinions about extrapolation and the mystery of the overlooked discontinuities*, International Journal of Forecasting, n°8, pp 575-582, 1992.
- [COL 92b] Collopy F., Armstrong J.S., *Rule based forecasting*, Management science, n°38, pp 1394-1414, 1992.
- [CON 85] Confection 2000, *Transport des vêtements sur cintres en "suspendu"*, n°57, pp 43-46, mai 1985.

- [COU 90] Coudurier J. F., *La simulation des flux de production*, Rapport d'étude CETIM, 1990.
- [CRO 94] Cross K.F., Feather J.J., Lynch R.L., *Corporate Renaissance : The Art of Reengineering*, Cambridge, MA: Blackwell, 1994.
- [DAL 74] Dalhart G., *Class seasonality - a new approach*, American Production and Inventory Control Society, Conference Proceedings, 1974.
- [DAL 87] Dalrymple D.J., *Sales forecasting practices : results from a United States survey*, International Journal of Forecasting, n°3, pp 379-392, 1987.
- [DAU 94] Daugherty J. P., Germain R., Dröge C., *Predicting EDI Technology Adoption in Logistics Management : The Influence of Context and Structure*, Logistics and Transportation Review, vol 31, n°4, pp 309-324, 1994.
- [DAV 79] Davis D.L., Bouldin D.W., *A cluster separation measure*, IEEE Transactions on Pattern Analysis and Machine Intelligence PAMI 1, pp 224-227, 1979.
- [DAV 92] David R., Alla H., *Du Grafcet aux réseaux de Pétri*, 2^{ème} édition revue et augmentée, Traité des nouvelles techniques (série automatique) Ed. Hermès, 1992.
- [DAV 94] David R., Alla H., *Petri Nets for Modelling of Dynamics Systems - a survey*, Automatica, vol 30, n°2, pp 175-203, February 1994.
- [DEF 94] DEFI, *La "Quick Response" de la filière textile aux Etats-Unis*, septembre 1994.
- [DEL 79] Delattre P., *classification optimale bi-critère*, rapport fac. Univ. Catholique de Mons, 1979.
- [DER 94] Dermagne J., *L'indispensable partenariat*, Industrie textile, n°1259, pp 23-24, novembre 1994.
- [DID 80] Diday E., *Optimisation en classification automatique*, INRIA, Rocquencourt, France, 1980.
- [DID 82] Diday E., Lemaire J., Pouget J., Testu F., *Eléments d'analyse de données*, Dunod, 1982.
- [DID 88] Diday E., (1988). The symbolic approach in clustering. In: *Classification and Related methods of data analysis* (H.H. Bock, (Ed. Elsevier)). North Holland, Amsterdam.
- [DID 96] Diday E., *Une introduction à l'analyse des données symboliques*, Ecole d'été, Lise-Ceremade, Université Paris 9 Dauphine, du 11-13 septembre 1996, Paris.
- [DOR 91] Dorchies J. C., *Perspectives de développement de l'approvisionnement de la grande distribution par l'industrie textile-habillement du Nord/Pas de Calais*, Rapport du GRIT / URIC, Octobre 1991.
- [DUD 73] Duda R. O., Hart P. E., *Pattern classification and scene analysis*. New York, Wiley, 1973.
- [DUN 74] Dunn J. C., *Well separated clusters and optimal fuzzy partitions*. J. Cybern., vol 4, pp 95-104, 1974.

- [EKE 89] Ekere N., Hammam R.G., *An evaluation of approaches to modelling and simulating manufacturing systems*, International Journal of production research, vol 27, n°4, pp 599-611, 1989.
- [EVE 74] Everitt B.S., *Cluster analysis*, John Wiley&Sons, Inc., New York, 1974.
- [FIL 83] Fildes R., *An evaluation of Bayesian forecasting*, Journal of Forecasting, n°2, pp 137-150, 1983.
- [FIL 95] Fildes R., Makridakis S., *The Impact of Empirical Accuracy Studies On Time Series Analysis and Forecasting*, International Statistical Review, 63, 3, pp 289-308, 1995.
- [FIZ 86] Fizazi H., Postaire J.G., *Classification optimale de petits échantillons par restauration des propriétés de convexité des fonctions de densité marginales de probabilité*, Congrès IASTED Int. Symp. On Identification and pattern recognition (RAI/IPAR), Toulouse, 1986.
- [FOU 85] Foucart T., *Analyse factorielle, programmation sur micro-ordinateurs*, 2ème édition, Masson, Paris, 1985.
- [GAR 81] Garcia-Cuevas G. J., Kirton T., Lowe P. H., *Demand forecasting for an integrated fashion business*, Clothing research journal, pp 3-15, 1981.
- [GAR 85] Gardner E.S., *Exponential Smoothing : the state of the art*, Journal of Forecasting, n°4, pp 1-28, 1985.
- [GEU 86] Geurts M.D., Kelly J.P., *Forecasting retail sales using alternative models*, International Journal of Forecasting, n°2, pp 261-272, 1986.
- [GOL 89] Goldberg D., *Genetic algorithms in Search Optimization and Machine Learning*, Addison-Wesley, Reading, MA, 1989.
- [GOL 95] Goldman S.L., Nagel R.N., Preiss K., *Agile Competitors and Virtual Organizations : Strategies for Enriching the Customer*, New York, NY : Van Nostrand Reinhold, 1995.
- [GOO 92] G. De Gooijer J., Kumar K., *Some recent developments in non-linear time series modelling, testing and forecasting*, International Journal of Forecasting, n°8, pp 135-156, 1992.
- [GOU 90] Gourieroux C., Monfort A., *Séries Temporelles et modèles dynamiques*, Economica, 1990.
- [GOV 75] Govaert G., *Classification automatique et distances adaptatives*, Thèse de 3ème cycle, université de Paris 6, 1975.
- [GOW 95] Gowda Chidananda K., Ravi T. V., *Divisive clustering of symbolic objects using the concepts of both similarity and dissimilarity*, Pattern recognition, vol. 28, n°8, pp 1277-1282, 1995.
- [GRA 80] Granger C.W.J., *Forecasting in business and economics*, New York, Academic Press, 1980.
- [GRE 95] Greenwell R.N., Angus J.E., Finck M., *Optimal Mutation Probability for Genetic Algorithms*, Math. Comput. Modelling, vol 21, n°8, pp 1-11, 1995.
- [GRO 72] Gross A. L., *A Monte Carlo study of the accuracy of a hierarchical grouping procedure*, Multivariate Behavior Research 7, pp 379-389, 1972.

- [GUF 78] Guftafson D.E., Kessel C.W., *Fuzzy clustering with a fuzzy covariance matrix*, Proc. IEEE, San Diego, pp 761-766, 1978.
- [GUN 83] Gunderson R.W., *An adaptative FCV clustering algorithm*, Int. J., Man-Machine Studies, n°19, pp 97-104, 1983.
- [HAC 90] Hackett R., *Investment in technology - the Service Sector Sinkhole?* Sloan Management Review, Winter, pp 97-103, 1990.
- [HAP 96a] Happiette M., Boussu F., Rabenasolo B., Zeng X., *L'approche de la classification dans les problèmes de prévision de ventes. Application à la filière Textile/Habillement/Distribution*, 5th Int. Conf. On CIMAT'96, France, Grenoble, pp 402-407, May 29-31, 1996.
- [HAP 96b] Happiette M., Rabenasolo B., Boussu F., *Sales partition for forecasting into textile distribution network*. IEEE International Conference on Systems, Man, and Cybernetics, Beijing, China, vol 4, pp 2868-2873, October 14-17 1996.
- [HAP 97] Happiette M., Boussu F., Rabenasolo B., Zeng X., *Clustering of symbolic objects : Decision aid for production planning*, 4th IFAC Workshop on Intelligent Manufacturing Systems, Seoul, Korea, pp 341-346, July 21-23, 1997.
- [HES 90] Hesser J., Männer R., *Towards an optimal mutation probability*, In Proceedings of the International Workshop Parallel Problem Solving from Nature, Springer Verlag, 1990.
- [HUG 95] McHugh P., Merli G., Wheeler W.A., *Beyond Business Process Reengineering : Towards the Holonic Enterprise*, Chichester, GB :John Wiley&Sons, 1995.
- [HUN 90] Hunter A., *Quick Response in Apparel Manufacturing*, The Textile Institute, Manchester, 1990.
- [HUN 92] Hunter A., King R.E., Nuttle H.L., *An apparel-supply system for QR retailing*, Journal of the Textile Institute, vol 83, n°3, 1992.
- [HUN 95] Hunter N.A., Valentino P., *Quick response - ten years later*, International Journal of Clothing Science and Technology, vol 7, n°4, pp 30-40, 1995.
- [HUN 96] Hunter A., King R.E., Nuttle H.L., *Evaluation of Traditional and Quick Response Retailing Procedures by Using a Stochastic Simulation Model*, Journal of the Textile Institute, vol 87, Part 2, n°1, 1996.
- [HUR 89] Hurel D., *Analyse et modélisation du processus boursier. Proposition d'une méthode de gestion de portefeuille*, Thèse de doctorat de l'université de Lille, le 6 janvier 1989.
- [HUS 85] Huss W.R., *Comparative analysis of company forecasts and advanced time series techniques in the electric utility industry*, International Journal of Forecasting, n°1, pp 217-239, 1985.
- [HUT 94] Hutchinson W., *RADAR (Radical Database Reduction)*, Center for retailing education and research, Working paper, University of Florida, 1994.
- [JAF 94] Jaffar J., Lassez J.L., *Constraint logic programming : a survey*, Journal of Logic Programming, Vol. 19/20, pp. 503-582, May-July 1994.
- [JAI 88] Jain A.K., Dubes R.C., *Algorithm for clustering data*, Prentice Hall, 1988.
- [JAM 78] Jambu M., *Classification automatique pour l'analyse des données*, Dunod, Paris, 1978.

- [JEN 91] Jensen K., *Coloured Pétri nets : a high level language for system design and analysis*, ROZENBERG G. (editions), Lecture notes in Computer Science 483, pp 342-416, Springer-Verlag, avril 1991.
- [JOH 88] Johnston J.J., *Econometric Methods*, McGraw Hill, Auckland, 1988.
- [JON 91] Jones D.R., Beltramo M.A., *Solving Partitioning problems with genetic algorithms*, in Proceedings of the 4th International Conference on Genetic Algorithms, Morgan Kaufmann Publishers, Los Altos, CA, pp 442-449, 1991.
- [JUA 82a] Juan J., *Le programme HIVOR de classification ascendante hiérarchique selon les voisins réciproques et le critère de la variance*, Les Cahiers de l'analyse des données, vol 7, n°2, pp 173-184, 1982.
- [JUA 82b] Juan J., *Programme de classification hiérarchique par l'algorithme de la recherche en chaîne des voisins réciproques*, Les cahiers de l'analyse des données, vol 7, n°2, pp 219-225, 1982.
- [KAM 85] Kamgar-Parsi B., Kanal L. N., *An improved branch and bound algorithm for computing k-nearest neighbors*, Pattern Recognition Letters 3, pp 7-12, 1985.
- [KEE 91] Keefe R., Haddock J., *Data driven generic simulators for flexible manufacturing systems*, International Journal of Production Research, vol 29, n°9, pp 1797-1810, 1991.
- [KIE 86] Kieffer J.P., *Les systèmes de production, leur conception et leur exploitation*, Thèse d'état, Aix-Marseille, 1986.
- [KIR 66] Kirby R.M., *A Comparison of Short and Medium Range Statistical Forecasting Methods*, Management Science vol 4, 1966.
- [KOE 88] Koehler A.B., Murphree E.S., *A comparison of results from state space forecasting with forecasts from the Makridakis Competition*, International Journal of Forecasting, n°4, pp 45-55, 1988.
- [KOU 93] Kouam A., *Approches connexionnistes pour la prévision des séries temporelles*, thèse de doctorat de l'université de Paris Sud, centre d'orsay, le 8 juillet 1993.
- [KRA 72] Krampf R.F., *The turning point problem in smoothing models*, Thèse de doctorat, Université de Cincinnati, 1972.
- [KSA 88] Kurt Salmon Associates, *Quick Response Implementation - Action Steps for Retailers, Manufacturers and Suppliers*, Technical Report, New York, NY, 1988.
- [KSA 94] KSA (Kurt Salmon Associates), *A Vision for the New Millennium*, I.T. Intelligence, 1994.
- [LAV 81] Lavoie R., *Statistique Appliquée*, les presses de l'université du Québec, 1981.
- [LAW 83] Lawrence M.J., *An exploration of some practical issues in the use of quantitative forecasting models*, Journal of Forecasting, n°1, pp 169-179, 1983.
- [LEB 95] Lebart L., Morineau A., Piron M., *Statistique exploratoire multidimensionnelle*, Dunod, 1995.
- [LEF 95] Lefort A., Yim P., *Modelling Hybrid Systems with Hypernets*, IEEE Inter. Conf. SMC, pp 1423-1429, Vancouver, Canada, 22-25/10/95.

- [LEN 93] Lenclud T., *Contribution à la conception d'un système intégrée de simulation des systèmes de production*. Thèse de doctorat, Université de Valenciennes, 1993.
- [LEN 95] Lenclud T., Tahon C., *Modélisation et évaluation des systèmes logistiques*, 1^{ère} rencontres internationales de la recherche en logistique, Marseille, 11-13 janvier 1995.
- [LEV 67] Levine A.H., *Forecasting Techniques*, Management Accounting, janvier 1967.
- [LEW 74] Lewandowski R., *La prévision à court terme*, Dunod, 1974.
- [LOV 94] Loveman G.W., *An assesment of the Productivity Impact of Information Technologies*, In Allen, T.J. and Scott Morton, M.S. (eds.), *Information Technology and the Corporation of the 1990s*, Oxford, Oxford UP, 1994.
- [MAH 36] Mahalanobis P.C., *On the generalized distance in Statistics*, Proc. Nat. Inst. Science India, vol.12, pp 49-55, 1936.
- [MAH 86] Mahmoud E., Rice G., Malhotra N., *Emerging issues in sales forecasting and decision support systems*, Journal of Academy of Marketing Science, n°16, pp 47-61, 1986.
- [MAK 79] Makridakis S., Hibon M., *Accuracy of forecasting : an empirical investigation with discussion*, Journal of the Royal Statistical Society, n°142, pp 97-145, 1979.
- [MAK 82] Makridakis S., Andersen A., Carbone R., Fildes R., Hibon M., Lewandowski R., Newton J., Parzen P., Winckler R., *The accuracy of extrapolation (time series) methods; results of a forecasting competition*, Journal of Forecasting, n°1, pp 111-153, 1982.
- [MAK 83] Makridakis S., Wheelwright S.S., McGee V.E., *Forecasting : methods and applications*, (2nd edition), New York, Wiley, 1983.
- [MAK 84] Makridakis S., Andersen A., Carbone R., Fildes R., Hibon M., Lewandowski R., Newton J., Parzen E., Winkler R., *The Forecasting Accuracy Major Time Series Methods*, Wiley, Chichester, 1984.
- [MAK 93] Makridakis S., Chatfield C., Hibon M., Lawrence M., Mills T., Ord K., Simmons L.F., *The M-2 Competition : A real-life judgemental based forecasting study*, International Journal of Forecasting, n°9, pp 5-29, 1993.
- [MAR 87] Martin F., *Méthodologie de modélisation et simulation de systèmes complexes décrits par réseaux de Pétri colorés*, Thèse de doctorat, INPG, 27 avril 1987.
- [MAS 90] Di Mascolo M., *Modélisation et évaluation de performances de systèmes de production gérés en Kanban*, thèse de doctorat, INPG, 1990.
- [MAT 87] Mathe H., Tixier D., *La logistique*, Que Sais-je ?, Ed PUF, 1987.
- [MEL 90] Mélard G., *Méthodes de prévision à court terme*, Ed. de l'université de Bruxelles, 1990.
- [MEN 84] Mentzer J.T., Cox J.E., *Familiarity, application, and performance of sales forecasting techniques*, Journal of Forecasting, n°3, pp 27-36, 1984.
- [MIC 92] Michalewicz Z., *Genetic algorithms + data structures = Evolution Programs*, Springer Verlag, 1992.

- [MIL 79] Milligan G. W., *Ultrametric hierarchical clustering algorithms*, Psychometrika 44, pp 343-346, 1979.
- [MIL 93] Miller T., Liberatore M., *Seasonal Exponential Smoothing with damped trends - An application for production planning*, International Journal of Forecasting, n°9, pp 509-515, 1993.
- [MOI 88] Le Moigne J. L., *La modélisation des systèmes complexes*, Afcet systèmes, Dunod, 1988.
- [NEW 74] Newbold P., Granger C.W.J., *Experience with forecasting univariate time series and the combination of forecasts with discussion*, Journal of the Royal Statistical Society, n°137, pp 131-165, 1974.
- [NOY 90] Noyes D., *Les réseaux de Petri et le modélisation modulaire des systèmes de production*, RAIRO APII, vol 24, n°6, pp 511-528, 1990.
- [RAB 97] Rabenasolo B., Happiette M., Boussu F., *Sales Forecasting under uncertain environment. Fuzzy classification in Textile Distribution*, IFAC-CIS 97 Workshop, Belfort, France, vol 3, pp 478-483, May 20-22 1997.
- [NUT 91] Nuttle H.L., King R.E., Hunter A., *A stochastic model of the apparel-retailing process for seasonal apparel*, Journal of the Textile Institute, vol 82, n°2, 1991.
- [PED 90] Pedgen C. D., Shannon R.E., Sadiowski T.E., *Introduction to simulation using SIMAN*, Ed. Mac Graw Hill, 1990.
- [PEG 69] Pegels C.C., *Exponential Smoothing : some new variations*, Management Science, n°12, pp 311-315, 1969.
- [PIC 90] Piccolo D., *A distance measure for classifying ARMA models*, Journal of Time Series, 11, vol 2, pp 153-164, 1990.
- [PIN 93] Pine B.J., *Mass Customization : The New Frontier in Business Competition*, Boston, MA : Harvard Bussiness School Press, 1993.
- [PON 93] Pons J., Chevalier P., *La logistique intégrée*, Ed Hermès, collection systèmes d'information, 1993.
- [PRO 95] Proth J.M., Xie X., *Les réseaux de Pétri pour la conception et la gestion des systèmes de production*, Ed. Masson, 1995.
- [QUE 96] Quelch J. A., Klein L. R., *The Internet and International Marketing*, Sloan Management Review, Spring, pp 60-75, 1996.
- [PET 77] Peterson J. L., *Petri Nets*, ACM Computing Surveys, vol 9, n°3, pp 223-252, september 1977.
- [PHA 93] G. Phanendra Babu, M. Narasimha Murty, *A near optimal initial seed value selection in K-means algorithm using a genetic algorithm*, Pattern Recognition Letters, n°14, pp 763-769, 1993.
- [RAI 71] Raine J.E., *Self adaptative forecasting considered*, Decision Sciences, avril 1971.
- [RHA 80] Rham C., *La classification hiérarchique ascendante selon la méthode des voisins réciproques*, Les cahiers de l'analyse des données, vol 5, pp 135-144, 1980.



- [RHO 97] Rhodes E., Carter R., *Impacts of globalization in textile and apparel retailing*, International ITAA symposium, Lyon, 10 - 12th July, 1997.
- [ROB 82] Roberts S.A., *A general class of Holt-Winters type forecasting models*, Management Science, n°28, pp 808-820, 1982.
- [ROB 91] Robinet P., *La prise de commande par EDI, Echanges de bons procédés*, Confection 2000, n°123, pp 80-89, juin 1991.
- [ROB 94] Robinet P., *L'EDI, élément clé du circuit court*, Confection 2000, n°154, pp 38-41, avril 1994.
- [ROU 85] Roux M., *Algorithmes de classification*, Masson, 1985.
- [RUS 69] Ruspini E.H., *A new approach to clustering*, Inform. and Control, vol 15, pp 22-32, 1969.
- [SAN 89] Sanders L., *L'analyse statistique des données en géographie*, G.I.P. RECLUS, Montpellier, 1989.
- [SAN 94] Sanders N.R., Mandrodt K.B., *Forecasting practices in US corporations : survey results*, Interfaces, 24 (2), pp 92-100, 1994.
- [SAU 84] Saulou J. Y., *Le tableau de bord du décideur - méthodologie de mise en place*, Les éditions d'organisations, Paris, 1984.
- [SCH 86] Schnaars S.P., *A comparison of extrapolation models on yearly sales forecasts*, International Journal of Forecasting, n°2, pp 71-85, 1986.
- [SCH 89] Schaffer J.D., Caruna R.A., Eshelman L.J., Das R., *A study of control parameters affecting on line performance of genetic algorithms for function optimization*, In Proceedings of the Third International Conference on Genetic Algorithms, pp. 51-60, Morgan Kaufmann, 1989.
- [SEL 84] Selim S. Z., Ismail M. A., *K-means type algorithms : a generalized convergence theorem and characterization of local optimality*, IEEE Trans. on pattern analysis and machine intelligence, vol. PAMI-6, n°1, january 1984.
- [SHA 92] Shaw C. T., King G. P., *Using cluster analysis to classify Time series*, Physica D Non linear Phenomena, n°58, pp 288-298, 1992.
- [SHI 86] Shingo S., *Maîtrise de la production et méthode Kanban. Le cas Toyota*, les éditions d'organisations, 1986.
- [SIN 96] Singletary E.P., Winchester S.C., *Beyond Mass Production : Analysis of the Emerging Manufacturing Transformation in the US Textile Industry*, Journal of the Textile Institute, vol 87, part 2, n°2, 1996.
- [SMI 94] Smith S. A., McIntyre S. H., Achabal D. D., *A two-stage Sales Forecasting Procedure Using Discounted Least Squares*, Journal of Marketing Research, vol XXXI, pp 44-56, February 1994.
- [SOK 58] Sokal R. R., Michener C. D., *A statistical method for evaluating systematic relationships*, Univ. Kansas Sci. Bull., t. 38, pp.1409-1438, 1958.

- [SOK 63] Sokal R.R., Sneath P.H.A., *Principles of numerical taxonomy*, W.H. Freeman and Cie, London, 1963.
- [SOR 48] Sørensen T., *A method of establishing groups of equal amplitude in plant sociology based on similarity of species content and its application to analyses of the vegetation on Danish commons*, Biol. Skr, 5 (4), pp 1-34, 1948.
- [SPA 84] Sparkes J.R., McHugh A.K., *Awareness and use of forecasting techniques in British industry*, Journal of Forecasting, n°3, pp 37-42, 1984.
- [SRI 94] Srinivasan K., Kekre S., Mukhopadhyay T., *Impact of Electronic Data Interchange Technology on JIT Shipments*, Management Science, vol 40, n°10, pp 1291-1304, 1994.
- [SYS 89] Syswerda G., *Uniform crossover in genetic algorithms*, in Proc. Third Int. Conf. On Genetic Algorithms, pp 2-9, 1989.
- [SZY 88] Szymankiewicz J., Donald J.MC., Turner K., *Solving business problems by simulation*, Mac Graw Hill, 1988.
- [TAR 89] Taranto G.M., *Sales forecasting practices : results from an Australian survey*, PhD thesis, University of N.S.W., 1989.
- [TAY 27] Taylor F. W., *Principes d'organisation scientifique du travail*, Paris, Dunod, 1927.
- [TEX 91] Texter Geriner P., Keith Ord J., *Automatic forecasting using explanatory variables : A comparative study*, International Journal of Forecasting, n°7, pp 127-140, 1991.
- [THI 93] Thiel D., *Management Industriel. Une approche par la simulation*, Economica, 1993.
- [THI 95] Thiel D., *Modélisation et simulation des systèmes de production*, congrès ENSAIT sur la logistique du 8 juin 1995.
- [THO 90] Thompson P. A., *An MSE statistic for comparing forecast accuracy across series*, International Journal of Forecasting, n°6, pp 219-227, 1990.
- [THU 92] Thurow L., *Head to Head : The Coming Economic Battle Among Japan, Europe and America*, London, Nicholas Brearley Publishing, 1992.
- [TON 90] Tong H., *Non-linear Time Series : A Dynamical System Approach*, Oxford University Press, Oxford, 1990.
- [VON 91] Von Laszewski G., *Intelligent Structural Operators for the K-Way Graph Partitioning Problem*, in Proceedings of the 4th International Conference on Genetic Algorithms, Morgan Kaufmann Publishers, Los Altos, CA, pp 45-52, 1991.
- [VRO 96] Vroman P., Happiette M., Rabenasolo B., *Dynamic Textile Sales Forecasting Model Using Fuzzy Theory*, CESA'96, Multiconference IMACS/IEEE, vol , n°, pp, Lille, 1996.
- [VRO 97] Vroman P., Happiette M., Rabenasolo B., *Fuzzy adaptation of Holt-Winters Model for Textile Sales Forecasting*, Soumis à la revue Journal of the Textile Institute, juin 1997.
- [WAR 63] Ward J.H., *Hierarchical grouping to optimize an objective function*, Journal of the American Statistical Association, n°58, pp 236-244, 1963.
- [WAT 87] Watson M.W., Pastuszek L.M., Cody E., *Forecasting commercial electricity sales*, Journal of Forecasting, n°6, pp 117-136, 1987.

- [WEB 96] Webby R., O'Connor M., *Judgemental and Statistical time series forecasting : a review of the literature*, International Journal of Forecasting, n°12, pp 91-118, 1996.
- [WHE 74] Wheelwright S.C., Makridakis S., *Choix et valeurs des méthodes de prévision*, Ed. d'organisation, Paris, 1974.
- [WHE 83] Wheelwright S.C., Makridakis S., *Méthodes de prévision pour la gestion*, Les éditions d'organisation, 1983.
- [WHI 89] Whitley D., Starkweather T., Fuquay D., *Scheduling problems and traveling salesmen : the genetic edge recombination operator*, in Proc. Third Int. Conf. On Genetic Algorithms, pp 116-123, 1989.
- [WIN 60] Winters P. R., *forecasting sales by exponentially weighted moving average*, Management Science, Vol 6, pp 324-342, 1960.
- [WIT 89] Whithycombe R., Forecasting with combined seasonal indices, International journal of forecasting, n°5, pp 547-552, 1989.
- [WIT 92] Witt S., Witt C., *Modelling and forecasting demand in tourism*, Academic Press, London, 1992.
- [WOM 92] Womack P., Jones T., Roos D., *Le système qui va changer le monde. Une analyse des industries automobiles mondiales dirigées par le MIT*, Dunod, 1992.
- [XIE 95] Xie X.L., Beni G., *A validity measure for fuzzy clustering*. IEEE Trans. on Pattern Analysis and Machine Intelligence, vol 13, n°8, pp 841-847, 1995.
- [YAN 93] Yang M.S., *A survey of fuzzy clustering*, Math. Comput. Modelling, vol 18, n°11, pp 1-6, 1993.
- [YAR 90] Yar M., Chatfield C., *Prediction intervals for the Holt-Winters forecasting procedure*, International Journal of Forecasting, n°6, pp127-137, 1990.
- [YIM 95] Yim P., Cadivel C., Lefort A., *Process Modelling in Information Systems by Format Nets*, IEEE Inter. Conf. SMC, pp 2263-2269, Vancouver, Canada, 22-25/10/95.
- [YOK 95] Yokum J. T., Armstrong J. S., *Beyond accuracy : Comparison of criteria used to select forecasting methods*, International Journal of Forecasting, n°11, pp 591-597, 1995.
- [ZAD 65] Zadeh L.A., *Fuzzy sets, Information and control*, n°8, pp 338-353, 1965.
- [ZEN 91] Zeng X., *Pilotage des processus dynamiques par échantillonnage statistique*, Thèse de Doctorat, Université de Lille 1, 1991.
- [ZEN 93] Zeng X., Vasseur C., *Recherche de la meilleure partition d'échantillons par scission-fusion, Diagnostic et sûreté de fonctionnement*, vol 3, n°2, pp 121-150, 1993.
- [ZHA 91] Zhang P., Engels C., Ethridge D., *Analyzing and Forecasting Weekly, Monthly and Quaterly Cotton Spot Prices ; A Time Series Analysis Approach*, 15th Cotton Economics&Marketing Conference, pp 354-359, Beltwide Cotton Conferences - 1991.

