



Outil d'aide à la décision dans le pilotage de plateforme logistique hospitalière : Mise en place d'un module de préconisation de commandes

Dhia Jomaa

► To cite this version:

Dhia Jomaa. Outil d'aide à la décision dans le pilotage de plateforme logistique hospitalière : Mise en place d'un module de préconisation de commandes. Génie logiciel [cs.SE]. université jean Monnet de saint-étienne, 2013. Français. NNT : . tel-01758331

HAL Id: tel-01758331

<https://theses.hal.science/tel-01758331>

Submitted on 4 Apr 2018

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

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



THÈSE

**En vue de l'obtention du
DOCTORAT DE L'UNIVERSITE JEAN MONNET DE SAINT ETIENNE**

**Délivré par :
Université Jean Monnet de Saint Etienne**

**Présentée par
Dhia Jomaa**

le 26 Novembre 2013

**Titre :
Outil d'aide à la décision dans le pilotage de plateforme
logistique hospitalière : Mise en place d'un module de
préconisation de commandes**

Ecole Doctorale Sciences, Ingénierie, Santé - ED SIS 488

**Unité de recherche :
Laboratoire d'Analyse des Signaux et des Processus Industriels – LASPI – EA 3059**

Jury :

Lionel Dupont	Professeur des Universités	Rapporteur
Michel Gourgand	Professeur des Universités	Rapporteur
Naoufel Cheikhrouhou	Professeur des Universités	
Djamal Simohand	Directeur R&D KLS	
Anne Meunier	Pharmacien des Hôpitaux HCL	
Thibaud Monteiro	Professeur des Universités	Directeur de thèse
Béatrix Besombes	Maître de conférences	Co-directeur de thèse

Remerciements

Je tiens à exprimer en premier lieu ma gratitude à mon directeur de thèse Thibaud Monteiro professeur de l'INSA de LYON. Je ne saurais jamais assez le remercier de son soutien, son aide et ses conseils avisés et pour le suivi attentif de mes travaux, ses encouragements, sa gentillesse et pour tout ce que j'ai appris grâce à lui. Merci pour les critiques qui m'ont permis de mener à bien ces travaux.

Je tiens ensuite à exprimer mes remerciements à ma co-directrice de thèse Béatrix Besombes. J'ai énormément appris à son contact, ses remarques pertinentes m'ont aidé à cadrer le sujet.

Je remercie Djamel Simohand, directeur R&D au sein de KLS, pour la qualité de l'encadrement et du soutien qu'il m'a accordés tout au long de ces trois années. Grâce à lui, je dispose d'une meilleure méthodologie ainsi que d'un meilleur recul par rapport à la recherche que j'ai effectuée.

Je remercie Gilbert Garcia, président directeur général de KLS, pour m'avoir accueilli dans sa société et de m'avoir ainsi permis d'effectuer ce projet de thèse.

Je tiens aussi à remercier tous les membres de KLS, véritable communauté d'experts. En particulier Abdelmoula Bechar, Jérôme Bidault, Benoît Lhôte et Julien Flocard.

Qu'il me soit permis de remercier également Gourgand Michel et Dupont Lionel qui m'ont fait l'honneur d'accepter d'évaluer ce travail, ainsi que Naoufel Cheikhrouhou et Anne Meunier pour avoir accepté d'être membres du Jury de soutenance.

Je remercie François Guillet, directeur du laboratoire LASPI, pour m'avoir accueilli dans son laboratoire et de m'avoir ainsi permis d'effectuer mes travaux de recherche.

Enfin, je rends hommage à ceux sans lesquels ce travail n'aurait pas pu être fait : à toute ma famille, à mes amis, à Marianne, Adel, Amine et Ines. Ce travail leur est dédié.

Dhia

Table des matières

CHAPITRE I.	INTRODUCTION	1
1	Evolution de la logistique hospitalière	1
2	Vers une chaîne logistique hospitalière : le cas de la pharmacie	2
3	Evolution du rôle des plateformes hospitalières	4
4	Evolution des outils informatiques d'aide à la décision pour la gestion des plateformes logistiques.....	7
5	Extension de la sphère d'application des WMS	9
5.1	Évolution fonctionnelle des WMS.....	9
5.2	La problématique de la préconisation des commandes : Module de gestion des approvisionnements.....	11
6	Problématique abordée : Développement d'un module de préconisation de commandes pour les plateformes pharmaceutiques.....	12
7	Organisation de la thèse	16
CHAPITRE II.	COMPOSANTES DU MODULE DE PRECONISATION PROPOSE.....	18
1	Module de préconisation de commandes existant	18
1.1	Système de prévision utilisé.....	20
1.2	Règles de gestion	20
1.3	Architecture informatique de support.....	23
2	Limites du système et pistes d'amélioration retenues.....	25
3	Structure et fonctions du système proposé.....	29
3.1	Objectifs du système proposé	29
3.2	Méthode	30
3.3	Structure du système.....	30
4	Conclusion.....	36

CHAPITRE III.	MODULE DE CLASSIFICATION DES ARTICLES DU STOCK.....	38
1	Intérêt de la classification des articles par une approche multicritères	38
2	Méthodes de classification des articles du stock.....	41
2.1	Méthode AHP : Processus de hiérarchisation analytique.....	42
2.2	Méthodes basées sur la programmation linéaire	45
2.3	Méthodes basées sur les méta-heuristiques et l'intelligence artificielle.....	52
2.4	Récapitulatif : Avantages et inconvénients des méthodes de classification présentées	56
3	Méthodes retenues et concept du module de classification des articles du stock.....	60
3.1	Critères de classification	60
3.2	Méthodes de classification retenues	61
3.3	Fonctionnement du module de classification des articles du stock	62
4	Conclusion.....	68
CHAPITRE IV.	MODULE DE PREVISION.....	69
1	Problématique de prévision : Revue de littérature sur les systèmes informatisés de prévision	69
1.1	L'importance des prévisions dans le contexte de la gestion de stock	69
1.2	Techniques de prévision de la littérature.....	71
1.3	Systèmes experts pour la prévision - Déploiement de la théorie de prévision dans les systèmes d'aide à la décision.....	80
1.4	Techniques retenues.....	85
2	Système de prévision proposé	87
2.1	Premier filtrage des données	88
2.2	Correction des valeurs extrêmes	90
2.3	Identification des typologies de consommation	93
2.4	Calcul des prévisions de consommation.....	99
2.5	Calcul des indicateurs de performance	105

3	Concept et interfaces développés.....	106
4	Conclusion.....	112
CHAPITRE V. ÉVALUATION DU SYSTEME : CAS D'APPLICATION.....		113
1	Cas des HCL : Hospices Civils de Lyon / Chaîne de réapprovisionnement.....	114
2	Quelles caractéristiques types pour la gestion des plateformes pharmaceutiques : Enquête auprès de centres hospitaliers	115
3	Données considérées (Plateformes + Historiques).....	117
4	Constitution des échantillons	119
5	Evaluation de la performance prévisionnelle.....	120
6	Evaluation de la gestion de stock	127
6.1	Politique de gestion adoptée.....	128
6.2	Résultats des simulations	131
7	Etude de cas d'un produit à profil complexe.....	137
8	Conclusion.....	139
CHAPITRE VI. CONCLUSIONS ET PERSPECTIVES		140
1	Conclusions	140
2	Perspectives.....	144
BIBLIOGRAPHIE		147

Liste des figures

Figure 1 : Description de la chaîne logistique pharmaceutique (Beretz, 2002).....	3
Figure 2:Différentes catégories des systèmes de gestion de la logistique	7
Figure 3: Composants du module de préconisation existant.....	19
Figure 4: Évolution du niveau de stock.....	21
Figure 5:Evolution du niveau de stock dans le cas de la politique sur seuil à niveau de recomplètement (s, S) Babai (2005)	22
Figure 6 : Architecture informatique en support au module de préconisation de commandes existant	24
Figure 7 : Composants du nouveau module de préconisation de commandes	31
Figure 8: Module de classification des articles	32
Figure 9: Module de prévision des consommations.....	33
Figure 10 : Architecture informatique déployée en support au nouveau module de préconisation de commandes	34
Figure 11: Coubre ABC	39
Figure 12 : Matrice bidimensionnelle de classification des articles du stock.....	41
Figure 13 : Structure hiérarchique des éléments influant la décision	43
Figure 14: Matrice de comparaison.....	44
Figure 15 : Agrégation et transformation des entrées dans un réseau de neurones	53
Figure 16 : Réseau de neurones à trois couches (Yu (2010))	54
Figure 17 : Réseau de neurones appliqué pour le problème de classification multicritères (Patrovi & Anandarajan (2002))	54
Figure 18: Processus de l'AHP	62
Figure 19 : Utilisation du module de classification	63
Figure 20 : Création d'un scénario ABC classique	64
Figure 21: Création d'un scénario multicritère automatique	65
Figure 22 : Création d'un scénario multicritère expert.....	66
Figure 23 : Application des scénarios de classification	67
Figure 24 : Classification de Pegels - Différentes variantes du lissage exponentiel.....	75
Figure 25 : Étapes de la méthode de Box-Jenkins (Makridakis et al. (1998))	77
Figure 26 : Système de prévision développé	88
Figure 27 : Valeurs extrêmes dans une série de données.....	91
Figure 28: Carte de contrôle pour l'identification des valeurs aberrantes.....	92

Figure 29 : Approximation par la droite des moindres carrés ordinaire.....	93
Figure 30 : Autocorrélogramme d'une série de données saisonnières.....	95
Figure 31 : Procédure d'identification de la saisonnalité	98
Figure 32 : Indicateur NF.....	103
Figure 33 : Indicateur AWS	104
Figure 34 : Module de prévision des consommations.....	107
Figure 35: Paramétrage de la procédure de prévision	108
Figure 36 : Performance prévisionnelle globale	109
Figure 37 : Options de filtrage des articles par typologie	110
Figure 38 : Affichage des articles saisonniers	111
Figure 39 : Chaîne logistique hospitalière CHU de Lyon	114
Figure 40 : Critères de classification énumérés	116
Figure 41 : Poids des différents critères.....	117
Figure 42 : Cas de 25_PE.....	125
Figure 43 : Cas de 32_PE.....	126
Figure 44 : Cas de 37_PE.....	126
Figure 45 : Taux de service (Pharmacie 25_PE).....	133
Figure 46 : Stock en volume (Pharmacie 25_PE)	133
Figure 47 : Stock en montant (Pharmacie 25_PE).....	134
Figure 48 : Taux de service (Pharmacie 24_MED)	135
Figure 49 : Stock en volume (Pharmacie 24_MED).....	135
Figure 50 : Stock en montant (Pharmacie 24_MED)	136
Figure 51 : Consommation historique de l'article AERIUS	137

Liste des tableaux

Tableau 1 : Travaux traitant de la classification des articles du stock	56
Tableau 2 : Plateformes pharmaceutiques considérées.....	118
Tableau 3 : Articles pilotes sur la base de la classification établie	120
Tableau 4 : échantillons finaux (Suite à l'élimination des cas atypiques).....	121
Tableau 5 : Typologies de consommation obtenues	122
Tableau 6 : Performances prévisionnelles obtenues	123
Tableau 7 : Performances des pharmacies	132
Tableau 8 : Performances prévisionnelles sur les 8 premiers mois de 2010	138
Tableau 9 : Taux de service atteint (Cas de l'AERIUS).....	138
Tableau 10 : Performance du stock (Cas de l'AERIUS)	139

Chapitre I.Introduction

1 Evolution de la logistique hospitalière

Les établissements hospitaliers se sont longtemps distingués des organisations industrielles par leur mode de fonctionnement centré sur le patient et le soin, réduisant leurs pratiques logistiques aux fonctions de transport des marchandises, des médicaments ou des patients. De par sa mission, principalement publique et sa vocation immatérielle, la recherche de la performance organisationnelle, objectif central dans le monde de l'entreprise, n'a pas été au centre des préoccupations de l'établissement hospitalier et un retard considérable, en terme de pratiques et d'outils informatiques décisionnels, s'est accumulé au fil du temps.

Le changement permanent du monde de la santé, le manque fréquent de personnel, la volonté de dégager les personnels soignants des tâches administratives ou encore les contraintes financières de plus en plus restrictives sont autant de facteurs qui ont rendu nécessaire une évolution profonde du métier de la logistique hospitalière. *Aptle(2001)* ou encore *Sampieri (2000)* ont considéré inéluctable une transformation radicale du mode organisationnel hospitalier afin de continuer l'amélioration de la qualité des prestations tout en supportant la contraction des budgets alloués.

Les considérations sur les budgets, longtemps délaissées au dépend de la qualité de soin, deviennent de plus en plus prioritaires. Les budgets affectés sont en effet de plus en plus serrés ce qui incite a fortiori à optimiser la gestion hospitalière par réduction des dépenses tout en étant assujetti à assurer la même qualité de soin au patient *Di Martinelly et al. (2005)*.

En France, cette évolution s'est accélérée au cours des deux dernières décennies, portée notamment, par la nécessité économique de rationalisation et de mutualisation des moyens impulsée par les tutelles à travers les différents plans Hôpital 2007, Hôpital 2010. Une prise de conscience de la part des acteurs de l'hôpital, autour des enjeux de la logistique pharmaceutique s'est imposée.

La pharmacie hospitalière constitue un élément majeur dans le support à l'activité de soin apporté aux malades. Elle alimente en médicaments et dispositifs médicaux les services de soins au sein des hôpitaux et les cliniques tout en assurant le réapprovisionnement auprès des laboratoires pharmaceutiques et des grossistes. *Landry et al. (2000)* ont estimé à 40% la part de la logistique dans les dépenses hospitalières annuelles dont, plus de la moitié est

imputable à la pharmacie hospitalière. Elle représente un poste de coût considérable, de sorte que l'optimisation de la logistique pharmaceutique devient un enjeu majeur pour les établissements hospitaliers

2 Vers une chaîne logistique hospitalière : le cas de la pharmacie

Durant les dix dernières années, un effort considérable a été entrepris par les responsables et décideurs hospitaliers afin de faire évoluer la logistique pharmaceutique vers la gestion de la chaîne logistique en s'inspirant des outils et approches utilisées dans le monde industriel. Cette évolution se traduit par une recherche de la performance globale de la chaîne et s'appuie sur une réorganisation des circuits de réapprovisionnement, de rationalisation des magasins de stockage et d'optimisation des circuits de dispensation vers les services tout en garantissant la qualité de délivrance et de soin aux patients.

Plusieurs chantiers logistiques se rapportant à la problématique de la "*Supply Chain*", en Europe et particulièrement en France, ont été initiés afin de mieux intégrer les flux et synchroniser entre acteurs en vue d'un meilleur réapprovisionnement et un rapprochement du fonctionnement industriel, fonctionnement ayant déjà fait ses preuves. Des plans de soutien à l'intégration logistique, aux regroupements et à la mutualisation des moyens hospitaliers ont fait l'objet de plans de l'état français, notamment le plan hôpital 2007 et le plan hôpital 2012. La pharmacie d'établissement est au centre de ces dispositifs.

Le regroupement des structures d'entreposage et de réapprovisionnement, souvent opéré en grande distribution et dans l'industrie, permet des gains de surface, la réduction des niveaux de stocks et une rationalisation des règles de gestion. Dans ce contexte évolutif, la pharmacie d'établissement est de moins en moins perçue comme un maillon isolé dans l'hôpital *Di Martinelly et al. (2005)*. Elle tend, désormais à faire partie intégrante d'un schéma logistique global, à l'image des réseaux de distributions classiques.

La structure de la chaîne logistique et l'évolution de la configuration du réseau de réapprovisionnement dans le domaine hospitalier ont été décrites dans certains travaux, dont celui de *Hassan (2006)*. Le schéma logistique hospitalier peut être résumé suivant la schématisation formalisée dans les travaux *Beretz (2002)* (Figure 1).

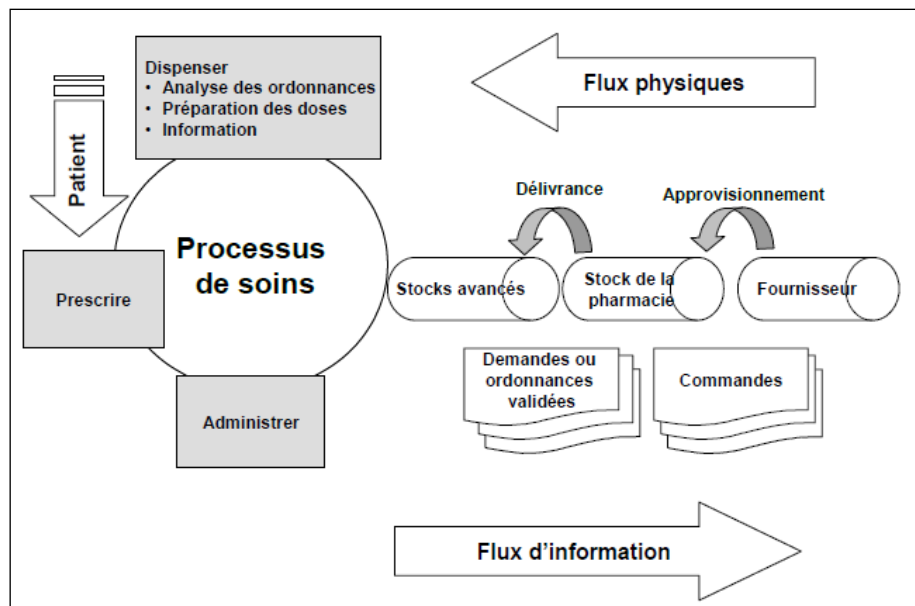


Figure 1 : Description de la chaîne logistique pharmaceutique (Beretz, 2002)

Le niveau le plus en amont de la chaîne est constitué par les fournisseurs. Ils sont présentés en général par de grands laboratoires assurant la mise à disposition d'un grand nombre de références de médicaments. Le niveau aval du réseau logistique comprend les établissements constitués de diverses unités de soin assurant la prise en charge des patients. Chacune des unités dispose d'un stock de produits pharmaceutiques (stock avancé), géré localement par le personnel de l'hôpital, notamment des infirmières, et s'approvisionnant auprès des pharmacies de l'établissement.

La pharmacie d'établissement (Stock de la pharmacie) constitue le maillon intermédiaire de la chaîne hospitalière. Dans le cas des gros établissements résultants du regroupement de plusieurs sites historiques, une ou plusieurs plateformes intermédiaires sont déployées afin de permettre l'alimentation fluide des hôpitaux en médicaments et dispositifs médicaux. Cet échelon est d'une importance majeure dans le processus de réapprovisionnement. Il doit assurer le maintien du stock en commandant auprès des fournisseurs afin d'être en capacité de répondre en permanence à la demande des hôpitaux et assurer les besoins des unités de soin. Comme évoqué précédemment, ce maillon est un poste de coût considérable dans les dépenses hospitalières annuelles et représente un grand potentiel d'optimisation des immobilisations. La gestion des stocks des pharmacies est cependant très complexe : elle l'est, en outre, par la multiplicité des références gérées, l'hétérogénéité des données logistiques, les volumes et conditionnements divers, les conditions de stockages spécifiques

(frigo, zones sécurisées pour les narcotiques, espaces stériles, ...), la gestion des dates de péremption, etc.

La multiplicité et la complexité des flux d'informations inhérents à la chaîne hospitalière pharmaceutique contribuent encore plus à la complexification de cette tâche. La difficulté de détention de l'information sur la vraie consommation des malades s'ajoute à la donne. Se situant en amont du processus de soin, les pharmacies d'établissement se basent aujourd'hui sur les commandes des unités de soins, lesquelles commandes ne reflètent pas nécessairement la vraie consommation des patients, mais plutôt les règles de réapprovisionnement des stocks avancés au niveau de ces unités. Cette contrainte usuelle, souvent rencontrée dans les réseaux logistiques classiques, peut entraîner un phénomène de distorsion de l'information connu sous l'appellation de l'effet *Bullwhip* et commenté dans plusieurs travaux de la littérature, dont ceux de *Zhang (2004)*, *Gaalman & Disney (2006)*, *Chandra & Grabis (2008)* ou encore *Duc et al. (2008)*. Une simple variation de la demande du client final peut entraîner une grande variation des demandes des échelons intermédiaires, polluant ainsi la vraie consommation et compliquant énormément la gestion et le dimensionnement des stocks.

La fonction logistique, longtemps absente du monde hospitalier, connaît un vrai bouleversement. Une émergence du profil du logisticien a lieu dans les établissements hospitaliers et dans les plateformes pharmaceutiques. Ce dernier devant travailler de plus en plus étroitement avec les pharmaciens et le corps médical. Pour créer les conditions favorables à cette mutation, l'utilisation d'outils de gestion informatisés devient indispensable. Le développement et la mise en place d'outils d'aide à la décision pour la gestion des plateformes pharmaceutiques paraissent inéluctables. L'évolution de la plateforme hospitalière et la complexification croissante des opérations d'entreposage, comprenant notamment le problème du réapprovisionnement et de l'ajustement du stock, œuvrent en faveur de la généralisation de ces outils au domaine hospitalier.

3 Evolution du rôle des plateformes hospitalières

Comme nous l'avons dit, dans la logistique pharmaceutique comme dans la grande distribution et dans l'industrie, les entrepôts et plateformes de distribution ont connu un essor considérable et jouent un rôle central dans le fonctionnement des opérations logistiques. D'une façon générale, le rôle d'une chaîne logistique consiste à mettre à

disposition des clients, les produits appropriés, au bon moment et dans les bonnes conditions : optimisation des coûts des opérations, respect des délais d'expédition auprès des clients et de la qualité des produits fournis. Ce rôle est, en grande partie, assuré par les plateformes de distribution, lesquelles structures permettent d'entretenir des opérations logistiques clés autour de l'expédition, de la labellisation des produits, des préparations pour les chargements ainsi que des opérations internes de picking, de transferts et d'optimisation du stock et des emplacements.

Les entrepôts ont été, longtemps, utilisés comme des centres de stockage intermédiaires permettant de mieux approcher les clients et les régions stratégiques afin de mieux maîtriser les coûts et accroître la réactivité de réponse aux clients et filiales. Cette fonction tampon, épargnant le recours aux intermédiaires et prestataires, permettait de mieux économiser sur les dépenses. De nos jours, l'entrepôt a atteint une grande maturité et son rôle s'étend au-delà de la fonction de stockage pour couvrir, de plus en plus, des opérations créatrices de valeurs ajoutées. *Richards (2011)* revient sur les principales opérations assurées par un entrepôt logistique. Six fonctions principales sont du domaine de l'entrepôt moderne :

- ***Le stockage*** : Fonction standard et commune, les entrepôts assurent, entre autres, le stockage des matières premières et des composants afin d'alimenter les usines, les manufacturiers, les magasins, les hôpitaux et établissements de santé, etc. La logique de regroupement et de mutualisation des moyens dans les pharmacies donne lieu à des immobilisations de grande envergure dont l'optimisation, en terme de valeur économique et en terme de taux de service, est un enjeu crucial.

- ***Centre d'assemblage, de traitements intermédiaires*** : Dans certains cas, les entrepôts sont utilisés pour personnaliser les produits avant l'envoi auprès des clients. Ces opérations, à valeur ajoutée, donnent énormément de flexibilité aux manufacturiers. Ainsi certains entrepôts pharmaceutiques peuvent se charger, entre autres, de préparer des piluliers dans le cadre de la dispensation nominative individuelle.

- ***Centre de transbordement*** : Les entrepôts récupèrent les produits massifiés auprès des fournisseurs et les éclatent en petits lots à destination d'une clientèle variée. Les plateformes pharmaceutiques assurent ainsi la desserte des établissements hospitaliers avec des formats de destination variés. En France, les contraintes de conditionnement imposées

par les laboratoires obligent les plateformes pharmaceutiques à mettre en place des unités de déblisterisation des médicaments pour permettre la dispensation individuelle nominative.

- *Centre de Cross - Docking* : Fonction centrale de la logistique moderne, le Cross - Docking consiste à faire véhiculer des produits à travers la chaîne logistique avec une grande réactivité. Certains entrepôts assurent l'identification des produits, leur consolidation et la distribution auprès des clients dans des délais très courts.

- *Centre de traitement des commandes* : entraînée par l'accroissement de l'e-Commerce, cette activité est de plus en plus supportée par les entrepôts, lesquels sont équipés pour gérer de larges volumes correspondant à des commandes mono-produit. Les plateformes pharmaceutiques s'orientent vers ce mode de fonctionnement en recevant, dans certains cas, directement l'ordonnance du patient (ex. : dispositif de Dispensation Journalière Individuelle et Nominative).

- *Logistique des retours* : Entraînées par des directives environnementales et une législation européenne de plus en plus rigoureuse, les entreprises accordent de nos jours plus de soin au retour des produits. Les entrepôts sont aménagés afin de supporter cette activité, en plein essor. Les pharmacies d'établissements reçoivent et gèrent des flux de retours de médicaments ou de produits périmés provenant des unités de soin.

Au regard de la mutation continue du métier de la logistique, l'entrepôt n'est alors plus considéré comme un simple centre d'immobilisation, dont la vocation exclusive est le stockage des produits, mais comme un système logistique plus complexe. Diverses opérations stratégiques font désormais partie du périmètre de l'entrepôt. Cette évolution n'est pas sans compliquer la donne aux gestionnaires, lesquels sont amenés à intégrer des opérations de plus en plus spécialisées, créatrices de forte valeur ajoutée, dans un environnement très concurrentiel et à gérer une multiplicité de flux physiques et informationnels. Dans cet environnement complexe et très contraint, les systèmes informatiques décisionnels et les systèmes de gestion ont joué un rôle crucial et constituent un élément essentiel de la compétitivité des entreprises ayant à gérer des opérations de logistique d'entrepôt. Une grande maturité a été atteinte dans ce domaine, due, entre autre, à une prise de conscience anticipée de la part des industriels et des groupes de distributions, de l'importance de tels outils et des avantages concurrentiels qu'ils peuvent apporter.

4 Evolution des outils informatiques d'aide à la décision pour la gestion des plateformes logistiques

L'utilisation des technologies de l'information et des systèmes informatiques constituent de nos jours un levier central dans l'optimisation des opérations logistiques au sein des plateformes de distribution. Les logisticiens et gestionnaires d'entrepôts se sont souvent appuyés sur diverses offres logicielles, en évolution et mutation continues, allant d'outils simples et répandus, comme les classeurs *Excel* ou des programmes faits maisons, à des outils très spécialisés comme les **ERP**, les **APS** ou encore les **WMS**.

Les établissements de santé, comme les entreprises, qui disposent de plateformes de distributions, s'appuient sur une multitude de logiciels pour couvrir les fonctions et les opérations de plus en plus diverses et propres aux entrepôts. Cette cohabitation logicielle, consistant à adopter plusieurs solutions informatiques à la fois, est souvent rencontrée en pratique. La sphère de fonctionnalités couvertes par les différentes offres logicielles ne permettant pas de répondre à la totalité des exigences des clients, ces derniers recourent assez couramment à l'interfaçage, présenté dans la figure 2, de différents outils pour couvrir la totalité de leur besoin.

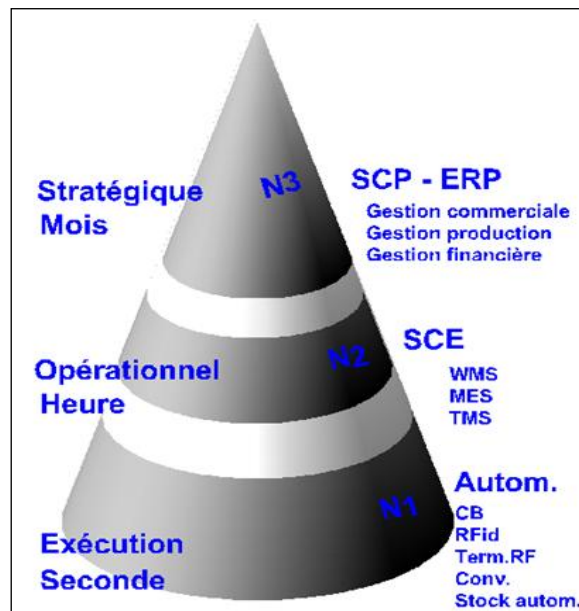


Figure 2: Différentes catégories des systèmes de gestion de la logistique

Dans le paysage, très riche, des outils de gestion informatisés, les systèmes de gestion d'entrepôt (connus sous l'acronyme anglo-saxon **WMS**) occupent une place centrale dans

l'optimisation du fonctionnement des entrepôts et dans l'amélioration de leur performance. Ces progiciels, largement adoptés par les industriels et la grande distribution, sont dédiés aux opérations de logistique interne, ils sont relativement récents dans leur mise en œuvre auprès des établissements de santé. La mise en place de ces outils peut être complexe pour certains établissements. Ce sont les grands centres hospitaliers qui se sont dotés en premiers de solutions de gestion d'entrepôts. Parmi ces établissements nous citons les Hospices Civils de Lyon, le CHU de Saint-Etienne, le CHU d'Avignon ou encore le CHU de Montpellier. Les plateformes desservant certains de ces établissements sont le résultat du regroupement de plusieurs pharmacies, fonctionnant auparavant indépendamment les uns des autres. Ces derniers ont présenté un intérêt pour des fonctions de traçabilité, pour l'optimisation des préparations, ainsi que le contrôle de machines de distribution automatique de médicaments.

L'utilisation des WMS devient, de nos jours, primordiale dans la gestion des plateformes de stockage en vue d'un fonctionnement optimisé. Cette gamme de logiciels permet une gestion de tous les flux physiques et informationnels au sein d'un entrepôt, depuis l'entrée des produits aux quais jusqu'à leur préparation pour envoi aux clients, en passant par toutes les opérations internes. Ceci s'étend généralement aux fonctions principales suivantes :

- Gestion de la réception des produits : comprenant la vérification des quantités reçues par rapport aux quantités prévues ainsi que les contrôles de la qualité des produits.
- Gestion d'entreposage : Optimisation de l'affectation des produits aux zones de stockage disponibles dans l'entrepôt.
- Préparation des commandes des clients : Définition des produits à choisir, opérations d'emballages, et fournitures des documents nécessaires comprenant les étiquettes adéquates.
- Fonctions de traçabilité des produits, fonctions d'inventaires de stock et fonctions de gestion et d'approvisionnement du stock.
- Fonctionnalités de *reporting* et de contrôle des performances de la plateforme.

Cette panoplie de fonctions, appuyées de plus en plus par des modules et fonctionnalités complémentaires introduits par les éditeurs, permet d'apporter plus d'efficacité et d'efficacité dans le fonctionnement quotidien des entrepôts. La traçabilité des produits,

assurée par ces systèmes, permet la localisation instantanée des produits et donc une réponse plus efficace aux clients. L'optimisation de l'affectation des produits aux emplacements de stock contribue aussi à l'optimisation des coûts et à l'amélioration des taux de service auprès des clients.

Les systèmes de gestion de l'entrepôt assurent ainsi la couverture d'une large sphère de fonctions propres au domaine de l'entreposage simplifiant le fonctionnement quotidien des plateformes. La cohabitation logicielle évoquée précédemment reste cependant souvent nécessaire et n'est pas épargnée par la simple installation d'un WMS. Selon les clients, et selon la complexité des opérations des entrepôts dont ils disposent, l'adoption d'autres outils de gestion ou de certains modules complémentaires s'avère souvent nécessaire. Les WMS sont en effet conçus pour fonctionner en toute autonomie ou dans un cadre plus global, intégrant des logiciels de gestion de tout genre.

Dans l'industrie, l'interfaçage avec les ERP, les TMS ou encore les logiciels de pilotage des armoires de stockage est chose courante dans les organisations industrielles et les entrepôts de distribution. Cette cohabitation, souvent nécessaire, n'est pas sans engendrer un coût financier et une complexité organisationnelle et informatique.

Du côté de l'hôpital, un retard considérable reste à rattraper sur les outils d'aide à la décision et notamment les systèmes de gestion d'entrepôts. La complexification de la logistique hospitalière et les contraintes budgétaires ont œuvré à rétrécir ce gap. La part importante réservée aux projets informatiques dans le plan hôpital 2012, 15% du montant des enveloppes régionales, témoigne en effet d'une prise de conscience manifeste dans ce domaine. La mise en place de *WMS* ne s'est cependant pas généralisée et seuls les grands établissements disposent de ces applicatifs. Les niveaux de performance des plateformes pharmaceutiques ne semblent pas atteindre celui des plateformes industrielles *Fustier (2010)*.

5 Extension de la sphère d'application des WMS

5.1 Évolution fonctionnelle des WMS

L'extension du périmètre des fonctionnalités couvertes par les WMS est au cœur de l'évolution récemment entreprise par les éditeurs de ces systèmes. Les logiciels de gestion d'entrepôt étant par nature génériques et destinés à divers types de clients, un effort accru est

orienté dans le sens de l'adaptation de ces outils pour répondre à différents secteurs d'activités et couvrir de multiples besoins.

Avec l'évolution de la logistique hospitalière, des WMS adaptés au monde de la santé ont vu le jour. Parmi les éditeurs proposant des solutions adaptées au monde hospitalier, nous citons la société a-Sis, Aldata Solution ou encore Manhattan Associates. Avec le progiciel Gildas Hospilog, KLS occupe une place importante sur le marché français des WMS hospitaliers.

Cette tendance témoigne d'une certaine maturité atteinte par les systèmes de gestion d'entrepôts, lesquels sont de plus en plus dotés de modules complémentaires étendant la sphère fonctionnelle au-delà des opérations classiques d'exécution et concurrençant certaines offres logicielles, longtemps considérées du domaine des ERP, des TMS ou même des APS. La complexité inhérente aux interfaçages de multiples programmes ainsi que le coût financier induit jouent également en faveur de cette tendance.

Dans le cadre de la mutation continue des WMS, des fonctions allant dans le sens de l'organisation des transports et des tournées de véhicules voient de plus en plus le jour. Située en aval par rapport à l'entrepôt, cette activité, longtemps monopolisée par les logiciels de gestion de transport (TMS), est désormais couverte par certains WMS. Des modules permettant la synchronisation de l'activité de l'entrepôt avec les tournées, l'organisation des expéditions ou encore la prise de rendez-vous deviennent indissociables de l'outil standard.

L'intégration de modules d'ordonnancement des commandes des clients est également à l'ordre du jour. Devant s'adapter aux changements de processus dans l'entrepôt et à l'introduction de nouveaux modes de préparation, il devient primordial de doter le WMS par un moteur d'ordonnancement. Des fonctions de *picking*, de *packing*, de ramasse multi clients ou de ramasses massifiées sont incluses. D'autres fonctions de mise en œuvre de *clusters* pour dispatcher des commandes, les assembler ou les associer selon certains critères sont également supportées.

L'autre fonction clé, qui se situe en amont de l'entrepôt, et qui intéresse également les éditeurs de WMS, est la gestion des approvisionnements. Connue sous l'appellation de préconisation des commandes, elle concerne les décisions sur les stocks détenus, les fréquences d'approvisionnement et les quantités à commander auprès des fournisseurs, afin d'optimiser le coût d'immobilisation tout en assurant un taux de service élevé. Nous nous

intéressons particulièrement, dans ce travail de thèse, à cette fonctionnalité. Nous revenons plus en détail sur la fonction préconisation des commandes.

5.2 La problématique de la préconisation des commandes : Module de gestion des approvisionnements

La préconisation des commandes est l'une des opérations auxquelles sont confrontés tous les entrepôts, qu'ils soient de nature industrielle ou qu'ils s'agissent de plateformes hospitalières. Cette tâche consiste à approvisionner les produits stockés par l'entrepôt auprès des fournisseurs afin d'être en capacité de répondre à la demande des clients. Cette fonction est considérée comme étant l'une des plus complexes parmi les opérations de gestion d'entrepôts. Si elle est bien gérée, cette fonction constitue en même temps, un levier considérable d'amélioration de performances et de gains financiers.

La difficulté consiste à ajuster au mieux le niveau de stock des différents articles en définissant le bon rythme d'approvisionnement, et les bonnes quantités de commandes, permettant d'opérer le bon compromis entre l'immobilisation dans l'entrepôt et le niveau de service client. Ceci se fait, souvent, dans un environnement contraint et dynamique. Des contraintes inhérentes aux produits stockés, aux rythmes de la consommation, aux fournisseurs et délais d'approvisionnements, aux capacités de stockage disponibles ainsi qu'au niveau d'exigences des clients peuvent être considérées.

Divers systèmes informatiques ont été développés pour traiter des préconisations de commandes et optimiser les stocks. Jusqu'à très récemment, ces modules faisaient généralement partie des systèmes ERP ou des modules des APS. Ces systèmes se basent généralement sur des méthodes de calculs confirmées, issues de la littérature logistique, permettant de calculer avec précision les quantités et fréquences de commandes. Des méthodes d'estimation de la demande et des règles de gestion de stock sont implémentées et permettent de réagir à l'évolution de consommation afin d'alerter sur les décisions judicieuses à prendre.

Jusqu'à présent, les fonctions d'optimisation du stock et d'anticipation de la consommation ne faisaient pas partie des fonctions couvertes par les WMS. Alors que les logiciels de gestion d'entrepôt assurent une traçabilité intégrale des produits stockés et des flux de consommation historiques, ces données sont rarement exploitées, par ce même logiciel, à des fins d'optimisation des réapprovisionnements.

Ce n'est qu'au cours de ces dernières années que les éditeurs de WMS commencent à évoluer par rapport à cette question en abordant la problématique de gestion de stock et des prévisions de consommation des articles. Cette évolution s'inscrit dans le cadre de l'élargissement de la sphère d'application couverte par les WMS pour intégrer des modules décisionnels.

Nous nous intéressons dans le cadre de ce travail de thèse à cette problématique. Nous traitons du cas d'un module de gestion de stock intégré dans un système de gestion d'entrepôt à destination de plateformes pharmaceutiques.

6 Problématique abordée : Développement d'un module de préconisation de commandes pour les plateformes pharmaceutiques

Le travail entrepris dans cette thèse se situe dans le contexte de l'amélioration d'un module de préconisation de commandes implémenté dans un système de gestion d'entrepôts pharmaceutique *Gildas Hospilog* développé par l'éditeur de logiciel *KLS Logisitic systems* pour répondre aux besoins spécifiques du domaine hospitalier. L'émergence de la logistique hospitalière et son évolution continue autour du circuit du médicament ont poussé la mise en place d'outils de gestion. Le pharmacien étant au centre de ce processus et garant du bon fonctionnement au sein de la pharmacie, il paraît de plus en plus indispensable de le doter d'outils informatiques pour l'aider dans cette fonction. Il s'agit, en effet, de lui permettre de mieux se focaliser sur les activités purement pharmaceutiques en automatisant au maximum les tâches logistiques fortement répétitives et sans grande valeur ajoutée au regard du métier de pharmacien.

KLS Logisitic systems est une société française, spécialisée dans le développement et la mise en place de logiciels de gestion d'entrepôts. Ayant une activité de plus de vingt ans, et ayant travaillé pour divers secteurs d'activités, l'entreprise couvre les besoins des entrepôts et, enrichit ainsi, au fur et à mesure son système de gestion. Prenant conscience de la mutation logistique du domaine hospitalier et des lacunes considérables en outils décisionnels, le système de gestion d'entrepôts hospitaliers *Gildas Hospilog* a été mis en place. Plusieurs plateformes hospitalières françaises ont été équipées par ce système, comprenant entre autres

structures, les Hospices Civils de Lyon, le CHU de Saint-Etienne ou encore le CHU d'Amiens.

Gildas Hospilog est un système de gestion d'entrepôt supportant l'ensemble des fonctionnalités classiques devant être proposées par un WMS standard : traçabilité des produits, gestion de la réception, gestion des zones, gestion des flux, etc. De par son implémentation dans diverses plateformes pharmaceutiques d'envergure, *KLS* est en continu confrontée à l'évolution de leurs besoins en logistique et à la nécessité de faire évoluer par conséquent son offre logicielle.

Dans ce contexte évolutif, des modules complémentaires ont dû être développés au fil du temps, en réponse à des besoins spécifiques du domaine hospitalier. *Gildas Hospilog* a atteint une certaine maturité fonctionnelle le situant au rang des principaux WMS hospitaliers du marché français et étendant sa sphère fonctionnelle au-delà du périmètre d'exécution, propres aux WMS classiques, pour couvrir d'autres opérations logistiques du domaine de l'entrepôt et du métier de l'hôpital.

La gestion de stock et des approvisionnements auprès des fournisseurs est l'une des fonctions ayant été anticipée par *KLS* et intégrée dans le logiciel *Gildas Hospilog*. Un module de préconisation de commandes a été mis en place pour traiter cette question et aider les pharmaciens à gérer leurs commandes de produits auprès des fournisseurs et laboratoires. Ce besoin en outils de réapprovisionnement a été exprimé initialement par les gestionnaires de plateformes pharmaceutiques qui ne disposaient pas des outils de gestion déployés dans l'industrie. L'intégration d'un module de préconisation de commandes dans le WMS s'avérait de plus en plus indispensable. Le besoin s'est exprimé en premier par les Hospices Civils de Lyon en début des années deux milles. *KLS* a finalement mis en place ce module en s'appuyant sur certains modèles classiques de la théorie de la gestion des stocks. Conscients de l'importance des approvisionnements dans l'hôpital, de la complexité des stocks pharmaceutiques et des apports considérables en performance qu'impliquerait l'optimisation de la pharmacie hospitalière, les décideurs de l'entreprise ont orienté un effort accru dans le sens de l'entretien et de l'amélioration continue de ce module.

Notre travail de thèse, financé par *KLS*, vient en appui à cette action afin de contribuer à accroître les performances et les fonctionnalités autour de la préconisation des commandes.

L'idée consiste à doter le **WMS** d'un module plus évolué, mieux à même d'intégrer la complexité inhérente à la plateforme hospitalière. Cette action s'inscrit également dans une volonté d'innovation visant à rétrécir le gap entre l'étendue des fonctionnalités proposées par les **ERP** et les **APS** et celles nouvellement supportées par les **WMS**.

La problématique de la préconisation des commandes peut couvrir un large spectre de réflexions. Nous abordons, dans ce travail, certaines de ses facettes et explorons quelques leviers d'amélioration. Une longue réflexion autour de ce sujet a été menée par les décideurs et ingénieurs de *KLS* le long du développement et de l'amélioration du module. Ce travail continu, enrichi par les retours des clients et un contact permanent avec les éditeurs de logiciels, a permis de mettre en exergue un certain nombre de pistes d'améliorations pertinentes. Ceci nous mène à répondre à certaines questions portant sur la définition, à partir de la littérature scientifique, d'un certain nombre d'outils et de modèles pertinents autour de l'optimisation des préconisations de commandes, d'évaluer leurs apports en performance et l'étendue de leurs adaptations pour le cas d'un système de gestion informatisé.

Nous avons, en début de cette thèse, procédé à une étude détaillée du module existant, ceci nous a permis de nous faire notre propre critique du système et de définir ainsi les principaux verrous de progrès qui nous paraissent prioritaires et atteignables. En confrontant notre synthèse aux visions et idées entretenues à *KLS*, nous avons posé un certain nombre de chantiers de recherches et de travail. Les principales questions que nous abordons peuvent être déployées en trois principaux axes de recherche :

1. La classification des articles du stock : La classification des articles du stock est une étape clé dans la gestion des approvisionnements. Elle permet de répartir les articles en différentes classes afin de définir des politiques de gestion adéquates. La majorité des systèmes logistiques permettent la classification des articles de stock sur la base du principe PARETO en utilisant, très souvent le critère du montant financier annuel. Cette méthode a révélé son insuffisance et a été longtemps critiquée pour son intégration exclusive de l'aspect financier, ce qui ne reflète pas nécessairement la réalité *Flores et al. (1992)*. Plusieurs travaux récents ont traité de la question de la classification des articles du stock et de la façon de l'améliorer

par intégration de l'approche multicritères, parmi ces travaux nous citons ceux de *Ramanathan (2006)*, *Zhou & Fan (2007)* ou encore *Cakir & Canbolat (2008)*.

Nous cherchons dans ce travail à enrichir le système par introduction de fonctionnalités de classification avancées. Nous recensons les travaux de la littérature qui ont traité de cette question et analysons les éléments pouvant constituer des points de progrès à notre module. Des questions autour des critères pertinents à considérer et des algorithmes de calculs à utiliser seront posées.

2. La prévision de consommation : La prévision de la consommation de l'entrepôt est une question centrale dans la gestion des réapprovisionnements. Le calcul des quantités de commandes à lancer auprès du fournisseur se fait sur la base d'une estimation de la consommation de l'entrepôt. Les systèmes de préconisation de commandes utilisent l'historique des consommations, qui est archivé dans la base de données, pour opérer ces calculs. La fiabilité du système de prévision est un facteur clé dans l'amélioration du système de gestion de stock. Nous traitons dans ce travail de la question des prévisions de consommation. Nous abordons un ensemble d'éléments autour de la fonction prévision permettant d'apporter de la valeur ajoutée à l'utilisateur du module. Nous nous posons la question de l'amélioration de la précision des calculs. Quelles méthodes de prévisions, se prêtant facilement au cadre d'un système expert, permettront d'améliorer l'estimation des consommations. Quelles informations autour de la prévision, quels indicateurs de performances, d'alertes, d'identification de tendances serait-il intéressant de remonter et constitueraient une information synthétique et pertinente pour aider dans la prise de décision.

3. Règles de gestion de stock : Les politiques de gestion de stock permettent de répondre aux deux questions classiques : Quand commander ? et combien commander ? dans le but d'optimiser le stock tout en assurant un bon niveau de service auprès du client.

Les systèmes de préconisation utilisent des politiques de gestion issues de la littérature pour opérer le calcul des quantités de commandes. Des méthodes classiques sur seuil ou de reapprovisionnement périodique sont souvent implémentées et permettent d'alerter sur la nécessité de lancer les approvisionnements. Le paramétrage des méthodes se fait sur la base des résultats issus du système de prévision. Nous traitons dans ce travail de la question des

politiques de réapprovisionnement et analysons l'apport du système de prévision sur la performance de stock.

7 Organisation de la thèse

Cette thèse est organisée en six chapitres. Dans le chapitre 2, suivant cette introduction, nous procédons à une présentation générale du système de préconisation de commandes proposé. Nous partons du module de préconisation existant et décortiquons ses différents composants. Cette analyse nous permet de mettre en vue les lacunes et insuffisances du système et de justifier ainsi les propositions que nous avons faites. Nous exposons, ensuite, les modules additionnels que nous avons développés ainsi que les fonctions couvertes. L'ensemble des modules développés seront abordés en détail dans la suite des chapitres.

Dans le chapitre 3, nous nous focalisons sur le module de classification des articles de stock. Il s'agit du premier module développé pour compléter le système de préconisation existant avec une fonction de classification permettant l'intégration de multiples critères lors du calcul des classes. Nous revenons sur certains travaux de la littérature ayant traité de la problématique de classification. Ceci nous permet de définir un certain nombre de méthodes multicritères pertinentes et de choisir celles qui nous paraissent les plus appropriées à l'implémentation dans un système informatique. Nous construisons ainsi un concept couvrant un certain champ de fonctionnalités de classification et de consultation autour de ces méthodes. Nous détaillons les différentes fonctionnalités développées.

Dans le chapitre 4, nous présentons le module de prévision des consommations. L'objectif de ce module est de doter le WMS d'un système de prévision performant pour calculer des prévisions plus justes et permettre de remonter des informations et alertes pertinentes et utiles aux utilisateurs. Ce module doit donc, à la fois permettre d'alimenter les calculs des préconisations par des estimations de consommations plus fiables, et servir de tableau de bord afin de communiquer aux utilisateurs des informations synthétiques et utiles sur l'évolution de la consommation des articles du stock. Nous nous appuyons sur les travaux de la littérature concernant les prévisions afin de bien définir les outils et modèles de calculs robustes, utiles et se prêtant facilement au cadre des systèmes experts. Des outils statistiques d'identification des tendances, de nettoyage des données et d'extrapolation des séries

temporelles seront ainsi abordés et un choix de techniques sera fait. Nous présentons en fin de chapitre le concept construit et développé.

Dans le chapitre 5, nous traitons d'un cas d'application et évaluons la performance du système développé en utilisant des données provenant de plateformes pharmaceutiques utilisant l'ancien système de préconisation de commandes implémenté dans le *WMS Gildas Hospilog*. Nous procédons, dans un premier temps, à une évaluation de la performance du système de prévision indépendamment des considérations sur le stock. Nous comparons les résultats produits par notre système de prévision à ceux du système de prévision initialement utilisé. Dans une deuxième phase nous évaluons l'impact sur le stock. Nous comparons les résultats obtenus en terme d'immobilisation et de rupture, selon le module de préconisation initial avec celui complété par les fonctions que nous avons développées.

Nous clôturons cette thèse par une conclusion générale reprenant les principaux résultats et apports de ce travail et suggérons des pistes d'améliorations ainsi que des perspectives d'évolutions.

Chapitre II. Composantes du module de préconisation proposé

Nous présentons dans ce chapitre la structure générale et les fonctions principales du module de préconisation de commandes que nous proposons. Les différents composants constituant le module seront développés et détaillés dans la suite des chapitres de la thèse. Nous adoptons la démarche suivante : nous partons d'un descriptif détaillé du module de préconisation actuellement implémenté, qui nous permettra de mettre à plat le système existant et de mieux identifier ses lacunes. Ce travail d'analyse permet de dégager des pistes d'amélioration et des axes de recherche potentiellement intéressants, de les confronter aux réflexions initialement entretenues par les décideurs et les ingénieurs de l'entreprise et d'en ressortir des propositions de fonctionnalités nouvelles pour compléter le module. Cette investigation a donné lieu à de nouvelles fonctionnalités et a permis de définir la nouvelle structure du module de préconisation de commandes.

1 Module de préconisation de commandes existant

Le module de préconisation de commandes existant permet d'assurer l'approvisionnement du stock de l'entrepôt auprès des fournisseurs pour permettre de répondre aux demandes des clients. Ce module a été initialement développé pour répondre aux besoins spécifiques du domaine hospitalier. Ceci s'est fait dans la perspective d'étendre les fonctions du WMS au-delà des opérations standards de traçabilité afin d'être en phase avec l'évolution continue et l'extension accrue de la sphère d'application de ces logiciels. La fonction d'approvisionnement est d'autant plus importante que l'ensemble des entrepôts pharmaceutiques sont confrontés à la problématique d'optimisation de leur stock.

Le module de préconisation existant s'appuie sur deux composants : Un système de prévision de la consommation de l'entrepôt, utilisant l'historique des consommations enregistré dans une base de données, et un système de gestion de stock s'appuyant sur deux politiques de gestion, une politique de reapprovisionnement périodique et une politique à point

de commande. Le système de prévision réalise un calcul périodique permettant de mettre à jour, mensuellement, une estimation de la consommation de l'entrepôt en ses différents produits stockés. L'historique de consommation est une information standard enregistrée par le WMS. Tous les produits étant tracés lors de l'opération de préparation pour l'envoi aux clients, les quantités envoyées sont donc enregistrées au niveau des tables historiques de la base de données.

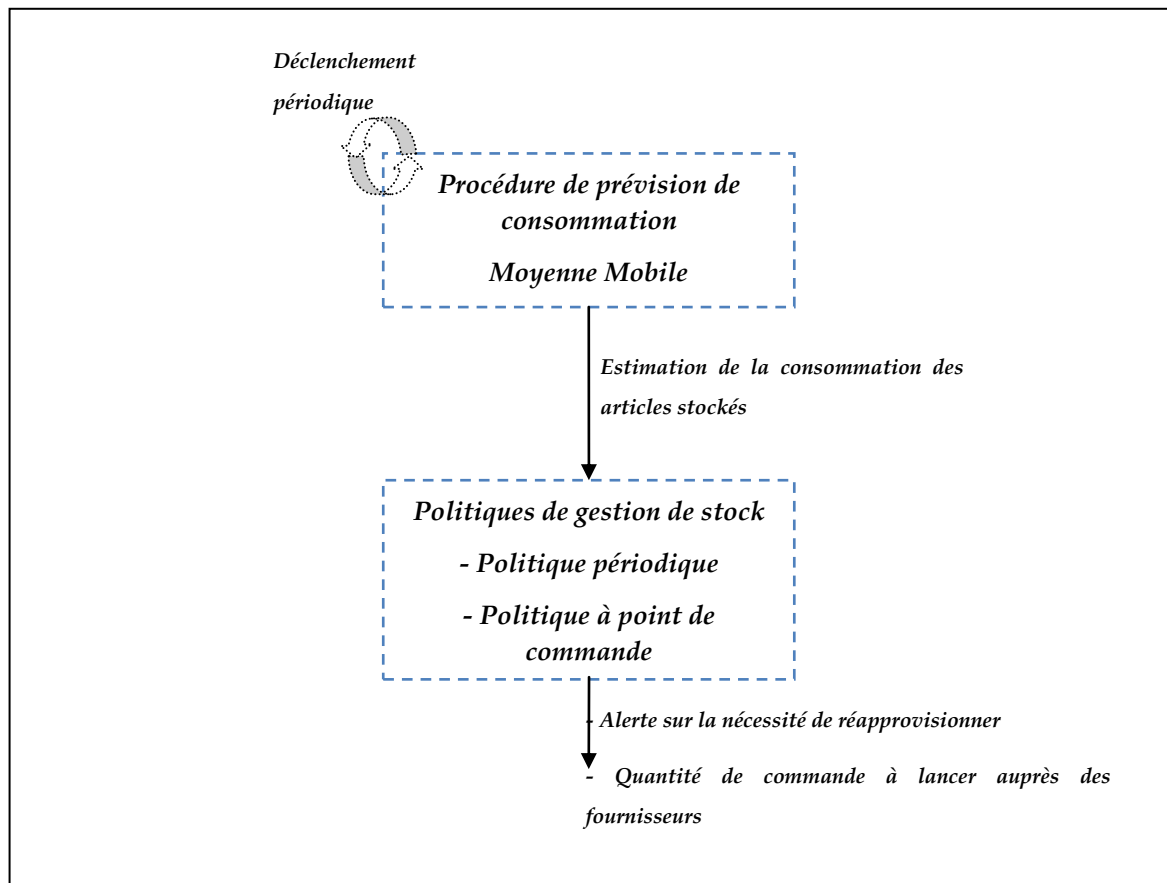


Figure 3: Composants du module de préconisation existant

L'estimation de la prévision calculée est utilisée par le système pour mettre à jour certains des paramètres des deux politiques de gestion (notamment le point de commande et le niveau de recomplètement). D'autres paramètres sont introduits manuellement par l'utilisateur lors de l'installation et du paramétrage du WMS. Sur la base de ce paramétrage et des estimations de consommation renvoyées, le rythme d'approvisionnement ainsi que les quantités de commandes seront définis. Nous revenons dans ce qui suit sur le détail des deux composants constituant l'actuel système de préconisation.

1.1 Système de prévision utilisé

Le module de préconisation est doté d'une procédure de calcul simple permettant d'estimer la consommation des articles stockés : une moyenne mobile sur l'historique annuel de consommation est appliquée. L'historique de consommation de l'entrepôt est une information clé stockée en continu dans la base de données par le WMS, une traçabilité complète de tous les mouvements des produits étant assurée en permanence par ce système. La procédure de calcul se déclenche, ainsi, mensuellement afin de mettre à jour l'estimation. À chaque calcul, le système opère une moyenne sur les douze derniers mois de l'historique. Cette moyenne glissante s'applique systématiquement sur tous les articles de stock indépendamment de toutes considérations sur les caractéristiques de l'article ou sur la typologie de sa consommation.

$$Prévision_{Article\ i} = \frac{\text{Mois courant} - \text{Mois courant} - 12}{12} Consommation_{Article\ i}$$

L'estimation de la consommation est utilisée par le système pour alimenter le deuxième composant et calculer les paramètres des politiques de gestion de stock. L'information est également remontée au niveau de la fiche de paramétrage des articles. L'utilisateur peut ainsi consulter cette information et la modifier à tout moment. Mis à part cette donnée, le système n'entretient aucune information autour de la fonction prévision. Il n'offre aucune option de consultation avancée, aucun affichage de graphiques d'évolution ou d'analyse autour des mesures de performances susceptibles de l'aider dans le suivi de l'évolution des produits de l'entrepôt pharmaceutique. La remontée d'informations pertinentes et ciblées, dans ce contexte, peut s'avérer d'une grande utilité afin d'aider à mieux appréhender le comportement des médicaments et simplifier la prise de décision.

La fonction prévisionnelle se limite dans le cas existant à l'alimentation des préconisations de commandes.

1.2 Règles de gestion

Le module de préconisation de commandes utilise deux politiques de gestion de stock classiques : La politique de reapprovisionnement périodique et la politique sur seuil de

commande. Ces deux techniques sont affectées à l'ensemble des articles du stock et vont déclencher une alerte au niveau de l'interface sur la nécessité de réapprovisionner le stock tout en préconisant une quantité à commander auprès du fournisseur.

1.2.1 Fonctionnement et paramétrage de la politique périodique implémentée

La politique périodique à niveau de reemplètement

Au début de chaque période de longueur T , si la position du stock descend en dessous d'une valeur donnée, appelée niveau de reemplètement et notée S , un ordre de réapprovisionnement est lancé de manière à ramener la position du stock à S . La commande est réceptionnée à l'issue du délai de réapprovisionnement L . Nous présentons dans la figure 4 l'évolution du stock suivant cette politique.

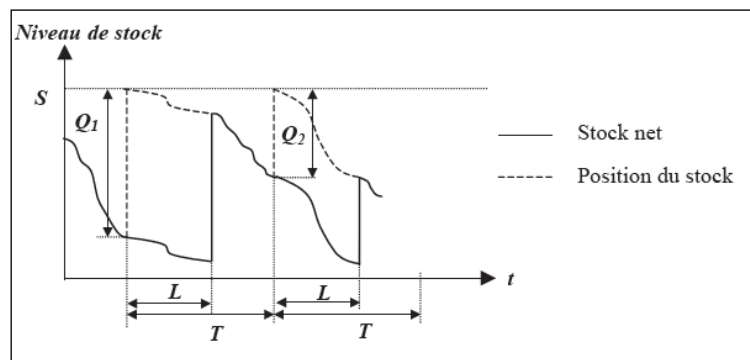


Figure 4: Évolution du niveau de stock
dans le cas de la politique périodique à niveau de reemplètement (T, S) Babai (2005)

Calcul des paramètres dans le cas de Gildas WM

Dans le cas du module de préconisation actuel, les paramètres de la politique sont pris en compte de la façon suivante :

- La périodicité de commande T n'est pas préconisée par le système, elle doit être introduite directement par l'utilisateur de l'application. Ainsi, lors de l'installation du WMS, l'utilisateur doit procéder, lors du paramétrage, à l'affectation d'une périodicité à l'ensemble des articles du stock. C'est au gestionnaire de définir, selon sa connaissance et sa propre estimation, une périodicité convenable pour chacun des articles. Aucune considération d'optimisation du stock n'est prise en compte par le système. Au regard du nombre très important de références à gérer dans une pharmacie, cette gestion manuelle et individuelle est clairement

un frein à toute mise en œuvre d'une quelconque optimisation des politiques de gestion des stocks.

- Le niveau de reapprovisionnement S est calculé par le système de façon à permettre de couvrir la période de temps comprenant la périodicité de commande T et le temps de réapprovisionnement L auprès du fournisseur. Ce niveau est augmenté par un stock de sécurité permettant de pallier les aléas. L'utilisateur doit également paramétrer ce niveau de sécurité en introduisant pour chaque article un nombre de jours de sécurité NbJ que le système traduit en quantité de sécurité. En se basant sur l'estimation de consommation de l'article, renvoyée par le premier composant, le niveau de reapprovisionnement est donc calculé de la façon suivante :

$$S = Conso * T * Coeff + Conso * L + Conso * NbJ$$

$Conso$ étant l'estimation de la consommation calculée par la moyenne mobile et $Coeff$ un coefficient de introduit par l'utilisateur.

1.2.2 Fonctionnement et paramétrage de la politique sur seuil implémentée

La politique sur seuil avec niveau de reapprovisionnement

Dans cette politique à suivi continu, dès que la position du stock descend en dessous du seuil de commande s , on reapprovisionne la position du stock jusqu'à un niveau de reapprovisionnement S . La commande à lancer est de taille variable. L'évolution du stock suivant cette politique est donnée par la figure 5.

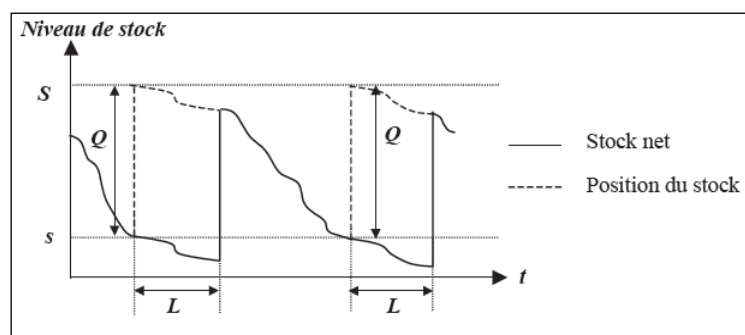


Figure 5: Evolution du niveau de stock dans le cas de la politique sur seuil à niveau de reapprovisionnement (s , S)
Babai (2005)

Calcul des paramètres dans le cas de Gildas WM

Dans le cas du module de préconisation actuel, les paramètres de la politique sont pris en compte de la façon suivante :

- Le seuil de commande s est calculé par le système de façon à permettre de couvrir le délai de réapprovisionnement auprès du fournisseur. Ce niveau est augmenté par le même stock de sécurité défini précédemment. Le seuil de commande est donc calculé de la façon suivante :

$$s = Conso * L + Conso * NbJ$$

- Le niveau de recombplètement S est calculé de la même façon que précédemment.

1.3 Architecture informatique de support

Nous revenons dans ce qui suit sur l'architecture informatique détaillée supportant le module de préconisation de commandes actuel. Cette analyse nous permettra de définir les composants informatiques complémentaires à développer ainsi que ceux, déjà existants, et qu'il convient d'utiliser et enrichir, afin d'assurer les nouvelles fonctionnalités que nous allons proposer.

La fonction préconisation de commandes fait partie des différents modules proposés par le WMS Gildas WM. Son fonctionnement est assuré par trois principaux composants : Une application PC supportant les interfaces graphiques du module, une DLL de calculs supportant des fonctions métiers nécessaires et une base de données Oracle permettant le stockage d'informations utiles au fonctionnement ainsi que l'exécution de certains traitements en complément à la DLL métier.

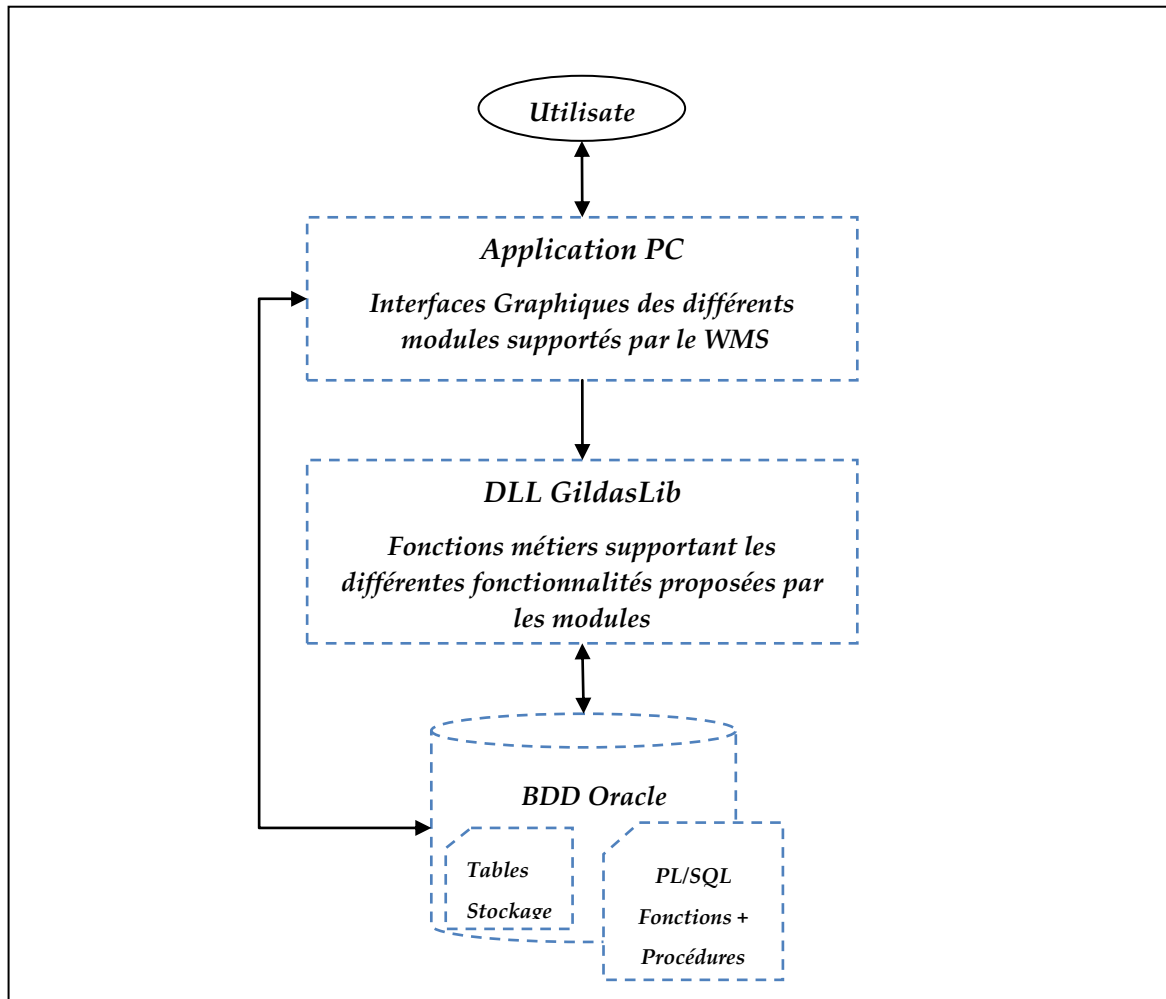


Figure 6 : Architecture informatique en support au module de préconisation de commandes existant

Application PC : L'application PC supporte la totalité des interfaces graphiques propres aux modules du WMS. L'utilisateur du système interagit avec l'applicatif via cette application. L'application PC est développée au moyen de l'environnement de développement **Borland C++ Builder**. Le module de préconisation est accessible via cette interface. Il est construit au moyen de composants Borland permettant à l'utilisateur d'opérer des consultations à travers des grilles, des panels, des labels ou d'exécuter des actions au moyen de boutons.

DLL métier : Il s'agit d'un fichier de bibliothèque logicielle contenant un ensemble de fonctions nécessaires au fonctionnement des différents modules du WMS et notamment le module de préconisation des commandes. Ces fonctions sont appelées à partir de l'application PC sur demande de l'utilisateur afin d'exécuter certains traitements.

Base de données : La base de données est un élément central dans le fonctionnement du WMS. Elle assure le stockage de toutes les informations utiles et nécessaires à mettre à disposition de l'utilisateur. Des tables spécifiques au module de préconisation sont utilisées. Au-delà de la fonction de stockage assurée par ce composant, la base de données utilisée assure aussi certains traitements en compléments aux fonctions de la DLL métier. Ceci se fait au moyen du langage structuré PL/SQL propre au système de gestion Oracle. Des fonctions et procédures de calculs sont stockées et se déclenchent pour lancer certains traitements. Le calcul de l'estimation prévisionnelle mensuelle est assuré par une procédure stockée. Un séquenceur (*un job oracle*) est programmé pour solliciter mensuellement cette procédure et mettre à jour les estimations.

2 Limites du système et pistes d'amélioration retenues

L'analyse effectuée ci-dessus nous permet de mettre en avant un certain nombre de limites qu'il convient d'améliorer et de compléter. Le module de préconisation actuel n'est doté d'aucune procédure de classification des articles du stock. La classification est une fonction centrale en gestion de stock et il est d'une importance majeure d'arriver à identifier les articles pilotes parmi l'ensemble des produits stockés. Les produits pilotes, appelées aussi articles phares, constituent en général autour de 20% des produits stockés et représentent une majeure partie du flux financier total (autour de 80% du montant financier annuel consommé). La classification se fait, dans la pratique courante, par des considérations exclusivement financières. Certains travaux de la littérature, notamment ceux de *Flores et al. (1992)* et *Cakir & Canbolat (2008)*, ont intégré, dans leur recherche de cette catégorie d'articles, plusieurs critères logistiques, en plus de l'attribut financier.

Cette approche multicritères peut s'avérer intéressante, notamment dans le cas d'une grande hétérogénéité des produits, ce qui est le cas des entrepôts pharmaceutiques, qui présentent une grande multiplicité de références dont les différences ne portent pas uniquement sur les coûts financiers. Des considérations plus pertinentes sur le choix des politiques de gestion ou des taux de services théoriques pourraient être définies sur la base des différentes classes. La fonction classification est, dans la plupart du temps, proposée dans les systèmes informatiques de gestion de stock.

Doter le système d'une fonctionnalité de classification nous paraît important, d'autant plus que cette problématique est de plus en plus abordée dans la littérature et précisément dans le

cadre d'implémentation dans les systèmes logistiques décisionnels. Plusieurs travaux, parmi lesquels ceux de *Ramanathan (2006)*, *Zhou & Fan (2007)*, ou encore *Yu (2010)*, ont traité de cette problématique dans le but d'aider les gestionnaires à mieux identifier les articles pilotes. Dans ce cadre, plusieurs techniques mathématiques, se prêtant facilement à une implémentation informatique, ont été testées et continuent d'être améliorées.

Le système de prévision utilisé actuellement applique une moyenne arithmétique à l'ensemble des articles stockés indépendamment de toutes autres considérations. Cette approche simpliste peut engendrer des estimations de consommation très dégradées. La moyenne arithmétique n'est pas adaptée au cas de consommations comprenant des profils tendanciels ou saisonniers, lesquels profils peuvent avoir lieu en pratique. Ce choix nous semble limitatif et il nous est paru nécessaire de retravailler sur la procédure d'estimation implémentée. La littérature sur les méthodes de prévision est en effet très riche et au-delà des travaux académiques sur la théorie des séries temporelles, plusieurs recherches ont traité de l'implémentation des techniques statistiques dans des outils de prévisions à destination des praticiens de la logistique et du marketing.

Au-delà du choix, discutable, des modèles d'estimation, la procédure de prévision actuellement déployée ne permet aucune interactivité avec l'utilisateur. Elle assure seulement l'alimentation périodique du module de préconisation des commandes pour la mise à jour des paramètres des politiques de gestion de stock. Ceci constitue une difficulté considérable du point de vue des systèmes décisionnels, lesquels systèmes sont centrés autour de la remontée d'informations pour aider l'utilisateur à la prise de décision. Au-delà du travail à entreprendre pour améliorer les précisions des calculs, il nous semble nécessaire de faire évoluer le concept du module afin de permettre la remontée de plus d'informations utiles à l'utilisateur. Ceci passera par la mise en place d'un tableau de bord prévisionnel concentrant une batterie d'informations utiles : des indicateurs de performances prévisionnels, des indicateurs de formes sur l'évolution des consommations, des tendances ainsi que des indicateurs d'alertes sur les comportements prévisionnels atypiques, etc.

Le système de préconisation utilisé s'appuie sur deux politiques de gestion de stock : politique périodique et politique sur seuil. Les paramètres des politiques sont calculés sur la base des prévisions de consommation et tiennent compte d'autres données que l'utilisateur est amené à renseigner. Aucune procédure d'optimisation ne permet d'ajuster ces attributs, l'utilisateur est ainsi appelé à définir, pour l'ensemble des articles, les périodicités qu'il estime appropriées. Il doit également apprécier les différents stocks de sécurité en définissant un

nombre de jours de couverture par article. Ceci se fait par une appréciation humaine, indépendamment de toute considération sur les caractéristiques du produit : prise en compte d'un taux de service cible défini, par exemple, sur la base de la classe d'importance du produit. A

En mettant à plat le module de préconisation de commandes existant et en confrontant notre analyse du système aux propositions et pistes initialement posées par les décideurs et ingénieurs de KLS, nous avons décidé d'orienter notre action au niveau de trois axes d'amélioration afin de faire émerger de nouvelles fonctionnalités et améliorer les performances du système existant :

Axe 1 : La classification des articles du stock : Module de classification

Le module actuel est dépourvu d'un système de classification des articles pour des fins d'optimisation de stock. Nous proposons de travailler sur un module de classification évolué permettant à l'utilisateur de définir avec précision la classe de chaque article. Ce module doit permettre à l'utilisateur d'intégrer différents critères dans sa classification, au-delà du critère financier classique, afin d'avoir des résultats plus réalistes et traduire ainsi plus justement les spécificités de la pharmacie hospitalière. Cette approche multicritères permettra également de s'adapter aux différents cas de figures d'entrepôts pouvant utiliser notre module. Le WMS étant à la base une solution générique devant s'adapter aux différents cas, il est d'un grand intérêt d'offrir des solutions adaptables aux différents profils de clients. Les considérations peuvent en effet changer d'un entrepôt à un autre.

Axe 2 : La prévision des consommations : Module de prévision

Le module de préconisation actuel utilise une routine de calcul appliquant une moyenne mobile qui est une méthode rudimentaire. Cette estimation est utilisée seulement pour alimenter le calcul des quantités de commandes sans offrir à l'utilisateur des fonctionnalités de consultation et d'information autour de la fonction prévision. L'unique information disponible étant une estimation de la consommation, consultable dans la fiche de paramétrage des articles.

Nous proposons de mettre en place un module de prévision à part entière en remplacement de la routine de calcul actuellement utilisée. Ce module doit nous permettre une évolution

sur deux fronts. Le premier étant l'amélioration des estimations de consommation, ce qui permettra d'alimenter le calcul des quantités de commandes moyennant des estimations plus précises. Le deuxième consiste à faire émerger la fonction prévision, fonction jusque-là absente du WMS. Ceci passe par la mise en place d'un tableau de bord prévisionnel, supportant des fonctionnalités de consultation des prévisions, des tendances d'évolution ainsi que des indicateurs d'évaluation des performances.

Le chantier de mise en place d'un système prévisionnel est chose novatrice, non seulement du point de vue du logiciel Gildas, mais plus généralement du point de vue de la famille des systèmes de gestion d'entrepôts. La sphère de fonctionnalités de ces applicatifs se limitait, jusqu'à très récemment, aux domaines de la traçabilité des produits et des opérations d'exécution. Ces systèmes sont souvent interfacés avec des modules complémentaires, notamment des modules de prévision, pour compléter certaines lacunes fonctionnelles. Cette cohabitation logicielle, souvent nécessaire, reste contraignante pour les clients, et au-delà de l'aspect financier, la multiplicité de logiciels au sein d'une même organisation est souvent créatrice de problèmes et conflits informatiques (problèmes d'interfaçage, problèmes de transferts de données entre application, etc.).

Une tendance croissante d'extension des fonctions des WMS par les éditeurs de logiciels, au-delà du périmètre des opérations d'exécution, est en cours. Cependant, autant que nous sachions, rares sont les éditeurs de WMS, du marché européen, qui ont doté leurs applicatifs par un module de prévision de consommation.

Axe 3 : Règles de gestion de stock

Le module de préconisation actuel est basé sur deux politiques de gestion dont les paramètres sont mis à jour par le système de prévision ainsi que par un paramétrage effectué manuellement par l'utilisateur. Le système de prévision utilisé a en effet un impact sur la performance du stock. L'amélioration de la précision des prévisions est, a fortiori, un levier d'action ayant un impact sur les approvisionnements. D'autres leviers d'actions pourraient être envisagés. Le choix de la politique adéquate et le calcul optimal des paramètres de ces politiques en font partie.

Nous nous limitons, dans le cadre de ce travail, à l'évaluation de l'impact de l'amélioration des prévisions sur la performance des stocks. Pour le faire, nous établirons un comparatif entre l'ancien système prévisionnel et le nôtre, couplés aux règles de gestion utilisées. Nous

n'établirons aucun développement informatique complémentaire au niveau de cette partie. Ceci s'inscrira dans la continuité de ce projet et fera l'objet de travaux futurs en complément ceux que nous avons entretenus.

3 Structure et fonctions du système proposé

3.1 Objectifs du système proposé

Comme présenté dans le paragraphe précédent, le module de préconisation de commandes existant présente un certain nombre de limites autant sur le plan des méthodes de calculs implémentées que sur le plan des fonctionnalités proposées aux utilisateurs.

Afin de combler ces insuffisances et apporter des améliorations au système, nous devons proposer des éléments de réponse aux questions précédemment posées et portant principalement sur la classification pertinente des articles du stock, sur la problématique de la prévision de consommation et sur l'estimation de la quantité de commandes optimale par utilisation des règles de gestion. La réponse à ces questions passera nécessairement par une étude approfondie et détaillée portant sur les méthodes scientifiques de classification multicritères, sur les modèles statistiques d'extrapolation des tendances et sur l'exploitation des estimations de consommation pour l'optimisation des niveaux de stock.

L'objectif principal de notre travail est donc de compléter le module existant par un certain nombre de briques fonctionnelles afin d'offrir à l'utilisateur :

- *Un outil de classification des articles de stock* : Outil d'analyse permettant de classer les articles suivant différentes démarches, comprenant la possibilité d'adoption d'une approche multicritères. Des classes d'importance A, B et C doivent pouvoir ainsi être affectées aux produits stockés. Ce résultat servira comme référence aux gestionnaires de l'entrepôt et pourra être exploité par le système dans la perspective d'un ajustement pertinent des politiques de gestion de stock.

- *Un outil de prévision des consommations* : Outil de prévision intelligent, permettant d'estimer périodiquement les consommations des articles stockés et offrant à l'utilisateur des moyens de consultation avancés sur des caractéristiques clés (évolution de tendance, prévision, consommation réelle, etc.) ainsi que sur des indicateurs de performances

prévisionnels. Ces calculs permettront aussi la mise à jour périodique des paramètres des politiques de gestion de stock.

3.2 Méthode

Pour aboutir aux objectifs décrits précédemment, et en partant de l'analyse approfondie de l'existant, notre démarche a consisté à définir, à partir de la littérature scientifique, un certain nombre de modèles et de techniques de calcul, ayant été validés et utilisés dans des problématiques similaires à celles que nous posons.

Notre choix tient compte à la fois de l'apport en performance de ces outils scientifiques, mais également de la faisabilité de l'implémentation informatique à partir des données et informations gérées par le système.

Nous prenons également en considération la facilité d'utilisation et de paramétrage de l'outil. Notre approche doit épargner, dans la mesure du possible, toute intervention humaine portant sur le paramétrage, afin de faciliter et simplifier l'utilisation. Nous considérons également la remontée d'informations et d'indicateurs d'alertes et de performance. Ceci est primordial du point de vue des systèmes experts et facilitera à l'utilisateur l'analyse et la prise de décision.

3.3 Structure du système

Nous présentons dans ce paragraphe l'architecture générale du système, enrichi par les modules supplémentaires que nous avons développés. Ces derniers seront présentés en détail dans la suite des chapitres de la thèse. Nous présentons d'abord les différentes fonctionnalités proposées (Apports fonctionnels) puis nous détaillons l'architecture informatique adoptée pour supporter ces nouveaux modules.

3.3.1 Apports fonctionnels

Le système que nous décrivons dans ce travail constitue un complément au module de préconisation des commandes implémenté dans le WMS *Gildas WM*. Ce système enrichit le fonctionnement du module actuel en permettant une classification des articles du stock sur

une base multicritères et en opérant des estimations de consommation par choix intelligent de techniques de prévision. Nous présentons dans la figure 7 la structure du système de préconisation de commandes enrichi par les modules proposés. Le schéma de la figure 7 permet de situer les nouvelles briques par rapport à l'ancienne configuration.

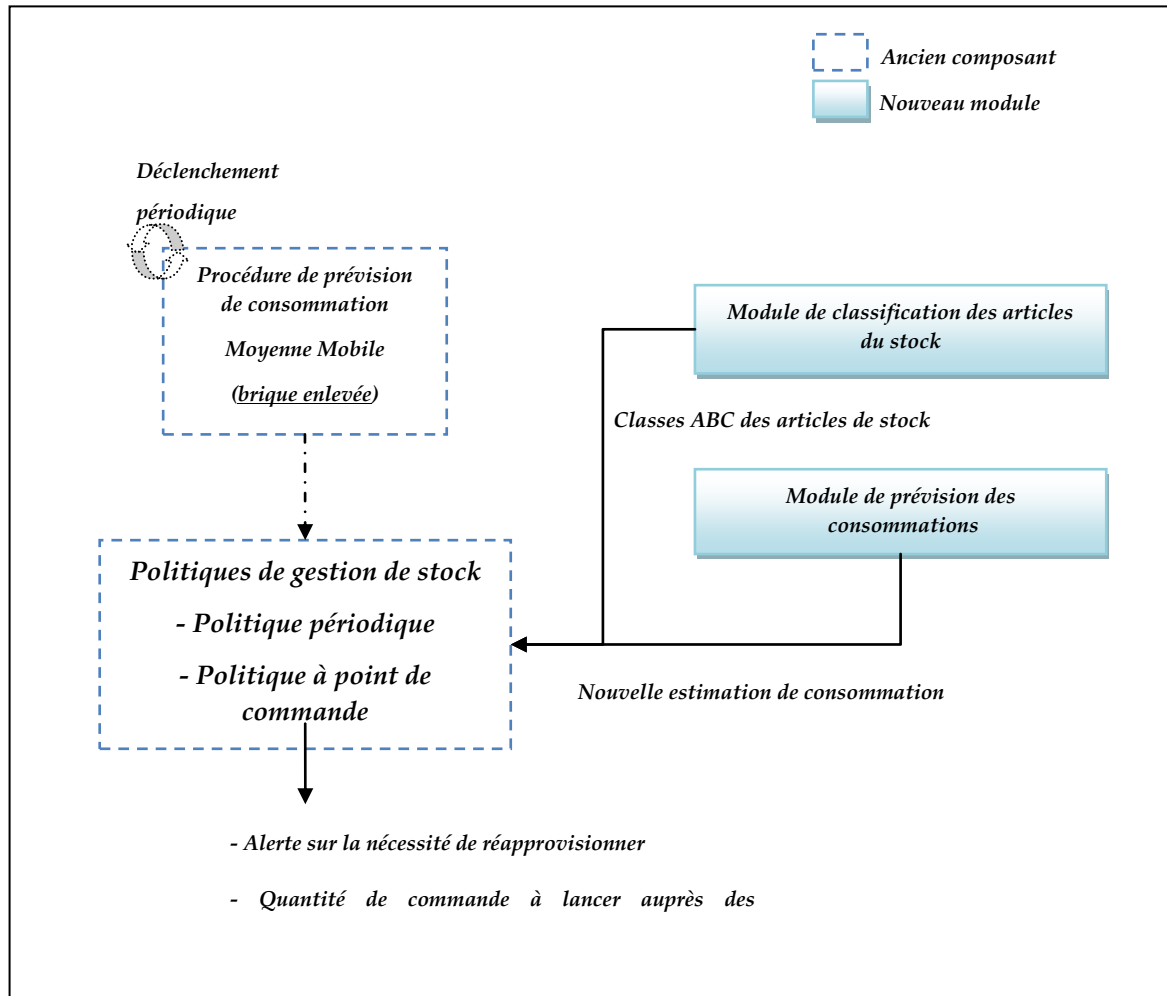


Figure 7 : Composants du nouveau module de préconisation de commandes

Deux modules d'appui viennent en support au système existant. Le module de classification multicritères des articles du stock et le module de prévision de consommation. Ces deux briques vont permettre à la fois d'alimenter le module de préconisation des commandes, en vue de meilleures performances de gestion, mais aussi feront l'objet de fonctionnalités à part entière pouvant être consultées, indépendamment des préconisations, et permettant de remonter à l'utilisateur un certain nombre d'informations et d'alertes dans le but de l'aider dans la prise de décision.

Module de classification

La fonction principale du module de classification est l'affectation de classes d'importances A, B ou C pour chacun des articles du stock. Ce module se déploie comme un outil d'aide à la décision interactif permettant à l'utilisateur de simuler plusieurs classifications et de les comparer afin de mettre à jour le plus judicieusement possible les classes d'importances relatives aux différents articles.

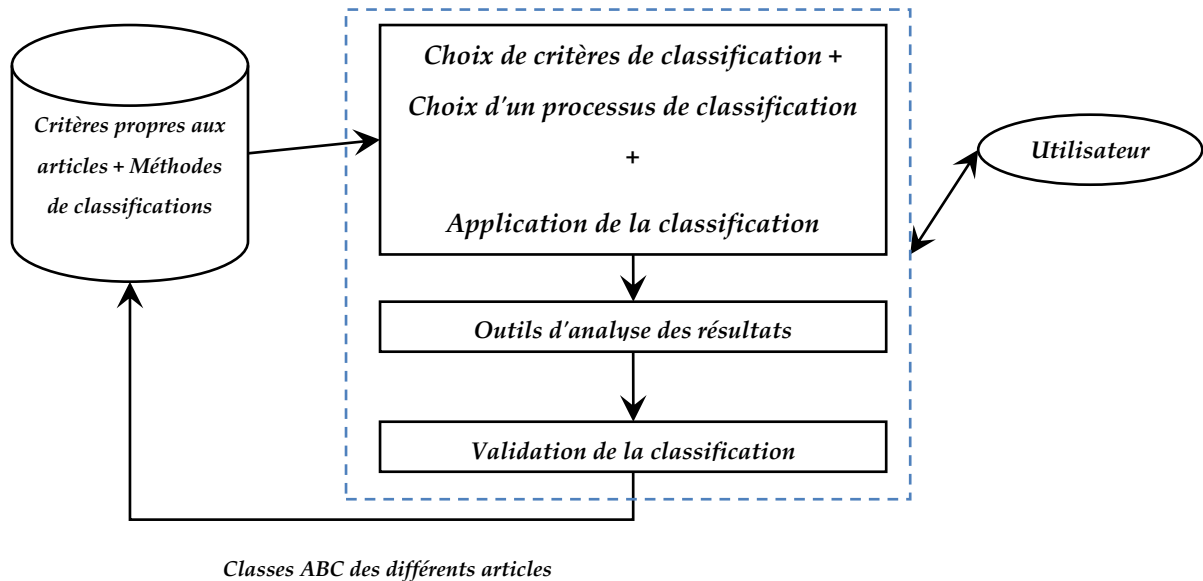


Figure 8: Module de classification des articles

L'utilisateur peut procéder de deux façons. La première est l'application d'une classification ABC usuelle utilisant le critère financier. La deuxième est l'application d'une démarche multicritères. Dans ce cas, l'utilisateur peut définir les critères qui lui semblent pertinents et qui correspondent le plus à son établissement. Le système en prendra compte dans la génération des classes. Une liste prédéfinie de critères est proposée et est sujette à évolution. Il est possible avant la validation, de lancer différents jeux de tests, avec différents critères puis de les comparer au moyen de différents outils d'analyse permettant de : Tracer les articles en communs relatifs aux différentes classes, calculer les proportions financières des différentes classes, etc.

Module de prévision

Le module de prévision des consommations permet d'estimer mensuellement la consommation prévisionnelle du mois en cours relative à tous les articles stockés. Cette estimation se base exclusivement sur l'historique de consommation. Une procédure dynamique permet d'analyser périodiquement les données et d'identifier certaines caractéristiques statistiques de forme. Un choix éclairé d'une technique de prévision, parmi une batterie de techniques implémentées, est opéré en fonction de cette analyse.

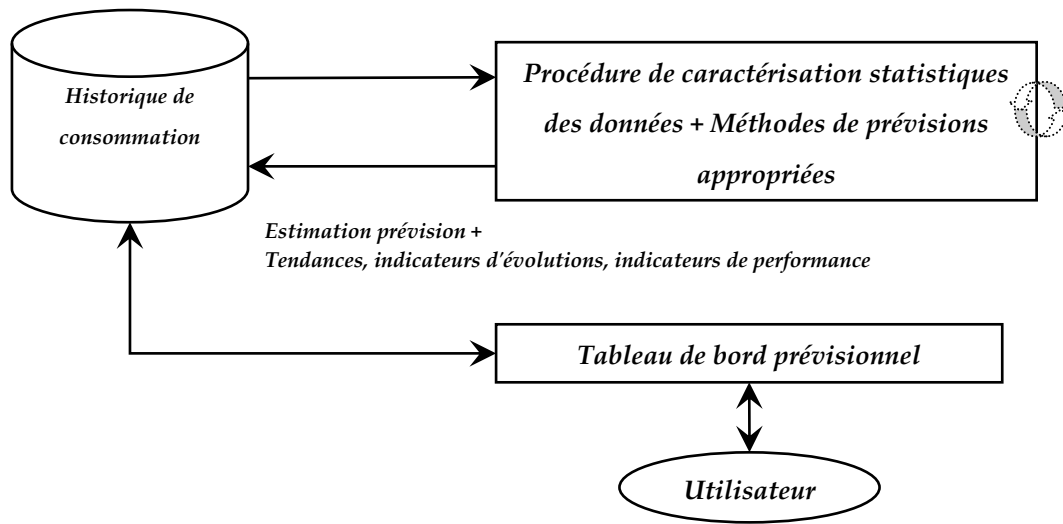


Figure 9: Module de prévision des consommations

L'utilisateur peut accéder aux résultats prévisionnels via un tableau de bord. Il s'agit d'une grille reprenant les informations les plus synthétiques qui traduisent l'évolution de consommation et des mesures de performances des estimations.

3.3.2 Architecture informatique déployée

Les fonctions proposées étant destinées à enrichir le WMS, il a fallu prendre compte du système existant afin de mettre en place les nouvelles fonctionnalités et s'adapter au mieux à l'architecture informatique existante. Au-delà de l'intégration de notre concept dans Gildas WM, notre choix a consisté à mettre en place des modules pouvant facilement se déployer en une application autonome, afin de pouvoir les proposer, indépendamment de l'offre WMS. Ce point nous semble très important et contribuera à élargir le champ d'application de nos modules. Cette considération a conditionné, en effet, l'architecture informatique à mettre en place.

Nous reprenons dans la figure 10 l'architecture détaillée du système informatique existant. Nous la complétons par les composants qu'il a fallu développer afin de faire évoluer le système. Nous commentons dans ce qui suit ces différents composants.

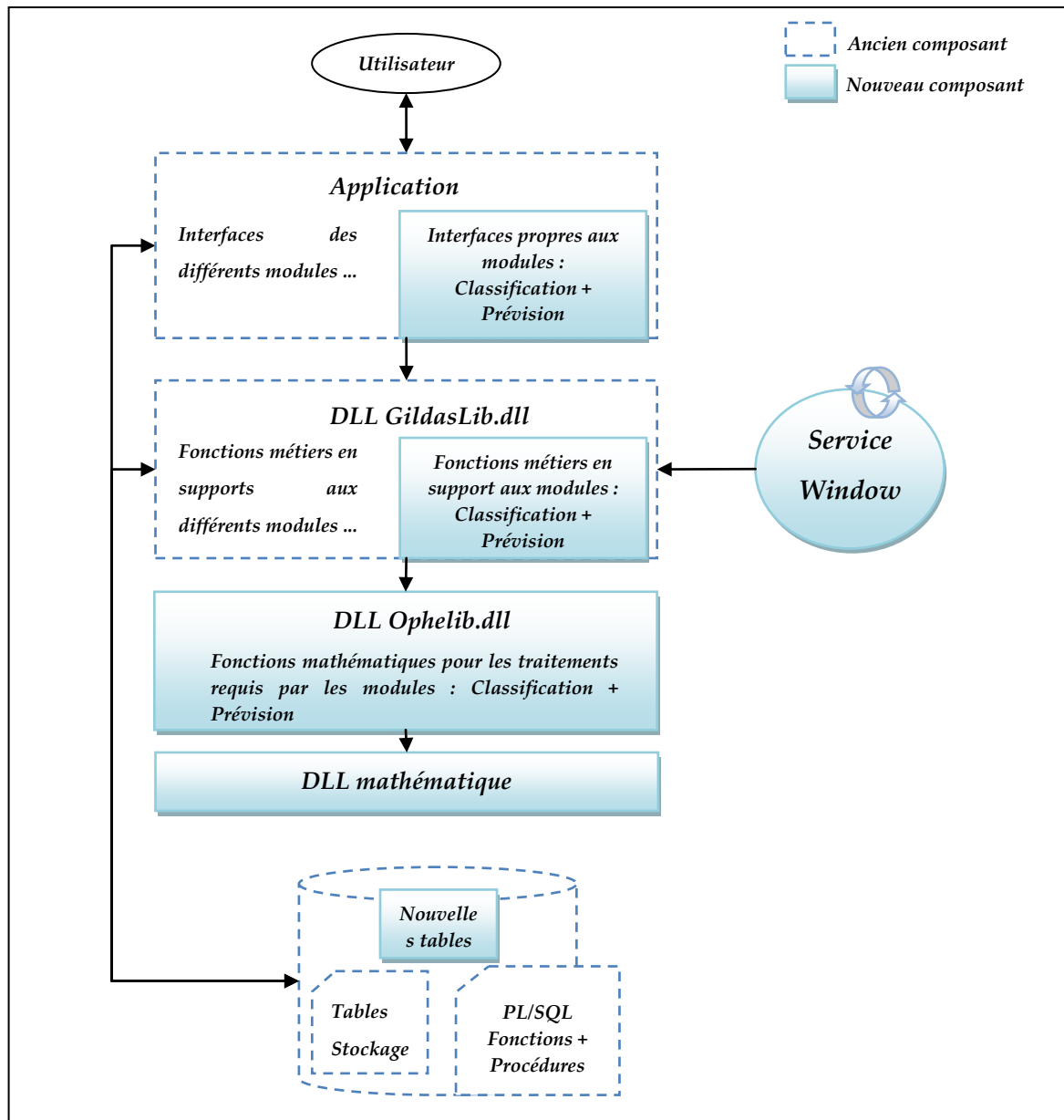


Figure 10 : Architecture informatique déployée en support au nouveau module de préconisation de commandes

Interface utilisateur

Nous utilisons la même interface graphique offerte par Gildas WMS afin de permettre à l'utilisateur d'accéder aux modules complémentaires que nous avons conçus. Des

développements informatiques complémentaires sur l'environnement *Borland C++ Builder* ont été réalisés afin de mettre à niveau le WMS. Ainsi, l'utilisateur peut accéder aux nouveaux modules à travers des interfaces graphiques comprenant des grilles de données, des boutons d'action et des graphiques.

Base de données

Les nouveaux modules ont introduit de nouvelles données devant être stockées afin d'être consultées par l'utilisateur à travers l'interface graphique : Historique des prévisions, indicateurs de performances sur les prévisions, les classes d'importances, etc., ainsi que des données intermédiaires nécessaires à certains calculs. Nous nous appuyons sur la base de données Oracle utilisée afin de supporter toutes ces informations. Pour le faire, une architecture complémentaire a été développée (Exclusivement des tables).

La base de données Oracle du WMS est également utilisée, dans des calculs additionnels (via le langage PL/SQL). Nous faisons le choix de ne pas utiliser les procédures et fonctions d'Oracle pour les traitements mathématiques. Nous consacrons la BDD Oracle exclusivement au stockage des données. L'option d'utilisation du langage PL/SQL pour effectuer des opérations mathématiques présente certains avantages. Oracle est en effet doté de composants très avancés comprenant un large spectre de fonctions mathématiques et statistiques et supportés par la licence du système. Ceci permet de se dispenser de l'achat de certaines bibliothèques mathématiques et allège le trafic au niveau du réseau, ce qui permet une certaine stabilité de l'application. L'inconvénient majeur de cette option réside dans la dépendance au système de gestion de base de données Oracle et complique par conséquent la génération d'applications autonomes transposables sur d'autres systèmes de gestion de BDD (SQL Server, etc.). En regard de cette contrainte nous avons choisi de nous émanciper de la BDD et de développer une DLL dédiée aux traitements mathématiques requis.

DLL mathématique

Nous développons un fichier DLL afin d'encapsuler les fonctions nécessaires aux traitements mathématiques. Ce fichier va supporter l'ensemble des opérations mathématiques et statistiques nécessaires au fonctionnement des deux modules ajoutés.

Cette DLL est développée sous *Borland C++ Builder* et s'appuie sur d'autres DLL mathématiques de plus bas niveau supportant des fonctions mathématiques de base prédéveloppées. Ces DLL sont téléchargées à partir d'internet et sont à licences non payantes.

Ces fonctions peuvent ainsi être appelées à partir de n'importe quel programme appelant, par liaison dynamique avec le fichier DLL. Ceci nous permet une plus grande flexibilité quant à l'adaptation de notre concept en vue de la génération d'une application autonome s'appuyant sur divers systèmes de gestion de base de données.

DLL intermédiaire (DLL des fonctions métier)

Cette DLL est déjà utilisée par l'application de gestion d'entrepôt et encapsule un certain nombre de fonctions métiers utilisées en support aux différents modules. Nous la complétons afin de permettre le déploiement de notre concept.

Cette DLL permet de synchroniser entre eux les différents composants de l'application. Elle assure la récupération des données à partir de la BDD, l'appel des fonctions de traitement à partir de la DLL mathématique et la mise à jour de la BDD. Elle est sollicitée, selon le cas, par le service Windows afin d'opérer des traitements périodiques, ou directement par l'utilisateur, via l'interface graphique, afin d'opérer d'autres types de calculs.

Service Windows

Le service Windows est un nouveau composant que nous avons développé afin d'entretenir les traitements périodiques. Dans l'ancien système, un mécanisme équivalent est assuré par la base de données Oracle, au moyen de 'Jobs', une sorte de compteurs, permettant de stimuler certains traitements à des périodicités paramétrables.

Ayant fait le choix de s'émanciper de la BDD, nous optons pour le développement d'un service afin d'assurer cette fonction. Ce composant, développé également sous *Borland C++*, va stimuler régulièrement, à des fréquences définies et paramétrables, la DLL intermédiaire afin d'opérer certains calculs et entretenir la mise à jour des données au niveau de la base. Le calcul mensuel des estimations de consommation est un traitement type entretenu par le service et devant se déclencher automatiquement en fin de chaque mois.

4 Conclusion

Dans ce chapitre nous avons procédé à une analyse détaillée du module de préconisation existant. Cette analyse nous a permis de définir les lacunes du système et d'avancer des éléments d'amélioration. Deux modules complémentaires ont été proposés : un module de classification des articles du stock ; et un module de prévision des consommations. Nous avons présenté l'architecture fonctionnelle et informatique détaillée en support à notre système.

Nous abordons, en détail, dans la suite de la thèse chacun de ces modules. Dans le chapitre suivant nous nous focalisons sur le module de classification multicritères des articles du stock.

Chapitre III. Module de classification des articles du stock

Dans ce chapitre, nous abordons le module de classification des articles du stock. Ce module permet de classer les articles de l'entrepôt suivant différentes méthodes incluant la possibilité d'une démarche multicritères. Nous revenons, en premier lieu, sur certains travaux de la littérature ayant traité de cette problématique. Ceci nous permettra de définir les méthodes et techniques de classification adoptées ainsi que l'étendue d'application dans le cadre des logiciels et des systèmes experts. Dans une deuxième phase, nous définissons les techniques et modèles de calculs retenus en justifiant notre choix. Enfin, nous présentons le concept imaginé et développé autour de ces modèles de calculs et nous détaillons les fonctionnalités proposées à l'utilisateur.

1 Intérêt de la classification des articles par une approche multicritères

Les entrepôts et plateformes logistiques assurent le stockage de composants, de produits divers et de matières premières afin de desservir les différents maillons de la chaîne logistique. Ces derniers peuvent correspondre à des plateformes intermédiaires, des distributeurs, des chaînes de production et de maintenance ou encore des magasins. Dans le cas de la chaîne hospitalière, les stocks peuvent être très élevés, allant jusqu'à plusieurs milliers de produits détenus par une même plateforme hospitalière.

Il est complexe, pour les gestionnaires, de suivre avec la même finesse et d'accorder une même importance à l'ensemble de tous les produits stockés. Il est ainsi important de pouvoir identifier, parmi la multitude des produits détenus, la classe des articles pilotes. Cette catégorie correspond en général à une petite fraction du stock (entre 10% et 20% des produits) et représente un taux financier élevé (autour de 80% du flux financier annuel de la plateforme). Un suivi rigoureux est généralement accordé aux produits appartenant à cette classe.

Il est courant, en pratique, de procéder, périodiquement, à cette priorisation dans le but de séparer les articles phares du restant des produits. L'effort à dépenser pour le suivi et le maintien du stock est souvent modulé par cette catégorisation. Différentes règles de gestion peuvent, en effet, être adoptées selon la classe d'importance du produit. Cette démarche est fondamentale. La bonne appréciation, ou non, des classes d'importance peut avoir des conséquences considérables sur le bon fonctionnement de l'organisation.

La méthode de classification ABC classique est parmi les techniques les plus utilisées, en pratique, dans la priorisation des articles du stock. Les produits sont répartis en trois classes d'importance : produits de la classe A, considérés comme étant les plus critiques, produits de la classe B, moyennement importants, et les produits de la classe C, considérés comme non prioritaires.

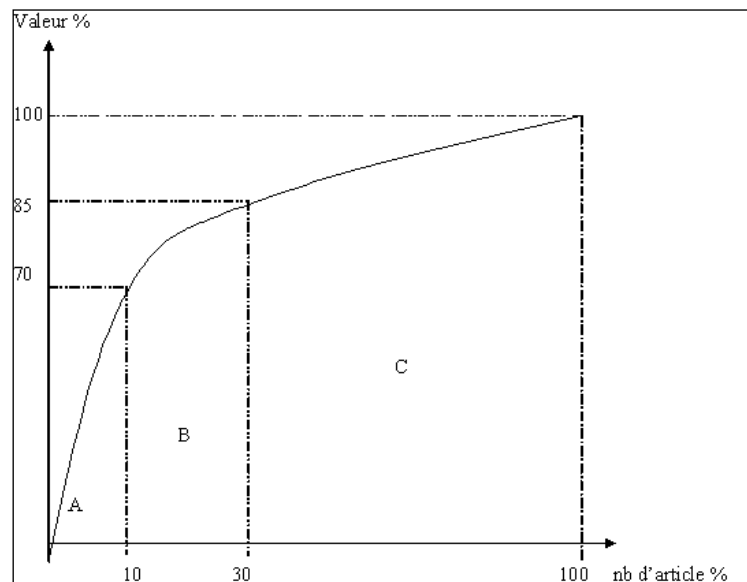


Figure 11: Coubre ABC

Cette répartition est établie, dans la plupart des cas, sur la base d'un critère financier. Les articles du stock sont, généralement, classés en sens décroissant de leur montant financier annuel, critère connu sous l'appellation anglo-saxonne d'ADU et représentant le produit de la consommation annuelle de l'article par son prix unitaire. La classe A, correspond en général à une faible proportion variant entre 10 % et 20 % du total des articles et représentant une grande proportion du total financier. Les produits de la classe C sont les plus nombreux (autour de 60 % du total), mais ont un faible poids financier. La proportion intermédiaire correspond aux articles de la classe B.

Cette méthode est très répandue auprès des gestionnaires et des pharmaciens. Elle est assez facile d'appréhension et très simple d'application. L'utilisation de classeur Excel permet facilement d'entretenir cette démarche. Elle est aussi souvent supportée par des modules de gestion de stock plus spécialisés.

Alors qu'elle est systématiquement appliquée en pratique, *Ramanathan (2006)* considère que l'approche de classification monocritère n'est pas toujours la mieux appropriée. Cette approche est justifiée dans le cas où les produits stockés présentent une certaine homogénéité et que la différence porte uniquement sur le critère financier. Ceci n'est pas forcément garanti dans le cas des produits pharmaceutiques. Les plateformes étant de plus en plus denses, desservant une grande variété de services de soins, les articles stockés présentent nécessairement une certaine disparité.

La méthode ABC classique a fait également l'objet de critiques dans les travaux de *Guvenir & Erel (1998)*, *Huiskonen (2001)* ou *Partovi & Anandarajan (2002)* lesquels ont reconnu que cette dernière ne constitue pas en général une bonne méthode de classification.

L'approche humaine, basée exclusivement sur l'expertise du gestionnaire, peut constituer une alternative à la classification ABC monocritère. Une démarche intuitive, s'appuyant sur la connaissance du manager, peut souvent s'avérer effective et pertinente, elle s'appuie couramment sur l'expertise des managers dans la gestion de certaines opérations. Ainsi dans les plateformes hospitalières, les pharmaciens ont coutume d'intégrer en plus des critères de coût ou d'ADU, des critères complémentaires comme le délai de réapprovisionnement, ou encore la criticité qui est liée à la rareté d'un article avec dans le cas des médicaments la possibilité ou non de le substituer à un produit générique.

Partovi & Anandarajan (2002) considèrent que cette approche est aussi limitée, le jugement humain, aussi avisé soit-il, n'est en aucun cas dépourvu de subjectivité et présente des risques d'incohérence.

En réponse à cette problématique, certains travaux ont considéré des approches de classification multicritères. Pour pallier les insuffisances du critère unique, plusieurs attributs ont été identifiés, comprenant en plus de l'aspect financier, le temps de

réapprovisionnement, la criticité de l'article, son obsolescence ou encore le quota de commande. Le principe consiste à intégrer l'effet de tous ces critères en un score global pour aboutir à une classification plus juste. Pour ce faire, différentes techniques de calcul, issues de la littérature multicritères, sont applicables. L'enjeu est de réussir à prendre compte de l'ensemble des critères tout en remédiant à l'approche ad hoc basée intégralement sur l'expertise de l'humain. Nous proposons dans ce qui suit de revenir sur les principaux travaux ayant traité de la problématique de classification des articles du stock. Ceci nous permettra de recenser des modèles de calcul, de les comparer afin de définir ceux qui répondent le mieux à notre problème.

2 Méthodes de classification des articles du stock

Flores & Whybark (1986, 1987) sont parmi les premiers à avoir traité de la problématique de classification multicritères des articles du stock. Leur approche est basée sur l'utilisation d'une matrice bidimensionnelle permettant l'intégration de deux critères (figure 12). Il s'agit d'affecter, à tous les articles du stock, un niveau d'importance (A, B ou C) en regard de chacun des deux critères définis. En croisant les deux dimensions de la matrice, neuf sous catégories de produits sont créées : famille 'AA', 'AB', 'AC', etc.

		Critère 1		
		A	B	C
Critère 2	A	AA	AB	AC
	B	BA	BB	BC
	C	CA	CB	CC

Figure 12 : Matrice bidimensionnelle de classification des articles du stock

Différentes politiques de gestion et pratiques de suivi peuvent ainsi être formulées et adoptées en fonction de l'appartenance de l'article aux catégories définies. Cette approche a permis, en effet, d'apporter une première réponse aux limitations de la démarche monocritère propre à la méthode ABC classique. Cependant, elle reste limitée et très mal

adaptée aux cas faisant intervenir plus de trois attributs. Ceci revient à manipuler des matrices multidimensionnelles afin d'intégrer la dimension multicritères. Cette situation est impraticable et aboutit à un grand nombre de familles de produits. L'autre limitation de cette matrice réside dans la pondération des attributs. Le croisement des deux niveaux suppose, en effet, une certaine parité en regard des deux critères choisis. Aucune distinction ne peut être attribuée dans le cas de cet outil.

Depuis les travaux de *Flores & Whybark (1986, 1987)*, plusieurs recherches ont suivi et traité de la problématique multicritères de classification des articles du stock. Différentes méthodes ont été adoptées, parmi lesquelles figurent des travaux utilisant des méta-heuristiques (Algorithmes génétiques) et des approches issues de l'intelligence artificielle (Réseaux de neurones). La programmation linéaire et les techniques d'optimisation ont fait également partie des méthodes explorées pour traiter de cette problématique. L'ensemble de ces approches constitue des modèles quantitatifs. Il peut être intéressant, lors de la classification des articles, de considérer à la fois des critères qualitatifs et quantitatifs, ceci peut donner lieu à des résultats plus réalistes. L'une des méthodes décisionnelles les plus utilisées et ayant été adaptée au problème de classification est la méthode AHP : '*Processus de hiérarchisation analytique*'. L'AHP permet d'intégrer, dans un même processus décisionnel, simultanément des critères qualitatifs et quantitatifs. Nous revenons dans ce qui suit sur ces trois lignées de méthodes et abordons certains travaux phares de la littérature.

2.1 Méthode AHP : Processus de hiérarchisation analytique

L'AHP est une méthode très réputée ayant été appliquée dans divers problèmes multicritères d'aide à la décision. Elle a été initialement introduite dans le travail de *Saaty (1977)*, pour aider les décideurs à intégrer plusieurs objectifs à la fois. Certains travaux ont considéré l'application de l'AHP dans le cadre de la classification multicritères des articles du stock. L'AHP permet de subdiviser un problème complexe en des sous-éléments suivant une structure hiérarchique, de traduire un jugement subjectif portant sur l'importance relative des différents éléments de la structure, en une valeur numérique en se basant sur un comparatif humain, puis de synthétiser le jugement en une mesure globale. *Flores et al. (1992)* résument la méthode AHP en trois principales étapes :

Étape 1 : Le décideur doit identifier l'ensemble des critères pouvant avoir un impact sur la décision qu'il doit prendre.

Étape 2 : Les critères doivent être structurés suivant une hiérarchie d'un ou de plusieurs niveaux. Certains critères, jugés de même nature, peuvent être regroupés pour constituer des sous-éléments d'un critère plus englobant.

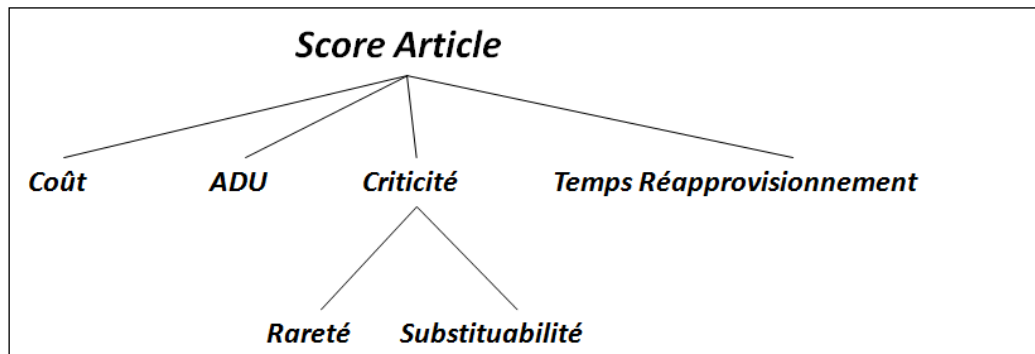


Figure 13 : Structure hiérarchique des éléments influant la décision

Étape 3 : Le décideur doit apprécier les différents critères de la hiérarchie. Les critères d'un même niveau sont comparés deux à deux par affectation d'une préférence. Cette étape permet de traduire le jugement humain en une mesure quantitative. *Saaty (1980)* a proposé d'exprimer la préférence par affectation d'une note allant de 1 à 9. Une matrice de comparaison symétrique est ainsi construite et permet de dériver, les poids relatifs aux différents critères ainsi qu'une mesure de la consistance du comparatif établi. La consistance du comparatif est une mesure intéressante qui traduit la cohérence de la comparaison établie par le manager. Cette métrique permet d'identifier les situations de comparaisons incohérentes, ceci peut souvent avoir lieu, surtout quand il s'agit de traiter de plusieurs critères. L'ensemble des poids obtenus est synthétisé en une formule reflétant l'impact de tous les niveaux hiérarchiques. Il s'agit enfin d'appliquer les valeurs des différents critères à cette formule pour obtenir le score final traduisant le mix multicritères. Ces valeurs doivent être normalisées avant l'application afin de les ramener à une échelle d'évaluation commune.

	<i>Coût</i>	<i>ADU</i>	<i>Criticité</i>	<i>Temps. R</i>
<i>Coût</i>	1	2	5	1/7
<i>ADU</i>	1/2	1	6	9
<i>Criticité</i>	1/5	1/6	1	7
<i>Temps. R</i>	7	1/9	1/7	1

Figure 14: Matrice de comparaison

L'AHP a été utilisée dans certains travaux pour aborder la problématique de classification multicritères des articles de stock. Parmi les premiers à l'avoir appliquée, nous citons le travail de *Flores et al. (1992)* ainsi que celui de *Partovi & Burton (1993)*. *Flores et al. (1992)* considèrent des données propres à un stock hospitalier. Ils reprennent la classification de 47 articles, ayant déjà été traités dans *Reid (1987)* au moyen de la classification ABC usuelle basée sur le critère financier. *Flores et al (1992)* introduisent quatre critères de classification : le coût unitaire moyen, le montant financier annuel, la criticité et le temps de réapprovisionnement de l'article auprès du fournisseur. Les articles considérés présentent une certaine disparité. Certains d'entre eux sont très coûteux et non critiques, d'autres sont très critiques, mais ayant des temps de réapprovisionnement très courts. Les résultats obtenus révèlent en effet une plus grande uniformité, en comparaison des résultats de la démarche classique.

Considérée comme référence, l'AHP a été reprise dans *Güvenir & Erel (1998)* ou plus récemment dans *Yu (2010)*. Ces travaux reprochent à la méthode une part de subjectivité introduite lors de l'intervention humaine dans le comparatif des critères. En alternative, ces deux travaux ont proposé l'application de méthodes issues du domaine de l'intelligence artificielle. Nous revenons, dans la suite de la thèse, sur cette lignée de modèles.

L'AHP a fait également l'objet d'implémentation dans les systèmes informatiques d'aide à la décision. *Cakir & Canbolat (2008)* ont développé un module informatique d'aide à la classification des articles de stock. Leur système est basé sur l'application d'une variante de la méthode AHP. L'utilisateur peut appliquer le processus de classification, développé ci-dessus, en suivant les étapes clés de la méthode : 1) Choix des critères pertinents à partir d'une liste évolutive, 2) comparatif deux à deux des critères choisis moyennant des variables

linguistiques et enfin, 3) évaluation des poids retournés par le système pour les appliquer aux produits du stock. À notre connaissance, ce travail constitue l'unique tentative d'application de l'AHP pour le cas d'un module spécifique à la problématique de classification multicritères des articles de stock. Par ailleurs, plusieurs logiciels supportant la méthode AHP existent, leur sphère d'application reste cependant généraliste. Le logiciel *Expert Choice* est parmi les plus réputés.

L'AHP constitue donc un moyen évolué pour combiner de multiples critères dans la classification des articles du stock afin d'aboutir à une classification plus réaliste. Cette méthode présente l'inconvénient de nécessiter l'intervention humaine ce qui peut entraîner de la subjectivité dans les résultats. Au-delà de cet inconvénient, l'AHP reste assez intéressante. Elle est plutôt simple d'appréhension par des gestionnaires et praticiens de la logistique. Son application est relativement simple et la sollicitation de l'humain n'est pas trop contraignante, ce dernier étant appelé pour établir la comparaison des critères, le reste des calculs pouvant se faire automatiquement par le système. L'implémentation fréquente de cette méthode dans des modules informatiques témoigne en effet de l'intérêt de son utilisation.

2.2 Méthodes basées sur la programmation linéaire

Une autre lignée d'articles ayant traité de la problématique de classification des produits du stock est basée sur les méthodes d'optimisation et de programmation linéaire. L'un des premiers travaux à avoir abordé cette démarche est celui de *Ramanathan (2006)*. *Ramanathan (2006)* propose une approche basée sur la résolution de programmes linéaires de maximisation afin de générer les scores propres aux articles stockés. Leur idée est similaire à une DEA (*Data Envelopment analysis*) et est relativement simple d'appréhension. En considérant N articles devant être répartis en 3 classes A, B et C suivant une répartition devant tenir compte de J critères et en supposant que tous les critères soient positivement reliés au niveau d'importance de l'article, *Ramanathan (2006)* propose la résolution, pour chaque article o , du programme linéaire (P1) suivant :

$$\begin{aligned}
 & \max_{j=1}^J \quad v_{oj} y_{oj} \\
 & \text{Sous Contraintes} \\
 & 1 \quad v_{oj} y_{ij} \leq 1, \quad i = 1, 2, \dots, N \\
 & \quad \quad \quad j=1 \\
 & 2 \quad v_{oj} \geq 0, \quad j = 1, 2, \dots, J
 \end{aligned}$$

Il s'agit de maximiser pour l'article o en cours, une fonction de score représentant l'agrégation des différentes valeurs des critères choisis y_{oj} modulées par les coefficients de pondérations v_{oj} . Le modèle est soumis aux contraintes de normalisation (1) assurant que l'application des poids v_{oj} , relatifs à l'article o , aux restants des articles du stock résulte en un score inférieur ou égal à l'unité. Les inéquations (2) constituent les contraintes d'intégrité du problème. De cette façon, chacun des articles va générer ses propres poids relatifs aux différents critères et permettant de maximiser son propre score.

Cette approche est assez intéressante. Elle permet, en effet, d'automatiser le processus de classification en éliminant toute intervention humaine. Ceci peut faciliter d'une part la procédure de classification, ce qui est intéressant du point de vue des systèmes experts, mais permet surtout d'éliminer la part de subjectivité liée à l'intervention de l'humain qu'on retrouve, notamment dans la méthode AHP. Le principe du modèle est relativement simple de compréhension par des praticiens de la logistique. Le travail de **Ramanathan (2006)** présente, cependant, certaines insuffisances et a été commenté et critiqué dans plusieurs travaux.

Ng(2007) reproche à cette approche de nécessiter le recours à une optimisation linéaire pour chacun des articles de stock. Ceci pourrait résulter en de longs temps de calcul quand il s'agit de traiter de milliers d'articles, ce qui est souvent le cas dans les plateformes hospitalières. Pour remédier à cette lacune, **Ng (2007)** propose de modifier le programme linéaire de **Ramanathan (2006)** en opérant certaines modifications et en changeant les contraintes du problème. Des traitements complémentaires sur le programme modifié ont enfin permis d'aboutir à une résolution analytique, évitant ainsi la procédure d'optimisation. Le résultat

des calculs est cependant différent de celui obtenu dans *Ramanathan (2006)*. Le score propre à chaque article o est calculé par la formule suivante.

$$\max_{j=1 \dots J} \left(\frac{1}{J} \sum_{k=1}^j y_{ok} \right)$$

L'apport de ce travail réside principalement dans l'alternative de calcul analytique. Les scores obtenus, conformément à la formule précédente, sont cependant indépendants des poids de l'article en question. *Hadi Venech (2010)* considère que ce résultat peut mener à la situation où un article est mal classé sans refléter sa vraie position par rapport aux produits stockés. De ce point de vue, le modèle de *Ramanathan (2006)* reste plus intéressant, d'autant plus que le problème du temps de calcul n'en est pas nécessairement un de notre point de vue. La classification des articles étant une opération peu fréquente (la reclassification se fait généralement à une fréquence annuelle), il n'est pas très contraignant d'avoir des temps longs, d'autant plus qu'il est toujours possible de lancer cette opération en parallèle à d'autres tâches ou durant des périodes creuses (le soir par exemple).

Zhou & Fan (2007) reviennent aussi sur le travail de *Ramanathan (2006)* en notant une autre lacune inhérente au modèle proposé. Le modèle de *Ramanathan (2006)* permet à chaque article de sélectionner ses propres poids de façon à maximiser son score global. Ceci peut entraîner la situation où un article, ayant une seule valeur dominante en terme d'un unique attribut, soit considéré comme étant un article de la classe A, indépendamment du restant des critères, lesquels peuvent avoir de faibles valeurs. Ceci peut donc entraîner une mauvaise appréciation de certains articles. Pour remédier à cette situation, *Zhou & Fan (2007)* proposent un complément au modèle de *Ramanathan (2006)* permettant d'atténuer l'effet de dominance du critère unique. Il s'agit dans leur cas de résoudre deux programmes linéaires pour chacun des articles du stock. Le premier étant le programme de maximisation (*P1*), le deuxième est un programme linéaire (*P2*) équivalent permettant de minimiser le score de l'article en regard des différents critères.

$$\begin{aligned} & J \\ \mathbf{P2} \quad & \min \sum_{j=1}^J v_{oj} y_{oj} \\ & \text{Sous Contraintes} \\ & \sum_{j=1}^J v_{oj} y_{ij} \geq 1, \quad i = 1, 2, \dots, N \end{aligned}$$

$$2 \quad v_{oj} \geq 0, j = 1, 2, \dots, J$$

Similairement au premier cas, ce programme permet de générer pour chaque article, les poids minimisant le score global, sous les contraintes de normalisation assurant que l'application de l'ensemble des poids au restant des articles renvoie un score supérieur ou égal à l'unité. De cette façon, chacun des articles ajustera ses propres poids de façon à générer le score le moins favorable. *Zhou & Fan (2007)* proposent ensuite de construire sur la base de ces deux résultats, le plus favorable et le moins favorable, un index composite combinant les deux extrêmes. En notant SF_o le score favorable de l'article o , SMF_o le score le moins favorable et en adoptant une constante arbitraire φ comprise entre 0 et 1, l'index composite s'écrit :

$$IC_o \varphi = \varphi \frac{SF_o - \min_{i=1..N} SF_i}{\max_{i=1..N} SF_i - \min_{i=1..N} SF_i} + (1 - \varphi) \frac{SMF_o - \min_{i=1..N} SMF_i}{\max_{i=1..N} SMF_i - \min_{i=1..N} SMF_i}$$

Cet index permet donc de faire un compromis entre les deux valeurs extrêmes et atténue ainsi les résultats renvoyés par le modèle de *Ramanathan (2006)*.

Chen (2011) considère que la démarche adoptée dans *Zhou & Fan (2007)* aboutit à des résultats plus raisonnables et plus réalistes. Cette approche est cependant toujours insuffisante. Du point de vue de *Chen (2011)*, il n'est pas tout à fait cohérent, de comparer des articles sur la base de scores obtenus à partir de poids différents, ce qui est le cas dans ces travaux antérieurs. De plus, l'indice composite introduit dans *Zhou & Fan (2007)*, et malgré l'effet d'atténuation qu'il introduit, ne permet pas nécessairement d'aboutir à une classification cohérente : D'une façon générale, un article donné, ayant une valeur élevée en terme d'un critère important et une faible valeur en terme d'un critère secondaire, devrait être priorisé par rapport à un article ayant une faible valeur en terme du critère important et une forte valeur en terme du critère secondaire *Chen (2011)*. L'indice composite ne permet pas cette distinction.

Pour combler les lacunes de ces travaux, *Chen (2011)* propose une extension des modèles précédents en adoptant une démarche qu'il intitule "Estimation par les paires pour la

classification multicritères". L'idée consiste à résoudre pour chacun des articles du stock le programme linéaire le plus favorable (**P1**) et le moins favorable (**P2**). La résolution permet de générer un ou plusieurs ensemble(s) de poids les plus favorables et un ou plusieurs ensemble(s) de poids les moins favorables. Il est en effet possible que la même valeur optimale ait plusieurs solutions. Dans le cas où le score le plus favorable (respectivement, le moins favorable) est atteint par plusieurs ensembles de poids, un choix particulier de la solution la plus favorable (respectivement, la moins favorable) est opéré. Ce choix se fait par sélection des poids représentant le moins de dissymétrie possible. Pour y parvenir, *Chen (2011)* propose en complément, pour le cas le plus favorable (respectivement le moins favorable), la résolution des programmes linéaires (**P1A**) et (**P1B**) (respectivement les programmes linéaires (**P2A**) et (**P2B**)) que nous présentons dans la suite.

Le principe du modèle (**P1A**) est le suivant : les trois premiers groupes de contraintes permettent de considérer l'ensemble des solutions optimales du modèle (**P1**) pour l'article o . Par ajout des deux derniers groupes, nous forçons les performances individuelles de chaque critère à être au dessus d'une valeur non négative α_o . L'objectif du programme étant la maximisation de cette dernière constante, les performances individuelles les plus élevées seront ainsi sélectionnées *Chen (2011)*.

$$\mathbf{P1A} \quad \max \alpha_o$$

Sous Contraintes

$$1 \quad \sum_{j=1}^J v_{oj} y_{ij} \leq 1, \quad i = 1, 2, \dots, N$$

$$2 \quad \sum_{j=1}^J v_{oj} y_{oj} = SF_o$$

$$3 \quad v_{oj} \geq 0, \quad j = 1, 2, \dots, J$$

$$4 \quad v_{oj} y_{oj} \geq \alpha_o, \quad j = 1, 2, \dots, J$$

$$5 \quad \alpha_o \geq 0$$

$$\mathbf{P1B} \quad \min \beta_o$$

Sous Contraintes

$$\begin{aligned}
 & \quad \quad \quad J \\
 1 \quad & \quad \quad v_{oj} y_{ij} \leq 1, \quad i = 1, 2, \dots, N \\
 & \quad \quad \quad j=1 \\
 & \quad \quad \quad J \\
 2 \quad & \quad \quad v_{oj} y_{oj} = SF_o \\
 & \quad \quad \quad j=1 \\
 3 \quad & v_{oj} \geq 0, j = 1, 2, \dots, J \\
 4 \quad & v_{oj} y_{oj} \geq \alpha_0^*, j = 1, 2, \dots, J; \alpha_0^* \text{ solution optimale de (P1A)} \\
 5 \quad & v_{oj} y_{oj} \leq \beta_0, j = 1, 2, \dots, J
 \end{aligned}$$

Le programme **(P1A)** est complété par le programme **(P1B)** dont le principe est le suivant : Les quatre premiers groupes de contraintes permettent de considérer l'ensemble des solutions optimales du modèle **(P1A)** pour l'article **o**. En ajoutant le dernier groupe de contraintes, nous forçons les performances individuelles de chaque critère à ne pas dépasser une borne positive β_0 . L'objectif du programme étant la minimisation de cette borne, en étant assujetti à l'optimalité de **(P1A)**, des poids individuels les moins dissymétriques seront ainsi sélectionnés *Chen (2011)*. La résolution de ces programmes complémentaires permet donc de définir les poids des critères les plus élevés tout en étant les moins dissymétriques. D'une façon analogue à celle présentée ci-dessus, les poids les moins favorables relatifs au programme **(P2)** sont affinés par résolution des programmes **(P2A)** et **(P2B)**.

$$\begin{aligned}
 & \quad \quad \quad \mathbf{P2A} \quad \max \gamma_0 \\
 & \quad \quad \quad \text{Sous Contraintes} \\
 & \quad \quad \quad J \\
 1 \quad & \quad \quad v_{oj} y_{ij} \geq 1, \quad i = 1, 2, \dots, N \\
 & \quad \quad \quad j=1 \\
 & \quad \quad \quad J \\
 2 \quad & \quad \quad v_{oj} y_{oj} = SMF_o \\
 & \quad \quad \quad j=1 \\
 3 \quad & v_{oj} \geq 0, j = 1, 2, \dots, J \\
 4 \quad & v_{oj} y_{oj} \geq \gamma_0, j = 1, 2, \dots, J \\
 5 \quad & \gamma_0 \geq 0
 \end{aligned}$$

$$\mathbf{P2B} \quad \min \vartheta_0$$

Sous Contraintes

$$1 \quad \sum_{j=1}^J v_{oj} y_{ij} \geq 1, \quad i = 1, 2, \dots, N$$

$$2 \quad \sum_{j=1}^J v_{oj} y_{oj} = SMF_o$$

$$3 \quad v_{oj} \geq 0, j = 1, 2, \dots, J$$

$$4 \quad v_{oj} y_{oj} \geq \gamma^*, j = 1, 2, \dots, J; \quad \gamma^* \text{ solution optimale de (P2A)}$$

$$5 \quad v_{oj} y_{oj} \leq \vartheta_0, j = 1, 2, \dots, J$$

Une fois les différents poids générés, **Chen (2011)** propose de calculer des poids moyens uniques au sens le plus favorable et des poids moyens uniques au sens le moins favorable. Ces deux ensembles sont utilisés pour le calcul des scores des articles et procurent ainsi une base de comparaison commune. Ainsi, le score final n'est pas généré par construction d'un indice composite. Ce dernier étant jugé subjectif, et le choix de la constante arbitraire, **Chen (2011)** procède par maximisation de l'écartement en terme des valeurs des critères. Une seule combinaison de constantes de lissage est donc générée ne nécessitant aucune appréciation humaine.

Les méthodes basées sur la programmation linéaire ont été utilisées pour traiter de la problématique de classification multicritères des articles du stock. En regard des éléments que nous avons discutés, ces techniques présentent certains avantages. Le processus de classification ne requiert aucune intervention du manager ce qui a le mérite d'enlever la part de subjectivité propre au jugement humain. Ceci nous paraît aussi être très intéressant dans le contexte de développement de systèmes informatiques experts où les utilisateurs cherchent une simplification significative des outils et fonctionnalités proposés. Les opérations d'optimisation linéaire requises par les modèles proposés ne nous paraissent pas contraignantes. De multiples bibliothèques d'optimisation mathématiques sont accessibles et permettent d'opérer ce genre de calcul. Parmi ces bibliothèques nous citons COIN-OR, GIPALS, CPLEX ou encore LP SOLVE. Le temps de traitement ne constitue pas non plus un grand enjeu de notre point de vue. La classification étant une fonction peu fréquente (à

lancer généralement annuellement), nous estimons acceptables les longs temps de calcul pouvant avoir lieu dans le cas de grandes plateformes de stockage, d'autant plus qu'il est toujours d'usage de déclencher ces traitements périodiques en parallèle ou pendant des périodes vacantes de non-utilisation du système.

2.3 Méthodes basées sur les méta-heuristiques et l'intelligence artificielle

Pour répondre à la problématique de classification des articles de stock et faciliter l'intégration de multiples critères, certains travaux ont adopté l'approche basée sur les méta-heuristiques, comme les algorithmes génétiques ou encore des méthodes issues de l'intelligence artificielle, comme les réseaux de neurones artificiels.

Certains travaux ont adopté un algorithme génétique pour traiter du problème de classification du stock. Inspirée par la génétique des populations, ces algorithmes ont été adaptés à divers problèmes scientifiques d'optimisation, ordonnancement, reconnaissance, etc. L'idée générale étant de faire évoluer une structure de connaissance (chromosomes) représentant des solutions potentielles au problème traité. Le principe peut être résumé en trois principales étapes. La première étape est la sélection : Analogue à un processus de sélection naturelle, elle permet d'identifier les individus les mieux à même à se reproduire et d'éliminer ceux qui sont inadaptés. La deuxième étape est l'enjambement. Des croisements de chromosomes s'opèrent à ce niveau afin de donner lieu à de nouvelles structures de chromosomes. La probabilité d'apparition d'un croisement entre deux chromosomes est un paramètre de l'algorithme génétique. La troisième étape est la mutation. Un gène peut être substitué à un autre avec une faible probabilité de mutation, variant généralement entre 0.001 et 0.01, afin de ne pas tomber dans une recherche aléatoire et conserver le principe de sélection et d'évolution. Les structures résultantes, sont évaluées et réintroduites dans la population. Ce processus itératif continue jusqu'à détection d'une structure satisfaisante dans la population ou atteinte d'une limite itérative.

Guvenir & Erel (1998) ont adapté cette technique pour traiter de la classification multicritères des articles de stock. Ils ont comparé les résultats d'une variante de l'algorithme génétique à ceux obtenus par la méthode AHP en évaluant, pour chacune des deux techniques, le

nombre d'articles en communs avec une classification étalon opérée par un gestionnaire, sur la base d'un échantillon test. L'algorithme génétique s'est avéré plus précis.

Le problème de la classification des articles a été repris dans *Partovi & Anandarajan (2002)* par application des réseaux de neurones. Les réseaux de neurones artificiels font partie des techniques de l'intelligence artificielle ayant été expérimentées dans multiples domaines scientifiques et problèmes d'aide à la décision. Cet outil permet de simuler le jugement du manager/décideur et de reproduire sa logique décisionnelle en s'appuyant sur un apprentissage par l'expérience. Un réseau de neurones consiste en une structure hiérarchique composée d'un ensemble de couches. Chacune d'entre elles est composée d'un certain nombre de neurones alimentant, à travers des synapses, les neurones du niveau aval. Des poids synaptiques sont associés aux différentes synapses afin de permettre la pondération des entrées de chaque couche. Il s'agit de l'équivalent d'une matrice de transformation multipliant le vecteur d'entrée. Une fonction d'activation est appliquée à la somme pondérée et permet d'introduire une non-linéarité. Ces fonctions, appelées généralement fonctions de seuillage, permettent, selon le cas, une excitation ou une inhibition du neurone. Dans le premier cas, l'information est transmise aux neurones adjacents, dans le cas contraire, l'information est atténuée. Ce fonctionnement constitue donc une sorte de membrane permettant de filtrer et de transmettre uniquement l'information quand elle est jugée importante.

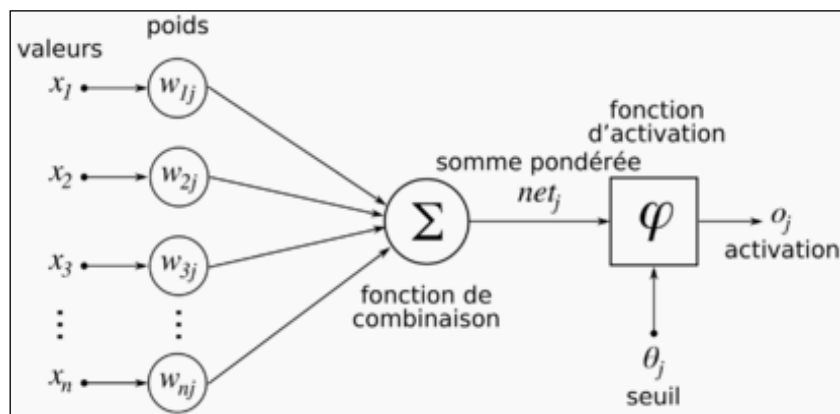


Figure 15 : Agrégation et transformation des entrées dans un réseau de neurones

Les réseaux de neurones se basent sur la notion d'apprentissage. Il s'agit d'un procédé par induction permettant de tirer une généralisation à partir d'observations limitées. Ceci se fait

par le moyen d'algorithmes d'entraînement permettant la modification des poids synaptiques en utilisant un jeu de données d'entrées. Une fois entraîné, le réseau est en capacité d'opérer une généralisation.

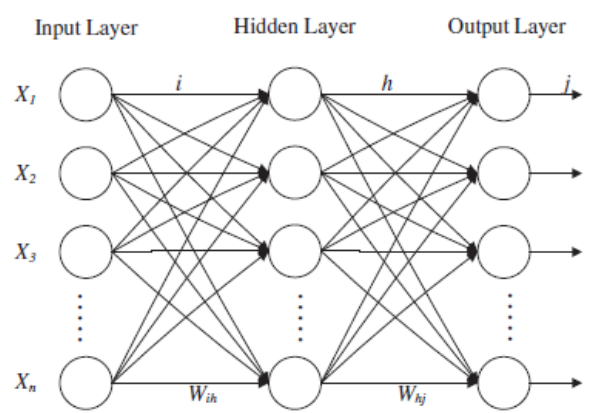


Figure 16 : Réseau de neurones à trois couches (Yu (2010))

Cette technique présente quelques avantages par rapport à certaines méthodes de régression ayant souvent été utilisées dans des champs d'application similaires. La première réside dans la capacité des réseaux de neurones à détecter et extraire la non-linéarité dans la relation entre les variables exogènes prises en considération. La deuxième est la capacité de cette méthodologie à générer des résultats précis indépendamment de la distribution des variables prises en compte.

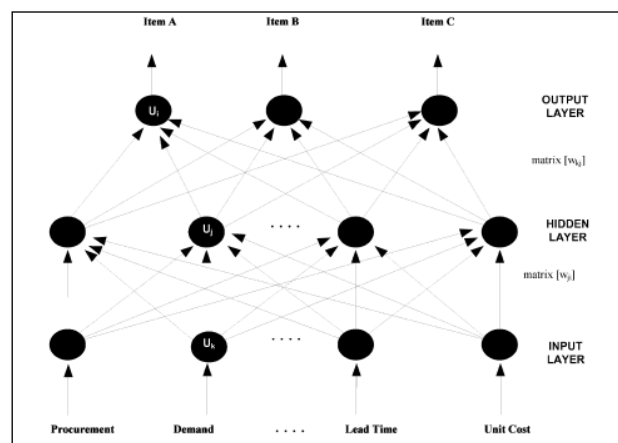


Figure 17 : Réseau de neurones appliqué pour le problème de classification multicritères (Partovi & Anandarajan (2002))

Considérant que les problèmes de classification des articles de stock sont non linéaires par nature, Partovi & Anandarajan (2002) ont considéré l'application d'un réseau de neurones

par introduction d'une couche intermédiaire. La sortie du problème étant continue et variant entre 0 et 1, les auteurs ont opté, conformément au travail de *Zahedi (1994)*, pour l'application d'une fonction de transfert sigmoïde, fonction non linéaire de la forme suivante.

$$f(w'x) = \frac{1}{1 + e^{-w'x}}$$

Deux algorithmes d'entraînement ont été considérés et comparés : la technique de rétro propagation du gradient et un algorithme génétique. Un échantillon de 96 articles, issus d'une entreprise pharmaceutique, a été considéré et quatre critères de classification ont été adoptés : le prix unitaire, le coût de commande, la demande annuelle et le temps de réapprovisionnement. L'échantillon a été divisé en deux sous parties : Une première classe, constituée de 50 produits étalons, ayant été classés au préalable par un gestionnaire, se basant sur sa propre connaissance, a servi pour l'entraînement du réseau. La deuxième partie constituée des 46 produits restants a été utilisée pour tester le réseau, une fois entraîné. Les deux algorithmes ont été ainsi confrontés et comparés également à la méthode **MDA**. La performance a été mesurée moyennant le niveau d'adéquation avec la classification étalon. Les réseaux de neurones se sont avérés plus précis et l'algorithme génétique a présenté les meilleures performances. Ce résultat a été déjà démontré et est conforme aux travaux entretenus dans *Doresey & Johnson (1998)*.

Les réseaux de neurones ont été appliqués plus récemment dans les travaux de *Yu (2010)* par adoption d'une fonction sigmoïde et utilisation de l'algorithme de rétro-propagation du gradient. *Yu (2010)* étend l'utilisation des outils de l'intelligence artificielle par application d'autres techniques usuelles et leur adaptation à la problématique de la classification multicritères. Ces méthodes comprennent, en plus du réseau de neurones, la technique des machines à vecteurs de support ainsi que la méthode des K - plus proches voisins.

Sur la base de 47 articles issus d'un stock pharmaceutique et en considération de quatre critères de classification, les mêmes que ceux utilisés précédemment, ces trois techniques ont été confrontées entre elles. Le comparatif s'est fait sur la base de quatre méthodes de benchmark : la démarche ABC classique telle que commentée dans *Reid (1987)*, la méthode AHP avec les poids considérés dans *Flores et al. (1992)*, le modèle d'optimisation présenté dans *Ramanathan (2006)* et le modèle d'optimisation adopté dans *Zhou & Fan (2007)*. La technique des machines à vecteurs de support s'est avérée la plus performante, elle permet d'aboutir aux résultats les plus précis, en regard des quatre méthodes étalons adoptées.

Les méthodes basées sur l'intelligence artificielle apportent une réponse intéressante à la problématique de classification multicritère des articles du stock. Plusieurs travaux ont en effet adopté cette voie et analysé leurs apports potentiels. Cependant, ces différentes méthodes nous semblent présenter certaines faiblesses. *Torabi et al. (2012)* estiment que ces méthodes ne sont pas tout à fait simples d'utilisation en pratique. Au-delà de la difficulté de vulgarisation de tels concepts auprès de managers et praticiens de la logistique, le processus de classification suivant une démarche d'apprentissage nous semble peu approprié au cas d'implémentation dans un système expert. Ceci peut paraître complexe et consommateur en temps. Les modules informatiques devant être de plus en plus simples et directs d'utilisation.

2.4 Récapitulatif : Avantages et inconvénients des méthodes de classification présentées

Nous reprenons dans le tableau 1, les principaux travaux ayant traité de la problématique multicritères des articles de stock, nous présentons les différentes techniques utilisées ainsi que les inconvénients et avantages de leur application en prenant en compte le cadre de développement de module informatique dans lequel nous nous situons.

Tableau 1 : Travaux traitant de la classification des articles du stock

Articles	Méthodologies de classification adoptées	Avantages	Inconvénients
<i>Flores et al. (1992)</i>	Méthode AHP	Simplicité de compréhension / Facilité d'implantation informatique	Méthode subjective (Subjectivité due au jugement humain)

Articles	Méthodologies de classification adoptées	Avantages	Inconvénients
<i>Guvenir & Erel (1998)</i>	Algorithme génétique	Algorithme génétique est plus précis que la méthode AHP (En comparant les deux méthodes sur la base d'une classification étalon effectuée par un expert)	Nécessité d'entraîner l'algorithme sur la base d'un échantillon étalon (Pas pratique pour le cas de systèmes experts)
<i>Partovi & Anandarajan(2002)</i>	Méthodes d'apprentissage : Réseaux de Neurones (Deux algorithmes d'entraînement)	Les réseaux de neurones sont plus précis que la méthode classique MDA (En comparant les deux méthodes sur la base d'une classification étalon effectuée par un expert)	Nécessité d'entraîner les réseaux de neurones sur la base d'un échantillon étalon (Pas pratique pour le cas de systèmes experts)
<i>Ramanathan (2006)</i>	Programmation linéaire : Maximisation d'une fonction score	Élimination de la subjectivité trouvée dans l'AHP / Méthode relativement simple de compréhension	Résolution requiert un solveur / Résolution d'un programme linéaire par article de stock (problématique pour le cas d'un grand nombre d'articles) / Certains articles, ayant un critère prépondérant, peuvent être classés, à tort, comme des articles de la classe A

Articles	Méthodologies de classification adoptées	Avantages	Inconvénients
<i>Zhou & Fan (2007)</i>	Programmation linéaire : optimisation de deux fonctions scores (Maximisation + Minimisation)	Élimination de la subjectivité trouvée dans l'AHP / Obtention d'un score plus équilibré que celui obtenu dans R. Ramanathan (2006) (Classification plus juste)	Nécessité d'utilisation de solveur (Résolution de deux programmes linéaires par article de stock)
<i>Ng (2007)</i>	Programmation linéaire : Maximisation d'une fonction score (Modification du programme linéaire proposé dans R. Ramanathan (2006))	Élimination d'une part de subjectivité trouvée dans l'AHP / La résolution du programme linéaire ne nécessite plus l'utilisation de solveur (Implantation informatique simple)	Un certain niveau de subjectivité (Introduit par les contraintes d'ordre supplémentaires) / La méthode peut entraîner une mauvaise classification des produits
<i>A. Hadi - Venech (2010)</i>	Programmation non linéaire : Maximisation d'une fonction score (Modification du programme proposée dans Ng (2007))	Élimination de la subjectivité trouvée dans l'AHP / Obtention d'un score dépendant des poids des critères et donc plus correct que celui obtenu dans Ng (2007)	Nécessité d'utilisation d'un solveur pour la résolution du programme non linéaire

Articles	Méthodologies de classification adoptées	Avantages	Inconvénients
<i>Yu (2010)</i>	Méthodes d'apprentissage : Réseaux de Neurones, Méthode SVM (Support Vector Machine), KNN (K - Nearest Neighbour)	Les trois méthodes d'apprentissage sont plus précises que la méthode classique MDA (En comparant les méthodes sur la base de quatre classifications étalons effectuées respectivement par une ABC classique, une classification AHP, le modèle de Ramanathan (2006) et le modèle de Ng (2007))	Nécessité d'entraîner les méthodes sur la base d'un échantillon étalon (Pas pratique pour le cas de systèmes experts)
<i>Chen (2011)</i>	Programmation linéaire : optimisation de deux fonctions scores et de 4 programmes complémentaires	Élimination de la subjectivité trouvée dans l'AHP / Obtention d'un score plus équilibré que celui obtenu dans Zhou & Fan (2007) (Classification plus juste)	Nécessité d'utilisation de solveur (Résolution de six programmes linéaires par article de stock)

La matrice bidimensionnelle présentée dans *Flores & Whybark (1986, 1987)* reste une méthode très rudimentaire et très limitée pour être appliquée dans le cadre d'un module informatisé de classification. L'AHP constitue, de notre point de vue, une alternative très intéressante. Cette technique a été souvent critiquée pour la part de subjectivité qu'elle introduit dans les résultats. Cependant, elle est très utilisée et a été implémentée dans plusieurs systèmes informatiques. La simplicité d'appréhension de la méthode, la simplicité de conceptualisation en une application interactive et la facilité d'implémentation informatique constituent des motivations quant à son implémentation. La comparaison deux à deux des critères étant l'unique étape nécessitant l'intervention humaine, le processus de comparaison reste relativement peu contraignant.

Les méthodes de l'intelligence artificielle nous paraissent peu appropriées au cas d'un module informatique de classification de stock à destination de pharmaciens et de praticiens de la logistique. Au-delà de leurs principes pouvant ne pas être directs et facilement appréhendables par des gestionnaires, l'application de ces techniques reste complexe.

L'entraînement des réseaux de neurones requiert au préalable la constitution d'un échantillon type ainsi que sa classification. Ce processus peut paraître long et fastidieux du point de vue des praticiens de la logistique qui sont généralement à la recherche d'outils simples supportant des démarches directes et peu contraignantes.

La famille des techniques de classification basées sur la programmation linéaire, nous semble d'un très grand intérêt. Ces méthodes présentent l'avantage d'éliminer complètement l'intervention humaine. À part le choix des critères, l'utilisateur est exempt de toute intervention. Ceci permet d'éliminer la subjectivité inhérente au jugement humain qu'on peut retrouver dans la méthode AHP. Ceci simplifie également le processus de classification, ce qui s'apprête parfaitement au cas d'un système de gestion informatisé. L'implémentation d'une approche de classification basée sur les techniques d'optimisation est aussi chose aisée. Les bibliothèques informatiques supportant les traitements d'optimisation sont abondantes. La contrainte du temps de calcul qui peut être long dans le cas de grandes plateformes de stockage supportant des milliers d'articles, n'en est pas une de notre point de vue.

3 Méthodes retenues et concept du module de classification des articles du stock

Nous présentons dans cette partie le module proposé de classification des articles. Notre module constitue un outil d'aide à la décision pour la classification des articles de stock avec la possibilité de prise en compte de l'aspect multicritères. Le module se base sur différents modèles de classification permettant d'adopter, soit une démarche monocritère classique, auquel cas le système se base sur le critère financier classique, soit une démarche multicritères, auquel cas l'utilisateur peut choisir parmi une liste les critères qu'il souhaite intégrer dans sa classification et définir la méthode d'intégration des critères.

3.1 Critères de classification

L'utilisateur va pouvoir choisir à partir d'une liste de critères prédéfinis ceux qui lui semblent les plus pertinents et qu'il souhaite intégrer dans la classification. De par la réalité industrielle de notre travail, les critères que nous avons fixés, et que le système propose, correspondent à ceux jugés pertinent, que l'on est capable de quantifier et dont la valeur est renseignée au niveau de la base de données. La base de données supportant le système de

gestion d'entrepôt est en effet conçue de façon à enregistrer un certain nombre d'informations logistiques utiles et nécessaires au fonctionnement des modules du WMS. Cette base est mise à niveau lors de l'installation du WMS, elle continue d'être alimentée au fur et à mesure, lors du fonctionnement de l'entrepôt.

En regardant de près les différentes tables constituant la base, nous avons défini les informations pouvant être récupérées pour être utilisées dans le module de classification. Actuellement, quatre critères sont accessibles :

- **Prix unitaire** : Renseigné dans la table des caractéristiques des articles, il s'agit du prix d'achat du produit auprès du fournisseur courant.

- **Consommation annuelle** : La consommation annuelle peut être reconstituée à partir de la table de l'historique des mouvements. À chaque livraison des services la quantité et la date correspondant aux articles et médicaments commandés sont enregistrées.

- **Temps de réapprovisionnement** : il correspond au délai s'écoulant entre le lancement de la commande par le WMS et l'arrivée des articles à l'entrepôt. Ce temps est négocié avec le fournisseur du produit et est enregistré dans la table des caractéristiques des produits.

- **Quota de commande** : C'est le multiple de commande que le fournisseur impose à l'entrepôt. Les quantités de commandes préconisées par le système sont systématiquement arrondies à ce multiple. Cette information est également enregistrée dans la table des caractéristiques des produits.

3.2 Méthodes de classification retenues

Sur la base des travaux de la littérature présentés et de l'investigation élaborée autour des différentes techniques de calcul, nous avons choisi de doter notre module par les techniques de classification suivantes :

- **La méthode de classification ABC classique** : Notre module permet d'adopter la méthode ABC classique. La méthode ABC classique reste une technique de classification référence que tous les pharmaciens logisticiens et les gestionnaires de stock peuvent utiliser. Ceci pourra se faire dans l'optique de la classification des articles du stock ou dans l'optique d'une

utilisation étalon, auquel cas les résultats serviront comme référence pour une comparaison avec des résultats issus de démarches multicritères.

- **La méthode AHP** : Notre module permet d'adopter la méthode AHP. C'est le premier processus de classification multicritères proposé à l'utilisateur. L'AHP se déploie suivant la démarche que nous avons commentée et que nous rappelons dans le schéma de la figure 18.

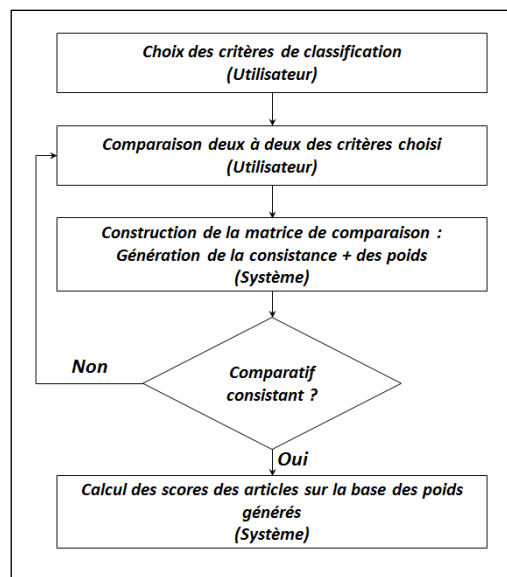


Figure 18: Processus de l'AHP

- **La méthode basée sur l'optimisation** : Notre module permet d'adopter la méthode basée sur les techniques d'optimisation linéaire présentées dans *Chen (2011)*. C'est le deuxième processus de classification multicritères proposé à l'utilisateur. Cette démarche de classification s'appuie sur la résolution des programmes linéaires présentés précédemment et épargne l'utilisateur de tout effort d'intervention. Il doit seulement définir, à partir de la liste des critères proposés, ceux qui lui semblent les plus pertinents et lancer le calcul.

3.3 Fonctionnement du module de classification des articles du stock

Le module de classification développé a été présenté dans *Jomaa et al. (2013. a)*. Le processus de classification proposé par notre module se déploie en quatre principales étapes comme schématisé dans la figure 19.

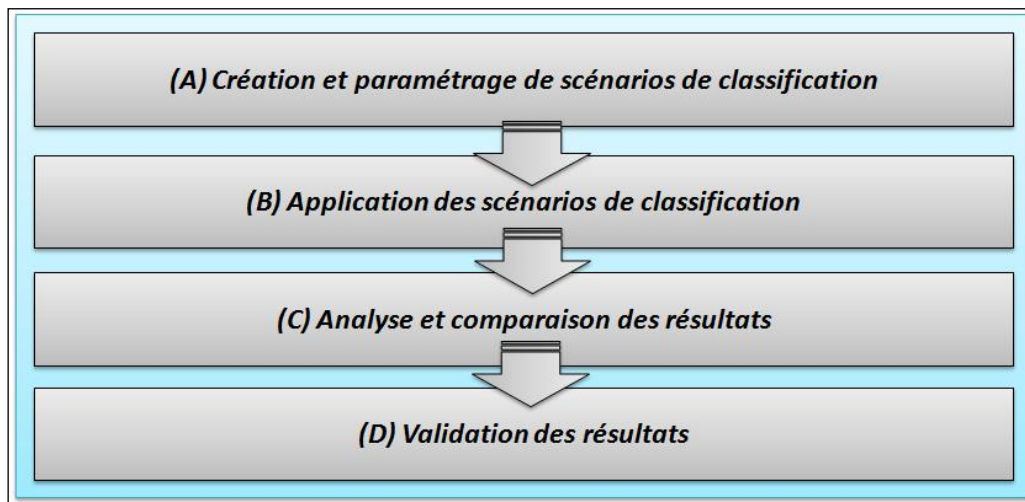


Figure 19 : Utilisation du module de classification

(A) Création et paramétrage de scénarios de classification

Le module permet à l'utilisateur de créer différents scénarios de classification, de les paramétrer et de les enregistrer afin de pouvoir les utiliser dans des classifications ultérieures. Un scénario de classification correspond à une combinaison d'une méthode de classification, choisie parmi les trois que le système permet d'appliquer, et un ensemble de critères que l'utilisateur choisit à partir d'une liste. Plusieurs combinaisons peuvent ainsi être créées :

- Adoption de la démarche classique : création de la méthode ABC monocritère.
- Adoption d'une démarche multicritères par appréciation, par l'utilisateur, des différents attributs de classification : création et paramétrage d'une méthode AHP.
- Adoption d'une démarche multicritères sans expression de préférences pour les attributs choisis : adoption de la méthode basée sur l'optimisation afin d'automatiser la pondération des critères définis.

Création d'un scénario ABC Classique

La méthode ABC classique constitue la première option possible. Aucun paramétrage n'est associé à ce choix. L'utilisateur peut seulement définir l'étendue de l'historique de consommation à prendre en compte lors des calculs, cet intervalle étant par défaut fixé à une année. Les calculs correspondant à ce scénario se feront systématiquement sur la base du

critère financier. Ce critère adopté correspond au produit de la consommation par le prix unitaire.

[21_MED] Création d'un scénario de classification

Création d'un scénario de classification

Nom de la classification
ABC MONO

Méthode de classification
☒ ABC Monocritère
☐ Multicritères Automatique
☐ Multicritères Expert

Période de calcul
Date de début 11/04/2012
Date de fin 11/04/2013

Etendue des classes
Classe A 0 % à 20 %
Classe B 20 % à 50 %
Classe C 50 % à 100 %

CLASSIFICATION ABC DES ARTICLES DU STOCK (CRITÈRE DE CLASSIFICATION : MONTANT ANNUEL CONSOMMÉ (EURO))

VALIDER QUITTER

Figure 20 : Création d'un scénario ABC classique

Création de scénarios multicritères automatiques

La deuxième option correspond à la création de scénarios de classification multicritères basés sur la méthode d'optimisation. Dans ce cas de figure, l'utilisateur est amené à définir, à partir de la liste, les critères qu'il souhaite intégrer dans sa classification (Interface de la figure 21). Ainsi, un premier scénario possible correspondrait au choix de cette méthode automatique, combinée à trois critères de classification. Un deuxième scénario potentiel correspondrait au choix de cette même méthode combinée à quatre critères de classification, etc.

Figure 21: Création d'un scénario multicritère automatique

Création de scénarios multicritères AHP

La troisième option possible correspond à la création de scénarios multicritères basés sur la méthode AHP. L'utilisateur choisit les critères qui lui conviennent à partir de la même liste. Les critères sont croisés deux à deux, par le système. L'utilisateur doit ainsi exprimer sa préférence sur la base de l'échelle d'importance proposée. Les poids des différents attributs ainsi que la consistance du comparatif, sont instantanément mis à jour, au fur et à mesure que le gestionnaire procède à la comparaison (Interface de la figure 22). Ceci simplifie énormément l'application du processus de l'AHP et donne une facilité d'intégration des informations synthétiques et utiles de cette approche. Plusieurs scénarios, autour de l'AHP, peuvent ainsi être créés, comprenant des critères différents ou intégrant les mêmes critères, mais avec des pondérations différentes.

[21_MED] Création d'un scenario de classification

Création d'un scenario de classification

Nom de la classification

AHP 1

Méthode de classification

☐ ABC Monocritère

☐ Multicritères Automatique

☒ Multicritères Expert

Période de calcul

Date de début : 10/05/2012

Date de fin : 10/05/2013

Etendue des classes

Classe A : 0 % à 20 %

Classe B : 20 % à 50 %

Classe C : 50 % à 100 %

CLASSIFICATION ABC MULTICRITÈRES DES ARTICLES DU STOCK (POSSIBILITÉ D'INTÉGRATION DE PLUSIEURS CRITÈRES) - MÉTHODE DE CLASSIFICATION EXPERTE (PRISE EN COMPTE DU JUGEMENT DU PHARMACIEN)

Choisir les critères de classification

CONSOMMATION	☒
TEMPS DE RÉAPPROVISIONNEMENT	☒
PRIX UNITAIRE	☒
QUOTA DE COMMANDE	☒

Poids des critères

CONSOMMATION	23%
TEMPS DE RÉAPPROVISIONNEMENT	15%
PRIX UNITAIRE	54%
QUOTA DE COMMANDE	6%

Consistance

☐

Comparaison des critères 2 à 2

CONSOMMATION	plus important que	TEMPS DE RÉAPPROVISIONNEMENT
CONSOMMATION	moins important que	PRIX UNITAIRE
CONSOMMATION	plus important que	QUOTA DE COMMANDE
TEMPS DE RÉAPPROVISIONNEMENT	beaucoup moins important que	PRIX UNITAIRE
TEMPS DE RÉAPPROVISIONNEMENT	beaucoup plus important que	QUOTA DE COMMANDE
PRIX UNITAIRE	plus important que	QUOTA DE COMMANDE

☒ VALIDER

☒ QUITTER

Figure 22 : Création d'un scénario multicritère expert

(B) Application des scénarios de classification

De multiples scénarios peuvent être créés et paramétrés comme présenté précédemment. Plusieurs questionnaires, d'une même structure, peuvent prendre part à ce processus en créant leurs propres modèles. Le système permet de stocker, de consulter, de modifier les différents scénarios et de les appliquer pour le calcul des classes ABC. L'application d'un scénario permet de générer les classes des articles selon la configuration du scénario en question, sans pour autant valider définitivement la classification. Le système permet d'appliquer trois scénarios à la fois. Il est ainsi possible de consulter les différents résultats, de les comparer entre eux ou par rapport à la classification courante, moyennant des outils d'analyse avancés, avant de mettre à jour les classes des articles stockés.

Classification / Criticité des articles									
Référence	Désignation	Classe actuelle	Score courant	Classe 1	Score 1	Classe 2	Score 2	Classe 3	
005881	AMIKACINE 250MG INJ FL - AMIKACI	B	0,001957	B	0,006408	B	0,001924	B	
000270	APOKINON 30MG INJ STYLO 3ML	B	0,001381	B	0,004442	B	0,001406	B	
000302	AMOXICILLINE/AC CLAVULANIQUE 1000MG IIA		0,023213	A	0,035833	A	0,02258	B	
000303	AMOXICILLINE/AC CLAVULANIQUE 2000MG IIA		0,005484	B	0,009156	A	0,005344	B	
002416	FLUMAZENIL 0,5MG INJ AMP 5ML	B	0,00218	B	0,00318	B	0,00215	B	
002427	FLUMAZENIL 1MG INJ AMP 10ML	B	0,002889	B	0,003937	B	0,002859	B	
002566	AMITRIPTYLINE 50MG INJ AMP 2ML - LAROX' B		0,001964	B	0,003688	B	0,001932	B	
004791	AMIKACINE 50MG INJ AM - AMIKACI	B	0,001199	B	0,006376	B	0,001197	B	

Scenari de classification									
Id classificati	Désignation	Créateur	Id méthode	Colonne de calcul	Date dernière modification	Date de validation	Numér	D	
13092001045	JF	JF	3		02/04/2013 11:30:34	02/04/2013 11:31:38	0	02	
13088001039	CLASSIFICATION MONO	DJ	1		29/03/2013 18:24:57		0	01	
13088001040	CLASS MULTI AUTO 1	DJ	2	CLASSE 1	29/03/2013 18:25:24		1	01	
13088001041	CLASS MULTI AUTO 2	DJ	2		29/03/2013 18:25:41		0	29	
13088001042	CLASS AHP 1	DJ	3	CLASSE 2	29/03/2013 18:26:24		2	01	
13088001043	CLASS AHP 2	DJ	3		29/03/2013 18:26:56	30/03/2013 17:07:11	0	01	
13089001044	SCÉNARIO CLASSI 1	DJ	2	CLASSE 3	30/03/2013 17:06:49		3	30	
13130001046	ABC MULTI 1	DJ	2		10/05/2013 14:47:09		0	01	
13130001047	AHP 1	DJ	3		10/05/2013 14:49:19		0	10	

Figure 23 : Application des scénarios de classification

(C) Analyse et comparaison des résultats

Le module permet une série de consultations avancées sur les résultats des différents scénarios. Ceci permet d'évaluer les résultats, de les confronter entre eux ou de les comparer avec la classification courante, afin d'avoir une appréciation globale avant de mettre à jour les classes des articles stockés. Il est ainsi possible d'évaluer le nombre d'articles en communs pour les différents scénarios et pour la classification courante, ce qui permet de tracer les produits transitant entre classes. Il est également possible de générer les proportions financières propres aux différentes classes d'un scénario donné. L'outil d'analyse est flexible et évolutif. Basé sur des requêtes SQL, il est possible de l'enrichir, par formulation de nouvelles requêtes afin de permettre la génération de statistiques particulières intéressantes tel ou tel gestionnaire.

(D) Validation des résultats

L'utilisateur peut enfin valider l'une d'entre les classifications lancées. La validation permet de mettre à jour les différentes classes propres aux articles suivant les résultats du scénario choisi.

4 Conclusion

Dans ce chapitre nous avons présenté le module de classification multicritères des articles du stock. En nous basant sur une revue détaillée des techniques multicritères issues de la littérature, nous avons défini celles qui nous paraissent s'apprêter le mieux à la classification des produits pharmaceutiques. La méthode AHP et une méthode basée sur la programmation linéaire ont été retenues.

Notre contribution est le développement d'un outil décisionnel d'aide à la classification multicritères autour de ces techniques, intégré au WMS. L'utilisateur peut choisir entre différents critères et opter pour plusieurs processus de classification. L'outil permet, à travers des fonctionnalités de tri et de consultations avancées d'analyser les différents résultats afin de valider la classification. L'intérêt de ce module est de permettre de bien cerner les articles pilotes du stock afin de mieux les analyser et les gérer. Nous nous intéressons dans le chapitre suivant au module de prévision des consommations.

Chapitre IV. Module de prévision

Dans ce chapitre, nous abordons le module de prévision des consommations. Ce module permet de calculer périodiquement les estimations de consommation des différents articles stockés par adoption d'une procédure dynamique s'adaptant à divers profils de consommations. Le module offre une panoplie de fonctions avancées de consultation et d'analyse permettant à l'utilisateur le suivi des évolutions ainsi que l'évaluation des performances prévisionnelles. Ce chapitre est organisé en trois sections.

Nous procédons, en premier lieu, à une revue synthétique de certains modèles de prévision couramment utilisés dans le domaine de la logistique. Nous présentons ensuite des travaux traitant du développement de systèmes prévisionnels experts et des travaux abordant la conception et la caractérisation de tels systèmes. En tenant compte de cette analyse et en intégrant les contraintes propres au contexte d'application dans lequel nous nous situons, nous définissons les modèles et techniques adaptés ainsi que le niveau d'automatisation des calculs qui nous semblent les plus adéquats.

Dans une deuxième section, nous détaillons l'architecture fonctionnelle développée en support à notre système de prévision. Nous abordons dans cette partie l'ensemble des procédures et modèles statistiques implémentés ainsi que l'interaction aboutissant aux estimations de consommations : procédure de catégorisation des données, procédure de nettoyage des valeurs aberrantes, procédure d'identification des tendances, modèles de prévision et paramétrage, calcul des indicateurs d'alerte et de performance.

En dernier lieu, nous présentons le concept développé autour de ces modèles de calculs et détaillons les fonctionnalités proposées à l'utilisateur.

1 Problématique de prévision : Revue de littérature sur les systèmes informatisés de prévision

1.1 L'importance des prévisions dans le contexte de la gestion de stock

La prévision de la demande ou de la consommation a été largement commentée dans la littérature. Cette pratique a été longuement considérée comme incontournable dans divers

domaines : économiques, logistiques, marketing, etc. Dans le contexte économique actuel, très perturbé et marqué par un accroissement de l'incertitude, la pratique prévisionnelle semble de plus en plus nécessaire. *Arinze (1992)* considère que la problématique de prévision reste d'actualité et que cette pratique est indispensable même en regard d'une configuration économique stable.

Dans le contexte de notre étude, la prévision des consommations revêt une importance particulière. Elle constitue l'un des leviers les plus conséquents, couramment déployés pour la maîtrise et la réduction des incertitudes propres aux opérations de planification et aux pratiques logistiques.

Plusieurs auteurs reviennent sur l'importance des prévisions, et notamment les prévisions à court et moyen terme, dans le contexte particulier de la gestion de stock et des réapprovisionnements. Parmi ces travaux, nous citons ceux de *Whybark (1972)*, *Winters (1960)*, *Silver et al. (1998)* ou encore *Pearce (1995)*. *Pearce (1995)* considère que la pratique des prévisions dans le contexte de la gestion de stock est assez particulière et comprend certaines caractéristiques propres. Le nombre élevé des produits stockés dans les entrepôts modernes, et notamment les plateformes hospitalières, constitue en effet une particularité. Ceci requiert une nécessaire automatisation des systèmes utilisés afin de permettre le calcul des prévisions, pour l'ensemble des produits, dans des temps maîtrisables tout en assurant une certaine fiabilité des calculs. L'outil informatique, le niveau d'automatisation ainsi que la puissance calculatoire paraissent fondamentaux. L'autre point souligné par *Pearce (1995)* est celui de la difficulté de choix entre les différentes techniques de prévisions. Le nombre d'articles en stock étant très élevé et les produits présentant une certaine disparité, les profils de consommation sont souvent multiples. L'utilisation d'une même procédure de prévision paraît clairement insuffisante et il est impératif d'adopter, dans ce cas, différentes techniques de prévision, afin de s'adapter au mieux aux typologies de consommation existantes. L'outil informatique d'aide à la décision, doté d'une certaine intelligence calculatoire, paraît encore une fois indispensable pour faciliter cette tâche et aider des gestionnaires, non nécessairement spécialisés ou familiarisés avec les techniques de prévision, dans l'estimation des consommations de leurs produits. La question des techniques de prévision à implémenter et l'automatisation de la sélection et des calculs est une question fondamentale.

Nous proposons, dans les sections suivantes, une revue synthétique des principales techniques de prévision couramment utilisées. Nous présentons, ensuite, certains des travaux ayant traité de la conception de systèmes prévisionnels experts. Nous définissons

enfin, sur la base de ces revues et en tenant compte de notre problème, les méthodes de calculs nous semblaient les plus appropriées ainsi que le niveau d'automatisation de notre système et permettant de répondre au mieux à notre besoin.

1.2 Techniques de prévision de la littérature

La prévision et son application dans les opérations logistiques sont parmi les sujets les plus commentés de la littérature. Nous proposons dans cette partie de faire une présentation concise et synthétique des principales familles et méthodes appliquées.

Dans le cadre de la problématique des prévisions, trois principales situations ont été identifiées par *Makridakis et al. (1998)* : une situation s'apprêtant à une approche prévisionnelle quantitative, pour laquelle des techniques de prévisions quantitatives ont été développées et adoptées, une situation s'apprêtant mieux à une démarche qualitative, pour laquelle des méthodes faisant intervenir l'expertise humaine sont utilisées et une troisième situation imprévisible.

Les méthodes de prévision quantitatives ont été développées pour répondre aux besoins de diverses disciplines. *Makridakis et al. (1998)* considère que cette classe de méthodes s'applique dans le cas où certaines conditions sont satisfaites :

- Des informations sur le passé doivent être disponibles ;
- Ces informations doivent être quantifiables : Données numériques ;
- Il est assumé que certaines caractéristiques du passé vont s'étendre dans le futur.

Dans le cadre de cette méthode, deux grandes familles d'approches peuvent encore être identifiées : les méthodes des séries temporelles et les méthodes explicatives.

1.2.1 Méthodes des séries temporelles

Les méthodes des séries temporelles s'appuient exclusivement sur les données historiques (consommation dans notre cas d'application) pour générer une estimation de la prévision. Les différentes méthodes développées considèrent la reconstitution de caractéristiques clés

propres à la série temporelle. La variation de la série est en effet caractérisée par certains attributs types que *Chatfield (1988)* a identifiés et résumés en quatre attributs :

- La tendance de la série ;
- L'effet saisonnier ;
- L'effet cyclique ;
- Autres fluctuations irrégulières non modélisables.

L'idée de cette approche consiste donc à modéliser ces caractéristiques afin de les extrapoler dans le futur. Il s'agit, en effet, de définir le modèle adéquat correspondant au mieux à la typologie d'évolution traitée. Pour cela, des techniques d'extrapolation, de complexités différentes, ont été construites, allant de la classe des méthodes du lissage exponentiel, relativement simple et facile de construction, à des modèles encore plus aboutis connus sous l'appellation des modèles ARMA, généralisant le lissage exponentiel et offrant un spectre de modélisation plus fin.

Méthodes du lissage exponentiel

Les méthodes du lissage exponentiel permettent de générer des prévisions sur la base d'une pondération des observations passées. Ces méthodes sont souvent perçues comme étant une extension de la technique de la moyenne mobile. Celle-ci permet d'appliquer une pondération uniforme de $1/k$ quand elle prend compte d'un horizon historique constitué de k observations.

$$\text{Moyenne Mobile} = \frac{\frac{\text{Mois courant}}{\text{Mois courant} - k} \text{Consommation}}{k}$$

Dans le cas du lissage exponentiel, le schéma de pondération est modifié de façon à favoriser les observations les plus récentes. Les poids introduits décroissent exponentiellement au fur et à mesure que les observations deviennent plus anciennes. Cette construction s'appuie sur l'hypothèse stipulant que les données les plus récentes permettent de mieux guider vers le futur, principe à la base de toutes les variantes du lissage exponentiel.

Exemple du lissage exponentiel simple

Le lissage exponentiel simple est le modèle le plus élémentaire de la famille des méthodes du lissage exponentiel. Couramment utilisée en pratique, cette technique permet de calculer la prévision d'une période future $T+1$ par ajustement de la prévision de la période passée T conformément à la formule ci-dessous.

$$P_{T+1} = P_T + \alpha(D_T - P_T)$$

La nouvelle prévision P_{T+1} est obtenue par addition de l'ancienne prévision P_T et d'une fraction α de l'erreur de prévision constatée, α étant une constante de lissage comprise entre 0 et 1. Une autre écriture couramment employée est celle de la formule suivante.

$$P_{T+1} = \alpha D_T + (1 - \alpha)P_T$$

La prévision de la période $T+1$ est la somme de la dernière observation (de la période T) pondérée par le coefficient de lissage α et de la dernière prévision pondérée par $1-\alpha$. Le développement de cette expression aboutit au schéma de pondération évoqué précédemment.

$$P_{T+1} = \alpha D_T + \alpha (1 - \alpha) D_{T-1} + \alpha (1 - \alpha)^2 D_{T-2} + \dots + \alpha (1 - \alpha)^{t-1} D_1 + \alpha (1 - \alpha)^t P_1$$

Le choix de la constante α permet de définir le niveau d'intégration des observations passées. Il est également à noter que cette technique n'est pas en capacité de bien reproduire les séries de données tendancielle, la prévision étant ajustée en continu par la dernière erreur constatée. D'où une utilisation accrue dans le cas de profils stationnaires.

Extensions du lissage exponentiel simple

Afin de prendre en compte des caractéristiques clés, comme la tendance et la composante saisonnière, des extensions au lissage exponentiel simple ont été élaborées. Ces extensions permettent de modéliser des séries historiques plus complexes.

Le modèle linéaire de Holt (**Holt (1957)**) permet d'introduire une estimation linéaire de la tendance, par adoption d'un schéma de pondération équivalent à celui utilisé précédemment. La prévision de la période future s'obtient ainsi par addition de deux composantes, estimées simultanément en fin de la période T : le niveau de la série de

données noté N et la composante tendancielle de la série notée b . Ce modèle se traduit par les équations présentées ci-dessous.

$$N_T = \alpha D_T + 1 - \alpha (N_{T-1} + b_{T-1})$$

$$b_T = \beta(N_T - N_{T-1}) + (1 - \beta)b_{T-1}$$

$$P_{T+1} = N_T + b_T$$

Le niveau N_T , calculé en fin de période T , s'obtient à partir de la dernière estimation de la pente b_{T-1} et de la dernière valeur lissée du niveau N_{T-1} , calculées en période $T-1$. La pente b_T permet de rattraper les décalages introduits dans les séries tendanciennes. L'équation tient compte du différentiel entre les deux dernières estimations de niveau ($N_T - N_{T-1}$). La pente est également lissée au moyen de la constante de lissage β . Cette construction est souvent appliquée à la place d'un lissage simple pour traiter des cas de séries de consommation tendanciennes. Cependant, elle reste insuffisante quand il s'agit de traiter de séries saisonnières. Une autre variante, extension du modèle de Holt, entretenue par **Winters (1960)** permet d'intégrer le phénomène saisonnier. Ce modèle est connu sous l'appellation de la méthode de **Holt-Winters** et introduit une troisième équation permettant l'estimation, en fin de période T , d'une composante saisonnière, moyennant le même mécanisme de lissage. La prévision de la période $T+1$ s'obtient en considérant trois composantes, conformément aux trois équations suivantes :

$$N_T = \alpha \frac{D_T}{S_{T-s}} + 1 - \alpha (N_{T-1} + b_{T-1})$$

$$b_T = \beta(N_T - N_{T-1}) + (1 - \beta)b_{T-1}$$

$$S_T = \gamma \frac{D}{N} + (1 - \gamma)S_{T-s}$$

$$P_{T+1} = N_T + b_T S_{T-s+1}$$

La troisième équation correspond à l'estimation d'un coefficient saisonnier, ratio de la demande en cours par la valeur lissée du niveau de la série. Ce ratio est modulé par un coefficient de lissage γ , conformément à la structure évoquée précédemment, en considération de la dernière estimation du coefficient saisonnier de la même saison S_{T-s} (décalé de $t - s$, s étant le décalage saisonnier de la série). Les deux premières équations correspondent aux estimations du niveau et de la pente. Le niveau étant opéré dans ce cas de figure, par rapport à la série désaisonnalisée, et non par rapport à la série brute.

Le schéma présenté dans le modèle précédent est multiplicatif, il est également possible de construire une autre variante par considération de coefficients saisonniers additifs. D'une façon plus générale, et en s'appuyant sur le même principe de pondération, il est possible d'aboutir à multiples variantes de la méthode du lissage exponentiel. Ceci se fait par introduction d'une tendance additive ou multiplicative, combinée à une saisonnalité additive ou multiplicative. Ce qui nous ramène, en considérant la totalité des combinaisons possibles, à potentiellement neuf variantes *Pegels (1969)*, offrant ainsi une large possibilité de modélisation de profils de consommation.

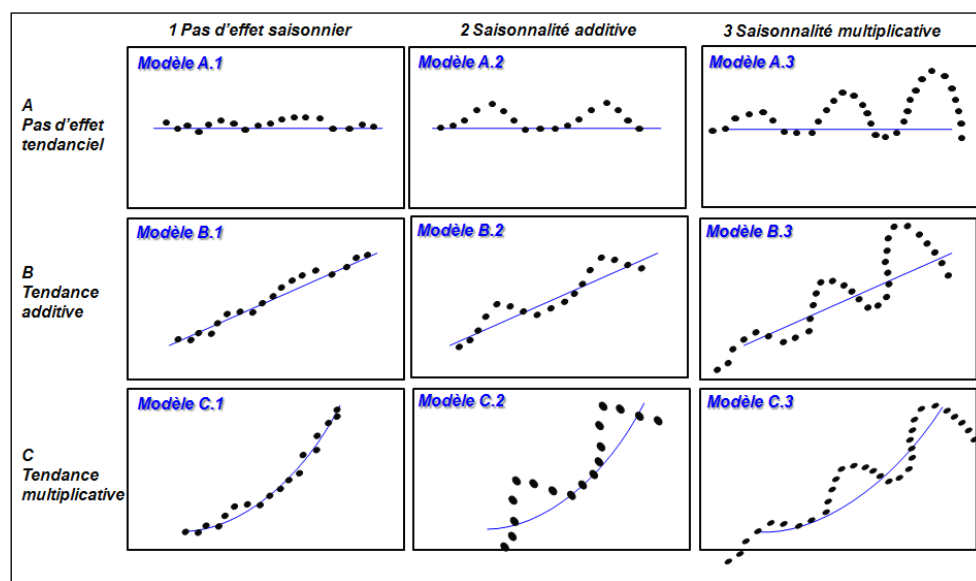


Figure 24 : Classification de Pegels - Différentes variantes du lissage exponentiel

Modèles ARIMA et méthodologie de Box – Jenkins

Les modèles ARMA ont été largement commentés dans la littérature et appliqués dans le domaine de l'analyse des séries temporelles et des prévisions. Comme dans le cas des techniques de lissage, cette approche de modélisation s'appuie, exclusivement sur l'analyse de la série de données. Cette approche est cependant plus complexe, mais pourrait aboutir à des modèles plus élaborés et donc potentiellement plus adaptés au cas de certaines séries de consommation. D'une façon générale, une demande historique D_t représentée par un processus ARMA peut être formulée de la façon suivante :

$$D_t = AR_t p + MA_t(q) + \varepsilon_t$$

$$AR_t = \alpha_1 D_{t-1} + \alpha_2 D_{t-2} + \dots + \alpha_p D_{t-p}$$

$$MA_t = \beta_1 \varepsilon_{t-1} + \beta_2 \varepsilon_{t-2} + \dots + \beta_q \varepsilon_{t-q}$$

$$D_t = \alpha_1 D_{t-1} + \alpha_2 D_{t-2} + \dots + \alpha_p D_{t-p} + \varepsilon_t + \beta_1 \varepsilon_{t-1} + \beta_2 \varepsilon_{t-2} + \dots + \beta_q \varepsilon_{t-q}$$

La série temporelle est exprimée comme étant la résultante de deux processus : Un processus **AR** (AutoRégressif) constitué par une combinaison linéaire finie des valeurs passées de la série et un processus **MA** (Moyenne Mobile) constitué d'une combinaison linéaire finie des valeurs passées d'un bruit blanc gaussien. Cette construction a été longuement analysée dans les travaux sur les séries de données univariées et a été considérée comme permettant la représentation d'une multitude de processus aléatoires stationnaires. *Anderson (1976)* estime qu'il est possible de reproduire un large spectre de séries chronologiques par adaptation de ces modèles. Ceci justifie l'intérêt pour l'application de cette approche dans divers domaines de modélisation et notamment les prévisions de consommations.

La construction des modèles **ARMA** s'appuie sur une analyse des fonctions statistiques d'autocorrélation et d'autocorrélations partielles propres à la série de données. Ces deux statistiques sont des indicateurs clés dans l'analyse des séries temporelles et traduisent le niveau de corrélation de la série de données avec elle-même à différents décalages de temps. Une investigation statistique approfondie autour de ces fonctions, que nous n'abordons pas dans cette thèse, permet de construire le modèle le plus adapté.

La condition de stationnarité évoquée ci-dessus est nécessaire quant à la bonne interprétation des autocorrélogrammes et la construction du bon modèle, ces derniers pouvant être pollués par une non-stationnarité de la série, due à l'existence de tendance ou de phénomènes saisonniers. Pour permettre l'intégration de ces composantes statistiques classiques, des extensions aux modèles ARMA ont été formulées comprenant les modèles ARIMA ou encore les modèles SARIMA (*Makridakis et al. (1998)*).

Ces variantes permettent d'une part, de rendre stationnaire un jeu de données, à la base non stationnaire (par effet d'une tendance, d'une saisonnalité ou de la combinaison des deux), au moyen de traitements statistiques (opération de différenciations) puis de rechercher le modèle ARMA correspondant à la série résultante. Les effets de tendance et de saisonnalité sont réintégrés dans les données par inversion des opérations de différenciations initialement appliquées. La finesse de modélisation, propre au modèle ARMA, est ainsi exploitée et étendue à diverses typologies.

Le processus d'application des modèles ARMA a été simplifié et formalisé par Box-Jenkins en une démarche, connue sous l'appellation de '*Méthode de Box Jenkins*', permettant la construction du modèle adéquat et son exploitation dans le calcul de la prévision.

Nous reprenons dans la figure 25 les principales étapes de cette démarche. La première étape, intitulée *Identification*, consiste à transformer la série, par opérations de différenciations, en une série stationnaire puis de définir, par interprétation des autocorrélogrammes, un modèle ARMA possible. La deuxième étape, intitulée *Estimation*, consiste à calculer les paramètres du modèle préconisé. Plusieurs techniques, que nous n'abordons pas, permettent le calcul des valeurs correspondant aux paramètres des processus AR et MA. Un test de résidus s'opère en fin de cette étape afin de confirmer la validité du choix opéré. Il est souvent nécessaire de repasser par la première étape afin d'opter pour une autre construction. La troisième et dernière étape, intitulée *Application*, consiste à appliquer le modèle construit dans l'estimation de l'évolution propre à un horizon futur.

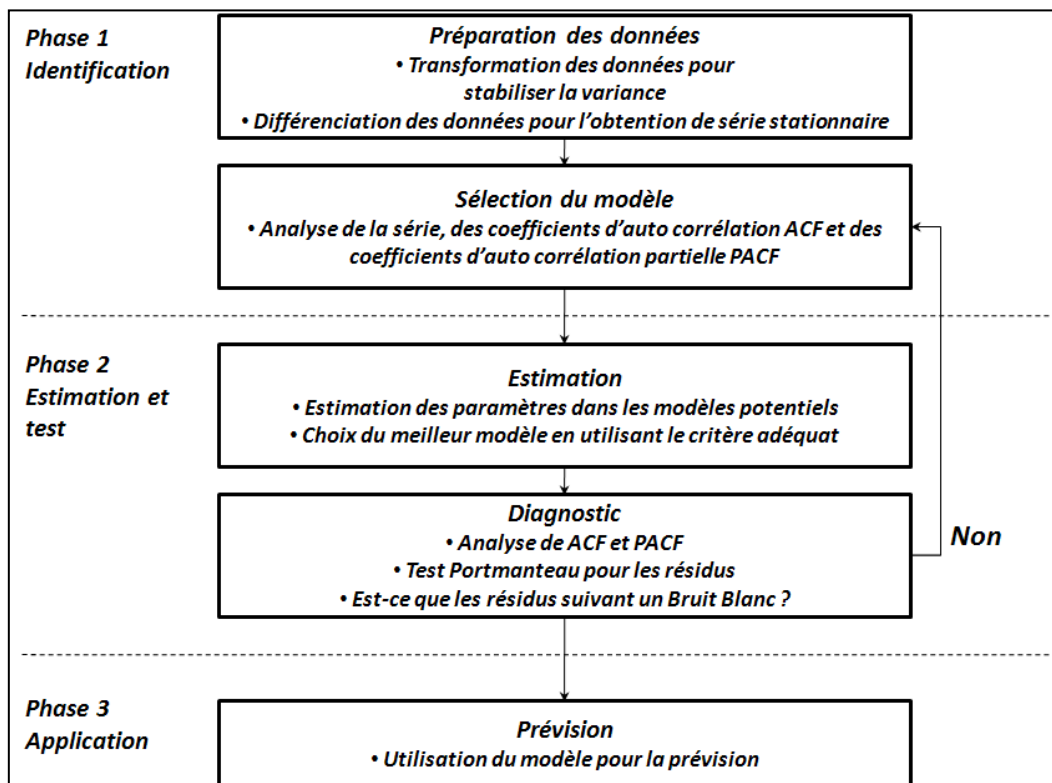


Figure 25 : Étapes de la méthode de Box-Jenkins (Makridakis et al. (1998))

La philosophie d'analyse et d'application de cette méthodologie peut être perçue comme équivalente à celle adoptée dans la première famille de méthodes. Dans le cas du lissage exponentiel, plusieurs variantes ont été développées afin de reproduire diverses typologies. Une analyse de l'évolution des consommations, moins formalisée que dans le cas des modèles ARMA, est à établir afin d'opter pour le modèle le plus intéressant. Nous retrouvons, en quelque sorte, une même démarche d'*identification - construction* dans les deux cas. Cette démarche doit permettre de traiter des caractéristiques statistiques communes, comprenant la tendance et la saisonnalité. La construction des modèles ARMA reste plus complexe, et requiert plus d'effort statistique, de test, de traitement et de construction. Les modèles en découlant sont également plus riches et permettent une modélisation plus fine.

1.2.2 Méthodes explicatives : l'approche causale

Les méthodes explicatives font partie des approches quantitatives appliquées dans les problèmes de modélisation et de prévision. Il s'agit, dans le cas de cette approche, de chercher à établir une relation entre l'évolution de la consommation (qui est la variable à prévoir, couramment appelée variable expliquée) et une ou plusieurs variables explicatives, autres que la série de données en elle-même, comme c'est le cas dans l'approche des séries temporelles présentée précédemment. La construction de ces modèles fait appel aux techniques de régression, dont la régression linéaire simple et la régression linéaire multiple. Dans le premier cas, la variable expliquée dépend d'une seule variable explicative donnant lieu à un modèle de la forme suivante :

$$D_t = a_0 + a_1X$$

Les coefficients a_0 et a_1 du modèle s'obtiennent, en général, par minimisation de critères d'écartement dont l'écart quadratique moyen. Dans le deuxième cas, la variable expliquée est reliée à plusieurs variables explicatives conformément au modèle suivant :

$$D_t = a_0 + a_1X_1 + \dots + a_nX_n$$

Makridakis et al. (1998) considèrent que ces deux cas, les plus communs de l'approche causale, font partie d'une démarche plus générale appelée la modélisation économétrique, permettant d'intégrer simultanément un ensemble d'équations de régressions multiples qui traduisent l'interdépendance entre les différentes variables. Toute la difficulté réside dans l'identification de l'étendue des interdépendances à intégrer dans la modélisation.

En pratique, plusieurs questions se posent lors de l'élaboration de ces modèles comprenant la définition des variables explicatives propres aux différentes équations, la détermination de la forme des fonctions (Linéaire, exponentielle, logarithmique, etc) ainsi que l'estimation des paramètres de ces dernières. Cependant, ces modèles complexes ne permettent pas nécessairement d'aboutir à des prévisions meilleures que celles obtenues moyennant les méthodes classiques des séries temporelles. Il s'agit plutôt d'un outil puissant pour l'évaluation et l'investigation de la dynamique des systèmes économiques plutôt qu'un outil prévisionnel *Makridakis et al. (1998)*.

1.2.3 Approche qualitative : Intervention humaine

Les jugements et appréciations humains sont souvent requis en entreprise dans divers domaines décisionnels. La prévision fait partie des opérations qui sont souvent entretenues par l'expertise métier. Plusieurs travaux de la littérature ont traité de l'efficacité du jugement humain dans la tâche prévisionnelle et de la portée de son intégration aux approches quantitatives précédemment discutées. *Pearce (1995)* revient sur certains d'entre eux, considérant qu'il est très difficile de définir la combinaison adéquate entre approches quantitative et qualitative, et que les résultats préconisés dans la littérature sont assez contradictoires. Des travaux comme ceux de *Armstrong (1985)* ou encore *Carbone (1983)* préconisent le non-ajustement des méthodes quantitatives par un jugement humain, leurs résultats ont démontré que des modifications opérées à la suite des calculs n'ont pas permis d'améliorer les précisions des prévisions. Dans d'autres travaux, les conclusions obtenues sont différentes et les précisions des calculs étaient améliorées après un ajustement statistique humain *Wolfe & Flores (1990)*. *Bunn & Wright (1991)* considèrent que le complément qualitatif peut être intéressant lorsque l'intervention humaine se fait par des experts ayant une bonne connaissance du contexte de l'application de la prévision. Cette réflexion a été confirmée dans *Sanders & Ritzman (2001)*, qui considèrent utile de faire intervenir l'humain du moment où ce dernier détient des informations complémentaires non

prises en compte par les modèles de calcul. Ces derniers précisent que dans le cas contraire, cette intervention risque de dégrader la qualité des résultats.

1.3 Systèmes experts pour la prévision - Déploiement de la théorie de prévision dans les systèmes d'aide à la décision

1.3.1 Systèmes experts prévisionnels

Nous nous intéressons dans cette thèse au déploiement de la théorie des prévisions dans un système informatisé. Nous revenons donc sur quelques travaux ayant traité de l'application et de l'implémentation des méthodes de prévisions dans des systèmes informatiques. Ceci nous permettra de mieux appréhender la structure générale de ce type de système, son évolution structurelle, les techniques et modèles souvent implémentés ainsi que l'étendue de l'implication de l'humain (Approche qualitative) dans le processus d'analyse à travers les fonctionnalités proposées. En nous appuyant sur ces éléments d'analyse, nous définissons les modèles de calculs ainsi que les caractéristiques de conception qui nous semblent les mieux appropriés au cas d'implémentation dans un WMS.

L'implémentation des techniques de prévision dans des systèmes informatiques à destination de praticiens de divers domaines, comprenant souvent les opérations logistiques, a fait l'objet de certains travaux de la littérature. *Jabbour et al. (1988)* ont considéré l'estimation de la consommation en électricité. Ils ont traité du cas d'un système informatisé calculant le taux de consommation électrique horaire sur un horizon de vingt-quatre heures. Leur système a permis de générer en quelques secondes des estimations qui se sont avérées, comparables, en termes de précision, à une approche purement humaine basée sur un jugement d'expert, mais requérant deux heures d'analyse. Le cas de la consommation électrique a été également considéré dans *Rahman & Bhatnagar (1988)*. Un système expert, utilisant une logique décisionnelle et un historique de données, a été mis en place, afin d'approcher la consommation électrique. Leur système renvoyait des résultats d'une précision légèrement supérieure à celle obtenue au moyen de techniques conventionnelles utilisées habituellement dans l'estimation de la consommation électrique.

Le système développé dans *Kwong & Chen (1988)* s'appuie sur un certain nombre de règles basées exclusivement sur des réponses et précisions que l'utilisateur est amené à fournir. Un fonctionnement similaire est retrouvé dans *Weitz (1986)*, un système prévisionnel intitulé Nostradamus a été développé. Ce dernier tient compte de l'intervention humaine, ainsi que de certaines règles permettant la recommandation d'une méthode de prévision appropriée au cas de la série traitée. Le processus reste long et requiert une interprétation de certaines caractéristiques statistiques, comprenant des autocorrélogrammes. Ceci suppose une certaine familiarité avec des concepts statistiques ainsi qu'une certaine expertise dans l'interprétation des séries temporelles ce qui peut être contraignant, lorsque les utilisateurs des systèmes prévisionnels ne sont pas nécessairement familiarisés avec les concepts statistiques.

Szmania & Surgent (1989) ont été parmi les premiers à avoir évoqué la nécessité d'éliminer, autant que possible, l'intervention humaine dans le processus prévisionnel afin d'accélérer le traitement de multiples séries en respect de contraintes temporelles de plus en plus restrictives. Leur problématique consistait à prévoir un grand nombre de séries de données, dans le cadre d'un réseau téléphonique, avec des temps de réponse très courts. Afin de répondre à cette situation, un système expert a été développé, utilisant une logique décisionnelle à base de règles (Si - Alors) permettant d'identifier certaines caractéristiques comprenant les valeurs accidentelles, les tendances, la saisonnalité puis d'appliquer un modèle prévisionnel issu de la famille des méthodes du lissage exponentiel.

L'élaboration de règles de sélection a été aussi évoquée dans *Collopy & Armstrong (1991)*, lesquels ont mené une enquête auprès de plusieurs experts de la prévision afin de définir des règles de choix à implémenter dans un système informatique. Aidé par l'intervention de l'utilisateur, leur système permet l'identification de caractéristiques statistiques, puis l'application de modèles prévisionnels, variantes du lissage exponentiel. Les auteurs estiment que l'amélioration de la procédure d'identification est toujours possible afin d'aller encore plus dans le sens de l'automatisation. Ceci peut se faire en effet, au moyen de procédures de tests statistiques facilement programmables.

Pearce (1995) et *Flore & Pearce (2000)* vont dans ce sens en développant et testant un système expert prévisionnel doté d'une procédure automatique d'identification de tendances et utilisé dans le cadre d'une plateforme de distribution. Leur système permet la correction des valeurs aberrantes, la détection des changements de niveau dans l'historique ainsi que l'identification de la tendance et de la saisonnalité. Des méthodes de tests statistiques ont été déployées afin de supporter ces différentes procédures. La prévision se fait dans leur cas, au

moyen de méthodes issues de la famille des séries temporelles, comprenant le lissage exponentiel simple et les méthodes saisonnières, et couvrant ainsi différentes typologies possibles.

1.3.2 Systèmes experts prévisionnels et niveau d'interactivité

Le niveau d'autonomie, devant être introduit dans un système prévisionnel par l'automatisation des calculs au moyen, notamment des méthodes quantitatives et des tests statistiques, reste une question difficile et un arbitrage assez complexe à opérer par les concepteurs de ces systèmes.

Il n'est pas anodin, dans certains cas, de faire intervenir l'utilisateur dans l'ajustement des résultats ou même dans le choix, voire le paramétrage, des techniques de calcul. La combinaison de l'approche quantitative et de l'intervention humaine a été souvent adoptée. Il est parfois intéressant de compléter le calcul prévisionnel opéré par le système par un jugement humain permettant de tenir compte d'autres connaissances auxquelles le système n'a pas eu accès. Cette combinaison doit être maniée avec beaucoup de prudence, car relève de l'apport en connaissance détenue par les praticiens et de sa pertinence. Pour cela, il faut prendre en considération le profil de l'utilisateur potentiel du logiciel et le contexte d'utilisation de l'applicatif. Ce niveau d'interactivité et de mise en participation de l'utilisateur dans un processus de prévision a été évoqué dans *Finlay & Marples (1997)*. Ces derniers font la distinction entre deux types de systèmes prévisionnels experts. Le premier type permettant à l'utilisateur une très grande interactivité en lui offrant la possibilité, entre autres, de choisir entre différentes techniques de prévision ou encore de procéder à leur paramétrage. Cette catégorie est désignée dans la littérature par '*Analytical - stochastic engineering*'. Dans la deuxième catégorie, le calcul prévisionnel est complètement automatisé, l'utilisateur n'intervient pas dans le choix ou l'ajustement des méthodes, des algorithmes de supports assurent cette tâche. Ce deuxième type est désigné dans la littérature par '*probabilistic - extrapolatory system*'. L'utilisateur peut intervenir dans l'ajustement des résultats, mais aucunement dans le choix et le paramétrage des routines de calcul.

Les progiciels de prévision professionnels proposés récemment sur le marché se situent entre ces deux catégories. Diverses offres logicielles sont disponibles (*FORECAST PRO, DEMAND SOLUTIONS, APERIA FORECASTER, PLANIPE, ETC.*), de niveau d'interactivité très variables, et allant de logiciels exclusivement spécialisés dans les prévisions à des offres

logistiques plus globales. Des fonctionnalités équivalentes sont retrouvées dans la plupart des cas et sont similaires à celles discutées dans la littérature sur les systèmes prévisionnels : nettoyage des données (Valeurs aberrantes), identification des tendances, méthodes de prévisions variées. Cependant, les techniques de prévisions proposées utilisent souvent des techniques de type lissage exponentiel.

1.3.3 Systèmes experts prévisionnels et conception optimale

La question de la conception optimale d'un programme informatique et du niveau d'interactivité adéquat qu'il convient d'introduire s'est souvent posée dans des travaux traitant de la conception des systèmes experts. Le cas particulier des systèmes prévisionnels est très peu commenté, *Adya & Lusk (2012)*. L'un des travaux ayant abordé cette problématique dans le but de définir les caractéristiques de conception aboutissant à un système de prévision optimal est celui de *Fildes & Goodwin (2006)*.

Fildes & Goodwin (2006) considèrent qu'un bon système de prévision doit assurer en priorité la tâche de stockage des données, le nettoyage des séries historiques, une constance dans la génération des estimations pour le cas de produits multiples (allant jusqu'à des milliers de produits) et la remontée des erreurs de prévisions. Afin d'aboutir à la structure logicielle la plus efficiente, intégrant ces différents points et assurant le compromis idéal entre les systèmes interactifs et les systèmes non interactifs, quatre dimensions principales, fondées à partir de la littérature, ont été avancées :

- *Manipulating effort -versus-perceived accuracy tradeoffs* : Des travaux de la littérature, sur la perception des technologies par les utilisateurs, insistent sur l'importance de la facilité d'utilisation afin d'améliorer l'acceptation de la solution technologique auprès des usagers *Davis (1989)*. *Todd & Benbasat (1999)* considèrent que l'automatisation des tâches est des plus pertinentes aux yeux des utilisateurs. Elle constitue la réponse idéale à cette question. Dans le cadre particulier d'un système prévisionnel, *Fildes & Goodwin (2006)* énumèrent un certain nombre de moyens permettant d'aller dans ce sens. L'automatisation des techniques de prévisions quantitatives avec la possibilité d'un ajustement qualitatif humain en fait partie. La facilité d'utilisation du système est aussi un point central. Ceci passerait, entre

autres, selon *Fildes & Beard (1992)* par le développement d'un outil permettant la correction de valeurs manquantes, la catégorisation des séries (nouveaux produits, demandes intermittentes, produits en fin de vie), l'introduction d'outils d'analyse statistiques, d'analyse des performances, des précisions prévisionnelles et comparaison des résultats ou encore la possibilité d'expérimentation instantanée des modèles de prévision.

- *Manipulating users' confidence* : l'étalonnage exprime l'étendue de la pondération humaine des résultats renvoyés par les calculs quantitatifs. Un mauvais étalonnage peut souvent avoir lieu et est, généralement généré par une confiance abusive accordée par les utilisateurs à leur perception subjective. La conception des systèmes experts prévisionnels peut, en effet, influencer ce paramètre. Une remontée d'informations non ciblée et peu synthétique pourrait résulter en un excès de confiance et entraîner par conséquent de mauvais calibrages. Le choix d'exprimer l'information sous forme numérique ou traduite en variable linguistique peut également avoir un impact sur la perception humaine et donc l'ajustement des résultats. Peu de travaux ont été réalisés dans l'analyse de l'impact de la conception des systèmes prévisionnels sur cette composante. *Fildes & Goodwin (2006)* reviennent sur les travaux de *Goodwin (2000)* ou *Lim & O'Connor(1995)*. Les premiers considèrent utile d'exiger auprès des prévisionnistes une justification de leur intervention dans le cas où un ajustement humain s'opère. Les derniers ont analysé cette configuration et ne constatent aucune atténuation après la mise en place de systèmes intégrant cette limitation.

- *Providing learning facilities* : Faciliter l'usage et l'apprentissage du système est une façon d'améliorer les performances de la tâche prévisionnelle. Ceci suppose un temps de manipulation et d'exploration du système par l'utilisateur qui pourrait être facilité par la mise en place de tutoriaux d'utilisation, ou d'options d'aide intégrée dans le logiciel. Un autre moyen qui permet de faciliter l'apprentissage est la remontée d'informations ou de feedback. *Benson & Onkal (1992)* considèrent que le feedback dans le cas des systèmes de prévision peut prendre diverses formes comprenant l'affichage de l'évolution historique des consommations, les précisions et performances des prévisions effectuées ou encore d'autres indicateurs statistiques : corrélation, variabilité, etc.

- *Fostering a sense of ownership of the forecasts* : L'une des façons d'améliorer l'acceptation d'un système prévisionnel, et donc d'améliorer nécessairement la performance prévisionnelle, est l'implication des utilisateurs dans la phase de conception du système. Ceci a été reporté dans certains travaux comprenant celui de *Alavi & Joachimsthaler (1992)* démontrant que l'implication de l'utilisateur est un facteur déterminant quant au succès du système développé. Le mauvais calibrage humain évoqué précédemment, et engendrant la dépréciation des résultats renvoyés par le système, pourrait être atténué en impliquant davantage l'utilisateur dans la phase de conception. Ce dernier étant impliqué dans le processus, plus de confiance serait accordée aux résultats. L'un des défis à remonter par les éditeurs de logiciel, et allant dans ce sens, consiste à réussir à impliquer davantage l'utilisateur dans le processus de prévision tout en l'orientant, en définitive, vers le choix de la technique la plus appropriée.

1.4 Techniques retenues

En regard de la revue de littérature élaborée et en tenant compte du contexte de notre application, nous présentons les choix que nous avons faits sur les caractéristiques fonctionnelles de notre système.

Le système étant développé en support à la gestion de plateformes pharmaceutiques, dont les produits gérés pouvant atteindre plusieurs milliers, il paraît évident de le doter d'une procédure de calcul quantitative avec un fort niveau d'autonomie dans la génération des calculs. Les gestionnaires ou pharmaciens utilisant une telle application n'ont, en général, aucune compétence particulière en interprétation des séries de données et dans les pratiques et méthodes de prévision. Le profil en question n'est pas non plus celui d'un prévisionniste. Les prévisionnistes n'ont pas nécessairement des compétences pointues en statistiques, mais sont familiarisés avec certains concepts et pratiques prévisionnelles leur permettant l'analyse des évolutions, la distinction des valeurs aberrantes, le choix et le paramétrage de certaines techniques en regard des profils de consommations analysées et l'interprétation de certaines statistiques élémentaires : moyenne, coefficient de variation, corrélations, etc. En considération de cette donnée, il nous paraît nécessaire d'épargner à l'utilisateur toute intervention au niveau du processus prévisionnel. La procédure de calcul doit permettre, en toute autonomie, l'estimation des consommations par adoption des modèles adéquats, et cela, sans demander aux gestionnaires de faire des analyses intermédiaires, des choix ou des

paramétrages de modèles. L'unique intervention envisagée, au niveau du processus prévisionnel, doit être la modification de la valeur finale de l'estimation. L'utilisateur doit être en capacité de faire cette modification. Nous revenons dans ce qui suit sur les modèles quantitatifs évoqués précédemment et regardons lesquels d'entre eux s'apprêtent au mieux à notre cadre : facilité de déploiement en un processus automatique et auto-adaptatif, facilité d'implémentation informatique, prise en compte de divers profils de consommation afin d'améliorer les performances.

Les modèles causaux et les méthodes économétriques ne nous semblent pas s'apprêter à notre cas d'application. Comme évoqué dans *Makridakis et al. (1998)*, l'inconvénient majeur de cette approche réside dans l'absence de règles bien formalisées permettant le déploiement de ces modèles dans diverses situations. La construction de tels modèles est en effet très dépendante du contexte d'utilisation. Notre système étant destiné à être utilisé par plusieurs clients, il paraît illusoire de tenter de mettre en œuvre un modèle générique s'adaptant aux diverses situations possibles, d'autant plus que ces modèles requièrent souvent un maintien et un suivi continu et ne sont pas de nature auto-adaptative. Des experts sont généralement sollicités dans l'analyse des données afin de définir les variables, générer les équations et estimer les paramètres optimaux.

Les méthodes basées sur les séries temporelles peuvent constituer une alternative intéressante dans notre cas d'application. Les approches évoquées précédemment, que ce soit celles des méthodes du lissage exponentiel ou de la méthode de *Box-Jenkins*, offrent à travers les diverses variantes, un large spectre de modélisation permettant de s'adapter à diverses typologies de consommation. La modélisation s'appuie exclusivement sur l'historique de consommation ce qui permet de simplifier la construction des modèles et leur implémentation dans une procédure dynamique autonome. L'historique de consommation est aussi une composante commune à toutes les plateformes équipées du WMS, des tables spécifiques sont supportées par la base de données standard et permettent la reconstitution de tous les flux de consommation depuis l'installation de l'application. Il s'agit de définir lesquelles des deux approches est la mieux appropriée, en regard des contraintes conditionnant notre cas d'application : amélioration des performances prévisionnelles, automatisation de la procédure et facilité d'implémentation informatique. Comme évoquée précédemment, la méthode de *Box-Jenkins* permet de couvrir un large spectre d'évolutions possibles tout en intégrant les cas des consommations tendanciellles et saisonnières. Le processus d'élaboration du modèle **ARMA**, le plus adapté à une série de données, reste

cependant très laborieux et requiert une multitude de procédures de transformations et de tests statistiques de validation. Dans la panoplie des progiciels de prévision trouvés sur le marché, rares sont ceux qui proposent des fonctionnalités autour du modèle ARMA. Des logiciels statistiques du domaine académique, notamment *SPSS*, simplifient la prévision au moyen de tels modèles. L'intervention de l'humain reste cependant nécessaire, notamment dans l'analyse des autocorrélogrammes pour la définition des ordres adéquats des processus. Une grande robustesse est donc exigée d'une procédure qui prétendrait automatiser intégralement cette méthodologie. À notre connaissance et depuis les travaux de *Box-Jenkins(1970)*, *Lu & Abourzik (2009)* sont les seuls à avoir proposé une automatisation intégrale de cette procédure d'estimation.

À cet égard, les variantes du lissage exponentiel restent plus simples d'exploitation, tout en donnant un spectre de modélisation assez large permettant de traiter des typologies de consommation typiques. Comme évoqué précédemment, des tests statistiques, relativement simples et facilement programmables, permettent de caractériser les attributs statistiques types (Tendance et saisonnalité) afin de mieux sélectionner les variantes à mettre en œuvre.

Sur la base de ces considérations, nous avons opté finalement pour l'implémentation de méthodes prévisionnelles issues de la famille des techniques du lissage exponentiel. Nous mettons en place une procédure dynamique, basée sur des tests statistiques de la série des consommations afin de faire le choix adéquat du modèle de prévision. Nous détaillons dans le paragraphe qui suit le fonctionnement de cette procédure : tests statistiques implémentés, variantes retenues, paramétrages des techniques de prévision.

2 Système de prévision proposé

Nous abordons dans ce chapitre l'architecture détaillée du système de prévision développé. Nous présentons l'ensemble des routines et modèles de calcul en support aux traitements prévisionnels. Nous reprenons dans le schéma de la figure 26, l'organigramme fonctionnel du système de prévision proposé.

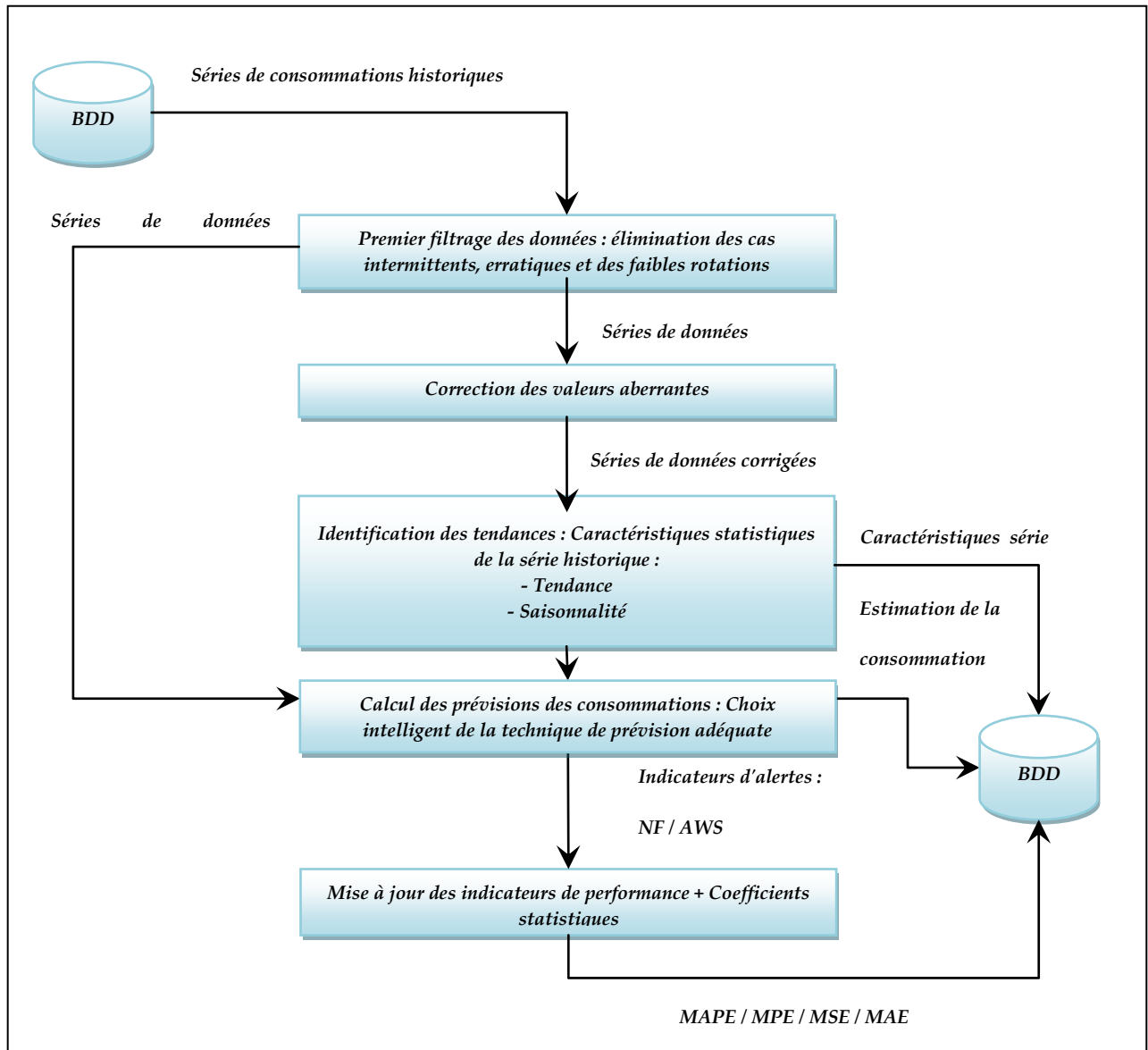


Figure 26 : Système de prévision développé

2.1 Premier filtrage des données

La multiplicité des articles gérés par les plateformes pharmaceutiques résulte en une grande variété de profils de consommation. Certains produits sont ainsi moins sollicités que d'autres ce qui peut entraîner des consommations atypiques : cas intermittents, à faibles rotation ou à forte variabilité. Il convient, en première étape, de séparer ces cas atypiques du restant des produits.

Le système opère un premier traitement des articles sur la base de trois attributs statistiques : la rotation, l'intermittence et le niveau de variabilité intrinsèque de la série historique. Cette première catégorisation des données est parfois adoptée dans les systèmes de prévision et de gestion de stock afin d'aider à la bonne affectation des techniques d'estimation de la

consommation ou le choix des politiques de gestion. *Boylan et al. (2008)* ont traité du cas d'un logiciel de gestion de stock opérant des calculs prévisionnels et choisissant des politiques de gestion sur la base de ces trois caractéristiques clés. Notre système va lancer cette procédure à un rythme annuel et procède aux calculs et classifications de la façon suivante :

- *L'intermittence* : Un produit est dit intermittent quand sa demande est non fréquente. Pour vérifier cette caractéristique, le système calcule le nombre de périodes à demande nulle en considération de la dernière année de l'historique. À partir d'une limite paramétrable, que nous fixons à 7 périodes comme dans *Boylan et al. (2008)*, la demande sera considérée comme étant intermittente. Ce choix reste arbitraire et peut être modifié.

- *La variabilité de la série* : L'un des attributs clés utilisés dans la mesure de la variabilité d'une série de données est le coefficient de variation CV qui est le rapport de l'écart type de la série par sa moyenne arithmétique, exprimé en général sous forme de pourcentage et donné par la formule suivante :

$$CV = \frac{\text{Ecart – Type (Demande)}}{\text{Moyenne (Demande)}} * 100$$

Nous appliquons, par défaut, la valeur limite et arbitraire, de 70% adoptée dans *Syntetos et al. (2005)*.

- *La rotation* : La définition des demandes à faibles rotations se fait aussi de façon arbitraire *Liang (1997)*. Dans le cas où la demande de la dernière année est en dessous d'un niveau (que nous fixons par défaut à 20 unités), nous considérons que le produit est à faible rotation.

Les trois caractéristiques sont recalculées et mises à jour annuellement. Il suffit que l'une d'entre elles soit vérifiée (demande identifiée comme étant intermittente et/ou demande identifiée comme étant erratique et/ou demande à faible rotation) pour que l'article soit isolé. Aucun traitement statistique complémentaire ne sera appliqué dans ce cas de figure, à part l'estimation de la consommation qui s'opère au moyen d'une moyenne mobile glissante d'ordre douze. Certains travaux de la littérature ont traité de la problématique d'estimation

des profils de consommation à faibles statistiques et notamment les cas intermittents. Des variantes du lissage exponentiel ont été adaptées à cette fin, nous citons notamment le modèle de *Syntetos & Boylan (2001)*. Certains travaux ont poussé l'analyse de l'adéquation de ces extensions en regard de caractéristiques statistiques incluant celles que nous avons choisies. L'adoption de la moyenne mobile nous permet de simplifier le mode opératoire pour le cas de ces profils spécifiques de consommation.

Par ailleurs, si la consommation historique est en dehors de ces trois cas (l'article en cours n'est ni intermittent, ni erratique et ni à faible rotation) un traitement statistique complémentaire sera appliqué avant de procéder à l'estimation de la consommation, laquelle se fera par un choix pertinent à partir d'une panoplie de techniques de prévision. Ce traitement complémentaire comprend une procédure de correction des valeurs extrêmes suivie de l'identification des caractéristiques clés de la série de données.

2.2 Correction des valeurs extrêmes

La procédure d'identification des valeurs aberrantes va permettre de détecter et de corriger les observations dont la probabilité d'avoir été produites par le même processus générant le restant des données est très faible. Les valeurs aberrantes résultent, en général, de certaines erreurs de saisie et d'enregistrement ou de certains facteurs influant la demande, mais dont la probabilité de reproduction est très faible. La correction de la série peut s'avérer nécessaire pour une bonne identification des caractéristiques statistiques, opérations qui vont suivre, et donc pour une bonne estimation de la prévision de la consommation, cette dernière étant modulée par la procédure d'identification.

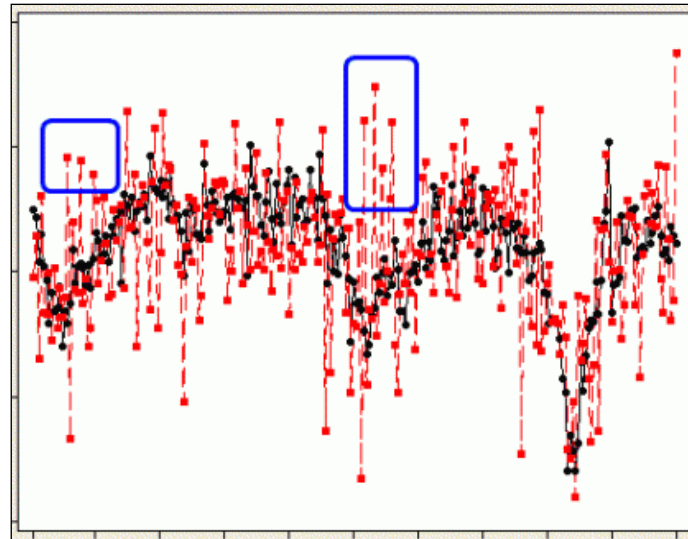


Figure 27 : Valeurs extrêmes dans une série de données

- *Cas de la tendance* : Comme nous le présenterons dans la section qui suit, l'identification de la tendance se fera par ajustement d'un modèle de données à la série de consommations historiques. La présence de valeurs extrêmes, peut en effet, induire une mauvaise interprétation du modèle.
- *Cas de la saisonnalité* : Comme nous le présenterons dans la section suivante, la saisonnalité sera identifiée moyennant une analyse de corrélation des observations de la série de données à certains décalages. La présence de valeurs aberrantes peut cacher cette corrélation.

La littérature sur l'identification des valeurs extrêmes est très abondante. Nous adoptons pour le cas de notre système, l'une des procédures les plus courantes et des plus simples d'utilisation : la méthode basée sur l'approche statistique. Couramment connue sous le nom de la méthode de 'la carte de contrôle', cette technique a été présentée, entre autres, dans **Han & Kamber (2001)** ou encore **Huang et al. (2006)**. L'idée consiste à faire une hypothèse sur la distribution de probabilité de la série de données en question, puis de procéder à l'identification des valeurs aberrantes moyennant un test de discordance. La distribution la plus couramment utilisée est la loi Normale. C'est la distribution que nous adoptons dans notre cas de figure. Le test statistique de discordance considère deux hypothèses. L'hypothèse H supposant que l'ensemble des données provient de la même distribution F et l'hypothèse alternative H supposant que les données proviennent d'un autre modèle de distribution G .

$$H : x_i \in F, i = 1, 2, \dots, n$$

$$H : x_i \in G, i = 1, 2, \dots, n$$

L'hypothèse H est acceptée ou réfutée moyennant une statistique de test et en considération d'un niveau de significativité α . Le test de discordance permet de vérifier si un élément x_i est significativement plus grand ou plus petit par rapport aux éléments de la distribution F . En considérant que la série de données X suit une loi normale de moyenne μ et d'écart type σ , $X \sim N(\mu, \sigma^2)$, la variable $X^* = (X - \mu)/\sigma$ suit une loi normale centrée réduite vérifiant $P \frac{X - \mu}{\sigma} = z_{\alpha/2}$. Un intervalle de confiance à $(1-\alpha)*100\%$ est de $[\mu - \sigma^*z_{\alpha/2}, \mu + \sigma^*z_{\alpha/2}]$, $z_{\alpha/2}$ étant la valeur critique correspondant au niveau de significativité $\alpha/2$. Selon cette équation, X appartient à l'intervalle $[\mu - \sigma^*z_{\alpha/2}, \mu + \sigma^*z_{\alpha/2}]$ avec un niveau de significativité α , ce qui signifie que la probabilité pour qu'un élément soit en dehors de l'intervalle est plus faible que $\alpha*100\%$. Une observation de la série historique se situant en dehors de cette plage constitue un évènement peu probable, l'observation en question sera considérée comme point aberrant.

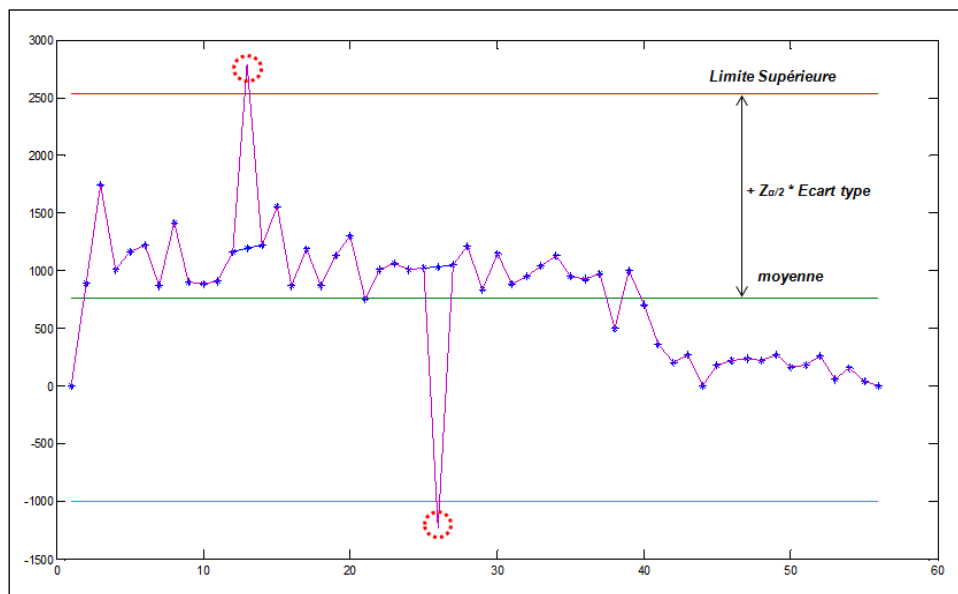


Figure 28: Carte de contrôle pour l'identification des valeurs aberrantes

La procédure d'identification implémentée est paramétrable. Il est possible à l'utilisateur de l'activer ou de la désactiver. Dans le cas où on choisit de l'activer, le système va appliquer mensuellement cette procédure et corriger systématiquement les valeurs identifiées comme étant extrêmes. La correction se fait en remplaçant la valeur jugée aberrante par l'observation

prévisionnelle calculée en début du mois. C'est l'une des techniques de correction adoptées dans la littérature et préconisées dans *Bourbonnais & Usumier (2013)*.

2.3 Identification des typologies de consommation

La procédure d'identification des tendances permet de caractériser les attributs clés de la série de consommations historiques : Tendance et saisonnalité. Cette caractérisation sert d'une part comme information, remontée auprès de l'utilisateur, et est utilisée d'autre part par le système pour opérer le choix pertinent de la technique de prévision la plus appropriée. La typologie de consommation d'un produit peut en effet varier au cours du temps, et par conséquent, une adaptation du modèle de prévision peut s'avérer nécessaire. Cette procédure est présentée dans *Jomaa et al. (2012)* et *Jomaa et al. (2013. b)*.

2.3.1 Identification de la tendance

Nous adoptons une règle de test équivalente à celle appliquée dans *Pearce (1995)*. La série de consommation historique est approchée par un modèle linéaire (Droite des moindres carrés ordinaire) et un test de significativité de la pente de la droite est effectué. Si le coefficient de la droite est significativement différent de zéro, la consommation est considérée comme étant tendancielle. Dans le cas contraire, la série n'est pas considérée comme tendancielle.

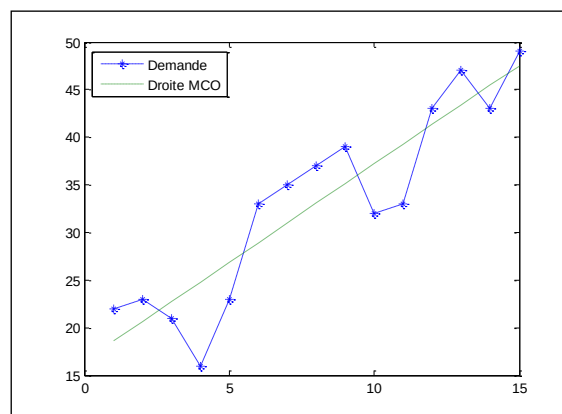


Figure 29 : Approximation par la droite des moindres carrés ordinaire

Le test de significativité de la pente consiste à vérifier l'influence réelle de la variable exogène (dans notre cas le temps) sur l'endogène (la consommation historique). Les deux hypothèses à confronter s'écrivent :

$$H_0 : a = 0$$

$$H_1 : a \neq 0$$

La statistique de test s'écrit :

$$t_a = \frac{a}{\sigma_a}$$

a étant l'estimation de la pente de la droite des moindres carrés approchant la consommation et σ étant la variance estimée du coefficient. Cette statistique suit une loi de *Student* à $(n - 2)$ degrés de liberté, n étant le nombre d'observations de la série historique de consommation. La région critique (de rejet de H_0) à un risque α donné s'écrit :

$$R.C : |t_a| > t_{1-\frac{\alpha}{2}}$$

$t_{1-\alpha/2}$ étant le quantile d'ordre $(1 - \alpha/2)$ de la loi de *Student*.

Le système calcule la statistique et opère le test R.C. Si l'hypothèse est vérifiée, H_0 est rejetée, la pente est significativement différente de zéro et donc la série est tendancielle. Dans le cas contraire, la série ne présente pas de tendance.

2.3.2 Identification de la saisonnalité

Coefficient d'autocorrélation

Makridakis et al. (1998) définissent la saisonnalité dans une série de données comme étant un motif qui se reproduit à des intervalles de temps fixes. Les ventes de mazout, par exemple, sont élevées pendant l'hiver et basses pendant l'été marquant ainsi une saisonnalité mensuelle. L'utilisation des antihistaminiques peut présenter une saisonnalité annuelle

marquée pendant le printemps et l'été et liée à la présence naturelle de pollen durant ces périodes.

Dans la littérature, une saisonnalité mensuelle correspond à un phénomène périodique se reproduisant chaque année à un décalage de douze mois. Le coefficient d'autocorrélation est une statistique clé souvent utilisée dans l'étude et l'identification des phénomènes périodiques. Nous considérons, dans le cas de notre système, une procédure d'identification basée sur cette statistique. Nous revenons en premier lieu sur la notion de corrélation avant de détailler la procédure en question.

Le coefficient d'autocorrélation d'une série temporelle à un décalage donné k permet de traduire le niveau de corrélation de la série avec la même série de données décalée de k observations. Pour une série de données D_t , la fonction d'autocorrélation $FAC(k)$ s'écrit :

$$FAC(k) = \frac{\sum_{t=k+1}^n (D_t - \bar{D})(D_{t-k} - \bar{D})}{\sum_{t=1}^n (D_t - \bar{D})^2}$$

$FAC(1)$ traduit la dépendance des éléments successifs de la série D , $FAC(2)$ traduit la dépendance des éléments de la série décalés de deux périodes de temps, et ainsi de suite. Une série de données de saisonnalité k , présente en général une valeur élevée du coefficient de corrélation à l'ordre k , traduisant la forte dépendance des observations à ce décalage.

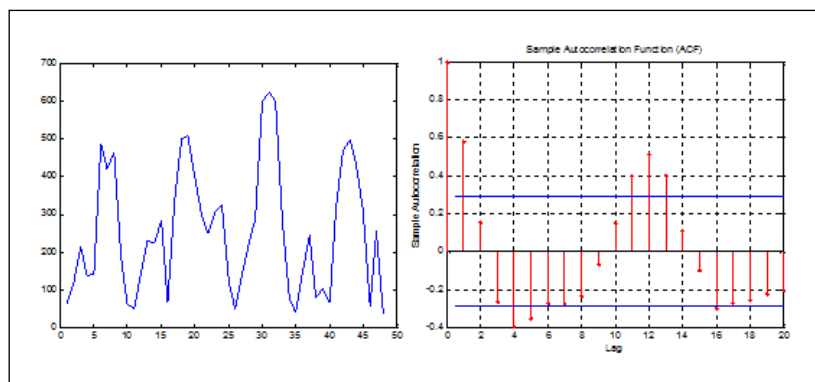


Figure 30 : Autocorrélogramme d'une série de données saisonnières

Les données que nous récupérons étant mensuelles, nous considérons lors de l'analyse de la saisonnalité un décalage de 12. Nous nous inspirons de la procédure évoquée dans *Pearce*

(1995), basée sur l'analyse du coefficient d'autocorrélation, et l'adaptions au cas de la maille mensuelle. Nous détaillons dans ce qui suit la procédure d'identification implémentée.

Procédure d'identification de la saisonnalité

Nous présentons dans l'organigramme de la figure 31, la procédure implémentée permettant de détecter l'existence de saisonnalité dans la série de consommation historique.

(a) Test de tendance et correction de la série

Le test de tendance, précédemment décrit, est systématiquement appliqué à la série historique. Dans le cas où la série présente une tendance, un premier traitement, consistant en une différenciation normale d'ordre 1, est appliqué afin d'ajuster la série. Les phénomènes tendanciels peuvent polluer les autocorrélogrammes des données, d'où la nécessité d'élimination de cette composante afin d'aboutir à des profils plus robustes. La différenciation est une technique statistique couramment appliquée permettant d'éliminer, à une certaine étendue, la non-stationnarité dans une série de données. La série résultante, appelée série différenciée, est obtenue par soustraction des éléments successifs de la série initiale.

$$D'_t = D_t - D_{t-1}$$

(b) Application du test de Ljung-Box

L'étape suivante consiste à opérer un test de *Ljung-Box* sur la série de données résultantes. (Différenciées ou non selon le cas). Le test de *Ljung-Box* fait partie de la famille connue sous l'appellation *Portmanteau tests*. Cette procédure permet d'évaluer la significativité statistique d'un ensemble de valeurs de la fonction d'autocorrélation d'une série temporelle. Ceci permet d'évaluer l'existence ou non d'une certaine dépendance dans les données. *Ljung & Box (1978)* considèrent un nombre arbitraire h de valeurs de la fonction d'autocorrélation d'une série donnée et construisent la statistique Q .

$$Q = n(n+2) \sum_{k=1}^h \frac{1}{n-k} \text{ACF}^2(k)$$

Cette variable a une distribution équivalente à une loi de Khi Deux. Il est démontré que si les données d'une série correspondent à un bruit blanc, la variable Q correspond à une loi de Khi Deux à h degrés de liberté. Si la valeur de Q se situe dans la région droite constituant les α % de la distribution de Khi Deux, il est conclu que les données représentent une certaine dépendance.

Si après l'application du test, il s'avère que la série ne présente pas de dépendance, le système conclut à l'inexistence d'un phénomène saisonnier. Dans le cas contraire, le système procède à la vérification de la douzième valeur de la fonction d'autocorrélation.

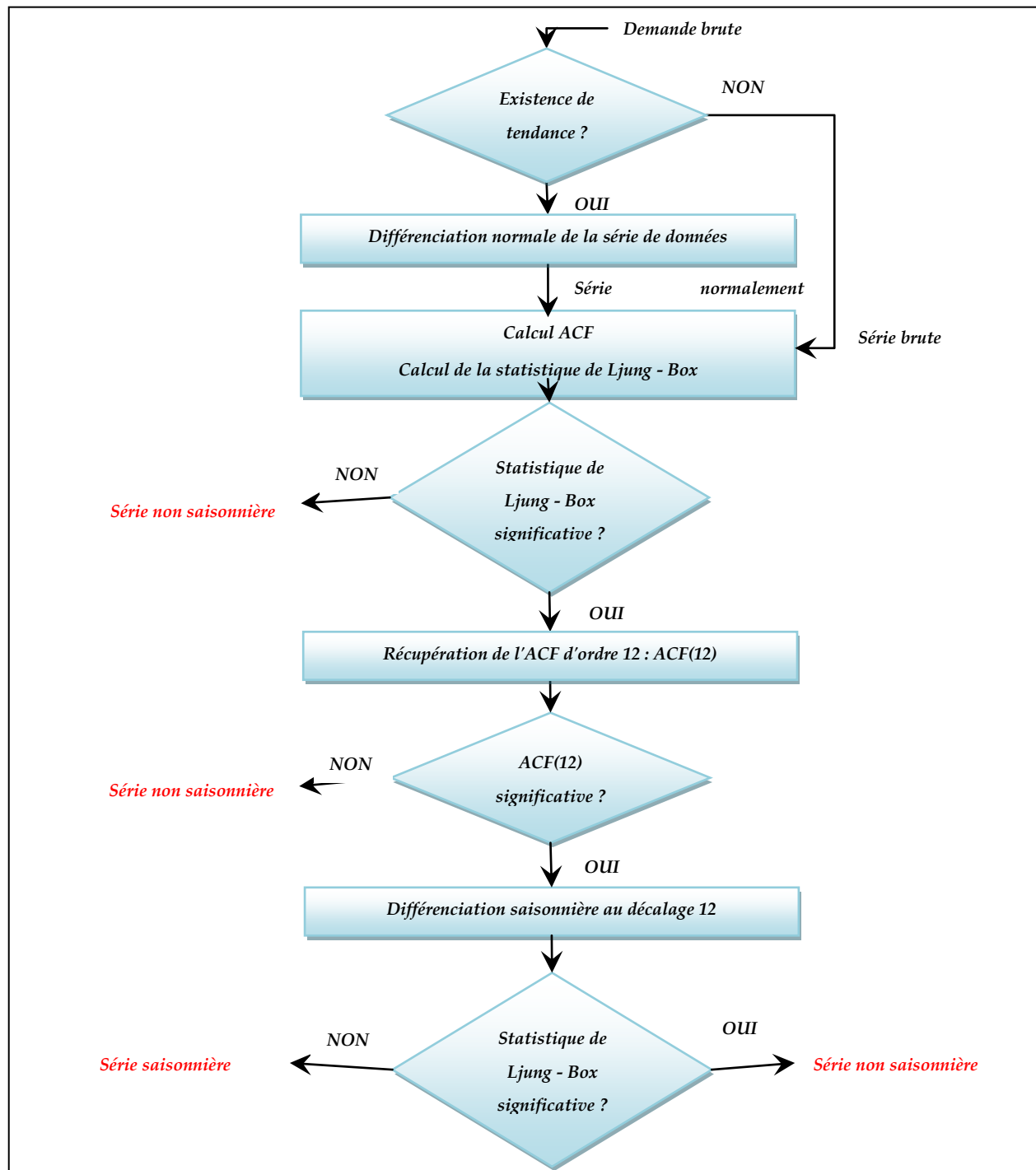


Figure 31 : Procédure d'identification de la saisonnalité

(c) Calcul et test de l'ACF à ordre 12

Dans le cas où les données sont dépendantes, le système va tester si la dépendance est due à un phénomène périodique mensuel. Ceci se fait par calcul et évaluation de la significativité statistique du coefficient d'autocorrélation d'ordre douze. Ce coefficient est comparé, en valeur absolue, à $\frac{1}{n}$, n étant le nombre d'observations de la série. S'il est en deçà de cette limite, le coefficient sera considéré comme statistiquement non significatif. Si le coefficient n'est pas significatif, le système conclut à l'inexistence d'un phénomène saisonnier. Si le coefficient est significativement différent de zéro, une différenciation saisonnière est appliquée à la série de données.

(d) Différenciation saisonnière et deuxième test de Ljung-Box

La différenciation saisonnière est une autre technique de traitement des séries temporelles permettant d'éliminer la non-stationnarité dans les séries de données par élimination des phénomènes périodiques. Une différenciation saisonnière d'ordre 12 consiste à appliquer à la série la différence entre ses éléments et les éléments décalés de 12 observations (une année). La série de données résultante se calcule comme suit :

$$D'_t = D_t - D_{t-12}$$

Si la consommation historique présentait réellement un phénomène saisonnier mensuel, l'application d'une différenciation d'ordre 12 éliminerait ce phénomène. Le système va calculer la série de données différenciée et réappliquer le test de dépendance de *Ljung-Box* sur la série résultante. Si le test est positif, cela signifie que l'opération de différenciation n'a pas éliminé la dépendance, supposée prévenir d'un phénomène saisonnier mensuel, dans ce cas le système conclut de l'inexistence d'un phénomène saisonnier. Dans le cas contraire, le système conclut qu'il existe un phénomène saisonnier mensuel.

2.4 Calcul des prévisions de consommation

2.4.1 Modèles de calcul implémentés

Suivant les considérations sur les attributs statistiques identifiés, le système va appliquer l'une d'entre les cinq techniques de prévision suivante.

Moyenne mobile

Une moyenne mobile d'ordre 12 est systématiquement appliquée pour le cas des articles à typologies atypiques et qui ont été isolés lors du premier filtrage des données.

Lissage exponentiel simple

Le lissage exponentiel simple est appliqué au cas des produits représentant des profils non tendanciels et non saisonniers. Les calculs se font conformément aux équations suivantes :

$$P_{t+1} = \alpha D_t + (1 - \alpha) P_t$$

$$P_{t+h} = P_{t+1}$$

P_{t+1} représente l'estimation de la consommation du mois (t+1) qui est calculée par le système en fin du mois t. D_t et P_t représentent respectivement la demande et la prévision du mois t qui vient de s'écouler. Le système initialise la prévision en considérant le premier mois de consommation :

$$P_1 = D_1$$

α est la constante de lissage, comprise entre 0 et 1 et ajustée par le système au moyen d'une procédure auto - adaptative que nous développons dans le paragraphe suivant.

Modèle linéaire de Holt

Le modèle linéaire de Holt est appliqué dans le cas de produits représentant une tendance. La prévision est obtenue par addition de deux composantes : une estimation du niveau de la consommation et une estimation de la pente. Le système applique cette technique conformément aux formules suivantes :

$$N_t = \alpha D_t + (1 - \alpha) (N_{t-1} + b_{t-1})$$

$$b_t = \beta (L_t - L_{t-1}) + (1 - \beta) b_{t-1}$$

$$P_{t+1} = N_t + b_t$$

$$P_{t+h} = N_t + h * b_t$$

L'initialisation des deux composantes se fait comme suit :

$$N_1 = D_1$$

$$b_1 = D_2 - D_1$$

Le calcul des constantes de lissages s'appuie sur la même procédure d'autorégulation.

Lissage exponentiel simple avec saisonnalité

Dans le cas des profils stationnaires présentant une saisonnalité, le système va appliquer un lissage simple en introduisant des coefficients de saisonnalité fixes. Les coefficients de saisonnalité sont calculés lors de l'identification des tendances, dans le cas où le profil en question s'avère saisonnier. Ils seront donc recalculés à chaque lancement de la procédure d'identification des tendances et non mis à jour mensuellement par une équation de lissage, comme c'est le cas dans le modèle de Holt-Winters.

Le système calcule pour ce genre de profil, la série de données débarrassée de sa composante saisonnière CSV, conformément à la formule suivante.

$$CSV_{mois\ t, année\ k} = \frac{D_{mois\ t, année\ k}}{S_{mois\ t}}$$

Un lissage simple est ensuite appliqué à la série de données conformément à la procédure présentée ci-dessus. Enfin, la prévision de la période suivante est estimée en réintroduisant le coefficient de saisonnalité du mois en question.

$$P_{mois\ t+1, année\ k} = CSV_{mois\ t+1, année\ k} * S_{mois\ t+1}$$

Les coefficients de saisonnalité sont calculés en utilisant l'approche du ratio à la moyenne mobile de la série *Makridakis et al. (1998)*. D'autres procédures de calcul ont été élaborées pour affiner l'estimation des coefficients saisonniers, la méthode *Census II*, commentée dans certains travaux, en est parmi les plus élaborées. *Makridakis et al. (1998)* a démontré que cette dernière n'aboutit pas nécessairement à des résultats prévisionnels plus précis.

Modèle linéaire de Holt avec saisonnalité

Ce modèle est appliqué pour le cas des profils présentant une certaine tendance et un phénomène saisonnier. La démarche est identique à celle du modèle précédent. Il s'agit seulement d'appliquer le modèle linéaire de Hot à la série désaisonnalisée, à la place du lissage exponentiel simple.

2.4.2 Calcul et ajustement des constantes de lissage

Nous revenons dans ce paragraphe sur la procédure d'ajustement des constantes de lissage des différents modèles utilisés. L'une des approches de calcul des coefficients de lissage, commentée dans la littérature et adoptée par certains éditeurs, est basée sur la recherche des valeurs permettant de minimiser un indicateur d'erreur. *Makridakis et al. (1998)* a présenté, pour le cas du lissage exponentiel simple, une démarche empirique basée sur la minimisation du carré moyen des écarts (*MSE*).

$$MSE = \frac{\sum_1^n (Demande - Pr\acute{e}vision)^2}{n}$$

Il s'agit dans ce cas de trouver la valeur optimale de la constante de lissage en recalculant l'indicateur *MSE* pour différents incréments du coefficient alpha allant de 0 à 1, le calcul étant fait en considérant une plage de données historiques. Des algorithmes d'optimisation non linéaires ont été également adoptés pour cette fin.

Cette procédure présente une insuffisance : l'optimalité étant atteinte sur des données historiques, il n'est pas garanti de continuer à satisfaire cette condition lors de l'introduction d'une nouvelle observation et calcul de prévision. Ceci constitue l'inconvénient majeur de cette approche. Nous choisissons, quant à nous, d'adopter une autre procédure de régulation, présentée par *Bourbonnais & Usunier (2013)* et basée sur l'utilisation de deux indicateurs : le *signal NF (Normal Forecast)* et le *signal AWS (Alert Warning Signal)*.

Signal NF (Normal Forecast)

Cet indicateur permet d'alerter sur une erreur de prévision, du mois courant, qui est anormalement élevée par rapports aux écarts de prévisions constatées au cours des derniers mois. Ceci peut être dû à l'effet d'une observation aberrante (Accidentelle), à une annulation brusque de la demande (qui pourrait correspondre à l'abandon brusque d'un article et qui se confirmerait dans les mois qui suivent) ou à tout autre phénomène non maîtrisable.

NF est l'équivalent d'une variable centrée réduite qui est calculée, à chaque période (Mois dans notre cas), par la formule suivante :

$$NF_t = \frac{EPS_t}{MAD_t}$$

EPS représentant un indicateur d'écart instantané et MAD étant l'indicateur de la valeur absolue moyenne des écarts. NF permet ainsi de traduire l'erreur observée à la période courante par rapport à la fluctuation moyenne de la série. Cet indicateur est en effet sensible aux observations anormales. Dans le cas où la consommation courante est anormalement élevée par rapport au reste de l'historique, l'erreur EPS sera élevée et de même l'indicateur NF.

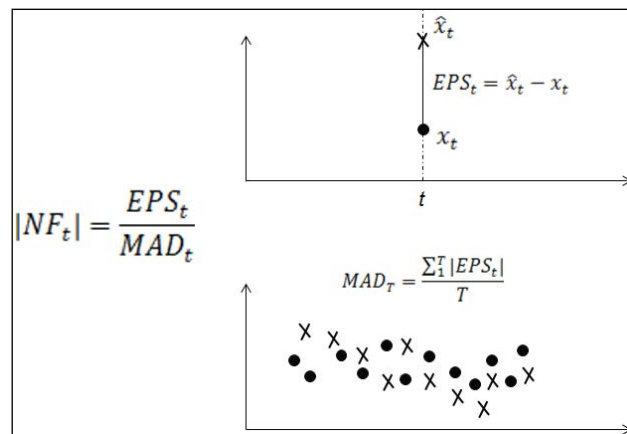


Figure 32 : Indicateur NF

Signal AWS

Cet indicateur permet d'alerter sur les cassures de tendances dans l'évolution de la consommation d'un article. Ceci peut traduire l'accroissement ou l'abandon progressif d'un produit ou encore la stabilisation de sa demande.

AWS est calculée suivant la formule suivante :

$$AWS_t = \frac{SUMEPS_t}{MAD_t}$$

SUMEPS représente un indicateur de la valeur cumulée des écarts (somme algébrique des écarts de prévision) calculé à chaque période sur la base de l'historique existant. Cet indicateur permet de renseigner sur un éventuel changement de tendance. Si les observations sont réparties régulièrement autour des prévisions, la valeur de SUMEPS est voisine à chaque instant de 0. Si elle croît en valeur, cela indique un changement de tendance ou de niveau.

L'indicateur NF étant spécifique aux valeurs anormales, une augmentation de NF supposerait une diminution de la constante de lissage (pour ne pas trop en tenir compte dans la prévision, l'observation étant accidentelle). L'indicateur AWS est propre aux changements de tendance, s'il augmente, il conviendrait d'élever la constante de lissage afin de rattraper l'inertie de la consommation, le changement n'étant pas accidentel, mais découlant d'une cassure de tendance. *Bourbonnais & Usunier (2013)* a proposé ensuite une procédure empirique permettant de réguler le choix de la constante de lissage alpha sur la base de ces deux considérations. Il convient de calculer le coefficient, en fin de période t, sur la base de l'équation empirique suivante :

$$\alpha_t = 0,7 - 0,1 * (NF_t + AWS_t)$$

Notre système calcule à chaque mois les valeurs de ces indicateurs et ajuste les constantes de lissages conformément à la procédure décrite ci-dessus.

Le coefficient de lissage *beta*, propre à la tendance, se déduit à partir du coefficient *alpha* en appliquant.

$$\beta_t = \alpha_t - 0,05$$

Il s'agit d'être légèrement moins réactif sur la tendance que sur la moyenne.

Le coefficient de lissage γ , pour le cas des modèles saisonniers, est fixé à 0,2. La rupture de tendance et les valeurs anormales n'affectent pas la composante saisonnière.

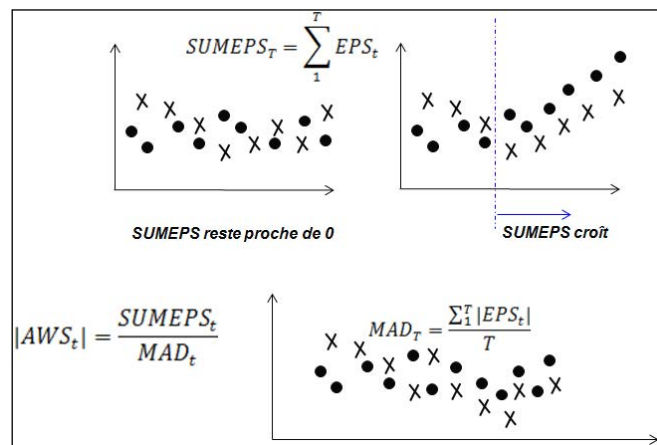


Figure 33 : Indicateur AWS

2.5 Calcul des indicateurs de performance

Le système met à jour mensuellement un certain nombre d'indicateurs de performance prévisionnels. Les calculs se font sur la dernière année de l'historique. Les indicateurs que nous avons retenus sont les plus couramment utilisés et que l'on retrouve dans la quasi totalité des travaux, lors de l'évaluation des performances prévisionnelles. Ces mêmes indicateurs sont souvent proposés dans les progiciels prévisionnels que nous avons évoqués précédemment.

MPE : Indicateur relatif qui traduit la moyenne des écarts entre les prévisions calculées par le système et les demandes réelles constatées en valeurs algébriques par rapport aux valeurs observées. Dans le cas de cet indicateur, les erreurs positives et négatives peuvent ainsi se compenser.

$$MPE = \frac{\sum_{i=1}^n (Prévision_i - Demande_i) / Demande_i}{n}$$

MAPE : Indicateur relatif plus adopté que le MPE, et traduisant la moyenne des écarts entre les prévisions calculées par le système et les demandes réelles constatées en valeurs absolues par rapport aux valeurs observées. Il se présente sous forme de pourcentage et constitue donc un indicateur pratique pour la comparaison.

$$MAPE = \frac{\sum_{i=1}^n |Prévision_i - Demande_i| / Demande_i}{n}$$

MAE : Indicateur absolu traduisant la moyenne arithmétique des valeurs absolues des écarts entre la prévision et la demande réelle.

$$MAE = \frac{\sum_{i=1}^n |Prévision_i - Demande_i|}{n}$$

MSE : Carré moyen des erreurs : Indicateur absolu traduisant la moyenne arithmétique des écarts entre les prévisions et les observations réelles.

$$MSE = \frac{\sum_{i=1}^n (Prévision_i - Demande_i)^2}{n}$$

3 Concept et interfaces développés

Nous présentons dans ce paragraphe le fonctionnement du module à travers les interfaces que nous avons développées. Nous mettons en vue l'étendue des fonctionnalités assurées en regard des différentes dimensions de conception pertinentes évoquées dans la littérature et formulées dans *Fildes & Goodwin (2006)* : **Niveau d'automatisation** (automatisation, correction des valeurs manquantes, outils d'analyse statistique, outils d'expérimentation instantanée), **le bon calibrage** (remontée d'information synthétique et très ciblée), la facilité d'utilisation à travers les facilités d'apprentissage (tutoriaux d'utilisation, tutoriaux d'aide, feedback d'utilisation), l'acceptation globale du système (facilité d'utilisation par rapport à des milliers d'articles, information la plus pertinente : précision des prévisions).

Interfaces du module de prévision

Pour aller dans le sens de la facilité et la simplification d'utilisation, notre module de prévision se présente sous forme d'une interface unique : un tableau de bord concentrant l'ensemble des données que nous avons jugées les plus synthétiques et les plus utiles à l'utilisateur (Figure 34). Cette interface se déploie en deux parties : une grille principale (partie supérieure de l'interface) et une grille de détail (partie inférieure de l'interface). Dans la grille principale, l'utilisateur va pouvoir consulter, pour les différents articles stockés : La classe de l'article (classe ABC multicritères en cours renvoyée par le module de classification des articles du stock), une prévision de consommation du mois courant, la vraie consommation mensuelle atteinte, un indice de confiance sur la prévision établie, la tendance de l'article, l'indicateur NF traduit en une information linguistique (existence ou non d'une erreur de prévision anormalement élevée) et l'indicateur AWS traduit en une information linguistique (existence ou non d'une cassure de tendance).

Dans la partie inférieure de l'interface, l'utilisateur va pouvoir consulter certains détails propres à l'article sélectionné dans la grille supérieure. Ces informations comprennent l'historique complet de la consommation et des prévisions établies ainsi qu'une visualisation (graphique) de l'évolution de ces deux données. Les indicateurs de précision propres à l'article sélectionné sont également renseignés (MAPE, MPE, MSE et MAE).

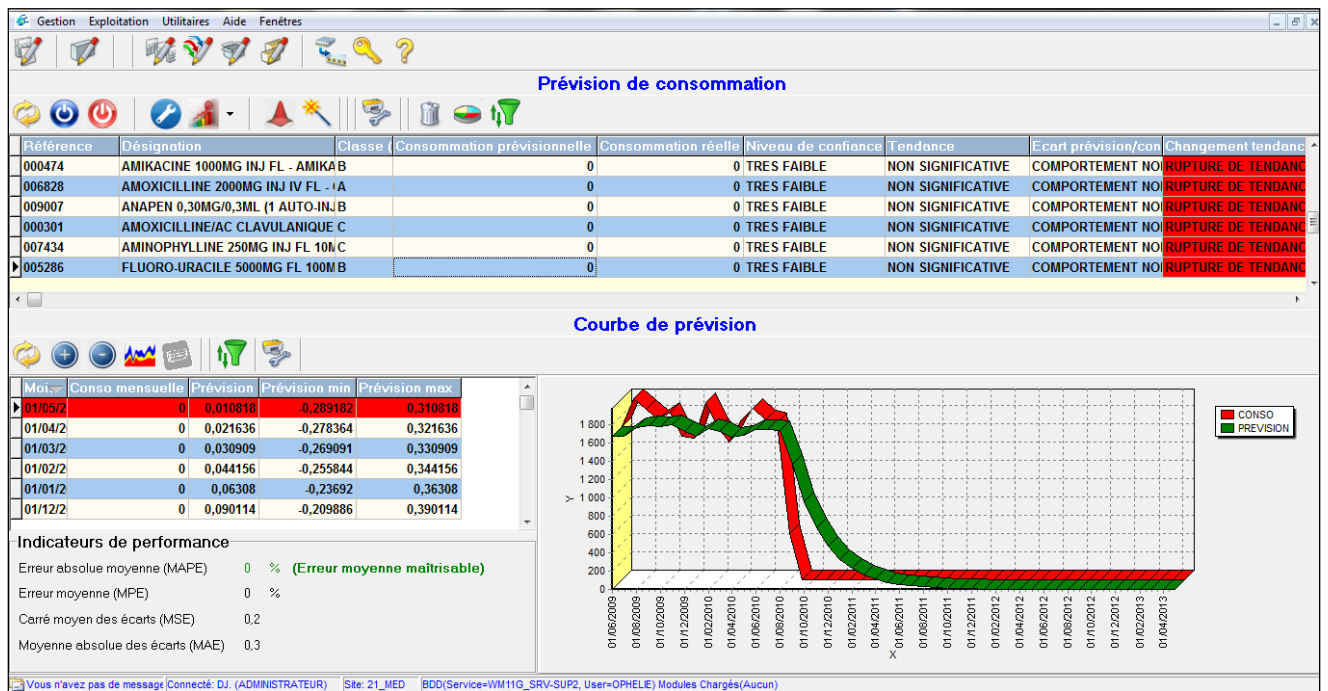


Figure 34 : Module de prévision des consommations

Calibrage / Intervention de l'humain

L'ensemble des données affichées est maintenu au niveau de la base et entretenu par le système, par le service Windows, conformément au fonctionnement que nous avons décrit dans le paragraphe précédent (paragraphe 4.2). Ce fonctionnement est modulé par un paramétrage portant sur la correction de la série et la fréquence d'identification des tendances. Un menu paramétrage permet à l'utilisateur d'activer ou de désactiver la procédure de correction des valeurs extrêmes et de définir, en nombre de mois, la périodicité de déclenchement de la procédure d'identification des tendances. Ces éléments constituent, en plus de la possibilité de modification de la prévision finale, les seules possibilités d'intervention de l'utilisateur.

Analyse des performances prévisionnelles

Le système permet une analyse de la performance prévisionnelle par article ou à un niveau plus global (sur la totalité des articles du stock). Les performances prévisionnelles par article sont exprimées par les quatre indicateurs affichés dans la grille de détail et traduisant la performance sur la dernière année de consommation. Le système traduit l'indicateur MAPE

en information linguistique binaire par comparaison de cet indicateur au coefficient de variation de la série de données. Si le MAPE est inférieur à la variation intrinsèque de la série, l'erreur de prévision est reportée comme étant maîtrisable, dans le cas contraire, l'erreur de prévision est considérée comme élevée.

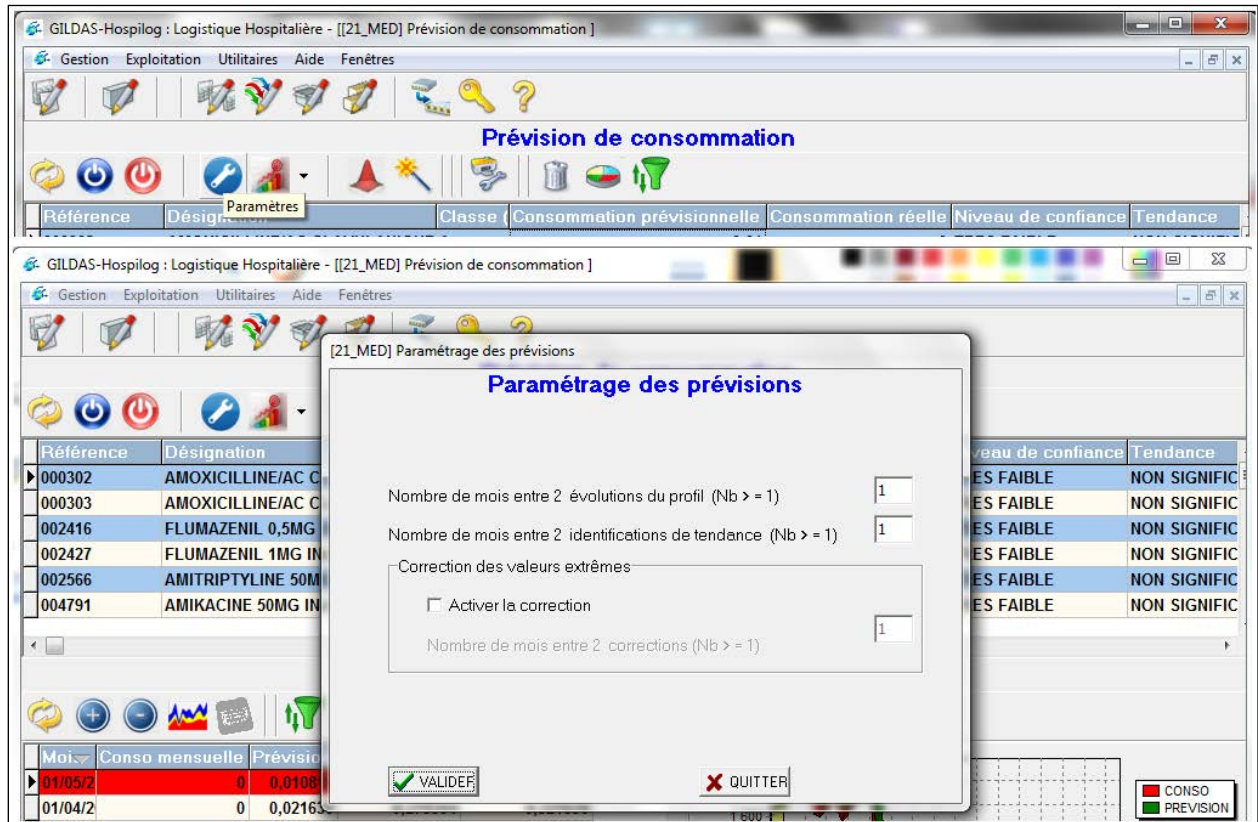


Figure 35: Paramétrage de la procédure de prévision

Une performance globale peut être consultée, elle se traduit au moyen d'une courbe d'évolution affichant simultanément l'erreur de prévision (MAPE) en abscisse et le pourcentage du chiffre d'affaires en ordonnée. Cette mesure de performance a été commentée dans *Bourbonnais & Usunier (2013)* et a l'intérêt d'afficher en un seul graphique la performance prévisionnelle pour tout le portefeuille du stock. Deux autres courbes d'évolutions théoriques sont affichées dans le même graphique. Construites sur la base de la variabilité intrinsèque des produits du stock (Coefficient de variation et $0,5 \times$ Coefficient de variation), ces courbes permettent de définir un objectif d'erreur optimal et un objectif d'erreur limite pour une portion donnée du chiffre d'affaires. L'utilisateur peut opérer cette

même analyse sur la base du flux de consommation ou sur la base des prix des différents produits.

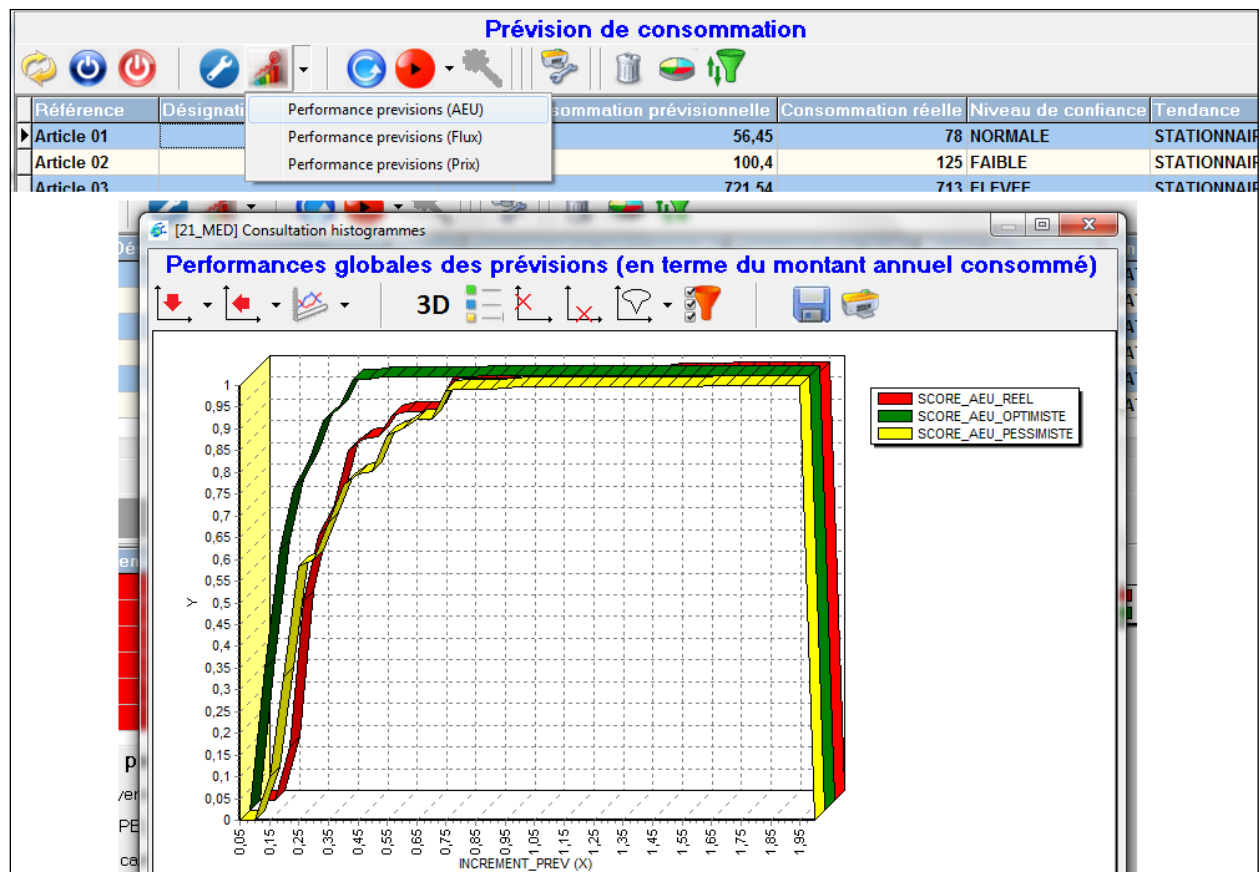


Figure 36 : Performance prévisionnelle globale

Facilité d'apprentissage / Ouverture du système

La facilité d'apprentissage du système est l'une des dimensions de conception œuvrant dans l'acceptation du système par les utilisateurs. L'ensemble des options d'affichage graphique et d'analyse de performances, dont nous avons doté notre système, est considéré par **Robert Fildes & Goodwin (2006)** comme étant des moyens de facilitation d'apprentissage. Pour poursuivre dans ce sens, notre système offre une option « **Simulation** » ouvrant le module de prévision à des données externes et permettant d'opérer des analyses prévisionnelles, instantanées, sur des séries de données importées à partir de fichiers Excel. Ce mode de fonctionnement permet l'application de l'ensemble des fonctionnalités d'analyse présenté précédemment, sur des données historiques externes. L'utilisateur va pouvoir dans ce cas, au moyen d'un bouton de commande, faire défiler le temps et expérimenter instantanément

l'ensemble des fonctionnalités du module. Ceci constitue un moyen puissant de prise en main de l'outil et d'appréhension de ses différentes composantes. Il constitue également une ouverture intéressante du système en généralisant l'application à des données externes.

Simplification de la manipulation de milliers d'articles

Pour faciliter la manipulation des différents articles, pouvant être comptés par milliers, notre système offre des facilités de tri et de filtrage permettant de se focaliser sur des groupes d'articles particuliers. Il est ainsi possible de limiter l'affichage sur les articles phares (classe A), d'opérer des analyses par typologie de consommation (cas des articles à baisse de consommation), ou de tracer les cas d'évolution atypiques en filtrant sur les indicateurs de formes (erreurs de prévisions anormalement élevées).

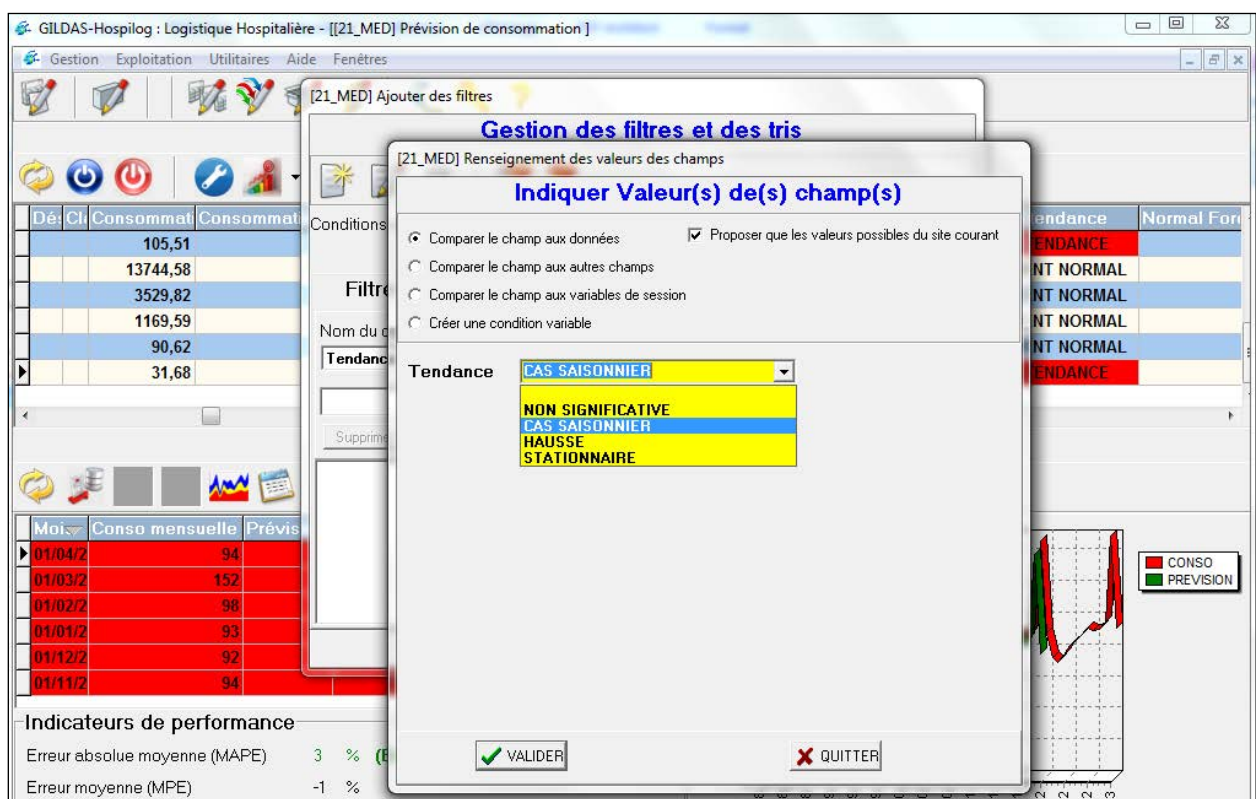


Figure 37 : Options de filtrage des articles par typologie

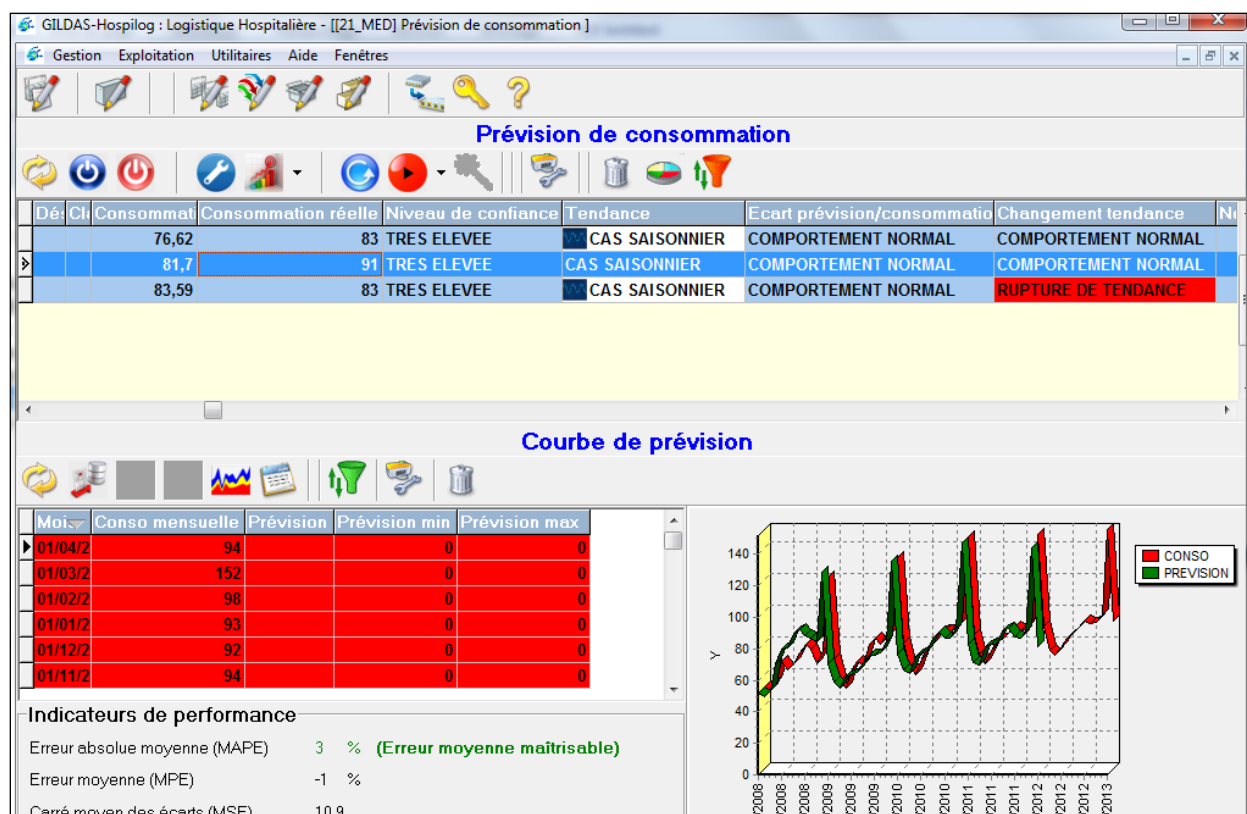


Figure 38 : Affichage des articles saisonniers

4 Conclusion

Dans ce chapitre, nous avons présenté le module de prévision de consommation développé. En se basant sur une revue de littérature des différentes techniques de prévision, nous avons opté pour l'implémentation de différentes variantes de la méthode du lissage exponentiel. Ces techniques de calcul présentent un bon compromis entre facilité d'implémentation informatique et robustesse de calcul.

La contribution de notre travail est de proposer un système auto-adaptatif basé sur une procédure d'identification des tendances. Ce système présente l'intérêt de s'adapter automatiquement à l'évolution dans le temps des comportements des produits et prend en compte le profil spécifique de la demande.

Les estimations issues de ce système permettent d'une part d'alimenter les calculs des quantités de commandes et d'autre part de remonter à l'utilisateur un tableau de bord récapitulatif reprenant des informations synthétiques sur les prévisions de consommation des différents produits de l'entrepôt.

Nous procédons dans le chapitre suivant à l'évaluation de notre système à travers une étude de cas propre aux Hospices Civils de Lyon.

Chapitre V. Évaluation du système : cas d'application

Nous procédons dans ce chapitre à une évaluation des performances du système développé à travers un cas d'application pratique. Pour le faire, nous proposons de comparer les résultats renvoyés par notre système, enrichi par le nouveau module de prévision, à ceux obtenus par adoption de l'ancienne configuration, à savoir l'estimation au moyen d'une moyenne mobile glissante.

Pour établir cette évaluation, nous considérons le cas de plateformes de distribution pharmaceutiques appartenant aux hospices civils de Lyon. Ce réseau logistique, s'appuyant sur des plateformes de distribution, permet de desservir les quatorze établissements des HCL en médicaments et dispositifs médicaux. *KLS* a équipé en 2006 l'ensemble de ces plateformes par le WMS *Gildas Hospilog*.

Le réapprovisionnement du stock, au sein de ces dépôts, s'opère au moyen du module de préconisation des commandes supporté par le *WMS*, décrit dans le deuxième chapitre. Afin d'établir nos analyses, nous avons extrait une copie de la base correspondant aux données historiques de fonctionnement de l'ensemble des plateformes depuis le démarrage de *Gildas Hospilog* en 2006. Après avoir calculé des estimations de consommation au moyen de l'ancien puis du nouveau système, nous comparerons les résultats prévisionnels renvoyés afin d'évaluer l'impact de ces estimations sur les niveaux de stock des plateformes et leurs taux de services.

Nous constituons, pour chacune des plateformes choisies, un échantillon d'articles tests sur lesquels nous réaliserons les simulations de performances des différentes méthodes de classification et de prévision proposées. Les échantillons correspondent aux produits de la classe A, obtenue au moyen du module de classification que nous avons développé, en adoptant une configuration multicritères tenant compte, autant que possible, de la spécificité de ce secteur. Une première évaluation se fera par rapport aux performances prévisionnelles, auquel cas nous comparerons les résultats retournés par chacun des deux systèmes, en fonction des indicateurs de précision. Une deuxième évaluation se fera par rapport au fonctionnement du stock : nous cherchons, par simulation, à évaluer les niveaux

niveaux (Figure 39) : des fournisseurs (Laboratoires médicaux), une pharmacie centrale, des plateformes intermédiaires (PUI) et les services hospitaliers (Clients finaux).

Les services de soin commandent auprès des pharmacies intermédiaires, lesquelles pharmacies commandent auprès de la plateforme centrale. La pharmacie centrale s'approvisionne auprès des laboratoires pharmaceutiques.

Le réapprovisionnement au niveau des deux maillons intermédiaires (PUIs et pharmacie centrale) se fait au moyen du module de préconisation de commandes. Ces différentes plateformes gèrent leurs stocks au moyen du module de préconisation de Gildas WM, basé sur la moyenne mobile. Nous proposons dans la suite d'évaluer les performances de notre système pour le cas de certaines de ces plateformes. Nous constituons des échantillons de tests, sur une base multicritères, et récupérons la classe des articles pilotes. L'évaluation se fera par rapport à ces échantillons constitués.

2 Quelles caractéristiques types pour la gestion des plateformes pharmaceutiques : Enquête auprès de centres hospitaliers

Nous proposons de recenser les critères caractéristiques de la gestion des plateformes pharmaceutiques. Ceci nous permettra d'établir une classification multicritères et de définir, plus justement, la classe des produits pilotes dans le contexte de la gestion de stock pharmaceutique. Divers attributs de classification ont été évoqués dans la littérature, cités notamment dans *Flores et al. (1992)*, *Ramanathan (2006)*, *Zhou & Fan (2007)*, etc. *Flores et al. (1992)* considèrent que le choix des attributs dépend fortement du type de l'activité de l'entreprise et que dans le cas particulier du secteur de la santé, la criticité du produit peut être plus importante que la considération financière.

Afin de définir les critères spécifiques au domaine de la gestion des plateformes hospitalières, nous avons consulté des pharmaciens responsables de plateformes pharmaceutiques afin d'avoir leur avis sur les attributs qu'ils considèrent comme les plus pertinents et qu'il convient d'intégrer dans la classification d'un dépôt de médicaments. Nous avons établi une liste d'attributs, à partir des travaux de la littérature sur la classification multicritère, et mené une enquête auprès de quarante pharmacies à usage intérieur correspondant à 18 centres hospitaliers répartis sur 15 régions françaises.

Nous avons réparti les critères recensés en trois typologies : les critères financiers (comprenant le prix unitaire, le coût de stockage et le coût de commande), les critères logistiques (comprenant la consommation annuelle du produit, le temps de réapprovisionnement auprès du fournisseur, la stockabilité, le quota de commande et la fréquence de commande) et les critères pharmaceutiques (comprenant la criticité du médicament, la substituabilité et la durée de vie du médicament).

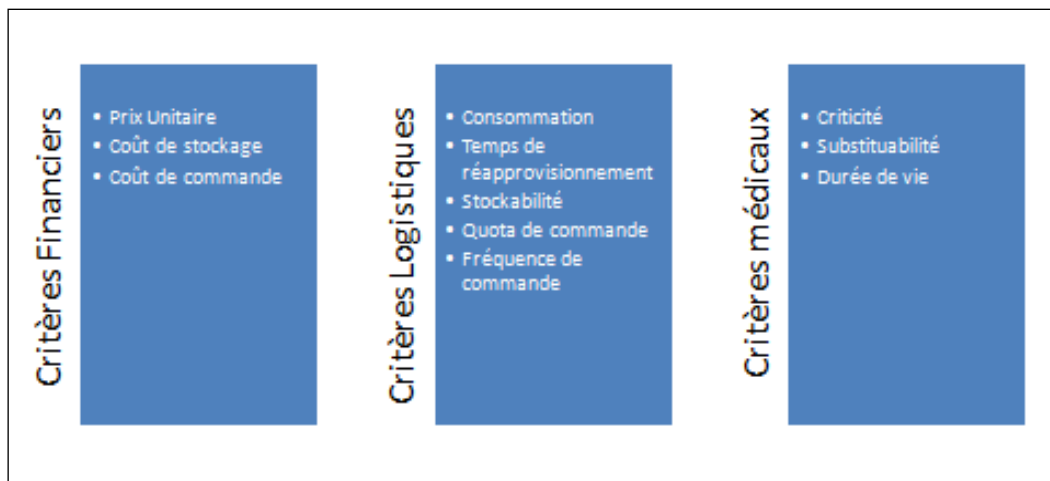


Figure 40 : Critères de classification énumérés

Nous avons demandé aux responsables sollicités d'évaluer l'importance de chacun des critères, en tenant compte du contexte de la gestion d'un stock hospitalier. Pour l'analyse de chacun des critères, nous avons proposé aux personnes interrogées d'évaluer de 0 à 5 chacun des critères selon son importance.

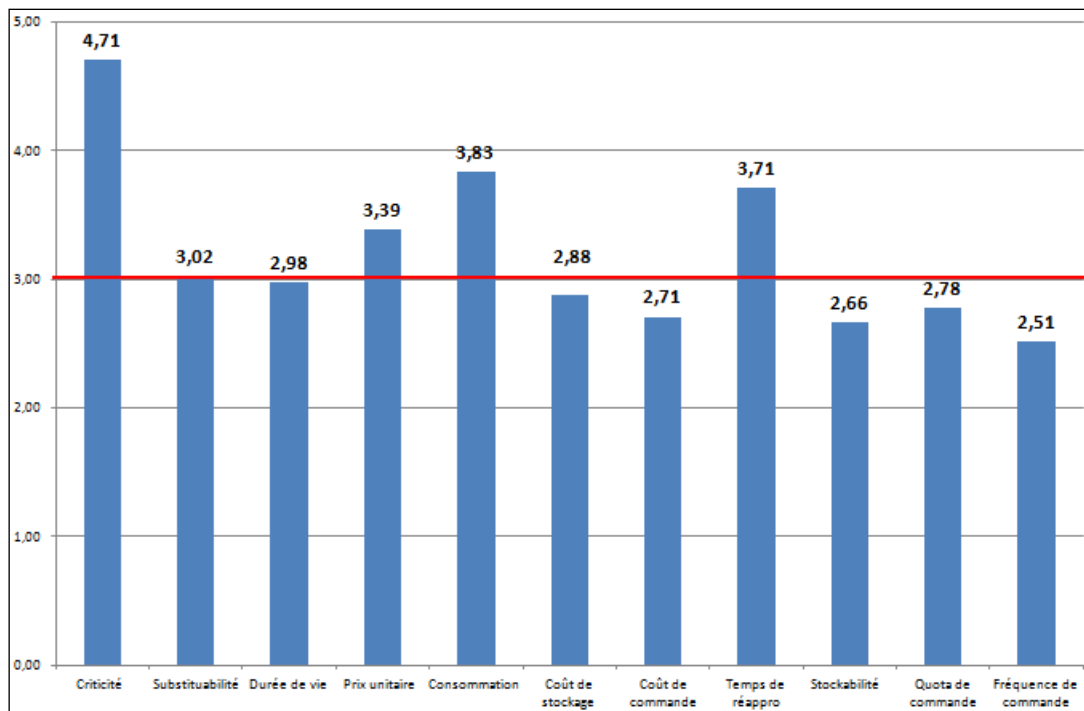


Figure 41 : Poids des différents critères

Nous reportons dans l'histogramme de la figure 41 les résultats moyens obtenus. La criticité du médicament se démarque nettement des autres critères (34 pharmaciens sur 40 lui ont attribué une note maximale). La criticité, telle que nous l'avons expliquée dans l'enquête, fait référence à l'impact de l'indisponibilité d'un médicament sur la santé du patient. Le temps de réapprovisionnement est également un critère important (parmi les 40 pharmaciens interrogés, 13 pharmaciens l'ont placé en première position, 12 en deuxième position et 13 en troisième position). L'enquête révèle aussi l'intérêt accordé à l'aspect financier classique, habituellement adopté en gestion de stock. Ceci se traduit par les deux critères 'Consommation' et 'Prix unitaire', qui se démarquent également du restant des attributs (pour la consommation, 10 pharmaciens lui ont affecté une note maximale et 18 la placent en deuxième position ; pour le prix, 8 pharmaciens le placent en première position, 11 en deuxième position et 15 en troisième position).

3 Données considérées (Plateformes + Historiques)

Nous avons récupéré auprès des hospices civils de Lyon une copie de la base de données correspondant au fonctionnement de certaines PUI depuis l'installation de Gildas WM en 2006 jusqu'au mois d'août 2010. En regardant de près la base de données, nous avons

sélectionné les plateformes pour lesquelles il était possible de reconstituer l'intégralité des données permettant d'opérer les simulations : Classification des articles, simulation de prévision de consommation et simulation du fonctionnement du stock. Les données requises correspondent aux :

- Articles de stock et certains attributs caractéristiques (Prix, Quota, fournisseurs) : Table Article dédié.
- Historique de consommation des différents articles : Tables Historique et Archives contenant l'ensemble des flux de l'entrepôt, notamment les flux de consommations.
- Historiques des commandes des clients (les hôpitaux) vers l'entrepôt : Tables Commandes historiques et Archives des commandes.

Nous présentons dans le tableau 2, les différents sites retenus, le nombre d'articles gérés (Médicaments ayant été consommés pendant l'année 2009), les fournisseurs assurant l'alimentation des stocks ainsi que les clients desservis. Les sites que nous considérons ne comprennent pas nécessairement la totalité des produits gérés par la plateforme. Une pharmacie peut en effet comprendre plusieurs sites, dédiés chacun à une famille ou à une typologie de produits.

Tableau 2 : Plateformes pharmaceutiques considérées

Site	Pharmacie	Références Médicaments stockés	Fournisseurs	Clients
20_MED	Pharmacie Hôtel Dieu	1302	Pharmacie centrale + Laboratoires médicaux	Hôpital Hôtel Dieu
24_MED	Pharmacie Croix Rousse	1640	Pharmacie centrale + Laboratoires médicaux	Hôpital de la Croix Rousse
25_PE	Pharmacie Renée Sabran	1998	Pharmacie centrale + Laboratoires	Hôpital Renée Sabran

Site	Pharmacie	Références Médicaments stockés	Fournisseurs	Clients
			médicaux	
28_PE	Pharmacie Charpennes	1237	Pharmacie centrale + Laboratoires médicaux	Hôpital gériatrique des Charpennes
32_PE	Pharmacie Pierre Garraud	1686	Pharmacie centrale + Laboratoires médicaux	Hôpital gériatrique Antoine Charial
37_PE	Pharmacie Antoine Charial	1377	Pharmacie centrale + Laboratoires médicaux	Hôpital gériatrique Antoine Charial

4 Constitution des échantillons

Afin de constituer des échantillons de tests, nous proposons de classer les médicaments suivant un scénario multicritères qui prend en considération les résultats de l'enquête réalisée auprès des pharmaciens. Les quatre critères principaux qui ressortent de l'enquête sont : la criticité, la consommation, le temps de réapprovisionnement auprès du fournisseur et le prix du médicament. Notre système ne permet pas, à ce niveau, la quantification de la criticité d'un médicament. Aucune information se rapportant à cette dimension n'est actuellement supportée par la base de données.

Nous optons donc pour l'intégration des trois critères restants. L'enquête réalisée auprès des professionnels a montré que ces trois critères ont une importance équivalente, nous proposons d'adopter un scénario de classification qui introduit les mêmes poids pour les trois attributs considérés. Ceci correspond à un scénario AHP, de poids uniformes de valeurs 33% et de consistance de comparaison nulle. Nous considérons pour chacun des sites l'échantillon correspondant à la classe A retournée par la classification adoptée. Nous présentons dans le tableau 3 la taille des échantillons, le taux en chiffre d'affaires par rapport à la totalité du portefeuille annuel ainsi que le pourcentage d'articles en communs avec ceux

d'une classe A obtenue par une classification ABC classique (basée sur le critère du montant financier annuel).

Tableau 3 : Articles pilotes sur la base de la classification établie

Pharmacie	Taille de l'échantillon	Pourcentage du CA total
20_MED	260	78,66%
24_MED	328	77 %
25_PE	400	81,77 %
28_PE	247	69,63 %
32_PE	337	84 %
37_PE	275	79,66 %

5 Evaluation de la performance prévisionnelle

Nous procédons à l'évaluation de la performance prévisionnelle en utilisant les échantillons précédemment constitués. Nous simulons le fonctionnement des deux systèmes (moyenne mobile et le nôtre) sur une période de temps et comparons les performances calculées sur cette période. La base de données utilisée comprend 4 années et 7 mois de consommation, que nous subdivisons en deux parties :

- Les quatre premières années (de Janvier 2006 jusqu'à Décembre 2009) représentent l'horizon de l'historique des consommations disponibles lors du démarrage des simulations. Il servira pour le premier filtrage des données ainsi que pour l'identification des tendances. Nous éliminons tous les produits identifiés, par le premier filtrage, comme étant atypiques (produits à faible rotation, produits erratiques et produits intermittents), ces derniers étant traités par notre système au moyen d'une moyenne mobile, il est inutile de les garder dans les échantillons de tests. Nous présentons dans le tableau 4, les échantillons obtenus à la suite

de l'élimination des cas atypiques ainsi que les taux financiers qu'ils représentent par rapport au portefeuille annuel.

- Les 7 derniers mois (de Janvier 2010 jusqu'à Août 2010) représentent l'horizon de simulation. Le calcul des prévisions se fera, pour chaque article, par adoption de la technique préconisée par le système lors de la phase d'identification des tendances. Les performances des calculs seront reportées par rapport à cette période.

Tableau 4 : échantillons finaux (Suite à l'élimination des cas atypiques)

Pharmacie	Echantillon initial	Echantillon final (obtenu suite à l'élimination des profils atypiques)	Taux du flux financier annuel
20_MED	260	223	69,53 %
24_MED	328	301	73,51 %
25_PE	400	378	75,6 %
28_PE	247	239	67,2 %
32_PE	337	289	79,2 %
37_PE	275	266	76,7 %

Nous présentons dans le tableau 5 la répartition des différents profils de consommation identifiés. Les profils stationnaires sont prépondérants (en dessus de 73% du total pour toutes les pharmacies). Les typologies tendancielle et saisonnières sont moins fréquentes (en dessous de 18% pour les cas tendanciels et moins de 8% pour les cas saisonniers), mais posent davantage de problèmes aux pharmaciens.

Tableau 5 : Typologies de consommation obtenues

Pharmacie	Echantillon considéré	Profils stationnaires	Profils tendanciels	Profils saisonniers
20_MED	223	180	31	12
24_MED	301	256	32	13
25_PE	378	278	69	31
28_PE	239	221	10	8
32_PE	289	233	50	6
37_PE	266	235	22	9

Pour l'évaluation de la performance prévisionnelle, nous adoptons deux indicateurs de mesures : un indicateur relatif (Le MAPE) et un indicateur absolu (Le MAE). Ces indicateurs ont été présentés dans le paragraphe 4 et font partie des métriques remontées par le module de prévision développé. Nous rappelons dans ce qui suit la signification et la construction de ces indicateurs : r^2

MAPE : Moyenne des écarts en valeur absolue par rapport aux valeurs observées. Il s'agit d'un pourcentage et donc un indicateur très pratique de comparaison.

$$MAPE_{Article\ k} = \frac{\sum_{i=Janvier\ 2010}^{Août\ 2010} |Demande_i - Prévision_i|}{Demande_i \cdot 7}$$

MAE : Moyenne arithmétique des erreurs absolues des écarts.

$$MAE_{Article\ k} = \frac{\sum_{i=Janvier\ 2010}^{Août\ 2010} |Demande_i - Prévision_i|}{7}$$

Ces mesures de performances sont effectuées au niveau du SKU (calculé par article). Afin d'avoir une appréciation de performance globale (sur l'ensemble des articles de l'échantillon), nous procédons, comme préconisé dans *Bourbonnais & Usunier (2013)*, au calcul de moyennes de performances pondérées par le flux financier de l'article. La littérature considère qu'il n'est pas opportun de comparer ou de moyenner des pourcentages d'erreurs. En effet, il est moins grave par exemple de se tromper de 10% sur un chiffre d'affaires de 100

K Euros que de se tromper de 5% sur un chiffre d'affaires de 10 000 *K Euros*. La pondération par le montant se fait conformément à la formule ci-dessous.

$$MAPE_{\text{moyen pondéré}} = \frac{\sum_{k=Article\ 1}^{Article\ n} Montant_{Article\ k} * MAPE_{Article\ k}}{\sum_{k=Article\ 1}^{Article\ k} Montant_{Article\ k}}$$

Nous reprenons dans le tableau 6, les résultats des différentes performances pour l'ensemble des pharmacies traitées.

Tableau 6 : Performances prévisionnelles obtenues

Pharmacie	Système utilisé	MAPE moyen pondéré (AEU)	MAE moyen pondéré (AEU)
20_MED	Ancien Système	67,28	29,66
	Système proposé	65,03	29,51
24_MED	Ancien Système	44,89	123,33
	Système proposé	43,10	118,84
25_PE	Ancien Système	52,02	87,62
	Système proposé	43,34	81,56
28_PE	Ancien Système	42,34	124,39
	Système proposé	44,95	136,67
32_PE	Ancien Système	55,64	164,27
	Système proposé	45,45	133,18
37_PE	Ancien Système	53,26	174,83
	Système proposé	49,38	167,40

Dans la majorité des cas testés, le nouveau système développé permet d'atteindre de meilleures performances prévisionnelles, en termes de mesures relatives et absolues:

Les meilleures performances sont atteintes au niveau des pharmacies 25_PE, 32_PE et 37_PE.

Des MAPE pondérés respectifs de 43,34 %, 45,45 % et 49,38 % sont réalisés. Ceci permet des gains respectifs de 8,76 %, 10,19 % et 3,88% par rapport aux performances retournées par

l'ancien système. Pour les pharmacies 20_MED et 36_MED les résultats sont de 65,03% et 43,1%, ce qui permet une amélioration de 2,25 % et 1,79 %.

Dans le cas de la pharmacie 28_PE, notre système a légèrement dégradé les résultats, la performance pondérée obtenue est de 44,95% contre une performance de 42,34 % dans le cas de l'ancien système, ce qui correspond à une régression de 2,61% par rapport à l'indicateur adopté.

Nous présentons ci-dessous les performances prévisionnelles de certaines pharmacies sous forme graphique. Nous adoptons la même représentation, discutée dans *Bourbonnais & Usunier (2013)* et implémentée dans notre système, visualisant la performance de tout le portefeuille de l'échantillon, par rapport à un objectif d'erreur cible calculé sur la base de la variation intrinsèque de la consommation historique : une enveloppe limite est définie sur la base du coefficient de variation de la série (Courbe bleue pointillée : Performance objectif) et de la moitié du coefficient de variation (Courbe bleue continue : Performance optimale). Les courbes correspondant au système développé (Courbes vertes) sont, pour une large plage du portefeuille considéré, en dessus des courbes propres à l'ancien système (Courbes rouges) ce qui traduit l'atteinte d'une meilleure précision pour un taux financier (Par rapport au portefeuille financier total de l'échantillon) plus élevé : sur environ 80% du portefeuille pour la pharmacie 25_PE et sur la totalité du portefeuille pour les pharmacies 32_PE et 37_PE.

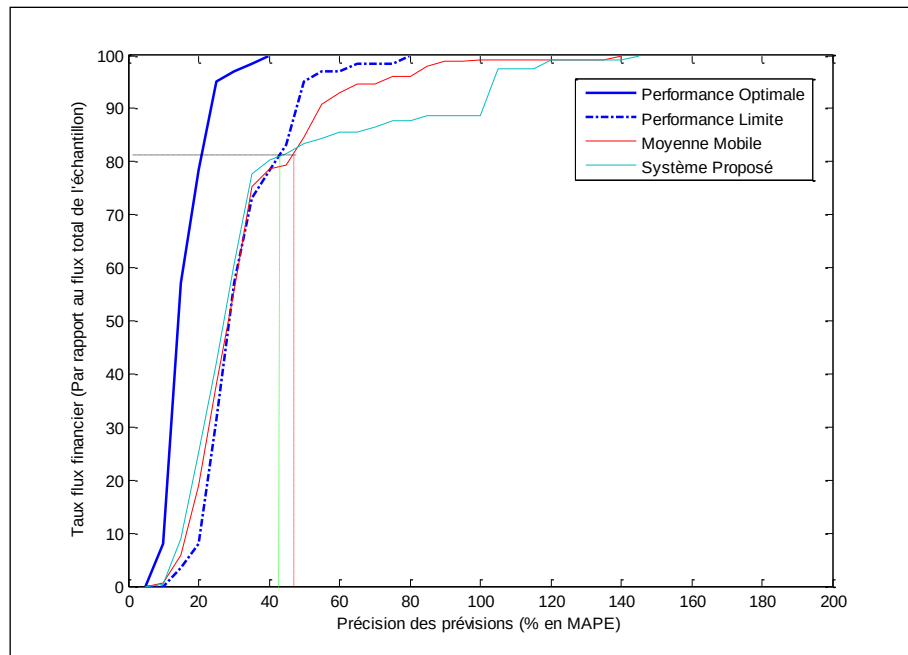


Figure 42 : Cas de 25_PE

Par ailleurs, et en comparaison avec les courbes théoriques, notre système permet de mieux coller à l'objectif de performance :

Pour la pharmacie 25_PE : Jusqu'à 80 % du portefeuille financier est prévu avec une précision supérieure à la limite théorique contre 78% pour la moyenne mobile.

Pour la pharmacie 32_PE : Jusqu'à 68 % du portefeuille financier est prévu avec une précision supérieure à la limite théorique contre 50% pour la moyenne mobile.

Pour la pharmacie 37_PE : Jusqu'à 65 % du portefeuille financier est prévu avec une précision supérieure à la limite théorique contre 55 % pour la moyenne mobile.

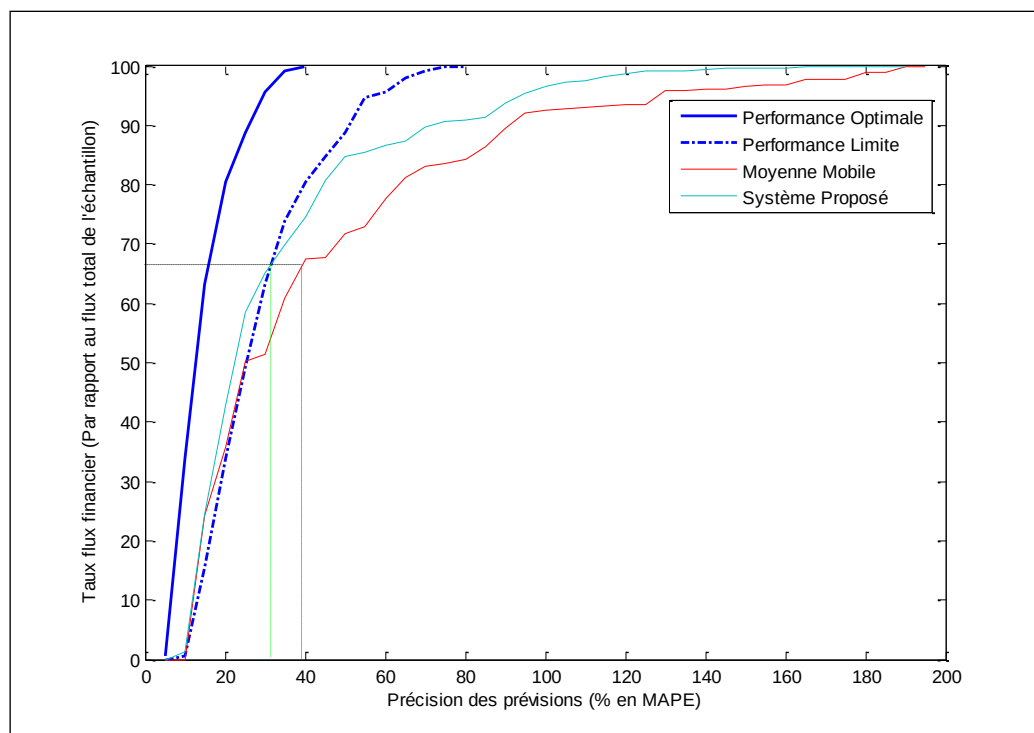


Figure 43 : Cas de 32_PE

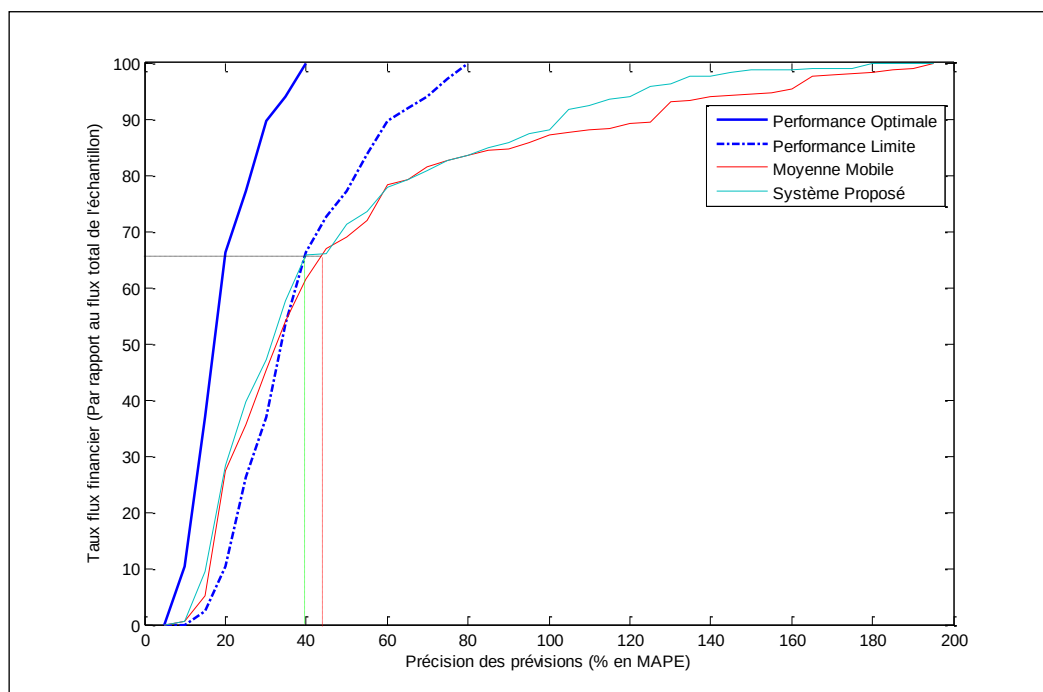


Figure 44 : Cas de 37_PE

6 Evaluation de la gestion de stock

Nous proposons, dans ce chapitre, d'étendre l'analyse des performances à l'évaluation de l'impact sur le fonctionnement du stock. Nous nous limitons dans le cadre de cette thèse à une première approche afin d'estimer les apports du système développé en termes de gains en immobilisation du stock et de taux de service rendu aux services de soins clients. Une réflexion plus approfondie et des analyses plus fines pourraient être poursuivies dans des travaux futurs.

Les travaux sur la gestion des chaînes hospitalières sont moins répandus que ceux portant sur les chaînes logistiques industrielles. Les contraintes sont cependant au moins aussi fortes puisqu'il est nécessaire de garantir un taux de service élevé auprès des unités de soin tout en réduisant les coûts de stockage des produits *Baboli et al. (2010)*. *Baboli et al. (2010)* revient sur certains travaux ayant traité de cette problématique : des questions portant sur la centralisation ou non des structures et des décisions ont été posées. La structure physique peut en effet être centralisée auquel cas, le fournisseur livre directement aux unités de soins sans passer par des pharmacies intermédiaires. La décision centralisée suppose que les différents maillons de la chaîne coordonnent afin de trouver un optimum de coût global pour toute la chaîne. Dans le cas de la décision décentralisée, chaque maillon essaie d'optimiser son propre coût indépendamment des autres membres. Ces diverses configurations ont été traitées dans certains travaux de la littérature en regard, notamment de coûts de transport et de stockage (*Axsater et al. (1999)*; *Ertogral et al. (2007)*) ou en considération de la nature de la demande : déterministe ou probabiliste (*Axsater (2003)*; *Ganeshan (1999)*).

La configuration actuelle aux HCL peut être perçue comme un fonctionnement (*Structure décentralisée, Décision décentralisée*). Les PUIs s'approvisionnent auprès de la pharmacie centrale et des laboratoires pharmaceutiques, moyennant le module de préconisation de *Gildas Hospilog*, afin de satisfaire les demandes des unités de soin tout en optimisant leurs propres coûts d'immobilisation. Comme présenté dans le deuxième chapitre, le module de préconisation de *Gildas Hospilog* s'appuie sur des politiques de gestion dont les paramètres (notamment le seuil de commande) sont calculés sur la base d'estimations prévisionnelles, opérées par moyenne mobile. Ce mode de gestion, est appelé dans la littérature '*gestion sur prévision*' *Babai (2005)*. Il consiste à mettre à jour les paramètres des politiques de gestion à la

fréquence de mise à jour des calculs prévisionnels, contrairement au mode de gestion classique dans lequel les paramètres des politiques de gestion sont statiques.

Pour évaluer les apports de notre système (en terme d'immobilisation et de taux de service), nous proposons de comparer par simulation, pour le cas de l'échelon PUI, l'impact de l'adoption d'une politique de gestion de stock (politique à point de commande) dans une configuration de gestion sur prévision en regard de l'ancienne configuration (Calcul des paramètres par moyenne mobile) puis de la nouvelle configuration (Calcul des paramètres par les nouvelles prévisions). La simulation se fera sur le même horizon de temps défini initialement (7 mois de l'année 2010) et les performances seront reportées par rapport à cet horizon : Immobilisation et taux de service. Nous revenons dans ce qui suit sur le détail des calculs adoptés.

6.1 Politique de gestion adoptée

L'ancien système s'appuie sur deux politiques de gestion, une politique périodique et une politique sur seuil à niveau de rechargement. Comme présenté dans le deuxième chapitre de la thèse, le choix des politiques s'opère manuellement et les paramètres sont calculés sur la base de la moyenne mobile et en tenant compte de données entrées directement par l'utilisateur : notamment le délai de réapprovisionnement et le niveau de sécurité objectif en nombre de jours de couverture. Nous proposons, pour ce cas d'application, d'adopter la politique de gestion (s, S) , politique sur seuil à niveau de rechargement. Nous introduirons respectivement les résultats prévisionnels de l'ancien puis du nouveau système en calculant autrement les paramètres : adoption de la notion du taux de service et de la notion de la quantité économique de commande.

La politique de gestion (s, S) a été commentée dans plusieurs travaux de la littérature notamment dans ceux de *Silver et al. (1998)*. Longtemps comparée aux autres techniques $((T, S), (s, Q)$ ou $(T, s, S))$, cette dernière s'est avérée plus performante : en terme de réduction d'immobilisation et de taux de service. *Liang(1997)* établit un comparatif concis et très synthétique de l'ensemble de ces politiques : L'inconvénient majeur des politiques périodiques $((T, S)$ et $(T, s, S))$ par rapport aux politiques continues $((S, Q)$ et $(s, S))$ est qu'elles nécessitent plus d'immobilisation pour assurer le même niveau de service, ce qui induit des coûts de stockage plus élevés. Il a été démontré que les coûts totaux engendrés par

l'adoption de la politique (s, S) ne dépassent pas ceux correspondant à la politique (s, Q). *Silver et al. (1998)*. Ces considérations font de la politique (s, S) l'une des méthodes les plus adoptées dans la pratique *Archibald & Silver (1978)*.

6.1.1 Calcul du point de commande

Nous procédons aux calculs des paramètres tels que commentés dans *Silver et al. (1998)*. L'estimation du point de commande se fait dans l'optique de trouver un compromis entre les coûts de stock et les coûts de rupture en regard d'un certain niveau de service théorique. Différents niveaux de service ont été commentés dans la littérature, nous adoptons la notion de niveau de service par cycle, il s'agit de l'un des plus utilisés. Cet indicateur, exprimé sous forme de pourcentage, mesure la probabilité de ne pas avoir de rupture de stock entre deux commandes successives (un cycle).

En fixant à $cslo$ le niveau de service cible, il s'agit de minimiser le coût de stockage tout en assurant une probabilité de non-rupture par cycle qui soit supérieure au niveau considéré. Le coût de stock et la qualité de service étant croissants en fonction du stock immobilisé, la résolution du problème considéré revient donc à trouver le plus faible stock satisfaisant le niveau requis, ce qui revient à résoudre l'équation suivante :

$$CSL = cslo$$

Il s'agit donc de trouver le seuil de commande qui satisfait l'équation ci-dessous (probabilité que la demande au cours du délai de réapprovisionnement soit supérieure au seuil est égale à $cslo$)

$$P \text{ Demande pendant délai réapprovisionnement} \leq s = cslo$$

En supposant que la demande suive une loi normale (hypothèse souvent adoptée) et que les délais de réapprovisionnement soient constants, nous aboutissons, à partir de l'équation précédente, à la formule du seuil de commande suivante :

$$s = \text{Demande}_{\text{Temps réappro}} + SS$$

$$s = \text{Demande}_{\text{Temps réappro}} + k * \sigma_{\text{Temps réappro}}$$

$Demande_{Temps\ réappro}$ étant l'estimation de la demande sur le temps de réapprovisionnement et $\sigma_{Temps\ réappro}$ l'écart type de l'erreur de prévision sur le même délai. Ces deux paramètres sont renvoyés par le système de prévision et mis à jour chaque mois. k est un facteur de sécurité dépendant du niveau de service cible choisi.

La consommation moyenne et l'erreur prévisionnelle seront approchées par le système de prévision. Le seuil de commande, et donc forcément la performance du stock, sont impactés par le système prévisionnel mis en place. Le seuil de commande sera ainsi mis à jour au même rythme de la mise à jour prévisionnelle (cadence mensuelle).

6.1.2 Calcul du niveau de reemplètement

Le niveau de reemplètement se calcule à partir du point de commande selon la formule suivante (*Liang (1997)*):

$$S = s + Q^*$$

Q^* étant la quantité économique de commande calculée suivant la formule de Wilson. Il s'agit d'un paramètre classique d'économie d'échelle établissant un compromis entre les coûts de commande et les coûts de stock.

$$Q^* = \frac{\sqrt{2 * A * m_D}}{h}$$

A : Coût fixe de commande ; m_D : Demande moyenne ; h : Coût de possession de stock (Unité de produit par unité de temps)

6.1.3 Calcul de la quantité de commande

À chaque fois que le niveau de stock atteint le seuil de commande, une quantité de commande est calculée suivant la formule suivante, et arrondie systématiquement au multiple de commande (Quota de commande) imposé par le fournisseur (principalement la pharmacie centrale).

$$Quantité\ de\ commande = S - Niveau\ de\ stock$$

Le niveau de stock à un instant donné s'obtient en comptabilisant les encours de réception et les quantités non satisfaites :

$$Niveau\ de\ stock = Stock\ existant + Encours - Ruptures$$

6.2 Résultats des simulations

Afin d'établir les simulations, nous avons extrait les données suivantes : l'historique des commandes des différents services auprès de la pharmacie à tester, les temps de réapprovisionnement de la pharmacie, les prix des différents articles stockés, les quotas de commandes (multiples des quantités de commandes imposés par les fournisseurs), les calculs prévisionnels établis sur la base d'une moyenne mobile et sur la base du système prévisionnel proposé (correspondant aux différents mois de l'horizon de simulation).

Nous avons reconstitué les stocks initiaux (correspondant au premier janvier 2010 (début de la simulation)) et nous avons considéré, comme dans *Hassan (2006)*, un coût de passation de commande à 2 euros et approché à 15 % le taux de possession de stock. Le coût de passation d'une commande comprend plusieurs composantes (Coût de télécommunication = 0,6 Euros, Coûts de fournitures administratives = 0,6 Euros, Coût par lignes de personnel = 0,75 Euros,...). Le coût de possession comprend les frais de gestion des stocks (coûts directs : loyers, frais d'entretien, services extérieurs, etc. ; coûts indirects : intervention du service informatique, service comptable, etc.) et est exprimé en pourcentage du prix du produit, variant généralement de 15 % jusqu'à 40 % selon le secteur d'activité de l'entreprise.

Nous mesurons, sur l'horizon de simulation considérée (7 mois de 2010), les indicateurs suivants : le taux de service par cycle réel, le stock moyen sur la période (souvent utilisé pour évaluer l'immobilisation des produits) et le coût de possession des articles (correspond au coût d'immobilisation du stock moyen disponible sur la période du temps).

Pour évaluer la performance globale, nous reportons le taux de service moyen sur les articles, le stock moyen total immobilisé ainsi que le coût de possession total correspondant. Nous présentons dans le tableau 7, pour le cas des pharmacies 25_PE et 24_MED, les résultats obtenus pour un taux de service cible de 80 %.

Le stock moyen cumulé est calculé suivant la formule suivante :

$$\text{Stock moyen cumulé} = \frac{\text{Stock Moyen}_{\text{Cycle } i, \text{Article } j}}{\text{Articles Cycles}}$$

Pour la pharmacie 25_PE, notre système renvoie un meilleur taux de service (77 % contre 72 % pour l'ancien système) pour un faible accroissement du stock immobilisé 0,73 %. Le coût d'immobilisation obtenu est presque le même (5625 pour l'ancien système contre 5623 pour le système proposé). Pour la pharmacie 24_MED, nous obtenons également un meilleur taux de service (81,2 % contre 78 %), ceci s'accompagne d'un accroissement maîtrisable du

stock immobilisé (un accroissement de 4,13 % en terme de volume et de 5,71 % en terme de coût).

Tableau 7 : Performances des pharmacies

Pharmacie	25_PE		24_MED	
Système de prévision	Ancien Système	Système proposé	Ancien Système	Système proposé
Taux de service réel atteint (%)	72,00	77,00	78,00	81,20
Stock moyen cumulé (Volume)	480388	483874	1360752	1416886
Coût immobilisation	5625	5623	17272	18258

Nous faisons varier dans ce qui suit le taux de service théorique (de 70 % jusqu'à 95 %) et reportons les résultats des simulations en terme de taux de service réel atteint, d'immobilisation en quantité et d'immobilisation financière.

Dans le cas de la pharmacie 25_PE (Figures 45, 46 et 47), le taux de service réel atteint par notre système est toujours en dessus du taux de service de l'ancien système. Ceci s'obtient avec le même niveau de stock immobilisé (Un accroissement d'environ seulement 1%).

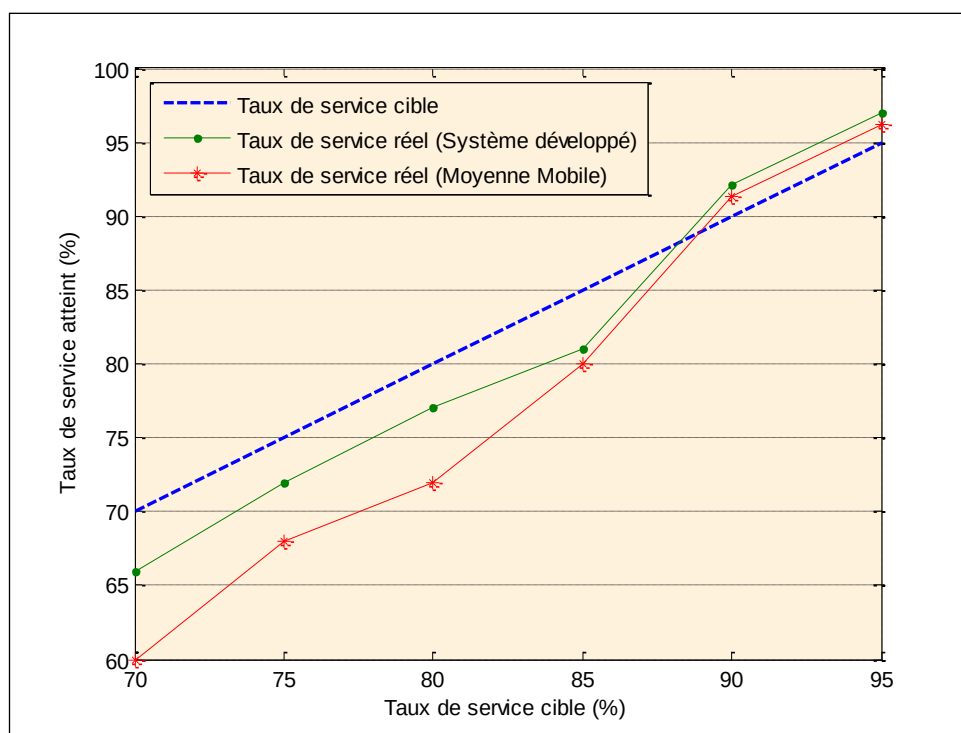


Figure 45 : Taux de service (Pharmacie 25_PE)

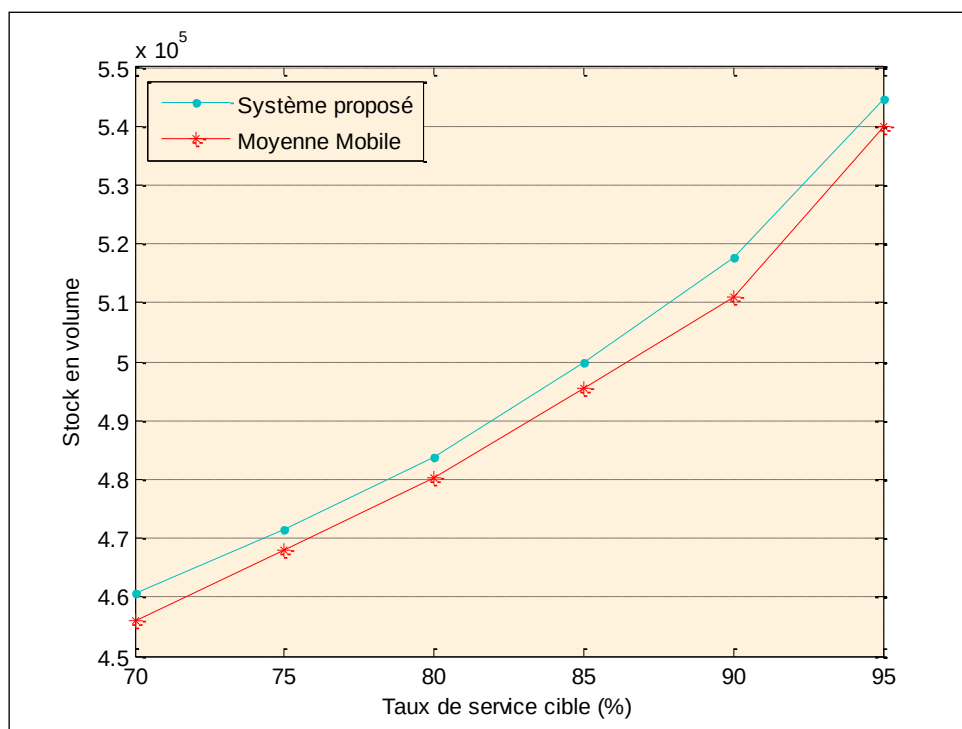


Figure 46 : Stock en volume (Pharmacie 25_PE)

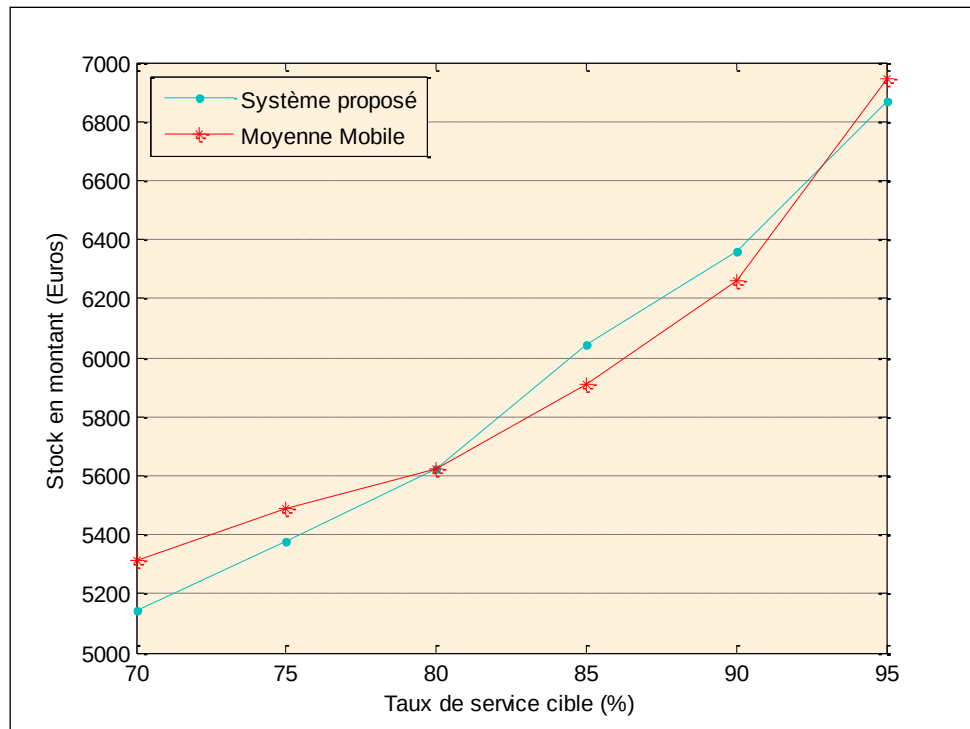


Figure 47 : Stock en montant (Pharmacie 25_PE)

Pour la pharmacie 24_MED (Figures 48, 49, 50), les taux de services de l'ancien système sont également en deçà de ceux renvoyés par le nouveau. Ceci se fait au dépend d'un accroissement du stock immobilisé qui reste maîtrisable (autour de 4% d'accroissement en volume, et de 5% d'accroissement financier).

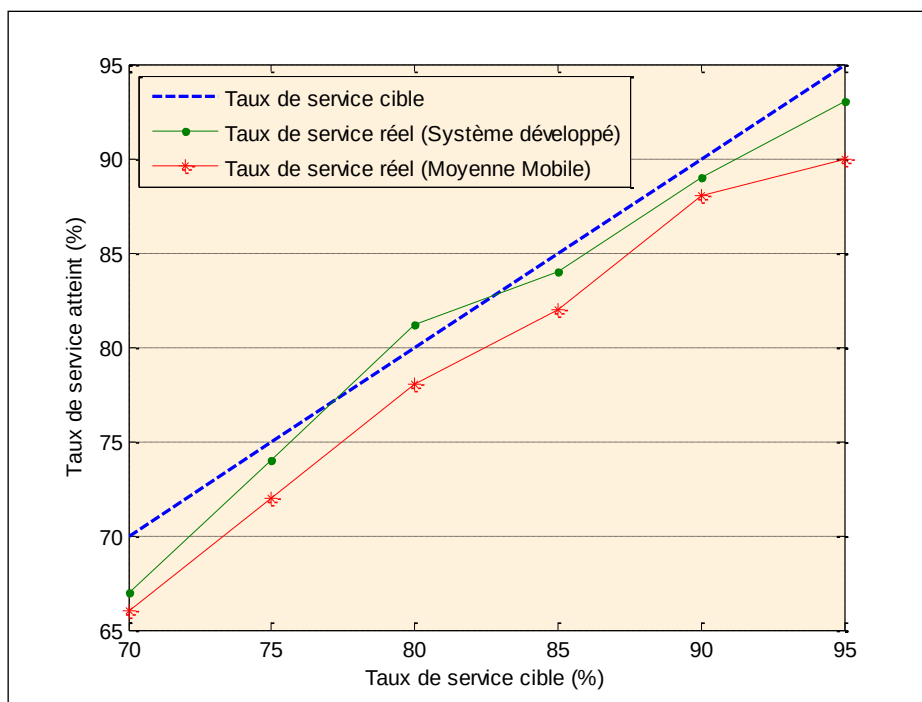


Figure 48 : Taux de service (Pharmacie 24_MED)

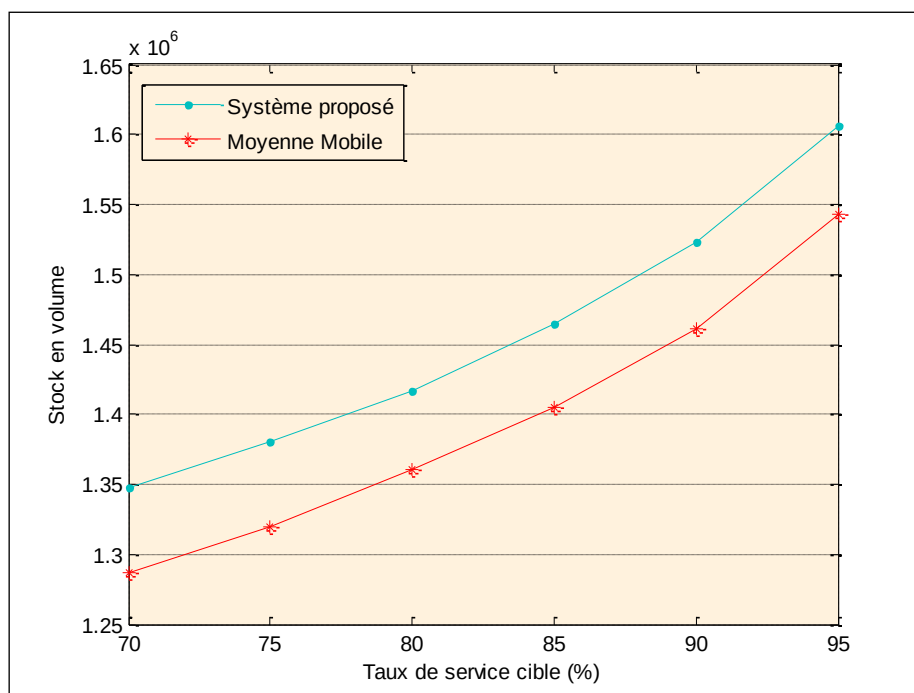


Figure 49 : Stock en volume (Pharmacie 24_MED)

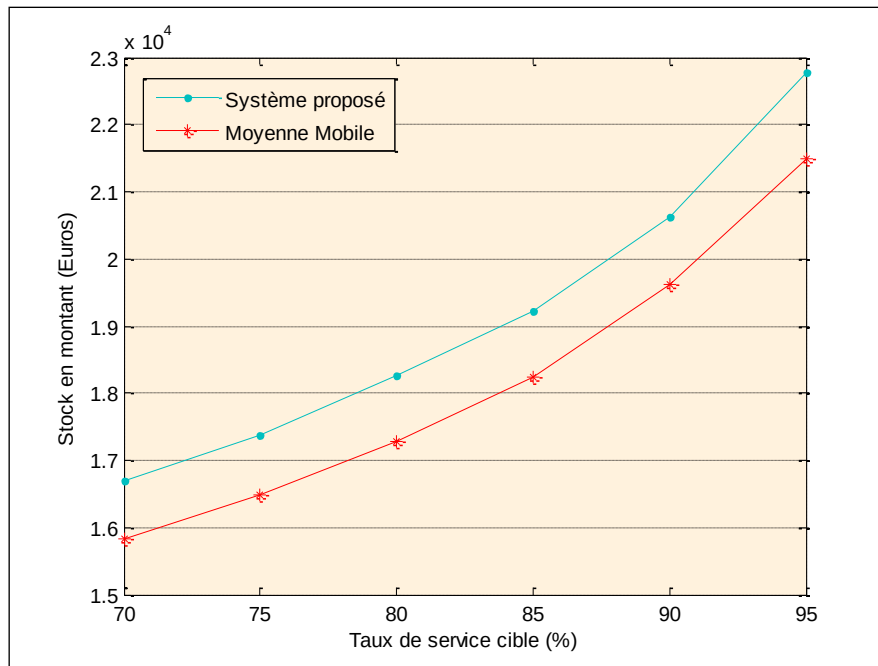


Figure 50 : Stock en montant (Pharmacie 24_MED)

L'analyse effectuée ci-dessus a été opérée sur des échantillons hétérogènes de la pharmacie (Contenant différents profils de consommation). Ceci nous a permis d'opérer une première évaluation globale du système.

Cependant, certains profils complexes sont ici agrégés avec de nombreux échantillons aux profils relativement stationnaires, pour lequel l'approche par moyenne mobile est déjà efficace. Ainsi, la valeur ajoutée, pouvant être apportée par notre système, par rapport à ces cas problématiques est masquée. Nous proposons d'illustrer, dans le paragraphe suivant, l'apport de notre système par l'étude de cas d'un produit complexe. En effet, ce sont ces produits, a priori imprévisible, qui posent le plus difficultés de gestion au quotidien en générant des ruptures récurrentes et de coûteuses demandes de réapprovisionnement en urgence.

7 Etude de cas d'un produit à profil complexe

Nous présentons dans ce paragraphe les performances renvoyées par notre système pour le cas d'un article non stationnaire. Nous considérons un produit identifié comme saisonnier par la procédure d'identification de tendance que nous avons implémentée. Nous opérons la même séquence de simulations réalisées précédemment, moyennant en premier lieu la moyenne mobile, puis le calcul recommandé par le système : (1) évaluation de la performance prévisionnelle sur les sept premiers mois de l'année 2010, (2) évaluation de la performance du stock.

Le produit identifié est l'AERIUS 5MG CPR géré au niveau de la plateforme 24_MED. Ce produit est indiqué pour le traitement de la rhinite allergique. Il s'agit d'une affection médicale bénigne secondaire à une hypersensibilisation à une substance étrangère dénommée allergène (pouvant être du pollen par exemple).

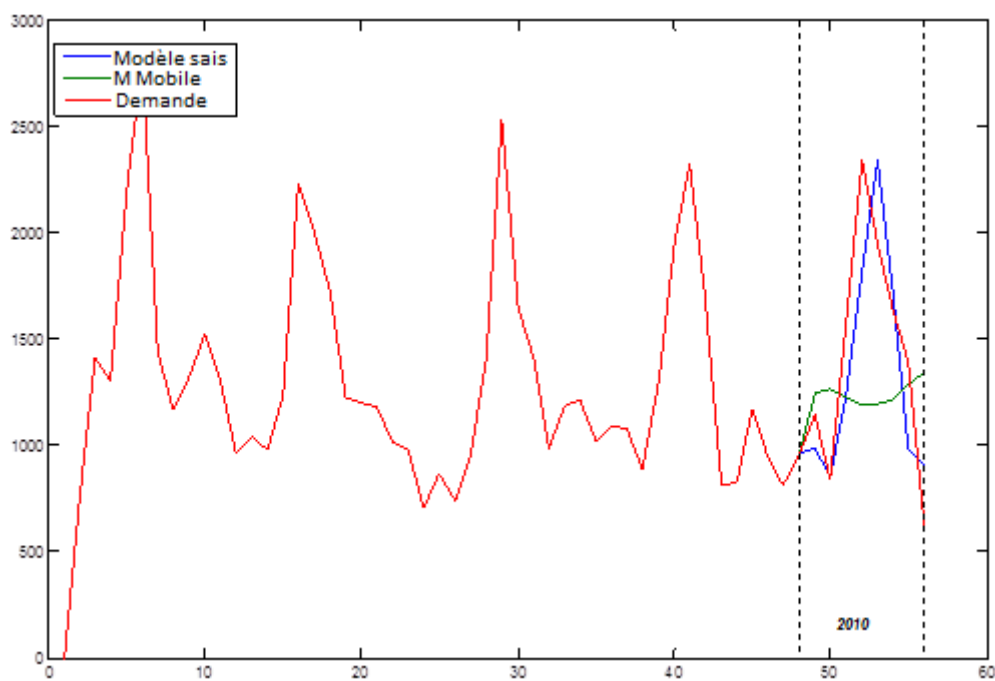


Figure 51 : Consommation historique de l'article AERIUS

Pour ce profil, le système détecte une saisonnalité, correspondant à un pic de consommation annuel se reproduisant autour du mois de mai. Le modèle appliqué est donc un modèle d'estimation saisonnier. Nous représentons dans la figure 51, le profil de consommation du produit retenu ainsi que les courbes d'estimations correspondant à une moyenne mobile et

au modèle saisonnier appliqué. Ce dernier permet en effet de mieux reproduire le pic de consommation.

Nous présentons dans le tableau 8, les performances prévisionnelles renvoyées par les deux méthodes appliquées.

Tableau 8 : Performances prévisionnelles sur les 8 premiers mois de 2010

	Ancien Système	Système Implémenté
Variation intrinsèque de la série	32,75%	
MAPE	38,25%	21,59%
MAE	472	233

Le système implémenté obtient de meilleures performances, pour les deux indicateurs de mesure d'erreur. Le pourcentage d'erreur obtenu par la moyenne mobile (38,25 %) est supérieur à la variation intrinsèque de la série (32,75 %). L'usage amène à considérer cette erreur comme trop importante pour que la prévision soit fiable. Le modèle saisonnier permet d'atteindre un pourcentage d'erreur de 21,59 %. Il est donc inférieur à la variation intrinsèque de la série ce qui le rend exploitable en terme de prévision. Par ailleurs, le nouveau système nous permet de diviser par deux l'erreur de prévision.

Tableau 9 : Taux de service atteint (Cas de l'AERIUS)

	Ancien Système	Nouveau Système
Taux de service cible (%)	80%	
Taux de service atteint (%)	65 %	78 %

Pour un taux de service cible de 80 %, le nouveau système permet d'atteindre un meilleur taux de service réel : 78 % contre 65 % renvoyé par l'ancien système.

Nous augmentons par la suite le taux de service cible, pour le cas de l'ancien système, afin d'atteindre un niveau de service réel comparable (Autour de 78 %) et mesurer le niveau de stock nécessaire à l'obtention de cette même qualité de service.

Tableau 10 : Performance du stock (Cas de l'AERIUS)

	Ancien Système	Nouveau Système
Taux de service atteint (%)	78 %	78 %
Stock moyen cumulé (Volume)	1243	1030

Pour garantir un niveau de service mesuré comparable, le nouveau système permet de réduire le stock moyen nécessaire de 20,68%. (le stock moyen en volume passe de 1030 à 1243).

8 Conclusion

Dans ce chapitre nous avons validé l'approche de la classification couplée à la prévision à travers l'étude du cas des Hospices Civils de Lyon. Des échantillons, issus de six plateformes pharmaceutiques, ont été testés avec l'ancien système et le système implémenté.

Une première évaluation, en terme de précision prévisionnelle a été effectuée. Le système implémenté permet d'atteindre de meilleures performances prévisionnelles pour le cas de cinq plateformes pharmaceutiques.

Une deuxième évaluation, en terme de performance de stock, a été aussi réalisée. Pour le cas des pharmacies testées, le système proposé permet d'atteindre des taux de service plus élevés pour des niveaux de stock comparables à ceux générés par l'ancien système. Par ailleurs, le système proposé réagit mieux par rapport à des profils de consommation plus complexes. Un même taux de service est atteint pour des immobilisations nettement plus faibles.

Chapitre VI. Conclusions et perspectives

1 Conclusions

Ce travail de thèse s'inscrit dans le cadre de l'amélioration d'un module de préconisation de commandes implémenté dans le système de gestion d'entrepôts pharmaceutique *Gildas Hospilog* développé et commercialisé par la société française *KLS logistic systems*. Le secteur de l'hôpital, et malgré l'évolution récente de la logistique hospitalière, manque considérablement d'outils d'aide à la décision. Les plateformes pharmaceutiques n'ont pas encore atteint le niveau de performance des entrepôts industriels classiques. On peut regretter que l'implémentation des systèmes de gestion d'entrepôts et des outils de gestion de stock ne soit pas assez généralisée au sein de ces structures.

Les fonctions de préconisation de commande et de gestion de stock ont longtemps été supportées par des modules intégrés dans des **ERP** ou des **APS** alors que les **WMS** étaient exclusivement dédiés aux opérations d'exécution et traçabilité des produits. L'évolution permanente des systèmes de gestion, due à la maturité de ces applicatifs et à l'évolution croissante des exigences métiers, notamment le monde de l'hôpital, a étendu le domaine d'application des **WMS** et leur périmètre fonctionnel. La gestion et l'optimisation des stocks, fonction centrale dans les pharmacies d'établissement, a toute sa place dans ce nouveau paysage fonctionnel.

Nous sommes partis du module existant, déjà implémenté et commercialisé, et avons procédé à une analyse approfondie de ces composantes. Cette analyse de l'existant, confrontée aux réflexions et idées entretenues par les décideurs et ingénieurs de l'entreprise KLS, nous a permis de définir les principaux chantiers de recherches à explorer et qui nous paraissaient générateurs de fonctionnalités nouvelles. Deux principales questions ont été définies : la problématique de classification des articles en stock, dans le but de générer le plus justement les articles pilotes et la problématique de l'estimation des consommations, dans le but de calculer, avec précision, des estimations de la consommation des articles afin d'alimenter le calcul des quantités de commandes et des différents paramètres des politiques de gestion.

À travers le premier volet, celui de la classification des produits, nous avons opéré une analyse de la littérature afin d'identifier les travaux ayant évoqué la classification des articles et les pratiques et méthodes appliquées. Des méthodes multicritères, issues de diverses disciplines, ont été avancées et adaptées afin de classer les articles de stock avec plus de justesse qu'avec une classification ABC classique. À cet égard, deux principales questions se sont posées, la première portant sur les critères les plus pertinents qu'il convient d'adopter et la deuxième sur les algorithmes de calcul et d'intégration des critères choisis. En analysant de près les algorithmes présentés dans la littérature, et en tenant en considération les contraintes inhérentes à notre problématique, nous avons choisi d'adopter deux méthodes multicritères, la méthode AHP et une méthode basée sur l'optimisation linéaire, en plus de la classification ABC classique. Un concept a été imaginé autour de ces éléments et un module de classification des articles a été développé. Ce module constitue un outil décisionnel d'aide à la classification des articles du stock. Il permet, par déploiement des méthodes considérées, d'appliquer ou bien une démarche monocritère classique, ou bien une démarche multicritères, auquel cas le choix de divers attributs, parmi ceux qui sont accessibles, et le choix de l'une d'entre les deux méthodes sont possibles.

Dans le deuxième volet, nous nous sommes intéressés aux prévisions de consommation de l'entrepôt, le système initialement utilisé appliquait une moyenne mobile pour l'ensemble des articles stockés, cette technique est très simple d'appréhension, mais peut entraîner des résultats dégradés notamment dans le cas de typologies de consommation présentant de la tendance ou un phénomène saisonnier.

Nous avons recensé diverses techniques de prévisions commentées dans la littérature et utilisées dans la pratique afin de définir des éléments de calcul nous permettant de mettre en place un système de prévision plus sophistiqué. Notre choix s'est porté sur les méthodes variantes du lissage exponentiel. Cette famille de techniques est relativement simple d'appréhension et très facile d'implémentation informatique. Elle offre, en même temps, des potentialités de modélisation à travers la prise en compte de caractéristiques clés comprenant les phénomènes tendanciels et saisonniers.

En nous appuyant sur des outils et tests statistiques, empruntés aux travaux sur les séries temporelles, nous avons mis en place un système dynamique auto-adaptatif permettant de

caractériser l'évolution de la consommation et de faire un choix intelligent de la technique de prévision la mieux adaptée au cas recensé.

Un module de prévision a été développé autour de ce concept et a permis de mettre en exergue la fonction prévision, fonction souvent absente des WMS. Au-delà de l'amélioration prévisionnelle pouvant être apportée par les modèles introduits, l'intérêt de ce module réside dans la remontée d'informations prévisionnelles, d'indicateurs d'alertes et performances ainsi que dans la visualisation de courbes d'évolutions et de tendances. Ces outils et modes de consultation constituent des éléments d'informations très utiles et synthétiques aidant dans la prise de décision.

Afin d'évaluer les performances du système développé, nous avons procédé à une étude de cas dans laquelle nous avons considéré des données propres aux hospices civils de Lyon. Cet établissement hospitalier, l'un des plus grands en France, assure le réapprovisionnement de ses différents services à travers une chaîne de plateformes logistiques, pharmacies d'établissement, gérée par l'ancien module de préconisation des commandes du *WMS Gildas Hospilog*.

Nous avons reconstitué des données historiques propres à certaines de ces pharmacies puis constitué des échantillons de tests par application d'un scénario de classification multicritères. En subdivisant l'historique de consommation en deux parties, un horizon d'initialisation et un horizon de simulation, nous avons généré des prévisions de consommation en nous appuyant respectivement sur l'ancien système puis sur le nôtre. Pour la majorité des pharmacies considérées, notre système permet l'amélioration des précisions prévisionnelles.

L'extension des analyses à l'évaluation du stock permet ainsi de montrer l'intérêt du système développé, un meilleur taux de service peut être atteint en regard d'un accroissement maîtrisable du stock. L'évaluation de l'impact sur le fonctionnement des stocks pourra être traitée plus en détail dans des travaux futurs.

Les deux briques fonctionnelles que nous avons développées (Module de classification et module de prévision) font désormais partie de l'offre standard du *WMS Gildas Hospilog* et ont été présentées lors de la trentième édition du salon européen de la logistique *SITL*, ayant

eu lieu à Paris en mars 2013. *SITL* est l'un des plus grands rendez-vous de la logistique et du transport en Europe et regroupe régulièrement des dizaines de milliers de professionnels de la logistique. Cette première exposition, dans un événement logistique d'envergure, nous a permis d'avoir une première appréciation de l'intérêt que pourraient accorder les professionnels à notre concept.

Un retour très positif a été recensé, notamment de la part d'entreprises industrielles, et exprimé particulièrement par l'équipe *R&D* de la société *Fives Cinetic*. *Fives Cinetic* fait partie du groupe mondial *Fives* qui conçoit et réalise des équipements de procédés, des lignes de production et des usines clés en main pour les plus grands acteurs mondiaux de divers secteurs. Il a entamé, depuis un certain temps, un partenariat avec *KLS*, afin de proposer des offres globales comprenant des systèmes de manutention et un *WMS*. Etant sensibles aux problématiques de gestion des flux et anticipation de la demande, cette équipe trouvait un grand intérêt à l'adoption du module dans divers contextes d'application. Une installation du module développé nous a été demandée, pour être testée au niveau de l'une de leurs plateformes de test située à Lyon.

Dans ce travail de thèse, deux principaux verrous scientifiques ont été identifiés. Le premier concerne la problématique de classification multicritères des articles de stock. Plusieurs travaux de recherche ont abordé cette question par adaptation de techniques de calcul multicritères. Les méthodes évoquées comprennent la méthode AHP, souvent appliquée dans les problèmes décisionnels multicritères, des techniques issues de l'intelligence artificielle ainsi que l'application de la programmation linéaire. *Cakir & Canbolat (2008)* ont présenté l'unique travail de la littérature ayant adapté une approche multicritères au cas d'un système de classification des articles du stock. Notre travail permet ainsi d'aller plus loin, en proposant un outil décisionnel permettant l'exploitation de deux approches multicritères.

Le deuxième verrou concerne la question des prévisions des consommations. La problématique des prévisions est largement commentée dans la littérature. Une lignée de travaux de la littérature a traité de la question de l'implémentation des techniques de prévision dans des systèmes experts. De par leur simplicité d'utilisation et leur robustesse, les méthodes issues du lissage exponentiel ont souvent été adoptées. Dans ce travail, nous puisons dans cette catégorie de méthodes. Une série d'outils et de tests statistiques, issus de la littérature sur l'analyse des séries temporelles, nous permet de compléter le concept et de mettre en place un système autonome et auto-adaptatif. Ces éléments nous permettent

d'apporter une première réponse intéressante à la problématique de la gestion de la chaîne logistique pharmaceutique.

2 Perspectives

Le travail entretenu dans cette thèse a permis d'enrichir le logiciel de gestion d'entrepôt pharmaceutique *Gildas Hospilog* en étendant son périmètre fonctionnel. Jugeant très positifs les résultats atteints et les premiers retours, une volonté d'aller encore plus loin dans ce travail est clairement exprimée par l'équipe dirigeante. L'amélioration de ce concept se fera sur deux fronts : (1) l'adaptation et la généralisation du module au restant des produits de *KLS* et (2) l'amélioration fonctionnelle du module.

(1) L'intégration du module aux restants des produits : *KLS* développe, en plus du système de gestion hospitalier *Gildas Hospilog*, des systèmes de gestion d'entrepôts standards à destination du monde industriel et de la grande distribution. Deux principales offres sont commercialisées : *Gildas WM* et *Gildas WE*. *Gildas WM* est l'équivalent de *Gildas Hospilog* mais constitue une offre plus générique à destination de divers secteurs. *Gildas WE* est dédié à la gestion d'entrepôts de petite taille, mais comprend la même logique fonctionnelle (Traçabilité, Zonage, etc.). Une mise à niveau de ces applicatifs est prévue, dans la continuité de notre travail, afin de les faire profiter de cette évolution fonctionnelle. Ceci peut s'avérer d'un grand intérêt, d'autant plus que le parc équipé est assez conséquent (plusieurs plateformes en France et en Suisse) et que les clients concernés expriment en continu des besoins de support plus pointus en gestion des stocks et en anticipation des commandes.

(2) Amélioration fonctionnelle du module : Le travail que nous avons réalisé a porté sur la partie classification et la partie prévision. Ces deux briques viennent en support aux calculs des quantités de commande opérés par les politiques de gestion implémentées. Cette dernière partie n'a pas été traitée et il nous paraît prioritaire, à ce niveau d'avancement, d'orienter dans le court terme, l'action autour des politiques de gestion. Plusieurs questions pourraient être traitées et donneraient lieu à d'autres évolutions fonctionnelles : l'affectation intelligente des politiques de gestion aux différents profils d'articles, l'introduction de la notion de taux de service, jusque-là absente de l'application, le calcul dynamique des paramètres des politiques (notamment le seuil de commande, les stocks de sécurité, etc.) sur la base de taux de service cible et bien d'autres points portant sur l'optimisation des stocks.

Sur le moyen et le long terme, d'autres chantiers portant sur des problématiques du domaine de l'entrepôt pourraient être envisagés. En plus de la fonction préconisation des commandes, des modules complémentaires ont été développés par *KLS* en supports à des activités et en réponse à l'évolution perpétuelle des systèmes de gestion d'entrepôts. Des questions portant sur la planification et l'ordonnancement des commandes, la définition intelligente des zones à travers l'entrepôt, l'affectation des caristes à travers les zones ou l'optimisation des plans de cueillettes à travers les différents emplacements pourraient être traitées et améliorées.

Dans ce travail de thèse, des modèles empruntés à la littérature scientifique ont été implémentés pour compléter un progiciel de gestion d'entrepôt. Le module de classification développé est basé sur deux techniques de classification multicritères. La première est la méthode AHP. Choisie pour sa simplicité d'application et de conceptualisation, cette dernière présente cependant l'inconvénient de la subjectivité liée au jugement humain. La deuxième technique adoptée est basée sur la programmation linéaire. Elle permet de pallier l'intervention humaine en permettant une classification automatique. Nous avons, pour notre part, appliqué ces résultats par implémentation dans un système expert. Un travail d'investigation sur la méthode utilisée pourrait être entrepris afin d'affiner les résultats de la classification multicritères.

Les résultats du processus de classification sont utilisés, à ce niveau d'avancement, à but informatif. L'utilisateur peut trier ces articles sur la base de ce critère, afin de mieux tracer les produits pilotes. L'ensemble des articles est cependant géré au moyen d'une même politique de gestion. Il est intéressant de pouvoir évaluer l'impact de l'adoption de différentes politiques de gestion, en fonction de la classe du produit, sur la performance de stock. Des travaux comme ceux de *Chu et al. (2008)* ou encore *Chen (2011)* ont abordé cette problématique.

Le module de prévision implémenté est basé sur des techniques extrapolatives. Ces modèles ont permis d'apporter un élément de réponse intéressant à notre problématique en intégrant plus de complexité dans les données et une facilité d'implémentation. Ces techniques exploitent exclusivement l'historique de consommation sans autre distinction par rapport au profil des produits traités. Il pourrait être intéressant d'investiguer l'intégration d'autres critères explicatifs, en plus de l'historique de consommation, et d'évaluer leur apport par rapport aux méthodes choisies.

Une des difficultés rencontrées dans le domaine hospitalier réside dans le fait que les données de consommation recueillies au niveau du WMS ne constituent pas l'information réelle sur la consommation finale des patients. Une distorsion de l'information a en effet lieu suite à l'agrégation de la consommation des services et la transmission le long de la chaîne logistique. L'évolution des Systèmes Informatiques Hospitalier incluant la gestion des dossiers patients et la prescription pharmaceutique par le médecin devrait permettre le déploiement de la dispensation nominative individuelle au sein des établissements. L'interfaçage de ces systèmes d'informations avec les WMS permettraient aux plateformes pharmaceutiques de travailler avec les consommations réelles des patients et de réduire les stocks . Il pourra alors être intéressant de tester les potentialités de notre système sur ces données plus fidèles.

Bibliographie

Adya, M. & Lusk, E. J. (2012). Designing Effective Forecasting Decision Support Systems: Aligning Task Complexity and Technology Support. Decision Support Systems. Book edited by Chiang Jao, ISBN 978-953-51-0799-6, Published: October 17, 2012 under CC BY 3.0 license.

Alavi,M. & Joachimsthaler,E.A. (1992). Revisiting DSS implementation research: a metaanalysis of the literature and suggestions for researchers. MIS Quarterly 16, 95-116.

Anderson, O. D. (1976). Time Series Analysis and Forecasting: The Box-Jenkins Approach. Butterworths, London.

APTEL O. and H. POURJALALI, " Improving activities and decreasing costs of logistics in hospitals A comparison of the U.S. and French hospitals ", The international Journal of Accounting, 2001, 36, 65-90.

Archibald, B. C. and Silver, E. A. 1978. (s, S) Policies under Continuous Review and Discrete Compound Poisson Demand. Management Science. Vol. 24, No. 9, pp. 889-909.

Arinze, B. Selecting appropriate forecasting models using machine learning methods. Department of management, Drexel University, Philadelphia, Pennsylvania, (1992), 1-20.

Armstrong, J.S. Long Range Forecasting. John Wiley and Sons, Inc., New York, 1985.

Axsater S. (2003) Approximate optimization of a two-level distribution inventory system, IJPE, 81-82, 545-553.

Axsater S, Zhang W.F. (1999) A joint replenishment policy for multi-echelon inventory control. IJPE, 59,243-250.

Babai, M. Z. (2005). POLITIQUES DE PILOTAGE DE FLUX DANS LES CHAÎNES LOGISTIQUES : IMPACT DE L'UTILISATION DES PREVISIONS SUR LA GESTION DE STOCKS. Thèse de doctorat.

BABOLI, B., MOYAUX, T., MEHRABI, A., MARC, N. (2010). Réorganisation des réseaux de distribution d'une chaîne logistique pharmaceutique aval : Comparaison des approches centralisée et décentralisée. GISEH 2010, Clermont-Ferrand, France, 2, 3 et 4 septembre 2010, 8 pages.

Benson, P. G. & Onkal, D. (1992). The effects of feedback and training on the performance of probability forecasters. *International Journal of Forecasting* 8, 559-573.

BERETZ L., "La logistique hospitalière: le point de vue de la pharmacie", Colloque "L'hôpital et la fonction logistique », Hôpital Expo, 20 mai 2002;

Bourbonnais, R., Usunier, J. C. (2013). Prévission des Ventes, Théorie et pratique. Economica, 5^{ème} édition.

Box, G. E. P and Jenkins, G. M (1970). Time series analysis: Forecasting and control, San Francisco : Holden-Day.

Boylan, J. E., Syntetos, A. A., Karakostas, G. C. Classification for forecasting and stock control: a case study. *Journal of the Operational Research Society* (2008) 59, 473 --481.

Bunn, B. L., and Wright, G. Interaction of judgmental and statistical forecasting: Issues and analysis. *Management science*. 37, 5 (1991), 501-518.

Cakir O, Canbolat MS (2008). A web-based decision support system for multi-criteria inventory classification using fuzzy AHP methodology. *Expert Sys. Appl.* 35 : 1367-1378

Carbone, R., Andersen, A., Corriveau, Y., and Corson, P. P. Comparing for different time series methods the value of technical expertise individualized analysis, and judgmental adjustment. *Management Science*. 29, 5 (1983), 559-566.

Chandra C, Grabis J. Inventory management with variable lead- time dependent procurement cost. *Omega* 2008; 36; 877 – 887

Chatfield, C. *The Analysis of Time Series: An Introduction*. Chapman and Hall, New York, 1988.

Chen, J.-X. (2011). Peer-estimation for multiple criteria ABC inventory classification. *Computers & Operations Research*.

Chu, C. W., Liang, G. S., Liao, C. T. (2008). Controlling inventory by combining ABC analysis and fuzzy classification. *Computers & Industrial Engineering* 55 (2008) 841–851.

Collopy, F., and Armstrong, J.S. Rule-based forecasting: Development and validation of an expert systems approach to combining time series extrapolations. Marketing Department, The Wharton School of the University of Pennsylvania, Philadelphia, Pennsylvania, (October 1991), 1-37

Davis,F.D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. *MIS Quarterly* 13, 319-340.

Di Martinelly, C. ; Guinet, A. ; Riane, F. : Chaîne logistique en milieu hospitalier : modélisation des processus de distribution de la pharmacie. 6e Congrès international de génie industriel – 7-10 juin 2005 – Besançon (France)

Duc, T. T. H.; Luong, H. T.; Kim, Y. 2008. A measure of bullwhip effect in supply chains with a mixed autoregressive-moving average demand process" . *European Journal of Operational Research* 187 : 243-256

Ertogral K., Darwish M. et Ben-Daya M. (2007) Production and shipment lot sizing in a vendor-buyer supply chain with transportation cost. *EJOR*, 176(3), 1592-1606.

Fildes,R. & Beard,C. (1992). Forecasting systems for production and inventory control. *International Journal of Operations and Production Management* 12, 4-27.

Fildes, R. & Goodwin, P. (2006). The Design Features of Forecasting Support Systems and their Effectiveness. *Decision Support Systems*, Volume 42, Issue 1, Pages 351-361.

Finlay, P.N. & Marples, C.G. (1997). A classification of management support systems. *Systems Practice* 10, 85-108.

Flores B.E., Olson D.L. & Dorai V.K. (1992) "Management of Multicriteria Inventory Classification",
Mathematical and Computer Modeling, 16(12), 71 - 82

Flores, B.E., Pearce, S.L. (2000). The use of an expert system in the M3 competition. *International Journal of Forecasting* 16, 485-496.

Flores B.E. and Whybark D.C., Multiple criteria ABC analysis, *International Journal of Operations and Production Management* 6 (3), 38-46 (1986).

Flores B.E. and Whybark D.C., Implementing multiple criteria ABC analysis, *Journal of Operations Management* 7 (1/2), 79-86 (1987).

Fustier, J. (2010). Les hôpitaux soignent leurs flux. *Supply Chain Magazine* N°48, 42-49.

Gaalman, G.J.C., Disney, S.M., 2006, State space investigation of bullwhip with ARMA(1,1) demand processes, *International Journal of Production Economics*, 104, 2, 327-339.

Ganeshan R. (1999) Managing supply chain inventory: A multiple retailer one warehouse multiple supplier model, *IJPE*, 59, 341-354.

Goodwin, P. (2000). Improving the voluntary integration of statistical forecasts and judgment. *International Journal of Forecasting* 16, 85-99.

Guvenir, H. A., & Erel, E. (1998). Multicriteria inventory classification using a genetic algorithm. *European Journal of Operational Research*, 105, 29-37.

Hadi-Vencheh, A. (2010). An improvement to multiple criteria ABC inventory classification. *European Journal of Operational Research*, 201(3), 962–965.

Han, J., & Kamber, M. (2001). *Data mining: Concepts and techniques*. Morgan Kaufmann Publishers.

Hassan, T. (2006). *Logistique hospitalière : organisation de la chaîne logistique pharmaceutique aval et optimisation des flux de consommables et des matériels à usage unique*. Thèse de doctorat.

Holt, C. C. (1957) Forecasting seasonal and trends by exponentially weighted moving averages, Office of Naval Research, Research Memorandum No. 52.

Huang, H., Lin, J., Chen, C., & Fan, M. (2006). Review of outlier detection. *Journal of Research on Computer Application*, 8, 8–13.

Huiskonen, J. (2001), "Maintenance spare parts logistics: Special characteristics and strategic choices", *International Journal of Production Economics*, vol. 71, no. 1-3, pp. 125-133

Jabbour, K., Riveros, J. F., Landsbergen, D., and Meyer, W. ALFA : Automated load forecasting assistant. *IEEE Transactions on Power Systems*. 3, 3 (1988), 908-914.

Jomaa, D., Monteiro, T., Besombes, B. (2012). Improvement of the inventory management module implemented in a pharmaceutical warehouse management system. ORAHS Conference. University of Twente.

Jomaa, D., Monteiro, T., Besombes, B. (2013, a). Development of an inventory classification module: Implementation in a warehouse management system. *International Conference on Industrial Engineering and System Management*.

Jomaa, D., Monteiro, T., Besombes, B. (2013, b). Design and development of a forecasting module: Case of a warehouse management system. MCPL Conference. International Federation of Automatic Control.

Kwong, K. K., and Cheng, D. A prototype microcomputer forecasting expert system. *Journal of Business Forecasting*. 7, 1 (1988), 21-27.

LANDRY S., BEAULIEU M., (2000), "Logistique hospitalière: un remède aux maux du secteur de la santé?", Groupe de recherche CHAINE, rapport n° 01-01, ISSN: 1485-5496

Liang, Y. (1997), ``The development of a knowledge- based inventory management system", PhD thesis, University of Salford.

Lim,J.S. & O'Connor,M. (1995). Judgemental Adjustment of Initial Forecasts: Its effectiveness and biases. *Journal of Behavioral Decision Making* 8, 149-168.

Lu, Y. & AbouRizk, S. M. (2009). Automated Box–Jenkins forecasting modelling. *Automation in Construction*. Volume 18, Issue 5, August 2009, Pages 547-558.

Makridakis, S. G., Wheelwright, S. C., Hyndman, R. (1998). *Forecasting: Methods and Applications*. 3rd Revised edition. John Wiley & Sons Inc.

Ng, W. L. (2007). A simple classifier for multiple criteria ABC analysis. *European Journal of Operational Research*, 177, 344–353.

Partovi, F. Y., & Anandarajan, M. (2002). Classifying inventory using artificial neural network approach. *Computers and Industrial Engineering*, 41, 389–404.

Partovi, F. Y., & Burton, J. (1993). Using the analytic hierarchy process for ABC analysis. *International Journal of Operations & Production Management*, 13(9), 29–44.

Pearce, S. L. (1995). Comparing an Expert System vs. supply, statistical process control, the use of bar codes in Traditional Approach to Forecasting Item Demand in a distribution. He does consulting and is principal in a Distribution Inventory Environment, Texas A and M software development firm dedicated to specific problems University, Unpublished Ph.D. Dissertation.

Pegels, C. C. (1969) Exponential forecasting: some new variations, *Management Science*, 12, No. 5, 311-315.

Rahman, S., and Bhatnagar, R. An expert system based algorithm for short term load forecast. *IEEE Transactions on Power Systems*. 3, 2 (1988), 392-399.

Ramanathan R. ABC inventory classification with multiple-criteria using weighted linear optimization. *Computers and Operations Research* 2006;33: 695–700

Reid R.A., The ABC method in hospital inventory management: A practical approach, *Production and Inventory Management* 28 (4), 67-70 (1987).

Richards, G. ; *Warehouse Management: A Complete Guide to Improving Efficiency and Minimizing Costs in the Modern Warehouse*. The Chartered institute of Logistics and Transport (UK).

Saaty T.L., A scaling method for priorities in hierarchical structures, *Journal of Mathematical Psychology* 15 (3), 234-281 (1977).

Saaty, T. L. (1980), "The Analytic Hierarchy Process", McGraw - Hill, New York.

Sampieri, N., Bongiovanni I., “Enjeux et perspectives des pratiques logistiques: pour une amélioration globale de la performance – le cas de l’hôpital public français”, RIRL 2000, les 3ème rencontres internationales de la recherche en logistique, 9-11 mai 2000, Trois-Rivières.

Sanders, N.R., & Ritzman, L.P. (2001). Judgmental adjustment of statistical forecasts. In. J.S. Armstrong (Ed.). *Principles of Forecasting*. Norwell:MA: Kluwer Academic Publishers, 405-416

Sexton, R. S., Doresey, R. E., & Johnson, J. D. (1998). Toward global optimization of neural networks: a comparison of the genetic algorithm and back propagation. *Decision support systems*, 22 (2), 171 - 185

Silver, E. A., Pyke, D. F., Peterson, R. (1998). Inventory Management and Production Planning and Scheduling. Third Edition. Wiley, John & Sons, Incorporated.

Syntetos, A. A., Boylan, J. E., Croston, J. D. On the categorization of demand patterns. Journal of the Operational Research Society (2005) 56, 495–503

Syntetos A.A., Boylan J.E., 2001, "On the bias of intermittent demand estimates" , International journal production economics, n.71, p. 457-466

Szmania, J., and Surgent, J. An application of an expert system approach to business forecasting. The journal of Business Forecasting Methods and Systems. 8, 1 (1989), 10-12.

Todd,P. & Benbasat,I. (1999). Evaluating the impact of DSS, cognitive effort, and incentives on strategy selection. Information Systems Research 10, 356-374.

Torabi, S.A., Hatefi, S.M., Saleck Pay, B. (2012). ABC inventory classification in the presence of both quantitative and qualitative criteria. Computers & Industrial Engineering. 63, 530–537

Weitz, R. R. NOSTRADAMUS a knowledge-based forecasting advisor. International Journal of Forecasting. 2, 3 (1986), 273-283.

Whybark, D. C. A comparison of adaptive forecasting techniques. The Logistics and Transportation Review. 8, 3 (1972), 13 - 26

Winters, P. R. Forecasting sales by exponentially weighted moving averages. Management science. 6, 3 (1960), 324 - 342

Wolfe, C., and Flores, B. Judgmental adjustment of earnings forecasts. Journal of Forecasting. 9 (1990), 389-405.

Yu, M.-C.(2010). Multi-criteria ABC analysis using artificial-intelligence-based classification techniques. Expert Systems with Applications,

Zahedi, F. (1994). Intelligent systems for business: expert systems and neural networks, California. Wadsworth Publishing, 1992.

Zhang, X. (2004) "The impact of forecasting methods on the bullwhip effect", International Journal of Production Economics, n° 88, pp.15–27.

Zhou P, Fan L. A note on multi-criteria ABC inventory classification using weighted linear optimization. European Journal of Operational Research 2007;182:1488–91