



AVERTISSEMENT

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

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

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

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

LIENS

Code de la Propriété Intellectuelle. articles L 122. 4

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

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

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

THESE

Présentée pour l'obtention du grade de
Docteur de l'Université de Lorraine
(mention automatique)

par

Mélanie NOYEL

Contrôle intégré du pilotage d'atelier et de la qualité
des produits.

Application à la société ACTA mobilier.

Soutenue publiquement le 10 novembre 2015 devant le jury composé de :

PR. PIERRE CASTAGNA

UNIVERSITE DE NANTES

RAPPORTEURS : PR. BERNARD GRABOT

ENIT TARBES

MDC HDR. MARIE-ANGE MANIER UNIVERSITE DE TECHNOLOGIE DE
BELFORT-MONTBELIARD

EXAMINATEURS : PR. PATRICK CHARPENTIER

UNIVERSITE DE LORRAINE

PR. ANDRE THOMAS (DIR)

UNIVERSITE DE LORRAINE

MDC HDR. PHILIPPE THOMAS
(Co-DIR)

UNIVERSITE DE LORRAINE

INVITÉ :

ALAIN GENET

PDG ACTA MOBILIER

REMERCIEMENTS

Du côté laboratoire, je tiens tout d'abord à remercier M. André Thomas, directeur de cette thèse, pour m'avoir fait confiance et permis de tenter l'aventure que représente la thèse CIFRE. Il aura su, au cours de ces trois années, me faire évoluer très positivement par ses remarques et conseils riches de sa sagesse que ce soit dans le domaine professionnel comme humain.

Je remercie ensuite M. Philippe Thomas qui aura su plus d'une fois me dépanner en me donnant des « coups de mains » appréciables dans les moments où je me retrouvais noyée sous le travail. C'est aussi grâce à lui que j'ai pu tant apprendre sur les réseaux de neurones dont il est un grand spécialiste.

Je remercie aussi M. Patrick Charpentier pour les échanges constructifs et les multiples prises de recul qui m'ont été possibles suite à nos discussions.

Du côté entreprise, je tiens évidemment à remercier en premier lieu M. Alain Genet, dirigeant d'Acta-Mobilier, pour la confiance qu'il a su m'accorder et la position fort intéressante qu'il m'a offerte aujourd'hui. J'en profite pour faire un point tout particulier sur cette entreprise dont la culture n'a rien à envier à quelques groupes de renommée mondiale. La politique mise en place qui a à cœur de valoriser l'humain en l'engageant dans tous les processus d'amélioration continue et de développement de l'entreprise est à la fois un pari risqué mais une opportunité énorme offerte à qui saura la saisir. Ce contexte fait de l'entreprise un vivier d'idées et de projets qui lui réussit et lui permet de tirer son épingle du jeu dans le contexte difficile actuel.

Le premier responsable technique avec qui j'ai eu l'occasion de travailler, M. Jean Baptiste Limoges, doit être tout spécialement remercié comme étant le co-créateur de POTER qui a évolué progressivement jusqu'à devenir un véritable MES pour l'entreprise. Il m'a permis de trouver une première légitimité dans l'entreprise.

Je tiens aussi à remercier mon proche collègue M. Thomas Brault pour les nombreuses discussions et réflexions autour du système d'information ainsi que le soutien moral dont il a fait preuve tout au long de mon travail. Je lui dois aussi en grande partie mon niveau actuel en matière de programmation.

Et je remercie aussi tout particulièrement le responsable technique actuel M. Maxime Villain qui a su avant tout me redonner confiance en mes travaux et me pousser à finir malgré la difficulté que représentait la simultanéité de cette fin de thèse avec mon nouvel

emploi de responsable Système d'Information. Il est parvenu à valoriser une grande partie de mes travaux en m'accompagnant efficacement dans le déploiement terrain.

Du côté plus personnel, je tiens à remercier en premier lieu M. Robert Pingot qui a mis à ma disposition tout ce dont je pouvais avoir besoin au cours de ces trois années, ma maman Mme. Catherine Noyel pour la confiance sans faille qu'elle a en moi et mon compagnon M. Mickael Pingot pour m'avoir supportée, encouragée et parfois même un peu obligée dans les moments les plus difficiles.

Bien sûr il y a encore beaucoup d'autres personnes qui ont apporté leur pierre à l'édifice. Je ne souhaite pas faire une énumération complète de peur d'oublier quelqu'un. Je préfère adresser un grand merci à tous, pour tout ce que vous avez pu faire pendant cette période pas toujours facile à vivre. Il est sûr que s'il avait manqué un seul d'entre vous, cette thèse ne serait très certainement pas celle qu'elle est aujourd'hui.

Table des matières

AVANT PROPOS	1
I Contexte industriel	3
I.1 Introduction et objectifs	5
I.2 Problématique industrielle	5
I.2.1 Présentation globale de l'entreprise	5
I.2.2 Méthodologie/démarche d'analyse du problème industriel	9
I.2.3 Analyse du problème industriel.....	9
I.2.4 Synthèse sur les problèmes industriels.....	20
I.3 Etat de l'art général.....	22
I.3.1 Maitrise de la qualité sur des processus instables	22
I.3.2 Pilotage des flux dans un contexte de production perturbé par les reprises	26
I.4 Articulation logique de la thèse	33
II Maîtrise de la qualité on-line sur des processus instables.....	37
II.1 Introduction	39
II.2 Etat de l'art	42
II.2.1 Les méthodes d'apprentissage (problèmes de classification)	42
II.2.2 La détection des évolutions de processus.....	55
II.2.3 Problématique scientifique spécifique	58
II.3 Proposition.....	59
II.3.1 Méthodologie générale.....	59
II.3.2 Conception du classificateur	60
II.3.3 Conception de l'observateur.....	65
II.4 Conclusion.....	70
III Evolution réactive des règles de pilotage en fonction de la saturation de l'outil de fabrication	73

III.1	Introduction.....	75
III.2	Etat de l’art.....	75
III.2.1	Introduction : impact de la non-qualité sur la perturbation des flux	75
III.2.2	Les indicateurs de contrôle de flux de la littérature.....	77
III.2.3	Les méthodes pour le basculement des règles de pilotage	79
III.2.4	Problématique spécifique	80
III.3	Proposition	80
III.3.1	Modèle de simulation	81
III.3.2	Intégration des variabilités dans le modèle.....	82
III.3.3	Validation du modèle.....	83
III.3.4	Un nouvel indicateur de non-qualité : le ratio d’opérations	84
III.3.5	Mise en évidence de la saturation rapide de l’atelier.....	84
III.3.6	Mise en évidence de la supériorité d’une règle par rapport à une autre dans un contexte donné	86
III.3.7	La détermination des seuils d’absorption et de saturation.....	87
III.3.8	Une cartographie visuelle	90
III.3.9	Les règles de pilotage associées	92
IV	Applications Industrielles.....	95
IV.1	Introduction.....	97
IV.2	Maîtrise de la qualité on-line sur des processus instables.....	98
IV.2.1	Justification du cas d’implémentation	98
IV.2.2	Description du poste de travail concerné.....	99
IV.2.3	Conception du classificateur.....	100
IV.2.4	Conception de l’observateur	120
IV.2.5	Résultats et discussions	123
IV.3	Evolution réactive des règles de pilotage en fonction de l’état de saturation de l’atelier	124

IV.3.1	Application et perspective d'implémentation.....	124
IV.3.2	Construction de la cartographie des états de saturation par clustering 125	
IV.3.3	Association des règles de pilotage.....	126
IV.3.4	Résultats et discussions	128
IV.4	Container intelligent.....	129
IV.4.1	Choix de la stratégie d'identification.....	129
IV.4.2	L'identification intuitive du « noyau élémentaire »	132
IV.4.3	La notion de container intelligent.....	135
IV.4.4	Résultats et discussions	137
IV.5	Ordonnanceur centralisé assisté par le produit	139
IV.5.1	Descriptif du problème	139
IV.5.2	Solution par programmation linéaire.....	141
IV.5.3	Une application du contrôle par le produit	143
IV.5.4	Résultats et discussions	144
V	Conclusion et perspectives	147
VI	Figures, tableaux et références	153
VI.1	Figures.....	154
VI.2	Tableaux.....	155
VI.3	Références.....	155
VII	ANNEXE	169

AVANT PROPOS

Le travail présenté dans cette thèse a été mené en collaboration avec l'entreprise Acta-Mobilier située à proximité d'Auxerre, spécialisée dans le laquage haut de gamme de façades destinées à l'agencement de stands, magasins, hôtels, cuisines, bureaux, etc... et le département ISET (Ingénierie des Systèmes Eco-Techniques) du CRAN (Centre de Recherche en Automatique de Nancy). L'objectif général des recherches de ce département est de définir des méthodologies, des modèles ou des outils pour disposer d'une représentation globale et cohérente d'un système (industriel) cible à étudier afin d'aider à la prise de décision optimale. Les travaux décrits dans cette thèse s'inscrivent parfaitement dans ce cadre et plus précisément dans l'une des 3 orientations choisies par le département, à savoir l'ingénierie du pilotage des flux physiques et leurs synchronisations avec leurs flux d'informations.

Ces travaux ont eu lieu dans le cadre d'une convention CIFRE (Conventions Industrielles de Formation par la Recherche), ce qui constitue une double expérience, à la fois industrielle et académique, particulièrement formatrice. Ce contrat a d'ailleurs débouché sur une embauche pour le poste de responsable Système d'Information de l'entreprise.

I CONTEXTE INDUSTRIEL

I.1 INTRODUCTION ET OBJECTIFS

L'objectif de cette première partie est avant tout de définir correctement la problématique industrielle et notamment son périmètre.

Nous décrivons, après avoir présenté l'entreprise et son fonctionnement, une méthodologie inspirée de l'ingénierie système pour analyser de manière structurée le problème industriel. Celui-ci se divise en deux branches distinctes et pourtant fortement liées, à savoir une problématique de maîtrise de la qualité sur des processus instables et une problématique de pilotage de flux dans un atelier dont la production est fortement perturbée.

Un état de l'art couvrant donc ces deux domaines sera détaillé afin de définir le cadre scientifique général de la thèse.

I.2 PROBLEMATIQUE INDUSTRIELLE

I.2.1 Présentation globale de l'entreprise



La société Acta-Mobilier est à la fois un agenceur des espaces de marque et un co-traitant industriel. Pour ce marché de l'agencement, de grandes marques telles que Chanel, Nespresso, Peugeot, Baume ou encore

Mercier lui font confiance pour réaliser leurs stands présentés sur des salons ou pour l'agencement de leurs magasins (Figure 1).



Figure 1 - Stand Peugeot fabriqué par Acta pour un mondial de l'auto

Pour le marché de la co-traitance, son savoir-faire lui permet de rester le leader français de la façade laquée haut de gamme pour assurer la production de portes de meubles de cuisines, salles de bains, banques d'accueil, de plateaux de bureaux pour des grands noms tels que Schmidt, Mobalpa, Vogica, Arthur Bonnet, Pyram ou encore Lapeyre. Depuis plusieurs années, l'entreprise se tourne vers l'export avec de nouveaux clients en Allemagne et en Autriche où l'exigence client est encore plus élevée au niveau de la qualité.

Elle produit en moyenne 3000 panneaux laqués de très haute qualité par semaine à partir de panneaux de fibres de moyenne densité (MDF). Elle maintient son statut de leader sur le marché grâce à ses innovations. Elle est par exemple à l'origine de la création des façades à « aspect béton » que la mode actuelle met très en avant mais elle invente aussi en permanence de nouveaux prototypes et de nouvelles finitions.



Figure 2 - Le service "couleur à la demande" permet de maintenir la compétitivité de l'entreprise

Comme toutes les entreprises depuis le 20^{ème} siècle, Acta-Mobilier est aussi soumise à la « personnalisation de masse » et propose donc à ses clients des produits de plus en plus personnalisés fabriqués à la commande. Ainsi ses clients peuvent, par exemple, choisir les dimensions de leurs pièces au millimètre près ainsi que la couleur suivant le référentiel de leur choix (Figure 2).

L'entreprise ne peut donc pas produire sur stock car les prévisions de produits finis sont impossibles, et elle ne déclenche la fabrication (voire l'approvisionnement en ce qui concerne les coloris de laque) qu'à la commande. C'est d'ailleurs dans cette perspective qu'elle envisage aujourd'hui de fabriquer elle-même ses matières premières suivant ses besoins, notamment pour les couleurs à la demande.

Ce choix stratégique a évidemment un impact conséquent sur le délai de livraison qui doit pourtant être affiché comme très court pour que l'entreprise puisse rester concurrentielle. D'autre part, ce mode de fonctionnement soumet l'entreprise à l'irrégularité et à l'imprévisibilité des demandes clients et l'oblige à gérer des flux de produits très divers. Cela implique pour l'entreprise d'avoir un personnel très flexible et de pouvoir être très réactive face aux éventuels aléas de production.

Outre le choix conséquent offert au client en matière de produits, Acta souhaite aussi offrir un gage de qualité. Dans ce domaine, elle obtient la certification « tñv » allemande choisie pour ses exigences et sa rigueur. Dans sa quête d'amélioration continue, elle est aussi certifiée ISO 9001, ISO 14001, OHSAS 18001 et implémente de nombreux processus Kaizen¹ sur le terrain. Ses quatre piliers sont la satisfaction client, le progrès continu, le fonctionnement par processus, et la responsabilisation et l'implication de chaque collaborateur.

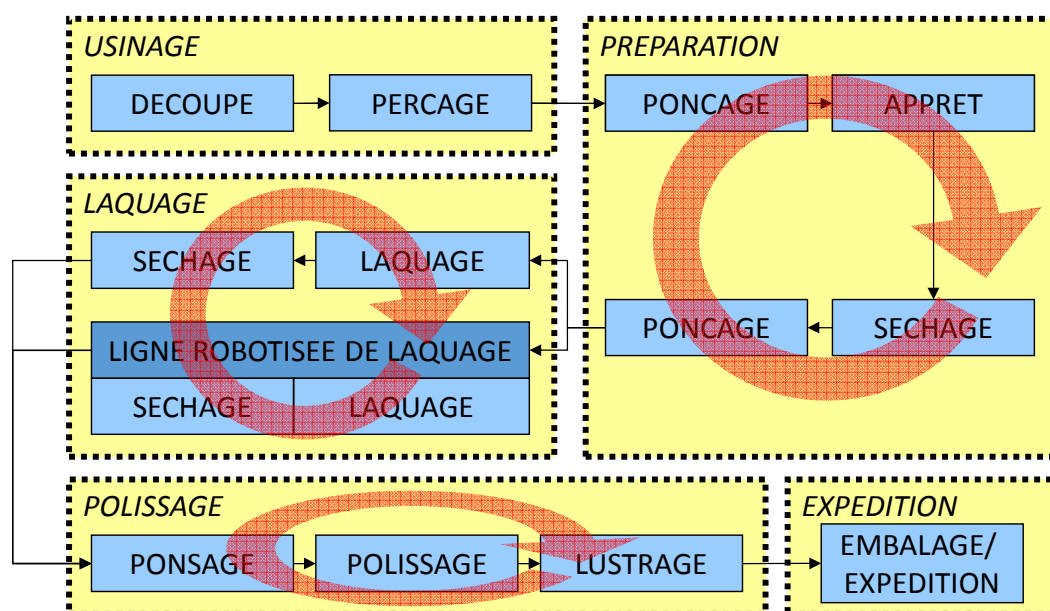


Figure 3 - Chaîne de production de l'entreprise Acta-Mobilier

¹ Le Kaizen est un état d'esprit partagé par tous les collaborateurs de l'entreprise à tous les niveaux hiérarchiques privilégiant, au fil des jours, les petites améliorations concrètes, simples et peu coûteuses.

L'activité de fabrication est divisée en 5 ateliers représentés sur la Figure 3. Le premier de ces ateliers est l'atelier usinage qui a pour objectif principal la découpe des pièces. Le type de panneaux le plus utilisé est principalement du MDF 19 mm mélaminé. Il s'agit d'un type de panneaux particulièrement adapté à la fabrication de meubles de cuisines et de salle de bains de par sa capacité à être mouluré comme le bois, ainsi qu'à être laqué (Voreux, 1987). Bien qu'il existe d'autres types de panneaux dans l'entreprise, nous suivrons, pour notre projet, uniquement le MDF car il représente le plus fort pourcentage d'utilisation pour les panneaux laqués. Les panneaux de dimensions non standards (approximativement 20% du flux) passent sur une scie de débit. Les autres se répartissent directement sur deux centres d'usinage qui fonctionnent en permanence en « deux-huit » ou en « trois-huit » suivant la demande client. Le code à barres assurant la majeure partie de la traçabilité est collé sur la pièce à la fin de cette étape de découpe. La possibilité d'erreur de l'opérateur à cette étape existe mais est très faible.

L'étape suivante est une étape de perçage automatique précédée d'une lecture des codes à barres permettant l'identification de la pièce.

Vient ensuite un ponçage manuel, ou par machine, qui précède l'apprêtage. Dans une même commande, il existe des pièces à apprêter sur une ou deux faces. Dans ce cas, les « apprêtées une face » doivent attendre les « apprêtées deux faces » pour rester dans le même lot.

Le laquage est soit manuel, ce qui oblige à poncer et laquer d'abord une face puis à poncer et laquer l'autre, soit automatique sur la ligne de laquage robotisée pour laquelle on doit poncer d'abord les deux faces.

Les panneaux passent ensuite sur une des deux machines de ponçage, puis sur une des deux machines de polissage, et enfin sur la machine de lustrage, ce qui constitue la chaîne de poly-lustrage. Toutefois il est possible que certaines pièces partent à l'emballage directement après le polissage sans passer par l'étape de poly-lustrage.

On a ensuite un point de contrôle qualité essentiellement visuel sous « cabine lumière » puis un filmage ou emballage carton.

Chacun des ateliers est géré par un responsable de production qui coordonne la production sans réel outil informatique. Face à un aléa, les décisions sont prises « au pied de la machine » grâce à des consensus entre responsables de production. Les plannings de

travail sont en réalité le reflet de l'ordre de livraison des commandes et la règle de pilotage mise en avant n'est autre que le traditionnel First In First Out (FIFO).

1.2.2 Méthodologie/démarche d'analyse du problème industriel

L'entreprise est particulièrement consciente de ses problèmes. Cependant, afin de cerner correctement ses exigences et ses besoins pour répondre correctement à ses attentes, une analyse globale a été faite. L'analyse du problème industriel doit être menée en adoptant un raisonnement le plus large possible afin de couvrir, si ce n'est la totalité, du moins la majeure partie des problèmes du contexte industriel. On peut s'inspirer de la méthode des « 5 Pourquoi » (Ohno & Rosen, 1988) en subdivisant chaque problème en sous-problèmes. Nous proposons la méthodologie présentée par le modèle de la Figure 4.

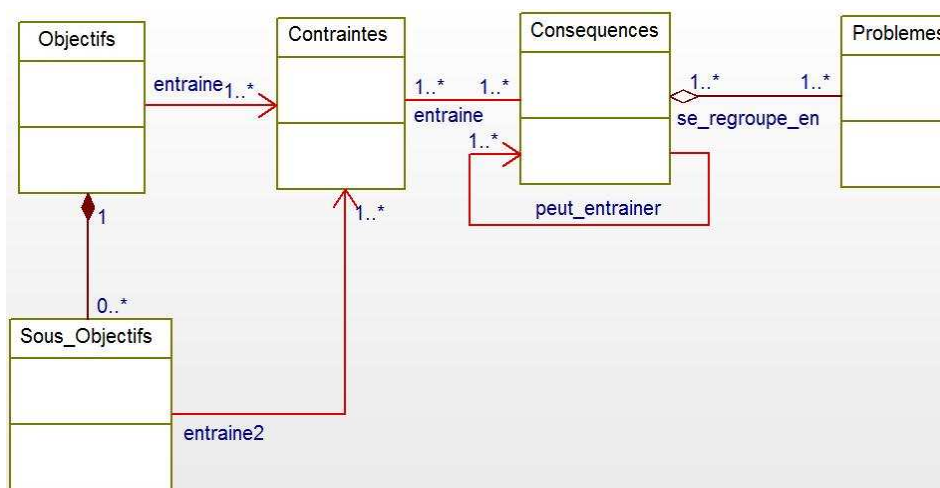


Figure 4 - Méthodologie d'analyse du problème industriel

Le premier niveau concerne les objectifs. Toute entreprise est conditionnée par ses objectifs définis dans sa stratégie, telle la compétitivité. Chacun des objectifs peut se suffire à lui-même ou comporter un ou plusieurs sous-objectifs. Chaque objectif ou sous-objectif entraîne ensuite des contraintes, chacune d'elles pouvant entraîner des conséquences, lesquelles permettront l'expression du problème industriel.

1.2.3 Analyse du problème industriel

La Figure 5 représente l'application directe de cette méthodologie pour la définition globale du problème industriel de l'entreprise Acta-Mobilier.

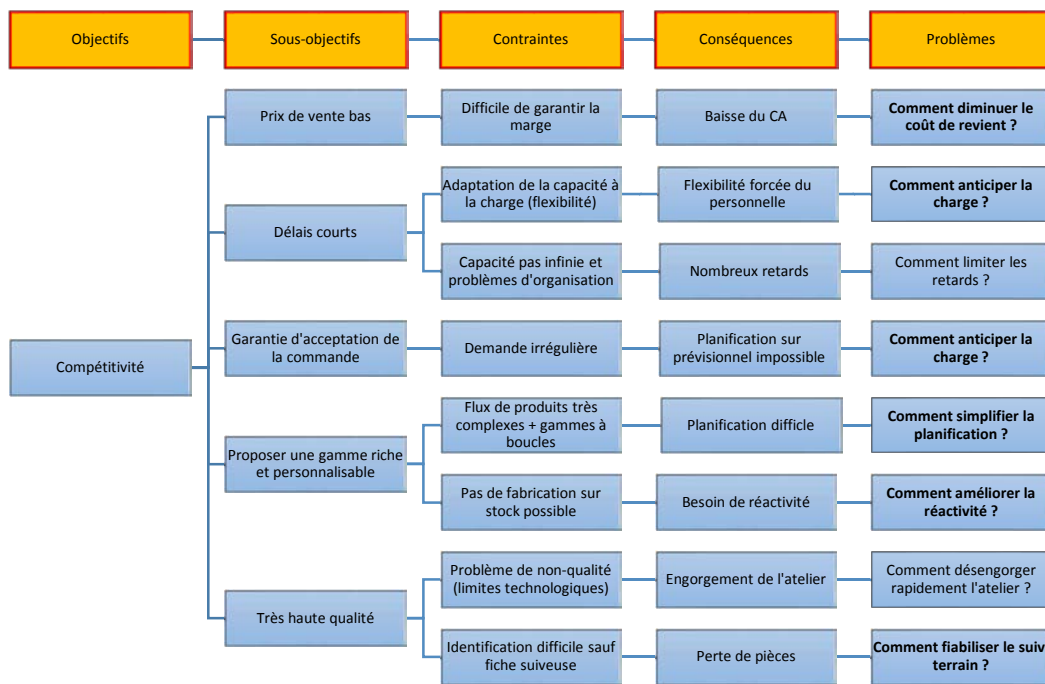


Figure 5 - Dérivation des exigences et contraintes pour définir le problème industriel

En partant de l'objectif principal qui est la compétitivité, il est possible d'en extraire 5 sous-objectifs qui permettront à l'entreprise de se démarquer face à la concurrence :

- fournir des produits à bas prix,
- présenter des délais de livraison très courts,
- garantir au client qu'on ne lui refusera jamais sa commande,
- proposer une gamme de produit riche et personnalisable,
- assurer un niveau de qualité très haut de gamme.

En transformant ces sous-objectifs en contraintes, puis en conséquences et enfin en problèmes, on obtient alors une liste de questions sur lesquelles travailler qui représentent le problème industriel. Les parties suivantes détaillent chacun de ces points.

1.2.3.1 Comment diminuer le coût de revient ?

Il y a divers moyens de minimiser le coût de revient (résumés dans Tableau 1) :

- Diminuer le nombre de réparations
- Conserver un stock minimal en évitant les ruptures, ce qui permet de limiter les dépenses de gestion de ce stock tout en évitant le surcoût d'une rupture en permettant un flux continu. La traçabilité peut avoir un impact très positif sur ce point (De Kok, Van Donselaar, & Van Woensel, 2008).
- Diminuer les stocks à l'actif circulant du bilan dégagerait une somme d'argent qui pourrait être réinvestie dans un projet de traçabilité à grande échelle sur l'entreprise mais nous n'aborderons pas ce point ici.

- Vérification de la disponibilité de chaque sous-composant avant lancement. Le phénomène de rupture de flux cité dans l'article (Joly, Frein, Gauthier, & Bernier, 2004) traitant de l'étude de l'impact des blocages sur le flux de production d'une usine terminale automobile se produit aussi couramment dans notre cas lors de rupture de stock de laque ou de panne du robot de laquage par exemple.

Si on connaît la composition exacte des stocks, on peut donc en tirer l'avantage de savoir avec exactitude si on est capable de tenir le planning de production. Par exemple, si le lot X de couleur xxx est prévu à l'usinage et que la couleur xxx n'est pas disponible avant 3 jours, ce n'est pas la peine de le lancer maintenant. Il serait plus intéressant de lancer à l'usinage le lot suivant dont le flux sera direct jusqu'à la fin de la production. Par ailleurs, la notion de stock de sécurité pourrait apporter une solution au problème de ruptures. Cependant, dans le cas de stocks très diversifiés (comme les couleurs à la carte) maintenir ces stocks pour chacune des couleurs induirait beaucoup de pertes financières.

Tableau 1 - Résumé des actions possibles vis-à-vis de la minimisation du coût de revient

Domaines	Moyens d'action
Qualité	Diminution du surcoût lié aux multiples réparations.
Gestion de stock	Eviter le surcoût d'une rupture des matières premières. Réinvestir l'argent dégagé par la diminution des stocks
Ordonnancement	Adapter le planning à la disponibilité des matières premières.

1.2.3.2 Comment anticiper la charge ?

L'entreprise garantit qu'aucune commande ne sera déclinée, quel que soit le nombre de commandes déjà passées par d'autres clients pour la même période. A ce jour, avant d'accepter une commande, l'entreprise n'a pas les moyens de vérifier s'il reste assez de disponibilité machine pour pouvoir la réaliser. Il faut donc adapter le temps de travail à la charge de travail (flexibilité). Ainsi, d'une semaine sur l'autre et comme nous l'avons déjà introduit, un poste peut passer d'un fonctionnement en 2/8 à un fonctionnement en 3/8 et *vice versa*. Il n'est donc pas possible actuellement de lisser la charge et le taux de charge est donc directement corrélé à la quantité commandée. Il pourrait être bon d'envisager une programmation à capacité finie et ainsi d'avertir le client que s'il commande aujourd'hui, le délai ne sera pas de 2 semaines mais de 3 à cause d'une demande élevée par rapport à d'habitude.

Les deux branches d'activité de l'entreprise présentent en plus une variation saisonnière très irrégulière. La variabilité de la charge est donc difficilement maîtrisable car totalement imprévisible. Il est donc particulièrement difficile de faire des prévisions dans ce contexte et les essais de prévisions actuels sont très peu concluants. Il est en effet possible d'atteindre des variations de demande de 30 à 40% d'une semaine sur l'autre et de +100% sur deux semaines.

La non-anticipation de la charge pousse l'entreprise à une flexibilité forcée du personnel. La solution envisagée par l'entreprise est de tenter d'avoir toujours 2 jours d'avance par rapport à la date d'expédition afin d'être moins pénalisée par les éventuels retards.

Toutefois, gagner du temps sur la production est certainement le point le plus facile à solutionner. En effet, il est bien connu qu'une pièce passe plus de temps à attendre entre les postes qu'à être transformée (temps à valeur ajoutée). Afin de nous rendre compte de l'effectivité de ce concept chez Acta Mobilier, en calculant les cadences à chaque poste et en ne tenant pas compte des postes de tri et des postes de séchage, il a pu être démontré que le temps de transformation réel était (action + manutention, en tenant compte des boucles présentées en Figure 6) de 40 min, soit 0,4 % du temps qu'elle passe dans l'entreprise. Le séchage au complet (cumul du séchage de toutes les couches qui constitue quand même un temps à valeur ajoutée) peut, quant à lui, représenter jusqu'à 50% du temps de production. Ce qui nous amène à un temps incompressible représentant seulement 50,5 % du temps total passé dans l'entreprise. 49,5 % de temps est donc perdu en attente, et est donc source d'amélioration potentielle.

Il serait donc possible, théoriquement, de diviser quasiment par deux le temps de production en travaillant en flux tendu. Pratiquement, il suffirait de réduire le temps d'attente des pièces de 20% (passer de 49,5% à 39,5%) pour atteindre le « J-2 ».

Tableau 2 - Résumé des actions possibles vis-à-vis de l'anticipation de la charge

Domaines	Moyens d'action
Information	Estimer la capacité restante afin de pouvoir donner un délai adapté lors de la commande en réalisant un suivi de l'atelier.
Ordonnancement	Se rapprocher du flux tendu pour réduire le temps de production avec une bonne gestion des règles de pilotage.
Statistiques des commandes	Sur les commandes Sur les variations de stock

Ainsi donc, ce manque d'anticipation sur la charge des semaines à venir impose à l'entreprise d'être particulièrement réactive en l'obligeant à considérer sa capacité de

production comme infinie. Les principales actions possibles dans ce cadre sont donc résumées par le Tableau 2.

I.2.3.3 Comment améliorer la réactivité ?

La nécessité de réaction (changer le planning établi) vient du fait que les paramètres de gestion de l'entreprise (niveaux des stocks, personnels absents, pannes machine...) ont changé par rapport au moment où la planification initiale (prédictive) a été faite. Pour pallier ce problème, il faut donc des informations reflétant la réalité des événements de production en temps réel. L'opérateur est le mieux placé pour observer cela. Il peut alors devenir acteur de la décision en ayant à sa disposition un module d'aide à la décision qui lui présenterait les meilleures alternatives possibles. Ce système lui permettrait de constater qu'il doit réagir et de choisir une « réaction » en connaissance de cause. C'est donc une question de « décision *juste à temps* ».

Le meilleur moment pour décider reste aussi à définir. Il peut être lié à la fréquence de mise à jour du « logiciel » ou à l'importance de changement des paramètres. Une planification globale pourrait être effectuée en amont de la production avec la volonté de faire un premier planning « au mieux », puis, de réagir « au fil de l'eau ». Le système serait alors réactif. La principale action possible dans ce cadre est résumée par le Tableau 3.

Tableau 3 - Résumé des actions possibles vis-à-vis du manque de réactivité

Domaines	Moyens d'action
Réordonnement réactif	Système de suivi du terrain (MES) et proposition des meilleures alternatives de réaction face à un problème de planification.

I.2.3.4 Comment simplifier la planification ?

La planification est très complexe. Elle prend en compte l'optimisation d'une multitude de critères dont le premier est certainement l'optimisation des temps de changement de production. Toutefois, les critères qui impactent le plus le système sont surtout la minimisation des chutes de panneaux à l'usinage et des restes de laques lors des changements de série au laquage. Ces deux points nous amènent à envisager au moins deux types de regroupements de pièces différents (d'abord par type de panneaux (matériau) au début, puis par couleur dans la suite du process) avec donc, au moins une étape de tri.

Actuellement, et principalement par souci de ne pas perdre de pièces, l'usinage est effectué par matériau et par client. Un lot représente donc par exemple toutes les pièces du client X en MDF mélaminé 19 mm. Arrivé à l'étape du tri avant laquage, le lot est reconstruit autrement. Le nouveau lot est alors la couleur xxx du client X en ne tenant plus compte du type de matériau. Un client peut avoir, par exemple, 10 couleurs différentes dans sa commande mais cette information est très variable.

Le problème du laquage réside dans le fait que le regroupement par couleur pour plusieurs clients n'est pas actuellement envisageable. Ceci conduit à utiliser plusieurs fois la même couleur dans la journée. D'autre part ces regroupements (couleur et matière), sont limités du fait des délais très courts. On pourrait au maximum regrouper sur la journée ce qui n'est pas toujours d'un très grand intérêt. Il est donc important de bien évaluer les avantages et les inconvénients, les gains et les contraintes de ces types de regroupements dans le cadre de ce projet.

Une particularité de l'entreprise visible sur la figure 3 est que toutes ses gammes de fabrication présentent des « boucles de retour de flux ». Il existe au moins deux types de boucles. Les boucles de production qui se produisent dès lors qu'une pièce doit être laquée à la fois sur la face et sur le dos. L'opération (apprêtage par exemple) doit alors être planifiée deux fois (minimum) avec entre les deux, une autre opération à durée déterminée (séchage par exemple). Ce cas se produit fréquemment dans l'atelier de préparation et dans l'atelier de laquage. Pour certaines pièces (avec les pièces bétons par exemple), il faut envisager une troisième boucle à cause des chants qui doivent aussi être transformés séparément. Ces boucles sont nécessaires dans la gamme. Le seul moyen d'y palier serait d'investir pour dupliquer le poste de charge concerné, ce qui représente un coût non négligeable et est donc inenvisageable.

En revanche, il est possible de réduire l'autre type de boucle lié principalement aux réparations. Le cas se produit si la pièce présente un défaut devant/pouvant être réparé. Ces réparations ne doivent pas être écartées de la réflexion relative au pilotage car elles représentent en moyenne 30% de la production de l'entreprise. Ce taux, dont l'importance peut choquer, est dû au fait que l'exigence qualité est très élevée et que la mise en évidence du poste de travail responsable de la création du défaut est difficile (et par conséquent, le problème est difficile à résoudre !). Par exemple, lors de l'usinage, le rayon sur les arêtes de la pièce a été usiné d'une dimension légèrement inférieure aux standards habituels à cause d'un léger dérèglement machine. Ce défaut peut être si peu visible qu'il restera

indétectable (dans l'état actuel de l'équipement de l'entreprise) à ce stade malgré le contrôle qualité en place et la pièce traversera donc les ateliers de préparation et de laquage sans que le défaut soit vu. Mais lors de l'atelier de polissage, ce défaut peut faire apparaître une « perce » qui se matérialisera par une bande plus claire tout le long de l'arête (là où l'épaisseur de laque va être plus faible à cause de la différence de rayon) et qui sera donc enfin détectée par l'opérateur. Cependant à ce stade, il est impossible pour l'opérateur de dire avec certitude si ce défaut est la conséquence d'un défaut de réglage machine lors de la création du rayon (atelier usinage) ou d'une épaisseur de laque trop fine (atelier laquage) ou d'un temps de séchage non adapté (atelier laquage)... En effet, les « perces » sont souvent des conséquences de multiples facteurs (variables et combinatoires) dépendant aussi de l'aspect de la pièce ou même de la couleur. La constatation du défaut et le traitement informatique via une éventuelle base de données des défauts ne permet donc que difficilement l'amélioration car le poste sur lequel agir n'est pas forcément facile à déterminer très rapidement.

Ces boucles de production et de réparation viennent perturber l'écoulement normal des produits, ce qui conduit à des ordres de fabrication qui se « doublent » malgré la logique simple de FIFO souhaitée initialement par l'entreprise pour le pilotage. Cette notion de vitesses différentes entre les lots est encore amplifiée par l'existence de « circuits courts (CC) ». Les pièces des lots CC peuvent être d'origines diverses. Soit la pièce a été perdue ou trop abimée lors de sa fabrication initiale et elle doit alors être recommencée depuis le début avec un délai beaucoup plus court afin qu'elle puisse tout de même partir avec le reste de la commande. Soit le client a modifié sa commande, ou réparé un oubli lors de sa commande initiale et l'entreprise souhaite lui offrir ce service client permettant la correction dans des délais raisonnables.

Enfin une autre particularité de l'entreprise est que, même si de nombreux ateliers sont articulés autour de machines automatisées (telles que les machines de découpe, le robot de laquage ou encore la ponceuse de chants), les opérations restent majoritairement très manuelles à cause de la variété, mais aussi à cause de la complexité des produits à fabriquer. C'est aussi ce savoir-faire sur les pièces complexes qui fait la force de l'entreprise par rapport aux concurrents plus automatisés. Le Tableau 4 donne quelques exemples des contraintes liées aux produits.

Tableau 4 - Exemples de contraintes liées aux produits

Produits	Contraintes	% de la production
Porte cintrée	Passage en machine interdit car réservé aux produits plans	2%
Façade béton	Application béton manuelle seulement	5%
Façade moulurée	Ponçage de la moulure à la main uniquement	6%
Poignées intégrées aux portes	Beaucoup de reprises manuelles liées au travail de la poignée	4%

Comme une grande partie des opérations reste donc manuelle, la production est fortement soumise aux facteurs humains qui induisent beaucoup de variabilité en termes de qualité et de productivité. Pour pallier à cette problématique, l'entreprise formalise des standards de travail et s'applique à les faire respecter en les faisant vivre et évoluer au rythme des changements de culture ou de production.

Une autre particularité vient encore s'ajouter à la complexité apparente de la planification. Il s'agit du fait qu'il existe de potentielles gammes alternatives en fonction du niveau de qualité produit.

Tableau 5 - Résumé des actions possibles vis-à-vis de la simplification de la planification

Domaines	Moyens d'action
Information	Récupérer les informations nécessaires à l'établissement concret des gammes.
Entité de gestion	Etablir une entité de gestion adaptée à la planification et au suivi terrain.
Ordonnancement	Gérer de manière réactive les perturbations liées aux réparations pour réaliser des plannings au plus près du besoin.

Dans le but de gérer correctement ces flux de produits complexes, il est donc nécessaire d'avoir une identification correcte de chacun des lots. Cependant, la particularité du produit rend l'identification pièce à pièce compliquée et contraint fortement le suivi terrain. Le Tableau 5 résume les différentes actions possibles pour simplifier la planification.

I.2.3.5 Comment fiabiliser le suivi terrain ?

Il n'est pas possible d'envisager des ré-ordonnancements sans outil de suivi de la production et ce suivi ne peut être réalisé que si les pièces ou lots sont correctement identifiés. Malgré le nombre conséquent de technologies d'identification (dont les technologies infotroniques) présentes sur le marché et analysées dans les travaux préliminaires à la thèse, aucune de celles-ci n'est la solution idéale et chacune présente au

moins un critère interdisant son choix (Noyel, Thomas, Thomas, & Beaupretre, 2012) comme montré dans le Tableau 6.

Tableau 6 - Analyse comparative de différentes technologies de traçabilité

	Utilisation consommables	Temps de marquage	Coût de marquage	Temps de marquage	Maintien de la qualité
Fiche suiveuse	-	-	-	--	++
	Les fiches suiveuses sont souvent « volantes » et trop nombreuses si l'on envisage une par pièce. Elles auraient tendance à se perdre ou à ne pas être consultées au profit de l'habitude.				
Etiquette adhésive code à barres	+	--	+	+	+
	Les codes à barres devraient être décollés de la surface à transformer pour être recollés de l'autre côté. Cette manipulation fastidieuse et qui doit pourtant être réalisée un nombre de fois important lors du processus de fabrication présente aussi l'inconvénient de dégrader le pouvoir d'accroche de l'adhésif. Les étiquettes ne collent plus et disparaissent dans les aspirations de poussière du process.				
Marquage laser	++	+	++	+	--
	Ce marquage nécessite un emplacement définitif sur la pièce (l'accord du client est souvent compromis) car il restera visible jusqu'à la fin. Certaines opérations dégradent le marquage (ponçage, ajout de matière telle que la laque), rendant obligatoire plusieurs marquages au cours de la gamme de fabrication avec, à chaque fois, une perte temporaire d'identification.				
Encre invisible	-	--	-	-	++
	L'encre ne survit pas au passage dans les fours de séchage. Il faut donc marquer plusieurs fois avec une perte temporaire d'identification.				
RFID	--	--	--	++	--
	L'insertion d'un tag, quel que soit le packaging (plastic, bois plastique, bois liquide, mdf tourné, etc...) dégrade la qualité du chant en laissant réapparaître la marque de l'insert après les tests de vieillissement.				
Caméra (couplé éventuellement à des lasers de dimensionnement)	++	++	++	--	++
	Solution plus coûteuse du fait du nombre conséquent de caméras à déployer tout au long de la chaîne de fabrication. Identification complexe compte tenu des faibles gradients de couleurs compliquant fortement la reconnaissance de pièces qui se ressemblent. Besoin de gérer plusieurs aspects différents pour une même pièce puisqu'elle évolue au fil de son processus de fabrication.				

Malgré l'attrait des technologies infotroniques qui peuvent permettre de connaître l'état exact du système de production, l'entreprise a dû se résigner jusqu'à ce jour, à conserver une technologie d'identification élémentaire : la fiche suiveuse ou carte d'identité avec suivi manuel. Elle consiste à mettre en place un système de carte d'identité, au format papier, qui suit physiquement la pièce qu'elle doit identifier tout au long du processus de

fabrication par l'intermédiaire des opérateurs. Chaque opérateur intervenant sur le produit est donc le garant de la bonne association entre pièce et carte d'identité à la sortie de son poste de travail. A cause de cela, c'est aussi cette technique de traçabilité qui procure le plus d'erreurs. Le principal inconvénient étant la multiplication du nombre de « feuilles volantes », il a été décidé d'identifier non pas la pièce mais le lot. De plus, la structure des lots change au fur et à mesure de la gamme de production et donc des regroupements, comme présenté précédemment.

La volonté de ne pas mélanger les pièces des clients est une conséquence de l'incapacité d'identifier pièce à pièce. Il est évident qu'une solution plus efficace d'identification serait la bienvenue pour pouvoir automatiser l'étape de tri, d'autant plus que les technologies infotroniques, malgré les quelques contraintes d'implémentation (Srivastava, 2004), (Sarac, 2010) sont un excellent moyen pour synchroniser les informations et les flux physiques en vue de satisfaire les exigences de flexibilité et d'adaptabilité citées plus haut (Thomas A. , 2009).

Les pertes de pièces sont donc une conséquence courante de cette identification sommaire dans un contexte de flux de pièces très complexes. En effet, les gammes traditionnelles présentent de nombreuses boucles comme cela est présenté sur la Figure 6.

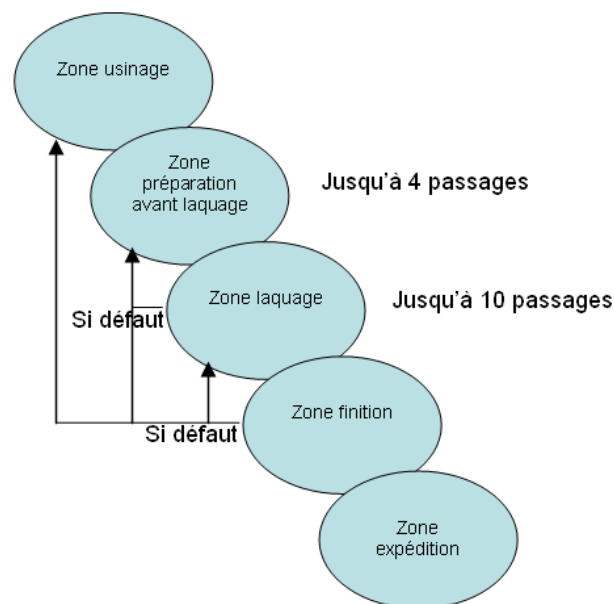


Figure 6 : Multiplicité des « boucles ».

Pour aller vers la solution du problème, il faudrait donc que les flux soient simplifiés, mais aussi formalisés (pour certaines pièces, on ne sait pas réellement quelle gamme a été suivie), fiabilisés (certaines pièces quittent le flux prévu sans que l'on sache pourquoi a

posteriori) et contrôlés (il est donc évidemment nécessaire, afin de prendre les meilleures décisions sur un horizon court terme, de connaître exactement l'état actuel de la production).

Ainsi nous avons vu que la perte de pièces peut provenir d'un problème de flux, mais elle est aussi intimement liée aux modifications conditionnelles des gammes. Pour éviter ce genre de problème, deux pistes sont à envisager. Premièrement une assistance au changement de lot en facilitant et fiabilisant le travail des opérateurs. Et deuxièmement, empêcher les erreurs en installant des détrompeurs à divers endroits stratégiques de la production. Le Tableau 7 résume les actions possibles pour fiabiliser le suivi de production.

Tableau 7 - Résumé des actions possibles vis-à-vis de la fiabilisation du suivi terrain

Domaines	Moyens d'action
Identification	Trouver un moyen d'identifier de manière fiable les pièces/lots.
Simplification des flux liés à la non-qualité	Réduire la complexité des flux en travaillant pour réduire la non-qualité
Entité de gestion	Etablir une entité de gestion adaptée à la planification et au suivi terrain ainsi qu'une assistance à la modification dynamique de ces lots si besoin.

1.2.3.6 Comment limiter le nombre de retards ?

Les retards sont courants dans l'entreprise. Ils ont des causes diverses. On peut citer parmi elles :

- Le délai très court associé à un manque de réactivité
- La charge qui peut parfois dépasser la capacité sans émettre d'alarme,
- L'engorgement de l'atelier à cause de la non-qualité,
- Les pièces perdues et retrouvées trop tardivement.

Tableau 8 - Résumé des actions possibles vis-à-vis de la limitation du nombre de retard

Domaines	Moyens d'action
Suivi terrain	Emettre des alertes suivant la position du lot par rapport à sa date d'expédition
Planification	Anticiper le retard potentiel en évaluant la charge dès la passation de la commande client
Qualité	Réduire les pertes de temps dues aux réparations
Complétude	Vérifier que les lots sont bien complets pour éviter la perte de pièces.

Les pièces non livrées car non terminées au moment de l'expédition sont appelées des « reliquats » et viennent à leur tour perturber la fabrication des pièces des semaines suivantes. Elles participent alors à l'engorgement de l'atelier. Le Tableau 8 résume les actions possibles pour limiter le nombre de retard de livraison.

I.2.3.7 Comment désengorger rapidement l'atelier ?

L'engorgement de l'atelier est le résultat de l'accumulation de pièces à cause des réparations ou des reliquats. Les flux de produits très perturbés tendent à augmenter l'encours et à paralyser petit à petit la production. Ce phénomène d'engorgement peut être assimilé à un effet « coup de fouet » cumulatif. En dessous d'un certain seuil d'engorgement, on constate peu de retard mais au-delà de ce seuil, le nombre de retards augmente très vite et il devient très difficile de le maîtriser. Cette méconnaissance de ce seuil ne permet à l'entreprise que de réagir après constatation du problème et non de l'anticiper. Les pistes de solution sont donc tournées vers deux moyens :

- Connaître le niveau d'engorgement (et l'anticiper) par un système d'indicateur,
- Développer des moyens d'action pour garder la situation sous contrôle.

Tableau 9 - Résumé des actions possibles vis-à-vis de l'engorgement de l'atelier

Domaines	Moyens d'action
Information	Trouver des indicateurs permettant de mettre en relief cet état d'engorgement de l'atelier et en faire une représentation permettant d'anticiper au lieu de réagir
Règles de pilotage	Evaluer l'impact des règles de pilotage sur cet engorgement et associer leur utilisation à des situations définies

Le Tableau 9 résume les actions possibles pour limiter l'engorgement de l'atelier.

I.2.4 Synthèse sur les problèmes industriels

Pour résumer tous les points précédents, nous pouvons représenter les diverses pistes de solutions par la Figure 7.



Figure 7 - Résumé des pistes de solutions groupées par axe de travail.

En définissant le périmètre de nos travaux, nous écartons les pistes de solution concernant les statistiques pour réaliser un prévisionnel car elles ne font pas parties de l'objet de la thèse. En conséquence, la réduction des stocks ainsi que les recherches pour éviter les ruptures ne feront pas, de même, parties de nos travaux. Afin d'apporter une solution aux problèmes industriels restants, nous envisageons de regrouper ces pistes de solution en quatre axes de travail distincts résumés dans le Tableau 10 (lequel est situé dans la partie I.4 présentant la structuration de la thèse).

Le premier axe consiste à agir directement sur la qualité pour diminuer le nombre de réparations et optimiser le traitement de ces non-qualités afin de limiter les perturbations et bouclages des flux. Les travaux sur la non-qualité ont toujours été une priorité pour l'entreprise. Le traitement prioritaire des pièces défectueuses permet d'assurer un taux de service qui reste remarquable au regard du pourcentage de pièces à réparer. Cependant ces traitements au cas par cas sont aussi parfois responsables des pertes de pièces qui empêchent souvent la livraison complète de la commande. Cet axe essentiel pour l'entreprise est un domaine scientifique intéressant qui relève de la maîtrise de la qualité sur les processus de fabrication instable. Nous aurons donc à l'investiguer.

Le second axe consiste à déterminer l'état de saturation de l'atelier pour faire apparaître « le bon moment » pour changer de règle de pilotage et éviter un engorgement. Il s'agit là aussi d'un domaine scientifique intéressant qui relève du pilotage des flux dans un contexte de production fortement perturbé par les reprises. Nous aurons aussi à l'investiguer.

Le troisième axe consiste à définir, identifier et suivre correctement une entité de gestion (que l'on peut appeler lot). Cette idée est associée dans nos travaux au concept de container intelligent. Bien que ce concept soit au cœur des recherches scientifiques actuelles, notre apport sur ce domaine relève plus de l'ingénierie.

Il en est de même pour le quatrième axe qui consiste à organiser les ré-ordonnancements de la production pour réagir aux problèmes de planification pouvant résulter de rupture d'approvisionnement de matières premières, de problèmes qualité ou de l'engorgement en affinant la connaissance des gammes. Le but est d'atteindre une gestion des flux de produits réactive mais toujours très simple afin de continuer de prendre les meilleures décisions très rapidement. Notre apport n'est donc là-aussi qu'applicatif et très peu scientifique.

I.3 ETAT DE L'ART GENERAL

L'état de l'art général présenté dans cette partie est donc focalisé sur les deux parties introduites précédemment et a pour objectif de cerner plus en détails le cadre scientifique de celles-ci et permettre de positionner les problématiques scientifiques spécifiques de la thèse.

I.3.1 Maitrise de la qualité sur des processus instables

I.3.1.1 Généralités

La maîtrise de la qualité est très souvent un des points prioritaires pour les entreprises, essentiellement en raison du fait que le niveau de *qualité produit* influence de manière conséquente la gestion de la production, comme nous l'avons vu précédemment, mais aussi car il s'agit souvent de l'exigence client principale. Très tôt, elle a été au cœur des réflexions et recherches avec par exemple, dès les années 50, la méthode COQ (Cost of quality) et le concept « d'usine fantôme » qui était vu comme un atelier parallèle à l'atelier officiel dont la mission était de réparer les défauts de l'usine officielle (Feigenbaum, 1951). Dans cette étude, l'usine fantôme nécessaire représente 40 % de la capacité de l'usine de production officielle. La première norme sur le sujet apparaît en France dès 1986 (X 50-126) pour évaluer les coûts de la non-qualité (Abouzahir, Gautier, & Gidel, 2003).

Parmi les 5 points de vue différents sur la qualité (Garvin, 1984), la définition la plus pertinente pour notre problème est « l'approche basée sur la fabrication ». Elle combine la conformité aux exigences client et le « faire juste du premier coup » dans le but de réduire les coûts. La notion de conformité est, elle aussi, définie comme « la mesure dans laquelle les caractéristiques de conception et d'exploitation d'un produit correspondent aux normes préétablies. ». La non-qualité peut donc s'exprimer sous 2 formes :

- Produit non-conforme aux exigences client.
- Produit conforme mais qui a nécessité une réparation pour l'être.

Analytiquement, un produit non-qualitatif peut être le résultat de causes diverses :

- Soit il s'agit d'une « erreur » de production, ce qui se produit régulièrement et aléatoirement lorsque les opérations sont très manuelles.
- Soit l'exigence qualité client devient supérieure à l'exigence habituelle ce qui peut survenir lors de la conquête de nouveaux marchés par exemple. Dans ce cas, les processus de production sont toujours gérés efficacement mais des travaux doivent être menés pour atteindre un niveau de qualité encore supérieur pour satisfaire les nouvelles exigences client.
- Soit le processus de fabrication est instable et dépendant d'une multitude de paramètres incontrôlables dont les interactions sont peu ou pas connues.
- Soit il s'agit du résultat de dérives d'un processus de fabrication ordinairement stable.
- Soit l'exigence qualité souhaitée est trop proche des limites technologiques de l'outil de production.

En général, le "zéro défaut" peut être obtenu en combinant deux approches :

- En optimisant les réglages initiaux de divers facteurs.
- Par la surveillance et la prévention des dérives.

Suivant ces deux considérations, une multitude de méthodes sont proposées dans la littérature pour réduire et maintenir sous contrôle le niveau de non-qualité.

1.3.1.2 Les approches les plus marquantes

a. Concepts

Les concepts prédominant du Lean Manufacturing touchent à la qualité (Lyonnet, 2010) avec notamment le « Total Quality Management » et le « Juste à temps » (Vollmann, Berry, & Whybark). Le Total Quality Management (TQM) est défini pour contrôler la qualité de la meilleure façon possible. L'APICS (Association for supply chain and operations management) définit le Total Quality Management comme « une approche de management pour obtenir le succès à long terme à travers la satisfaction client. La TQM est basée sur la participation de tous les membres de l'organisation à l'amélioration des processus, des

biens, des services et de la culture dans laquelle ils travaillent. » Cette définition rejoint celle du “Juste à temps” que la même institution présente comme « une philosophie de fabrication basée sur l’élimination planifiée de tous les gaspillages et sur l’amélioration continue de la productivité. L’un des éléments principaux du “Juste à temps” est d’améliorer la qualité jusqu’au « zéro défaut ». Au sens large, il s’applique à toutes les formes de fabrication ».

La méthode Kaizen, contrairement aux démarches d’innovation de rupture, est aussi très efficace car elle favorise les évolutions graduelles de qualité en réalisant des efforts continus sur les processus de production et en impliquant directement chaque partie prenante. Un point sur les 4 qui définissent l’idéologie Kaizen fait référence uniquement à la qualité avec le but d’éliminer complètement les défauts.

La norme ISO 9001 traite du contrôle des processus et de l’amélioration continue.

b. Outils

La première approche présente des outils simples et très efficaces. Il s’agit des 7 outils de base de la qualité (diagramme d’Ishikawa, feuilles de relevés, cartes de contrôle, histogramme, diagramme de Pareto, diagramme de corrélation, graphiques). Dans cette approche, les pièces sont contrôlées *a posteriori* à la fin de l’opération et les propositions d’amélioration sont faites pour les pièces suivantes en utilisant les connaissances d’experts.

La roue de Deming, plus communément connue sous le nom de PDCA (Plan – Do – Check – Act), décrit les différentes étapes d’un processus de progrès continu où une cinquième étape “Observation” est parfois ajoutée en amont (OPDCA) dans le but de souligner l’importance de cette étape de mesure et d’analyse avant d’agir.

La méthode Statistical Process Control (SPC) suppose que, comme les produits sont fabriqués par des processus de fabrication dont les comportements fluctuent au cours du temps et tendent à devenir instables, voir désorganisés, les entreprises ont besoin d’approches préventives pour garantir et maintenir un niveau de régularité donné pour chaque processus grâce à des systèmes de surveillance adaptés (Oakland, 2007) tels que les plans d’expériences, les diagrammes de Pareto ou encore les cartes de contrôles.

L’inconvénient majeur de certaines de ces approches est qu’elles sont principalement réactives : on constate le défaut, puis on agit pour qu’il ne se reproduise plus. Dans le cadre de processus très instables et évolutifs au cours du temps, des méthodes positionnées plus en amont de la génération du défaut sont requises.

c. Une méthode proactive : Taguchi's Quality Engineering

La méthode Taguchi's Quality Engineering (TQE) (Taguchi, 1989) a été initialement étudiée pour les applications industrielles, elle est la première à proposer une anticipation de la qualité avant le travail direct sur le produit, à travers l'ajustement et le contrôle des paramètres influant sur la qualité. Avec des plans d'expériences optimaux (Optimal Experimental Design – OED), elle permet le réglage optimal de la machine en fonction d'un objectif de qualité donné. Cette méthode aboutit généralement à une amélioration notoire du niveau de qualité (Duffua, Khursheed, & Noman, 2004). Elle a été spécialement conçue pour permettre de grandes économies de temps et d'argent en réduisant considérablement le nombre d'expériences nécessaires (grâce aux plans fractionnaires) pour tester les différents facteurs influents et en organisant l'ordre des expériences pour minimiser le nombre de réglages nécessaires. Cependant, même avec cette méthode à bas coût, il arrive que les décideurs refusent de lancer des expériences, principalement en raison de trois facteurs :

- La nécessité de temps. Il est parfois impossible de réserver le temps nécessaire pour faire les expériences. Par exemple, sur une machine goulot d'étranglement où le taux de charge est optimisé, immobiliser la machine signifie une perte de production avec un impact direct sur le taux de livraison.
- La nécessité de matériel. Des expériences sur une machine consomment des produits semi-finis et/ou des matières premières. Le coût de ces produits doit être pris en compte pour calculer le coût de l'expérience. Dans les cas de produits très coûteux ou de temps opératoires conséquents, il est éventuellement possible de diffuser le plan d'expériences sur le planning de production (PDP) si les modalités des facteurs garantissent un niveau de non-qualité moindre n'affectant pas la livraison client.
- Le réglage des facteurs techniques obtenu est optimal à la fin du plan d'expériences, mais cette solution peut ne pas être assez robuste (même si un plan produit a été mis en œuvre et a pris en compte des facteurs de bruit) à l'évolution du processus de production (modification de la machine ou du processus de fabrication, évolution conséquente des conditions météo comme le passage de l'été à l'hiver, etc...) parce que cette approche est hors ligne par nature.

Ces 3 points se traduisent souvent par l'abandon de véritables modèles expérimentaux au profit des petits changements "tests" dont l'impact réel est souvent apprécié plus que mesuré. En outre, la méthode Taguchi a un autre inconvénient important qui est la difficulté de trouver le réglage optimal de chacun des facteurs dans le cas où ils sont trop nombreux, si beaucoup sont non contrôlables, ou s'il y a trop d'interactions. Par exemple, si la température est un facteur important, il peut être difficile de simultanément tenir compte

de sa gamme complète de variation et de garantir un réglage optimal. D'autre part, un OED devrait être déployé sur chaque poste de travail générateur de défauts, ce qui n'est pas forcément concevable.

Toutes ces méthodes de maîtrise de la qualité, qu'elles soient réactives ou proactives, peuvent efficacement éliminer les causes principales de variabilité sur le système de production (Apley & Shi, 2001), mais elles peuvent aussi se révéler infructueuses si un des paramètres influençant la qualité du produit n'est pas pressenti par les experts et donc non pris en compte dans l'étude. Les résultats peuvent donc être soit peu améliorants, soit non robustes dans le temps. Si on admet que les experts ont une connaissance globale du problème et n'omettent pas de paramètres influents, le dernier inconvénient majeur est la non-réactivité. En effet, ces méthodes ne permettent pas toujours l'adaptation à une évolution du système de production. En résumé, les solutions trouvées, si l'on recherche de la robustesse seront généralement sub-optimales, ce qui est globalement très contraignant pour notre problématique ou nous recherchons à la fois l'optimalité et la robustesse à cause des nombreuses évolutions de processus et d'environnement souvent incontrôlables. Une solution apparente réside dans les méthodes « on-line » qui permettent le suivi dynamique des paramètres et variables en vue d'une réadaptation automatique dès que nécessaire.

I.3.1.3 Synthèse et cadre scientifique relatifs à la maîtrise de la qualité

Les techniques traditionnelles de maîtrise de la qualité « off-line » ne sont pas adaptées à notre environnement changeant et nous mènent à nous tourner vers des techniques « on-line ». L'utilisation de machines d'apprentissages pour maîtriser les paramètres de processus influant sur la qualité est ainsi une piste que nous avons investiguée.

I.3.2 Pilotage des flux dans un contexte de production perturbé par les reprises

Le pilotage des flux est aujourd'hui un sujet sur lequel la recherche investit beaucoup car il s'agit d'un point de plus en plus crucial pour les entreprises dont les processus de fabrication se complexifient au fur et à mesure de l'évolution des exigences client et de la diversification des produits. Les incertitudes provenant du système de production (maintenances curatives, ruptures d'approvisionnement matière, problèmes qualité) ou de paramètres externes (variabilité des demandes client par exemple) rendent le problème difficile à résoudre, que ce soit d'un point de vue prédictif (prévision de planning généraux ou ordonnancement sur une fenêtre de temps à venir), ou réactif (ordonnancement

dynamique déclenché suite à un événement imprévu). Il faut donc fournir aux industriels des modèles et des méthodes qui peuvent assurer à la fois une performance globale et une capacité à s'adapter et à réagir face aux événements perturbateurs dont la probabilité d'occurrence augmente de plus en plus en proposant des estimations claires des états possibles de leur système de production.

Le problème de pilotage des flux peut être abordé avec des solutions locales ayant des performances à court terme qui doivent être menées de manière itérative, ce qui correspond à des méthodes telles que le Kaizen ou le Lean. Les industriels préfèrent aujourd'hui des systèmes de contrôle qui fournissent des solutions satisfaisantes, adaptables et robustes plutôt que des solutions optimales qui sous-entendent d'atteindre des compromis difficiles (Thomas, Trentesaux, & Valckenaers, 2012). Il existe dans la littérature plusieurs approches possibles concernant les besoins industriels cités ci-dessus :

- Centralisée/Prédictive : ordonnancement centralisé/planning réalisé par l'ERP et implémenté par un MES. Cette approche est considérée comme optimale tant que la modélisation du système reste réaliste et déterministe. Dès lors qu'un paramètre dévie significativement (ex : panne), l'application du planning, si celui-ci reste viable, conduit à une gestion médiocre. Par ailleurs, dans le cas où le ré-ordonnancement est incontournable, le temps nécessaire est souvent en incohérence totale avec la réactivité recherchée.
- Proactive/Réactive : L'objectif de cette approche est de pallier au manque de robustesse de l'approche classique en gardant sous contrôle la nervosité induite par trop de ré-ordonnancement successifs suite à des événements perturbateurs (Herrera, 2011). Ces techniques utilisent généralement la redondance (temporelle ou orientée ressource), les méthodes probabilistes, les méthodes contingentes qui dressent plusieurs plannings possibles et basculent de l'un vers l'autre en fonction des perturbations possibles et des fonctions objectif qui intègrent des critères de robustesse.

1.3.2.1 Règles traditionnelles de pilotage des flux

Il existe de nombreuses règles de pilotage des flux produits dans la littérature. Chacune peut se révéler plus ou moins efficace suivant le domaine du problème étudié. Les plus connues et utilisées en entreprises manufacturières sont les suivantes :

- FIFO – Premier entré, premier sorti
- LIFO – Dernier entré, premier sorti
- SPT – Priorité au produit/lot dont le temps de fabrication est le plus court
- LPT – Priorité au produit/lot dont le temps de fabrication est le plus long
- SRPT – Priorité au produit/lot dont le temps de fabrication restant est le plus court

- LRPT – Priorité au produit/lot dont le temps de fabrication restant est le plus long
- EDD – Priorité au produit dont la date d'échéance est la plus courte
- CR – Ratio critique – Priorité au produit/lot dont le ratio $\frac{\text{temps de fabrication restant}}{\text{temps restant avant la date d'échéance}}$ est le plus haut.

Suivant la règle choisie, l'impact sur le flux de produits en cours de fabrication va être différent. Par exemple, tandis que la règle SPT favorisera la réalisation des petites tâches peu consommatrices de temps, la règle EDD favorisera la réalisation des lots dont la date de fin est la plus proche, indifféremment du temps gamme nécessaire pour le finir.

Lorsque le taux de reprises est important, les règles de pilotage les plus adaptées en fonction des situations ne sont pas forcément les mêmes. FIFO, SPT, EDD et CR sont des règles très couramment utilisées dans les ateliers de fabrication et sont souvent considérées comme des benchmarks des recherches par simulation (Kuhl & Laubisch, 2004) lorsqu'on fait varier le taux de reprises. De nombreux travaux dans le domaine des semi-conducteurs ont montré que dans le cas de reprises et de très gros volumes de production une règle spécifique comme FSVCT (Fluctuation Smoothing of Variance of Cycle Time) peut donner de meilleurs résultats (Mittler & Schoening, 1999). Pour autant, ce facteur « volume de production » et le % de reprises concerné sont très différents de ceux de notre contexte.

Ainsi il semble que généralement, dans les ateliers où il y a de nombreuses reprises, les règles SPT, EDD et FIFO sont les plus souvent utilisées, même si SPT semble être la plus efficace.

1.3.2.2 Système de décision adapté

Dans un contexte fortement perturbé par les reprises, les systèmes de décision centralisés sont donc peu indiqués. En effet, bien qu'ils présentent des qualités de robustesse et de stabilité, ils affichent surtout des faiblesses en matière de réactivité et d'adaptabilité (Thomas A. , De la planification au pilotage pour les chaines logistiques, 2004), ce qui est fortement requis pour palier toutes les perturbations engendrées par les reprises qui nécessitent une réévaluation complète de l'état du système. Les systèmes de décision distribués peuvent, en travaillant sur des optima locaux (par atelier par exemple ou même par poste), détecter très vite une perturbation et y trouver rapidement une solution car déterminer un optimum local nécessite moins de réflexion que pour déterminer un optimum global. Par contre, le vieil adage « la somme des optimums locaux n'est pas égale à

l'optimum global » qu'on peut retrouver dans la théorie des contraintes (Goldratt & Cox, 1992) induit de tirer profit des deux systèmes. Comme présenté sur la Figure 8, lorsqu'on souhaite faire du contrôle réactif, on peut se tourner vers des systèmes « hybrides », centralisés/distribués, dans lesquels une partie de la décision est réalisée de façon locale et ce dans un cadre défini de façon globale. L'information utile est alors restreinte et traitée de manière plus locale (Herrera, 2011).

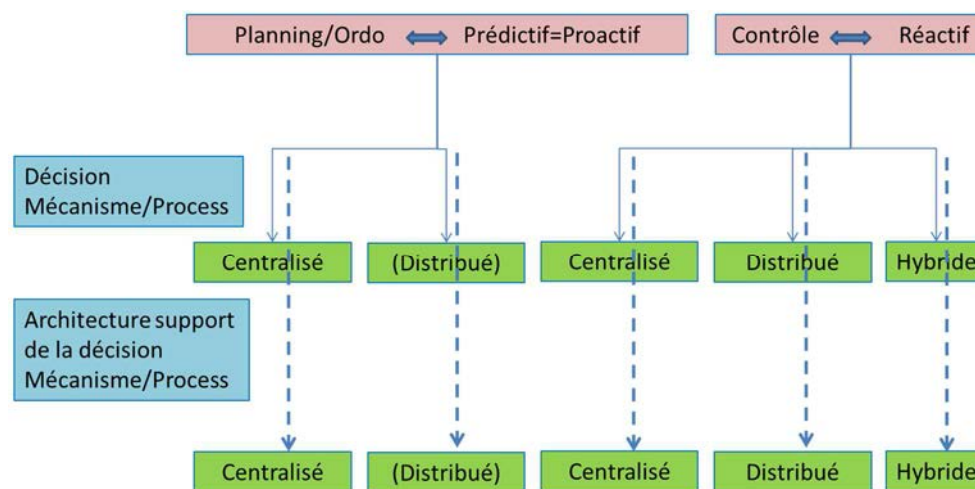


Figure 8 – Mécanismes et architectures de décision

Or, la distribution de la décision, même de manière partielle, nécessite dans un premier temps de parfaire la synchronisation des flux physiques et informationnels. En effet, la prise de décision réactive ne peut être envisagée que si l'état exact du système est connu.

La distribution de la décision peut s'effectuer suivant 3 axes :

- La distribution spatiale. Le degré de distribution spatial s'étend d'un extrême (système centralisé, toutes les décisions sont prises au même endroit) à l'autre (système hyper-distribué ou chaque poste a sa propre cellule de décision). Là encore la notion d'hybride peut être intéressante en allouant la capacité de décision à différents niveaux. A l'échelle d'un l'atelier parfois, ce qui lui conférerait la capacité de s'auto-optimiser, ou à l'échelle du poste de travail pour les goulots structurels ou conjoncturels (Thomas & Charpentier, 2005), à qui il serait judicieux d'attribuer leur propre cellule de décision.
- La distribution temporelle. Elle correspond schématiquement à l'intervalle de temps entre les prises de décision de la cellule. Elle dépend des quantités d'informations qui devront circuler (temps d'exécution) mais aussi de la pertinence car il n'est pas forcément utile d'utiliser des fréquences trop courtes. La réévaluation systématique n'est pas forcément indiquée, surtout dans le cadre de systèmes contrôlés par le produit (que nous définirons plus loin).
- La distribution physique. D'après les études faites sur la réflexion humaine, lors de la prise de décision on retrouve des notions de collaboration et de coopération entre les « actants », à savoir toutes entités prenant part, de près ou

de loin, à la décision (Grosjean & Robichaud, 2010). Dans le cas de décision d'ordonnancement, les actants peuvent être définis comme l'ensemble des entités comprenant les pièces en cours de fabrication, les opérateurs, l'état des machines, les prévisions...

I.3.2.3 Système contrôlé par le produit, la place des technologies infotroniques

C'est dans cette démarche de distribution de la décision en quête de réactivité que s'inscrivent les travaux sur les Systèmes Contrôlés par le Produit². Même si H. Van Brussel et al. abordent déjà l'idée en 1998, le concept se formalise en 2003 (Morel, Panetto, Zaremba, & Mayer, 2003) sur les fondements des systèmes holoniques (Van Brussel, Wyns, Valckenaers, Bongaerts, & Peeters, 1998) avec le principe de combiner avec plus de flexibilité les modes de pilotage centralisés avec les modes de pilotage distribués en tenant compte des capacités du produit à jouer un rôle actif de synchronisation des échanges entre différents systèmes d'entreprise centralisés (ERP) et distribués dans le système physique (automates, CNC (computer numerical control), ...). En effet dans un contexte où les flux de production sont très anarchiques et changeant en fonction du volume de non-qualité, il devient nécessaire d'avoir l'information en temps réel lié à chaque « granulat de gestion » (élément de gestion : produit, groupe de produits, ressource, ...). Les voies classiques de contrôle centralisé doivent donc être abandonnées au profit du côté réactif offert par ce concept de contrôle par le produit. Dans ce domaine, les travaux d'El Haouzi (El Haouzi, 2008) et de Klein (Klein, 2008) qui présentent une méthodologie de conception et d'intégration de SCP ont pu démontrer la pertinence de ces systèmes en production industrielle. La connaissance de l'état exact du système peut être obtenue « sur le terrain » grâce à des outils tels que les technologies infotroniques. En plus de permettre une traçabilité simple et sécurisée des produits sur la chaîne logistique (Srivastava, 2004), la communauté de l'Intelligent Manufacturing System affirme que coupler les produits (agents) avec ce type de technologie (telle que la RFID) peut permettre de satisfaire les nouvelles exigences de flexibilité et d'adaptabilité dont les entreprises ont besoin face aux contraintes actuelles (Thierry, Thomas, & Bel, 2008). Les techniques d'identification sont un prérequis pour la distribution de l'intelligence comme pour la simple traçabilité dans le sens où elles permettent de rendre le produit « intelligent ». En effet, dans notre

² Le contrôle par le produit est à envisager au sens large et s'applique dès lors que le produit devient garant d'au moins une information lui permettant d'interagir avec d'autres éléments de son environnement.

interprétation du concept SCP, les produits deviennent des « holons » composés d'une partie physique et d'une partie informationnelle qui leur donnent la capacité de participer à des décisions locales concernant leur propre fabrication. Chaque produit devient alors « agent » du système et constitue le lien matériel entre les flux physiques et d'informations. Les technologies d'auto-ID (Mac Farlane, Sarma, Chirn, Wonga, & Ashton, 2003) illustrent ce principe en attribuant à chaque produit un identifiant unique. Même en déportant physiquement les capacités de mémorisation, de communication et de prise de décision sur des systèmes externes aux produits, c'est tout de même le produit qui est le garant de l'information. Il est alors possible d'assurer le suivi et le contrôle de chacun et, donc, de permettre de nouvelles méthodes de décision centrées sur les produits eux-mêmes. (Meyer, Främling, & Holmström, 2009) dressent un état des lieux des divers concepts concernant les produits intelligents et proposent notamment une classification de l'intelligence en trois dimensions représentée en Figure 9 : le niveau, la localisation et le degré d'agrégation.

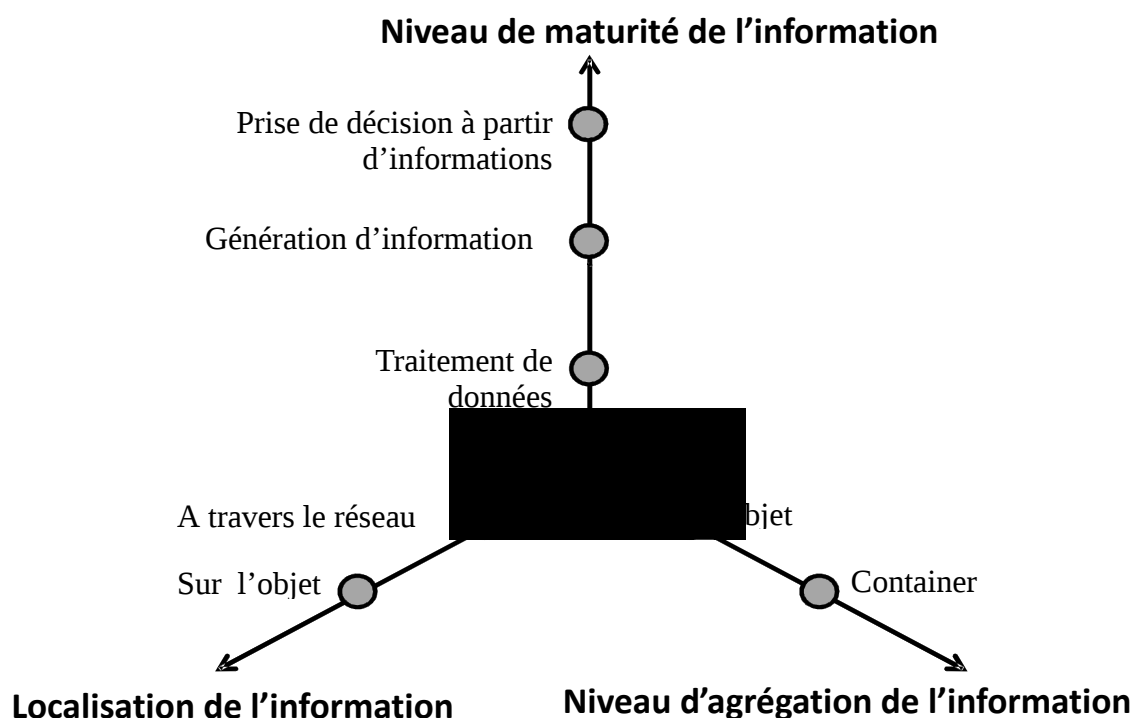


Figure 9 – Classification de l'intelligence en trois dimensions

Grâce à cette classification, la notion de système de produits intelligents s'étend donc à des produits contenant simplement leurs identifiants leur permettant d'être reconnus par l'intermédiaire d'une base de données, aux produits embarquant directement un module

décisionnel et étant capables de se gérer, seuls ou assemblés, sans passer par l'intermédiaire d'une base de données. Evidemment, la position souhaitée du système dans cet espace tridimensionnel apporte des contraintes au niveau du support de l'information liée à la pièce. En effet, alors qu'un code à barres pourrait suffire dans le premier cas, un micro-processeur peut être requis dans le second cas. Ces contraintes ont aussi une répercussion financière.

I.3.2.4 Synthèse et cadre scientifique relatifs au pilotage des flux par le produit dans un environnement changeant

Dans un contexte de production fortement et aléatoirement perturbé par les reprises (comme c'est le cas dans l'entreprise Acta Mobilier), l'emploi de la règle SPT semble être la plus utilisée parmi les règles traditionnelles de gestion des flux telle que le FIFO ou l'EDD. Mais cette règle adaptée à une utilisation locale ne permet pas une optimisation et un contrôle des flux à l'échelle de l'entreprise. Les systèmes de décision centralisée sont généralement les mieux placés pour réaliser cette optimisation globale. Cependant, ils sont généralement trop lents et trop déconnectés du « terrain » et sont donc par nature des systèmes à forte inertie ne permettant pas les adaptations réactives aux divers problèmes qualité rencontrés. A l'opposé, les systèmes de décision distribuée sont très réactifs et donc quelques fois plutôt instables. D'autre part, ils ne peuvent pas toujours garantir l'optimalité comme nous l'avons introduit précédemment. Toutefois, dans les conditions de production décrites dans l'introduction, le besoin de réactivité est primordial et ce type de système de décision devient alors intéressant. Dans ce contexte, les systèmes contrôlés par le produit peuvent allier la réactivité en déportant une grande partie de la décision au niveau terrain au moment opportun et conserver de la robustesse en respectant des règles gérées à un niveau supérieur. De multiples façons « d'hybrider » ses deux types de décisions peuvent être envisagées. Cette thèse propose un système de décision hybride permettant l'adaptation prédictive aux perturbations qualité en contrôlant les règles de pilotage à utiliser. Il s'agit donc d'une solution de pilotage distribuée avec des règles d'ordonnancement adaptées à la situation (paramètres et variables) à partir d'informations dynamiques portées par le produit.

I.4 ARTICULATION LOGIQUE DE LA THESE

Conformément à la synthèse des problèmes industriels présentée dans la partie I.2.3 et rappelée dans le tableau 10, nous proposons d'articuler nos travaux de la manière suivante. Les parties II et III correspondent aux deux apports scientifiques de la thèse. En partie II, la prévision de la non-qualité dans un contexte de flux fortement perturbés par les reprises qui aboutira à la mise en place d'un système d'aide à la décision de production, évolutif dans le temps, permettant de garantir un niveau de qualité acceptable. La partie III traitera de l'état de perturbation résultant de la non-qualité et des boucles de production et proposera une manière de faire évoluer les règles de pilotage en réaction à cet état de perturbation en s'appuyant sur la mise en place de cartographies dynamiques permettant d'éviter les états de saturation de l'atelier.

Tableau 10 - Résumé des axes de travail

Problèmes industriels	Domaines scientifiques	Applications
Augmentation de qualité pour rester concurrentiel	Maîtrise de la qualité sur des processus instable (Partie II)	Machines d'apprentissage (Partie IV.1)
Engorgement de l'atelier ayant pour conséquences de multiples retards	Pilotage des flux dans un contexte fortement perturbé par les reprises (Partie III)	Cartographie dynamique associée aux règles de pilotage les plus adaptées (Partie IV.3)
Suivi terrain de pièces à identification unitaire difficile		Container intelligent et MES dédié (Partie IV.4)
Planning et réordonnancement dynamique rapide		Ordonnanceur centralisé assisté par le produit (Partie IV.5)

Nous avons fait le choix de positionner l'apport scientifique de la thèse uniquement sur ces deux parties et même si, comme nous l'avons dit précédemment, le contrôle par le produit et la mise en œuvre du réordonnancement dynamique restent importants, notre apport scientifique restant limité sur ces parties, nous les aborderons uniquement dans la partie IV dédiée à l'application. Ainsi, cette partie IV sera plus volumineuse et décrira l'implémentation dans l'entreprise en 4 chapitres. Les chapitres 1 et 2 reprendront respectivement les propositions des deux parties « scientifiques ». Le chapitre 3 concernera l'apport de « containers intelligents » dans le cadre d'un système contrôlés par le produit sur un poste de travail stratégique (goulot) fortement contraint. Et le chapitre 4 se concentrera sur la mise en œuvre du ré-ordonnancement dynamique rapide tout en maintenant le niveau de qualité requis sans devoir faire appel à des approches d'ordonnancement analytiques optimales trop chronophages pour une utilisation en mode réactif.

Ces 4 chapitres sont présentés comme 4 outils qui ont été implémentés dans l'entreprise. Ils peuvent éventuellement être exploités indépendamment les uns des autres suivant les problématiques rencontrées dans les entreprises mais sont interconnectés dans la solution finale adaptée à l'entreprise Acta-Mobilier. Chacun des outils est donc apporté comme une brique d'un système global de supervision et d'aide à la décision présenté sur la Figure 10 qui sera donc lui-même un système de systèmes capable d'échanger des informations et de trouver un consensus à partir de différents points de vue.

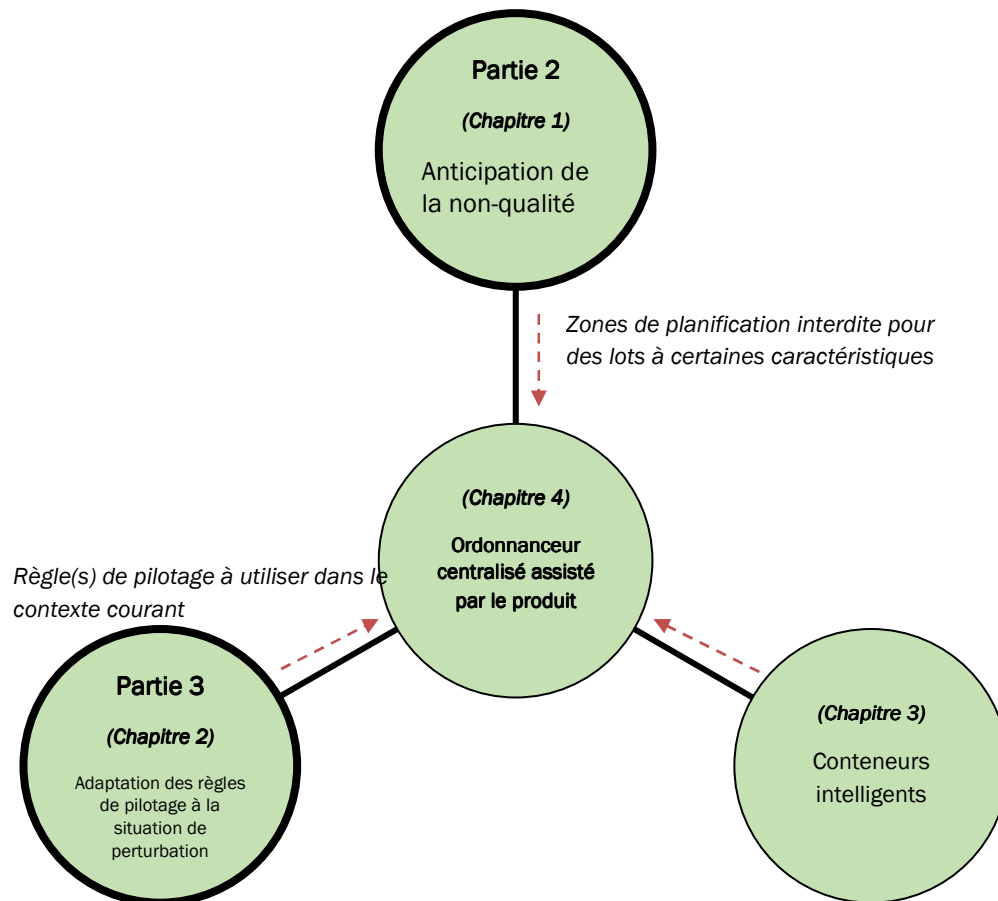
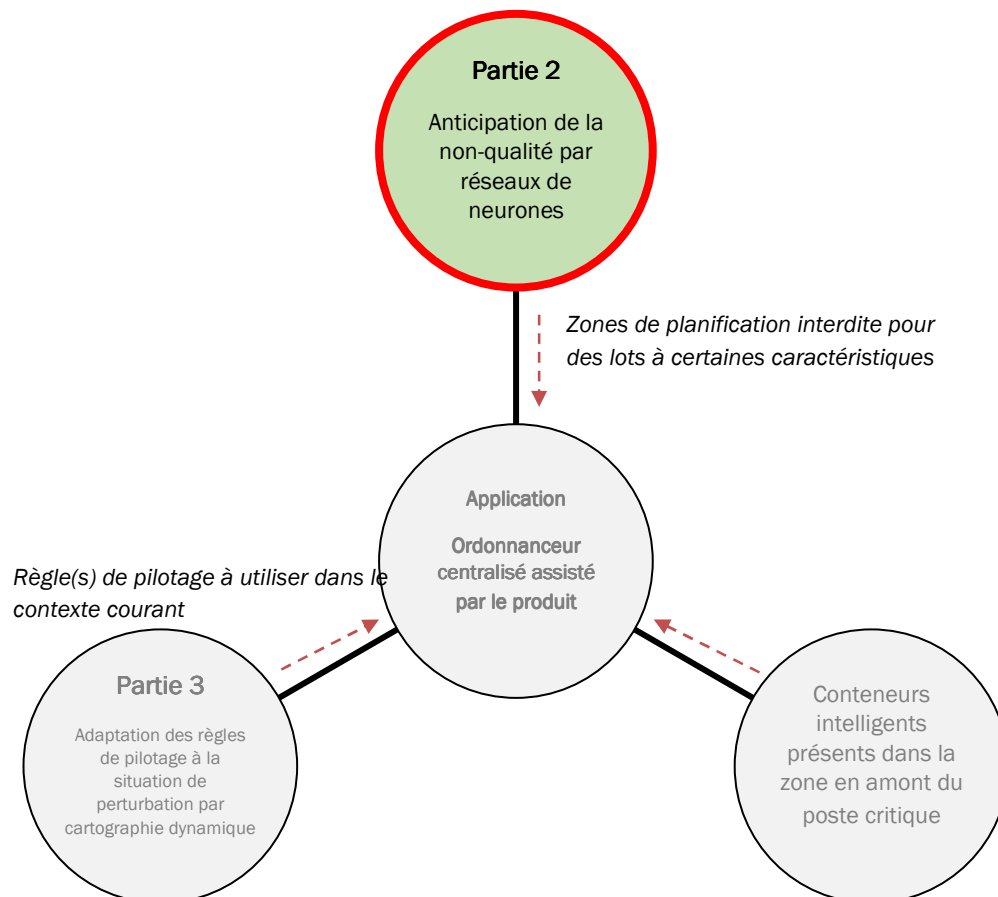


Figure 10 – Interconnexion des outils et corrélation avec l'organisation structurelle du document.

II MAITRISE DE LA QUALITE ON-LINE SUR DES PROCESSUS INSTABLES



II.1 INTRODUCTION

Cette partie traite de la maîtrise de la qualité de produits obtenus par des processus instables. D'après le problème industriel et l'état de l'art général sur ce sujet présenté en partie I, les outils classiques de maîtrise de la qualité peuvent s'avérer insuffisants dans le contexte de systèmes fonctionnant aux limites technologiques. L'objectif est donc de proposer un dispositif de prévision de la qualité fournissant à l'opérateur des prévisions fiables et adaptées à l'état actuel du système.

Une solution réside dans un système de prévision (Figure 11) construit en parallèle du système physique anticipant les comportements de celui-ci de façon aussi fidèle que possible. Ce système est capable de mesurer différentes entrées du système physique et de prédire la sortie de ce même système physique avant de comparer ses prévisions avec la réalité. Le système de prévision considéré doit donc être capable de qualifier les données de sortie et il peut donc être utilisé dans le but d'évaluer une décision prise en amont, c'est donc un problème de classification³. Le modèle de comportement ainsi proposé est donc obtenu à partir des données réelles par une méthode d'apprentissage.

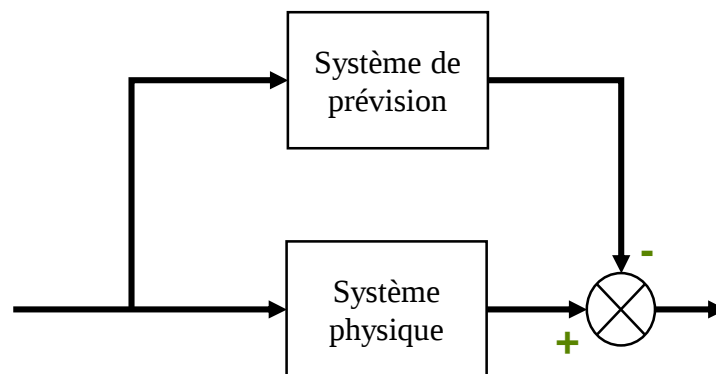


Figure 11 - Relation entre le système de prévision et le système physique.

Par ailleurs, dans le domaine de « l'industrie intelligente », les systèmes industriels sont perçus comme changeant dans un environnement qui évolue lui-aussi. Les qualificatifs qu'on leur donne sont systèmes évolutifs, dynamiques ou encore non-stationnaires. Ainsi,

³ On distingue trois types de problèmes solubles avec une méthode d'apprentissage automatique supervisée et lorsque l'ensemble des valeurs de sortie est fini, on parle d'un problème de classification

maintenir un niveau de qualité donné ne se résume plus à trouver et imposer le réglage optimal des facteurs influents (ou des standards de travail pour les opérations manuelles) pour une robustesse donnée sur un temps donné, mais plutôt à signaler les évolutions de réglage (ou de comportement) nécessaires pour garantir ce niveau de qualité en fonction des facteurs auxquels le système de production doit se soumettre. Réagir « en ligne » est donc la meilleure solution afin de permettre de suivre la dynamique des événements. Les systèmes « en ligne » sont souvent reconnus comme essentiels dans toutes les applications industrielles où les décisions doivent être prises en temps réel (Karp, 1992). Grâce à son niveau de connectivité, le système « en ligne » est capable d'obtenir, d'interpréter et de répondre aux informations reçues. Dans l'idéal, il peut même être capable de se synchroniser lui-même avec le comportement du système réel. Beaucoup de travaux soulignent l'intérêt d'approches « en ligne » dans divers secteurs de recherche tels que, par exemple le plasma (Refke, Barbezat, & Loch, 2001), la transformation d'énergie (Stirl & Skrzypek, 2003) ou encore les CD (compact disc) (Clifford & Duncan, 2002)... Dans le domaine de la qualité, on peut aussi citer les études sur l'impact des facteurs environnementaux (température, humidité...) dans les processus de moulage sous pression (Thomas, Suhner, Meuteler, & Brachotte, 2004), ou sur l'impact de l'usure de l'outil sur l'état de surface lors de l'usinage (Sick, 2002).

La nouvelle approche qui met l'accent sur l'apprentissage avec une évolution de la structure des modèles informatiques en ligne se traduit par le terme « Evolving Intelligent Systems » (EIS) (Angelov & Kasabov, 2005) à ne pas confondre avec les algorithmes évolutionnaires qui comportent la notion d'évolution au fil des générations successives. Les EIS sont traditionnellement associés à des flux de données récupérés en ligne et le plus souvent en temps réel sur des modes de fonctionnement. Ils peuvent être perçus comme des systèmes intelligents adaptatifs de complexité de calcul relativement faible. En raison de la multitude de méthodes adaptatives, évolutives et dynamiques, les EIS représentent une part importante dans le domaine de l'apprentissage basé sur les données dans des environnements non stationnaires (Sayed-Mouchawed & Lughofer, 2012). On peut notamment citer l'article de (Koksal, Batmaz, & Testik, 2011) qui propose un état de l'art sur l'utilisation des machines d'apprentissage pour des problématiques liées à la qualité.

L'adaptation automatique en ligne est donc couramment envisagée pour les techniques d'apprentissage mais elle n'est pas toujours nécessaire car elle présente aussi quelques inconvénients. Sarle (Sarle, 2002) a montré que l'apprentissage en ligne est généralement

plus difficile à réaliser et moins fiable que l'apprentissage hors-ligne. En effet, les méthodes « en ligne » sont souvent mises de côté, justement en raison du manque de ressources ou d'expertise lors de la construction du modèle, l'objectif étant d'obtenir un comportement proche du comportement réel dans toutes les situations envisageables.

L'adaptabilité peut aussi être obtenue de manière différente avec des réapprentissages successifs hors-ligne. La méthode hors-ligne permet en particulier d'exploiter de nombreux algorithmes permettant d'éviter les minimums locaux et les problèmes de surapprentissage. Une validation croisée sur des données de validation peut aussi être effectuée, ce qui n'est pas possible avec les apprentissages en ligne.

Ce type de système est à la fois en ligne lors de la consultation des données d'entrée et de la comparaison des données de sortie, mais conserve cependant une caractéristique de réapprentissage hors-ligne. Le modèle résultant est donc statique et nécessite des étapes de réapprentissage pour s'adapter à des évolutions du système physique qu'il surveille. Il existe deux types d'évolutions détectables au niveau d'un système (Raquel & João, 2009) :

- La notion de « shift » (virage) est associée aux changements brusques.
- La notion de « drift » (dérive) est associée aux changements graduels.

La dérive est plus difficile à détecter et est souvent confondue avec le bruit. C'est pourtant ce type d'évolution qu'il faut mettre en évidence si l'on souhaite conserver le comportement du modèle le plus proche possible du comportement réel. Afin de pouvoir traiter les dérives, 3 étapes sont généralement proposées :

- Surveillance
- Synchronisation
- Diagnostic

La surveillance se fait généralement en comparant les prévisions avec les résultats réels. Il est possible de suivre des informations concernant, par exemple, le nombre de mauvaises classifications, ou la fréquence de mauvaises classifications.

La synchronisation peut se faire de diverses manières. Comme cette synchronisation est souvent consommatrice de temps (la révision du modèle pouvant prendre de quelques minutes à plusieurs jours), la fréquence de synchronisation doit être correctement étudiée. 3 possibilités sont traditionnellement considérées :

- synchronisation périodique (toutes les heures ou une fois par semaine par exemple),
- synchronisation déclenchée sur événement (arrivée d'une information fournie par un outil connecté, sollicitation par un opérateur),

- synchronisation basée sur la détection statistique de dérives.

Les « blind methods » permettent des adaptations sans aucun mécanisme de détection de dérives et sont donc particulièrement compatibles avec les synchronisations périodiques, tandis que les « méthodes averties » peuvent tenir compte des données récentes présentant la dérive et correspondent donc mieux aux synchronisations basées sur la détection statistique de la dérive.

Quelle que soit l'option choisie, il faut toujours différencier le bruit dans les données et les changements réels afin d'atteindre un compromis entre la robustesse du modèle vis-à-vis des facteurs de bruit et la flexibilité dans le suivi des dérives (Lughofer & Angelov, 2011).

II.2 ETAT DE L'ART

II.2.1 Les méthodes d'apprentissage (problèmes de classification)

Lorsque l'on parle de méthodes d'apprentissage, on parle d'induction (apprentissage par expérience). Il s'agit de techniques d'apprentissage automatiques où l'on cherche à produire automatiquement des règles à partir d'une base de données d'apprentissage contenant des données (entrées/sorties) traitées et validées. Il existe beaucoup d'exemples d'utilisation connus tels que la reconnaissance de formes, d'écriture, vocale... Dans chacun des cas, l'objectif est d'apprendre au système à reconnaître un pattern X (un carré par exemple) en lui montrant lors de l'apprentissage suffisamment de représentations différentes de ce même pattern (carrés de différentes couleurs, différentes tailles, différentes position...) pour qu'il puisse estimer si un nouveau pattern inconnu qui lui serait présenté peut être classé comme pattern X ou non.

Il existe deux grandes approches de classification par apprentissage, la classification supervisée et la classification non supervisée. Pour la classification non supervisée, également dénommée clustering, l'objectif est de regrouper des données d'entrée en un nombre réduit d'ensembles appelés clusters. Un état de l'art sur ces problèmes est présenté dans (Jain, Murty, & Flynn, 1999). L'inconvénient est que l'on sait que les données incluses dans un cluster sont similaires, mais on ne sait pas à quel label/étiquette elles correspondent. L'apprentissage supervisé exploite lui une collection d'exemples pré-classés, c'est-à-dire dont on dispose déjà des labels les rattachant à une classe prédéterminée. Cet ensemble d'exemples est utilisé pour construire un modèle liant les données aux labels, ce modèle

étant ensuite utilisé pour classer de nouveaux exemples dont le label est cette fois-ci inconnu (Jain et al. 1999). Une taxonomie des problèmes de classification est présentée sur la Figure 12 qui inclut la taxonomie des problèmes de clustering proposée par Jain et al. (1999).

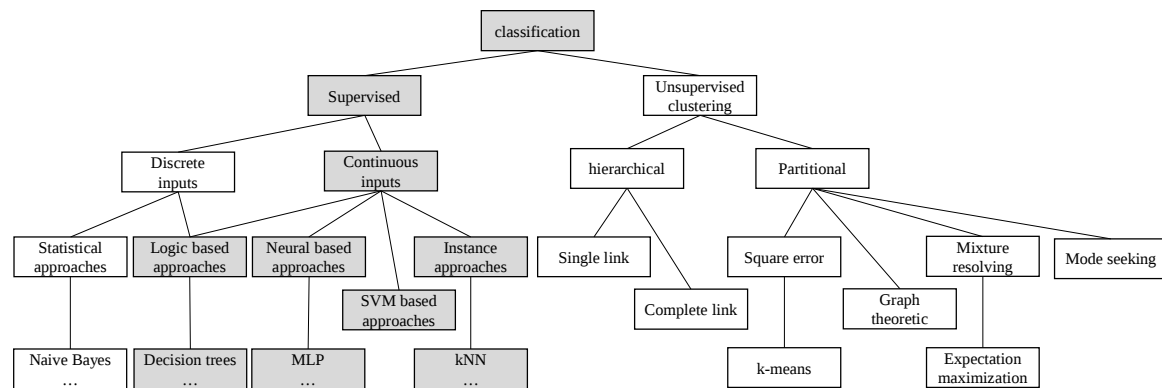


Figure 12 - Taxonomie des approches de classification (Jain et al. 1999)

D'après la nature de notre problème, les méthodes non-supervisées ne sont pas indiquées (on ne s'intéresse donc pas à la branche de droite de la Figure 12). Parmi les méthodes supervisées, seules celles relatives à des entrées continues nous concernent car nos facteurs d'entrée sont continus (température, grammage, etc...). Il reste donc 4 approches potentielles à exploiter : les approches basées sur la logique, les réseaux de neurones, les SVM et les approches basées sur les instances. Par conséquent notre état de l'art s'attachera à étudier et analyser la pertinence des 4 types d'outils suivants : arbre de décision, perceptron multicouches, support vector machine et k-plus proche voisin.

II.2.1.1 k-plus proches voisins (kpp)

kpp est un algorithme d'apprentissage « non paramétré paresseux » qui utilise des bases de données dans lesquelles tous les points sont séparés dans plusieurs classes pour prédire la classification de nouveaux échantillons (Cover & Hart, 1967). Cette classification est réalisée en déterminant la classe d'appartenance d'un échantillon inconnu par le vote à la majorité de ses k plus proches voisins dans la population d'apprentissage. La précision des kpp dépend de la métrique utilisée pour calculer la distance entre les échantillons. Différentes métriques de distances ont été proposées dans le but de déterminer ces voisins. La plus classique est la distance Euclidienne. Malheureusement, la distance Euclidienne ne prend pas en compte les irrégularités statistiques (dans notre cas industriel nos lots de

données en présentent de nombreuses) qui devraient être estimées à partir de la population d'échantillons déjà classés ayant servi pour l'apprentissage (Weinberg & Saul, 2009). La distance Mahalanobis prend en compte les données considérées comme présentant beaucoup de dispersion en leur accordant moins d'importance que pour les données sans dispersion. Cette méthode permet de réduire les conséquences des données bruitées. Par exemple, dans le domaine des calibrations multi-variables, la distance Mahalanobis est utilisée pour détecter les valeurs aberrantes (Roussew & Leroy, 1987). Donc cette distance est plus précise dans le cas de données réelles qui sont généralement bruitées et polluées par les valeurs aberrantes.

Le choix de k est critique. Plus k est petit, plus le bruit va prendre de l'importance. A l'inverse, les grandes valeurs de k conduisent à un temps de calcul de plus en plus long et à un modèle moins flexible (James, Witten, Hastie, & Tibshirani, 2013). Le choix de k est généralement réalisé par validation croisée. Le Tableau 11 résume les principaux avantages et limitation des kpp.

Tableau 11 - Avantages et limites de la méthode Kpp

Avantages	Fonctionne bien avec des données bruitées ou polluées par des valeurs aberrantes
Limites	La précision dépend de la distance utilisée et le temps de calcul du choix de k Le modèle est figé et ne s'adaptera pas aux évolutions du système qu'il représente

II.2.1.2 Arbres de décision (DT)

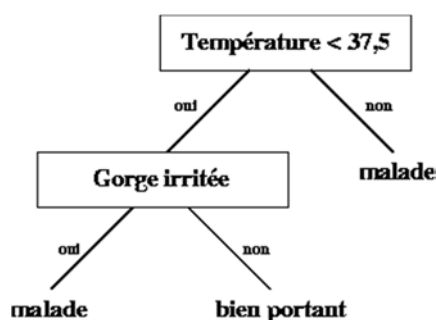


Figure 13 - Exemple d'arbre de décision

Les arbres de décision (Decision Trees DT) sont des modèles séquentiels qui combinent de manière logique une séquence de tests très simples. Chaque nœud représente une caractéristique d'une instance en cours de classification et chaque branche qui découle de ce nœud représente une valeur possible pour cette caractéristique. Les instances sont donc classées en démarrant par le nœud racine et sont

triées le long des différentes branches sur la base de leurs caractéristiques (Kotsiantis S. , 2007).

Cette méthode est très populaire en statistique et en automatique grâce à la lisibilité évidente du modèle, à sa capacité à sélectionner automatiquement les variables

discriminantes dans un fichier de données contenant un très grand nombre de variables potentiellement intéressantes et surtout à sa rapidité dans le classement des données.

De multiples variantes ont été développées afin d'améliorer encore les capacités des DT telles que les forêts d'arbres décisionnels par exemple qui effectuent un apprentissage sur de multiples DT dont l'apprentissage a été réalisé sur des sous-ensembles de données légèrement différents. Cet algorithme combine alors de multiples avantages dont une classification rapide de haute précision, la gestion d'un très grand nombre de variables d'entrées, la possibilité de traiter les cas même avec une grande partie de données manquantes toujours avec précision grâce à des estimations ou encore la possibilité de déterminer la proximité des cas utiles pour le clustering. L'avantage principal des DT résulte de leur structure récursive de tests logiques desquels il est possible d'extraire de la connaissance directement interprétable.

De nombreux algorithmes pour les DT ont été développés. On peut citer par exemple le C4.5 (Quinlan J. , 1996), l'arbre de classification et régression CART (Breiman, Friedman, Stone, & Olshen, 1984), SPRINT (Shafer, Agrawal, & Mehta, 1996), ASHT (Bifet et al. 2009), ou encore SLIQ (Mehta, Agrawal, & Rissanen, 1996). Outre CART, les algorithmes les plus populaires sont ID3 et ses extensions plus récentes C4.5 et C5.0 (Quinlan J. , 1993). Si C5.0 est très similaire à CART, ses précédentes versions (ID3 et C4.5) sont plus limitées et, en particulier, n'incluent pas de processus de « pruning » (Tibshirani & Friedman, 2009) qui simplifie le modèle en supprimant les branches non significatives.

Cependant, il a été démontré que les arbres de décision se comportent généralement mal dans le cas de données bruitées (Patel & Panchal, 2012). Dans ces conditions particulières pourtant courantes en industrie, ils sont souvent sujets à des problèmes d'overfitting (R. Segal, 2004). De plus, le processus de construction est déterministe et nécessite une procédure de "bagging" (Breiman L. , 1996) qui va permettre de construire divers jeux de données d'apprentissage différents à partir d'un seul, assurant ainsi la diversité entre les modèles. Enfin, ces types de modèles ne s'adaptent pas naturellement aux évolutions du système qu'ils représentent ou de l'environnement. Dans ces cas, une nouvelle modélisation complète doit être relancée. Le Tableau 12 résume les principaux avantages et limitation des DT.

Tableau 12 - Avantages et limites des arbres de décision

Avantages	Lisibilité logique du modèle grâce à sa structure récursive Rapide
------------------	---

Limites	Gère un grand nombre de variables
	Les forêts permettent de classer des données incomplètes
	Mauvais comportement sur des données bruitées
	Sujets à de l'« overfitting »
	Nécessiter de réaliser du « bagging »
	Pas d'adaptation aux évolutions du système qu'ils modélisent sans modélisation complète.

II.2.1.3 Machines à vecteurs de support (SVM)

Les « machines à vecteurs de support » (ou séparateurs à vastes marges – Support Vector Machine en anglais) sont eux aussi des classificateurs linéaires développés à partir de la théorie statistique de l'apprentissage « théorie de Vapnik-Chervonenks » (Vapnik, 1999). La principale différence qui distingue les SVM des autres machines d'apprentissage est la non existence de minimum locaux (Cristianini & Shawe-Taylor, 2000). Les SVM peuvent être vus comme une extension des machines d'apprentissage linéaires qui effectuent une séparation linéaire de l'espace d'entrée.

Dans le cas de données linéairement séparables, la première idée clé des SVM est de trouver la frontière de séparation qui maximise la marge maximale, c'est à dire, à maximiser la distance entre la frontière séparant deux classes, d'une part, et les éléments les plus proches de la frontière appartenant à chacune des deux classes, d'autre part. Ces éléments les plus proches de la frontière seront appelés vecteurs de support.

Cependant, la plupart des cas réels d'application nécessitent plus qu'une simple séparation linéaire de l'espace des entrées. Dans le cas de données non-linéairement séparables, la deuxième idée clé des SVM est de se ramener à un problème linéairement séparable en projetant l'espace de représentation des données d'entrées en un espace de plus grande dimension, dans lequel il existe une frontière de séparation linéaire (Christianini and Shawe-Taylor 2000). Dans le but de s'affranchir du problème de la définition mathématique de cette projection (qui peut s'avérer impossible dans de nombreux cas), Cette projection s'effectue en utilisant des noyaux. Différents types de noyaux ont été proposés par le passé : linéaire, polynomial, sigmoïdal ou radial basis function (RBF). (Hsu, Chang, & Lin, 2003) préconisent d'utiliser des noyaux RBF qui incluent moins d'hyperparamètres que les noyaux polynomiaux. De plus (Keerthi & Lin, 2003) ont prouvé que les noyaux linéaires sont un cas particulier des noyaux RBF et Lin and Lin (2003) ont montré que certains choix de paramètres des noyaux RBF permettent de retrouver les noyaux sigmoïdaux.

Par ailleurs, dans le cas de classes non séparables (chevauchement de classes), l'algorithme d'optimisation utilisé pour déterminer la frontière optimale doit être à marge souple qui exploitent des variables ressorts (Cortes & Vapnik, 1995).

Le principal avantage des SVM comparativement aux autres machines d'apprentissage réside dans l'absence de minima locaux. Par contre, de même que les réseaux de neurones, un modèle SVM sera un modèle de type boîte noire. De plus, la performance de l'algorithme d'apprentissage sera assez sensible au réglage (par l'utilisateur) de plusieurs paramètres du modèle tel que le poids des variables ressorts (Cortes & Vapnik, 1995) dans le critère à optimiser ou la largeur du noyau dans le cas de noyau RBF. Le Tableau 13 résume les principaux avantages et limitation des SVM.

Tableau 13 - Avantages et limites des machines à vecteurs de support (SVM)

Avantages	Absence de minima locaux
Limites	Apparaît comme « boîte noire » pour l'utilisateur Très sensible au réglage des paramètres Le modèle est figé et ne s'adaptera pas aux évolutions du système qu'il représente

II.2.1.4 Réseaux de neurones (NN)

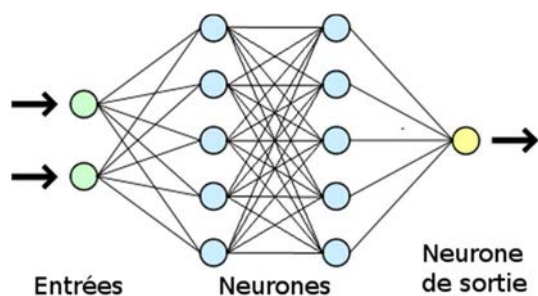


Figure 14 - Exemple de réseau de neurones

Un réseau de neurones est une structure informatique directement calquée sur les neurones biologiques du cerveau où les neurones d'entrées (vert sur la Figure 14) vont être stimulés en fonction de l'environnement de manière à activer ou non des synapses qui conduiront jusqu'à

l'activation des neurones cachés (bleu) puis d'un ou plusieurs neurones de sortie (jaune).

L'ancêtre de tous les réseaux de neurones est le perceptron (Rosenblatt, 1959) qui s'appuie sur une modélisation particulière d'un neurone, le neurone formel (McCulloch & Pitts, 1943). De ce perceptron découle tous les types de réseaux de neurones encore utilisés à ce jour. Les réseaux de neurones sont donc tous un ensemble de neurones formels⁴

⁴ A noter qu'une troisième génération de réseaux de neurones basée sur un nouveau type de neurone formel, le neurone impulsionnel est étudié dans le cadre principalement de travaux de modélisation en neurosciences (Thorpe, 2012).

associés en couches et fonctionnant en parallèle. Il est possible de classer ces réseaux en fonction de leur architecture d'une part (Herauld & Jutten, 1994):

- Les réseaux « feed-forward » ; l'information se propage de l'entrée vers la sortie sans retour en arrière (perceptron multicouches MLP, réseaux à fonction radiale (RBF)).
- Les réseaux « feed-back » ; ce sont des réseaux récurrents (cartes auto-organisatrices de Kohonen (SOM), réseaux de Hopfield, adaptive resonance theory (ART), réseaux de neurones multicouches sigmoïdaux récurrents (RMLP)).

Et en fonction de leur mode d'apprentissage d'autre part (Dreyfus & al., 2002) :

- Apprentissage supervisé : on dispose d'un ensemble de données dont le résultat associé est connu. Cet ensemble est utilisé lors de l'apprentissage pour déterminer quand le réseau a raison et quand il se trompe (MLP, RBF, RMLP).
- Apprentissage non supervisé : on ne dispose pas de la connaissance de la solution pour un ensemble de données et on va chercher à regrouper ces données en classes selon des critères de ressemblance inconnus *a priori* (SOM, ART).
- Sans apprentissage : Hopfield.

Ces différents types de réseaux vont donc être amenés à avoir différents types d'applications.

- Les réseaux de neurones bouclés sans apprentissage (Hopfield) sont exploités pour l'optimisation combinatoire.
- Les réseaux de neurones bouclés à apprentissage non supervisé (SOM, ART) permettent d'effectuer de l'analyse et visualisation de données en général, et du clustering en particulier.
- Les réseaux de neurones bouclés à apprentissage supervisé (RMLP) sont utiles pour la modélisation dynamique et la commande de processus.
- Les réseaux de neurones non bouclés à apprentissage supervisé (MLP, RBF) servent pour traiter des problèmes de modélisation statique, d'approximation de fonction et de discrimination (classification). Ce sont donc ces derniers qui nous intéressent ici car ce sont les seuls qui permettent de traiter efficacement un problème de classification.

La principale différence existant entre les MLP et les RBF est relative à la fonction d'activation exploitée par les neurones cachés.

Dans un RBF, la fonction d'activation associée aux neurones sera une fonction radiale de base et le principe d'un tel réseau sera de calculer la distance euclidienne entre le vecteur d'entrée et les centres de ces fonctions radiales. L'apprentissage d'un modèle RBF consiste à déterminer son architecture (le nombre de neurones cachés) et à fixer la valeur des paramètres (les centres des fonctions radiales, leur écart type, et les poids). Dans la plupart des applications, pour simplifier l'apprentissage, le nombre de fonctions radiales est

déterminé par validation croisée et essais-erreurs et les centres et écart types des fonctions radiales sont fixés (Viennet, 2006). Les RBF sont cependant peu adaptés au traitement de problèmes de grandes dimensions (Viennet, 2006).

Dans un MLP, la fonction d'activation associée aux neurones sera de type sigmoïdale (sigmoïde ou tangente hyperbolique) et le principe d'un tel réseau est d'effectuer la somme pondérée de ses entrées transformées par les fonctions d'activations des neurones cachés. L'apprentissage d'un tel réseau revient à déterminer les paramètres (poids et biais) liant les neurones d'une couche à ceux de la couche suivante. Cet apprentissage se fait généralement par optimisation d'un critère en exploitant l'algorithme de rétro propagation du gradient (McClelland & Rumelhart, 1986).

Le principal avantage des réseaux de neurones (MLP et RBF) est que ce sont des approximateurs universels (Hornik, 1991). De plus, vu que la sortie d'un MLP n'est pas linéaire par rapport aux poids du réseau, un MLP sera un approximateur parcimonieux⁵ (Barron, 1993), (Dreyfus & al., 2002). Pour qu'un RBF devienne un approximateur parcimonieux, il est nécessaire de rendre les centres et écart-types des fonctions radiales variables ce qui complexifie grandement l'apprentissage (Dreyfus & al., 2002).

Les principales limites des modèles neuronaux sont que le modèle résultant est un modèle « boîte noire » duquel on ne peut pas extraire de connaissance et que leurs algorithmes d'apprentissage effectuent une recherche locale d'optimum ce qui les rends sensibles à leur initialisation. Le Tableau 14 résume les principaux avantages et limitations des MLP.

Tableau 14 - Avantages et limites des réseaux de neurones

Avantages	Approximateur universel et parcimonieux Le modèle peut s'adapter aux évolutions du système
Limites	Apparaît comme « boîte noire » pour l'utilisateur Très sensible à l'initialisation Le modèle peut s'adapter

II.2.1.5 Ensembles de classificateurs

Les ensembles de classificateurs, aussi connus sous les noms de « committees of learners », « mixture of expert », « multiple classifier systems », ou encore « classifier ensembles » représentent un large domaine de recherche sur les dernières décennies. Il

⁵ La « parcimonie » exprime le fait que le MLP nécessite un faible nombre de paramètres ajustables pour réaliser correctement sa tâche.

s'agit d'un groupe de classificateurs qui vont fonctionner simultanément en parallèle et dont on va combiner les réponses comme présenté sur la Figure 15.

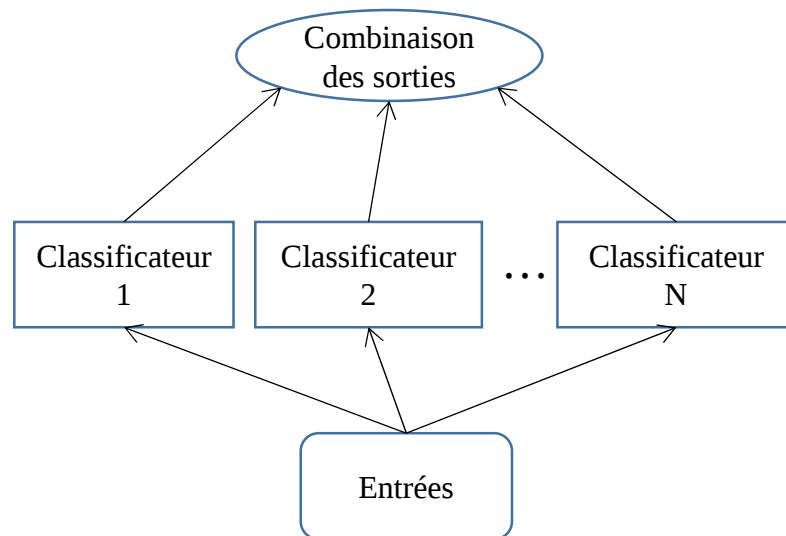


Figure 15 - Exemple d'ensemble de classificateurs

Au sens large, ces ensembles permettent d'obtenir plusieurs points de vue et donc d'éviter les cas où la structure particulière du modèle aura conduit à une aberration. Avec les machines d'apprentissage, les méthodes des ensembles permettent d'obtenir de meilleures performances en prédiction que ce qui aurait été obtenu par un seul des constituants de l'ensemble.

Les ensembles peuvent être utilisés à plusieurs niveaux. Dans notre cas industriel, deux niveaux au moins sont possibles :

- Pour la prévision de la qualité avec éventuellement une population de différents types de classificateurs
- Pour la détection de la dérive

Nous allons nous focaliser sur le premier niveau, c'est-à-dire, chercher à améliorer les résultats de classification en associant plusieurs classificateurs.

Faire fonctionner un ensemble de classificateurs requiert inévitablement plus de travail que de ne faire fonctionner qu'un seul classificateur, c'est pourquoi on trouve plus généralement des ensembles d'algorithmes rapides comme les arbres de décision par exemple. Mais il existe cependant de nombreuses autres méthodes basées sur les ensembles dont quelques-unes sont présentées dans les paragraphes suivants.

Le but des ensembles d'apprentissage est de construire une collection de classificateurs individuels qui soient divers mais précis. Nous pouvons alors développer une méthode de classification très précise en combinant les décisions des classificateurs individuels de

l'ensemble en utilisant des méthodes de votes (Dietterich, 2000). Un ensemble de classificateurs peut être construit à 4 niveaux différents (Kuncheva, Combining pattern classifiers: methods and algorithms, 2004) :

- Les données,
- Les caractéristiques (Ho, 1998),
- Les classificateurs (Kuncheva, 2002),
- Les niveaux de combinaison (Kuncheva, 2002).

Le principe sous-jacent aux ensembles de classificateurs est présenté sur la Figure 15. La création d'un ensemble de classificateurs repose sur 2 étapes : la génération de multiples classificateurs et leur fusion (Dai, 2013). Ceci conduit à deux questions :

- Combien de classificateurs sont requis ?
- Comment les combiner ?

La clef pour concevoir un ensemble performant repose sur la diversité des classificateurs. Les 4 algorithmes (Kotsiantis S. , 2011) les plus populaires pour augmenter la diversité sont :

- Bagging : réalisation de k jeux de données d'apprentissage différents de même taille que le jeu de données d'apprentissage initial par tirage aléatoire avec remise des exemples dans le jeu d'apprentissage initial (Breiman L. , 1996),
- Boosting : construction de l'ensemble classificateur en modifiant la distribution du jeu de données d'apprentissage en fonction de la précision des classificateurs précédemment appris (Freund & Schapire, 1997),
- Rotation Forest : utilisation d'une décomposition en composante principales pour extraire les caractéristiques pour chacun des classificateurs individuels (Rodriguez, Kuncheva, & Alonso, 2006),
- Random Subspace : l'apprentissage de chaque classificateur se fait sur un sous espace du domaine d'entrée défini aléatoirement (Ho, 1998).

Les méthodes "bagging" et "random subspace" sont plus robustes que les autres lorsque les données sont bruitées (Dietterich, 2000), (Kotsiantis S. , 2011). Comme les données sur lesquelles nous travaillons sont issues du domaine industriel et fortement bruitées, nous utiliserons l'approche « bagging ».

a. Critère de sélection des classificateurs individuels

Différents critères de sélection ont été proposés dans la littérature. La méthode la plus simple qui se révèle également fiable et robuste est d'utiliser la meilleure performance individuelle. Elle est généralement préférée pour les applications industrielles (Ruta & Gabrys, 2005). L'erreur individuelle minimum (MIE) est un indicateur qui représente le

taux d'erreurs minimum des classificateurs individuels et promeut une stratégie de sélection du meilleur classificateur. Cette stratégie est définie par :

$$MIE = \min_j \left(\frac{1}{n} \sum_{i=1}^n e_j(i) \right) \quad (1)$$

où $e_j(i)$ représente l'erreur de classification du classificateur j pour la donnée i .

Une autre approche pour la sélection des classificateurs est de considérer la diversité des classificateurs individuels même si cette notion n'a pas réellement de définition acceptée de tous et est donc souvent difficile à mesurer explicitement. Différentes mesures de diversité ont été proposées comme les Q statistiques, la corrélation, le désaccord, la faute double, l'entropie des votes, l'index de difficulté, la variance Kohavi-Wolpert, l'interrater agreement, ou encore la diversité généralisée. De nombreux auteurs ont testé et comparé ces mesures en utilisant différents exemples (Kuncheva & Whitaker, 2003), (Tang, Suganthan, & Yao, 2006), (Aksela & Laaksonen, 2006), (Bi, 2012). (Kuncheva & Whitaker, 2003) montrèrent que toutes ces mesures donnent des résultats sensiblement équivalents. De plus, les mesures se sont révélées très corrélées les unes aux autres. De ce fait, le choix de l'une ou l'autre mesure de diversité proposée importe peu et nous utiliserons la mesure de « faute double » (DF) qui présente l'avantage de mettre en avant la proportion de données mal classées par les deux classificateurs comparées, ce qui peut se rapprocher de la notion de mauvaise classification d'un ensemble de classificateur. Cette mesure fut proposée par (Giacinto & Roli, 2001) pour construire une matrice de diversité en comparant les classificateurs par paires. Elle sélectionne les classificateurs qui sont les moins corrélés en évaluant la proportion des cas qui ont été mal classés par les deux classificateurs en question.

$$DF_{i,j} = \frac{N^{00}}{N^{11} + N^{10} + N^{01} + N^{00}} \quad (2)$$

où i et j représentent les deux classificateurs, et N^{ab} est le nombre de sorties correctes du classificateur i quand $a=1$ (ou incorrectes si $a = 0$) et le nombre de sorties correctes du classificateur j quand $b=1$ (ou incorrectes si $b = 0$). N^{00} représente le nombre de sorties incorrectes simultanément pour les deux classificateurs. Plus cette valeur est faible, plus la diversité entre les deux classificateurs considérée est grande.

b. Processus de sélection des classificateurs individuels

La sélection des classificateurs est un problème qui a été abordé par de nombreux auteurs (Ruta & Gabrys, 2005), (Hernandez-Lobato & Hernandez-Lobato, 2013), (Dai, 2013)).

Deux approches peuvent être exploitées (Ruta & Gabrys, 2005):

- Sélection statique des classificateurs. La sélection optimale est déterminée sur le groupe de données utilisé pour la validation, puis fixée et utilisée pour classer les nouvelles données.
- Sélection dynamique des classificateurs. La sélection est faite “en ligne” pendant la classification et se base sur les performances d’apprentissage et les différents paramètres des données en cours de classification.

Cette deuxième approche peut s’avérer très gourmande en temps de calcul.

Différentes stratégies, plus ou moins complexes, ont été proposées pour sélectionner les classificateurs pour un ensemble (Kim & Oh, 2008), (Ko, Sabourin, & Britto Jr, 2008), (Tsoumakas, Partalas, & Vlahavas, An ensemble pruning primer, 2009), (Yang, 2011), (Guo & Boukir, 2013), (Soto, Melin, & Castillo, 2013). Nous n’en présenterons que deux qui sont simples, basées l’une, sur la précision, l’autre, sur la diversité des classificateurs.

La première est donc basée sur l’utilisation de la précision des classificateurs individuels donnée par l’équation 2. Les classificateurs sont ajoutés à l’ensemble de manière séquentielle en allant du plus précis au moins précis. La structure choisie sera celle qui fournira les meilleurs résultats de classification.

La seconde stratégie est basée sur une approche de sélection par la précision et la diversité (SAD – Selection by Accuracy and Diversity) proposée par (Yang, 2011). C’est une stratégie récursive simple qui consiste à suivre les étapes suivantes :

1. Evaluer la précision de chacun des classificateurs en utilisant les données de validation.
2. Choisir le classificateur le plus précis pour réaliser un vote.
3. Calculer la diversité entre l’ensemble de classificateurs et les classificateurs individuels restants en utilisant la mesure double-faute.
4. Choisir le classificateur avec la plus forte diversité et l’ajouter à l’ensemble de classificateurs pour construire un nouvel ensemble.
5. Evaluer la performance de ce nouvel ensemble de classificateurs. Si tous les classificateurs ne sont pas utilisés, répéter les étapes 4 à 6, autrement, comparer tous les ensembles de classificateurs et sélectionner le meilleur.

c. La fusion des classificateurs

Les classificateurs sont généralement fusionnés à travers un vote majoritaire. (Kuncheva & Whitaker, 2003) ont étudié le problème du manque d’indépendance entre les classificateurs ce qui peut limiter les bénéfices attendus d’un ensemble de classificateurs.

Le vote majoritaire est généralement effectué sur les sorties binaires des classificateurs individuels.

Cependant, certains types de classificateurs tels que les MLP ou les SVM ne se contentent pas de donner une sortie binaire mais fournissent une valeur réelle qui peut être assimilée à la fonction d'appartenance à une classe que l'on retrouve dans la logique floue. Cette valeur peut être comparée à un seuil dans le but de déterminer la classe de l'exemple considéré. Par exemple, avec le seuil classique de 0,5, une valeur de sortie de 0,49 correspondra à l'association de l'exemple considéré à la classe 0 quand une valeur de 0,51 permettra d'associer l'exemple à la classe 1. Cette approche conduit à une perte d'informations. Aussi pour ces classificateurs, une deuxième approche de fusion sera de construire la moyenne des sorties réelles des différents classificateurs individuels qui sera dans un second temps comparée au seuil considéré dans le but d'associer l'exemple à la classe considérée.

Enfin, une autre approche pour fusionner les différents classificateurs est de traiter ce problème comme un processus d'apprentissage à part entière. Wozniak a ainsi proposé d'effectuer la fusion en effectuant l'apprentissage d'un MLP (Wozniak, Experiments with trained and untrained fusers, 2007) ou encore en exploitant un algorithme évolutionnaire (Wozniak, 2009). Aussi un MLP peut être utilisé pour effectuer cette tâche de fusion, la sélection des classificateurs individuels à inclure dans l'ensemble étant obtenue par l'intermédiaire d'un processus de « pruning » des entrées (élagage des classificateurs individuels non pertinents pour le modèle global).

d. Choix du classificateur final

Deux approches peuvent être exploitées pour construire le classificateur final :

- Sélectionner le classificateur unique qui donne les meilleurs résultats,
- Construire un ensemble classificateur en associant un processus de sélection et de fusion de classificateurs.

L'objectif de notre travail n'étant pas focalisé sur la construction des ensembles, nous n'allons présenter qu'un seul exemple de chacune de ces deux approches.

Les ensembles classificateurs sont souvent construits en utilisant un seul type de classificateurs. On trouvera ainsi des ensembles neuronaux (Zhou, Wu, & Tang, 2002), (Windeatt, 2005), des ensembles SVM (Hajek & Olej, 2010), (Haghighi, Vahedian, & Yazdi, 2011), (Li, et al., 2012), des ensembles kpp (Kim & Oh, 2008), (Ko, Sabourin, & BrittoJr, 2008)), ou des ensembles arbres de décision (Dietterich, 2000), (Tsoumakas,

Katakis, & Vlahavas, 2004), (Soto, Melin, & Castillo, 2013). Peu d'applications ont utilisé conjointement différents types de classificateurs (Aksela & Laaksonen, 2006), (Yang, 2011). (Wozniak, Grana, & Corchado, 2014) ont proposé un état de l'art concernant les classificateurs multiples qui donne une vision du spectre d'applications de ces approches incluant la télédétection, la sécurité informatique, la banque, la médecine et les systèmes d'aide à la décision. L'application de cette approche au problème de contrôle de la production en général et du contrôle de la qualité en particulier n'a pas été étudiée. De plus l'impact de la combinaison de divers classificateurs sur la performance de l'ensemble classificateur est généralement non évalué. Le Tableau 15 résume les principaux avantages et limitations des ensembles classificateurs.

Tableau 15 - Avantages et limites des ensembles de classificateurs

Avantages	Amélioration des performances par rapport à un classificateur seul Existence de méthodes pour traiter les données bruitées
Limites	Plus chronophage en création comme en utilisation qu'un seul classificateur

Toutes les méthodes et outils abordés présentent donc des avantages pour traiter notre problème, mais ils ont aussi tous le même inconvénient, à savoir, ils ne permettent pas ou mal de suivre les dérives des process et ainsi de permettre de déterminer le « bon moment » où il faudra faire un réapprentissage.

II.2.2 La détection des évolutions de processus

Les systèmes de production sont des systèmes évolutifs évoluant dans un environnement changeant. Du fait de ces évolutions, provenant soit du système, soit de l'environnement, il est nécessaire de rendre le système de surveillance de la qualité évolutif. Pour ce faire il faut déterminer quand et à partir de quand une telle évolution rend le modèle inopérant et implique donc une réadaptation de ce modèle. D'un point de vue technique de la machine d'apprentissage, deux raisons peuvent engendrer une dérive de comportement.

La première raison est qu'au moins l'une des données d'entrée est en dehors des zones de fonctionnement normales du système de prévision. Dans cette situation, le modèle n'est pas capable de donner une réponse correcte. Par exemple, sur la Figure 16, les prévisions sont possibles dans les cas 1 et 2 mais pas en 3 (facteur 5 hors des bornes) ni 4 (facteur 2 hors des bornes). La raison pour laquelle les valeurs des facteurs se retrouvent en dehors

des bornes peut être le résultat d'une dérive (usure progressive d'un outil) ou d'un brusque changement (par exemple une chute de pression due à une panne de compresseur).

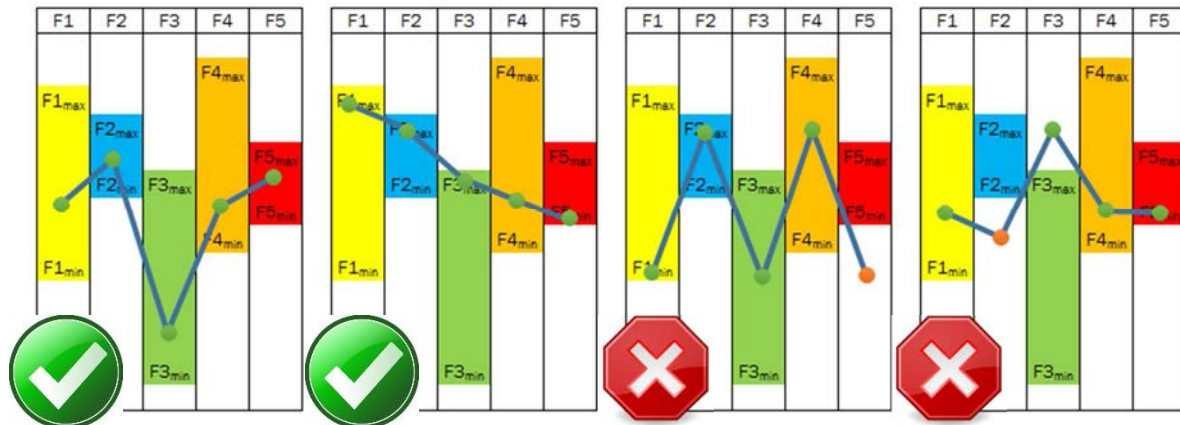


Figure 16 – Situations où les prévisions sont possibles ou impossibles.

La seconde raison concerne les modifications non-contrôlées de comportement du système réel. En effet, en changeant un paramètre (volontairement ou involontairement) qui n'est pas une entrée du système de prévision, il se peut que le comportement du système réel s'en trouve quand même affecté. Une autre raison encore peut provenir d'un changement du process lui-même (déplacement de capteur, changement de composant ...).

La capacité à vérifier les hypothèses émises en restant informé de la réalité est un prérequis pour la détection des dérives. Ceci implique que le système informatique puisse collecter les informations à la sortie du système réel. En comparant ces nouvelles données avec les hypothèses du système de prévision, il est possible alors de suivre les écarts entre la théorie et la réalité. Cette valeur sera appelée le taux d'erreur. L'évolution de ce taux de mauvaises classifications peut aussi être liée à des facteurs de bruit, considérés comme normaux dans l'industrie, et qui rendent donc difficile la détection de la dérive. Les approches suivantes ont été proposées comme détecteur de dérive.

II.2.2.1 Les cartes de contrôle

Les cartes de contrôle sont particulièrement adaptées au contrôle dynamique de séries de données temporelles (Tague, 2004). Cet outil est utile pour déterminer statistiquement si la variation est une variation « normale sous contrôle » ou s'il s'agit d'une dérive. En effet, même sous contrôle, il est connu qu'il existe approximativement une probabilité de 0.27% qu'une donnée sorte des bornes 3-sigma (Pareto). Dans l'idéal, la rencontre d'un seul de ces points isolés ne devrait pas engendrer de resynchronisation. En revanche,

plusieurs points détectés montrent la présence d'une dérive dont la cause n'est peut-être pas encore connue.

II.2.2.2 DDM et EDDM

Ces deux méthodes sont détaillées dans (Baena-Garcia, et al., 2006). La DDM (Drift Detection Method) surveille le nombre d'erreurs produites par le système et suppose que le taux d'erreurs suit une loi binominale. Si la distribution des données est supposée stationnaire, le taux d'erreurs diminue avec l'augmentation du nombre de données au fur et à mesure que l'algorithme d'apprentissage améliore son apprentissage. Une augmentation significative du taux d'erreurs lors de l'apprentissage implique donc un changement dans la distribution générant les données. Le principal inconvénient de cette méthode est la détection difficile des dérives lentes.

EDDM (Early Drift Detection Method) considère cette fois le nombre de données se trouvant entre des erreurs de classification consécutives. Comme la distribution des données est toujours supposée stationnaire, cette distance augmente avec le nombre de données. Une diminution de la distance implique donc une dérive. Cette méthode présente l'inconvénient que, en fonction des valeurs de ses paramètres, elle est adaptée à la détection, soit des dérives lentes, soit des dérives rapides ; mais ne parviendra donc pas à être correctement efficace si l'objectif est de détecter sur un même jeu de données les deux types de dérives.

II.2.2.3 EDIST (Error DISTance based approach)

Cette méthode présentée par (Khamassi & Sayed-Mouchawed, 2014) est dérivée de EDDM mais en utilisant un test d'hypothèse statistique afin de comparer les distributions des distances d'erreurs de classification sur deux fenêtres de données différentes. Si la différence dépasse un seuil, lui-même ajusté de manière adaptative par la méthode, une dérive est alors détectée. EDIST est donc globalement plus performant que DDM et EDDM.

II.2.2.4 PHT

Le test de Page Hinkley (Hinkley, 1971) est une technique d'analyse séquentielle utilisée typiquement pour détecter les changements dans le traitement du signal. Elle est couramment utilisée et se retrouve dans beaucoup de benchmarks tel que (Sebastiao & Gama, 2009). Il permet la détection efficace de changements dans le comportement

normal d'un processus modélisé. Le test est conçu pour détecter un changement dans la moyenne d'un signal Gaussien. Il considère une variable cumulative définie comme la différence cumulée entre les valeurs observées et leur moyenne au moment du test. La valeur de cette variable est comparée à sa valeur minimale et lorsque la différence entre les deux est supérieure à un seuil donné, le saut de moyenne est détecté. L'intérêt de cette approche réside dans le fait qu'elle permet également d'estimer l'instant de ce saut.

Les avantages et inconvénients des méthodes présentées dans ces paragraphes sont résumés dans le Tableau 16.

Tableau 16 - Quelques méthodes de détection de dérives

Méthodes	Avantages	Inconvénients
Cartes de contrôles	Adaptées au contrôle dynamique de séries de données temporelles. Solution visuelle.	
DDM	Indépendante de l'algorithme d'apprentissage	Détection difficile des dérives lentes.
EDDM		Adaptée soit aux dérives lentes, soit aux dérives rapides suivant le réglage des paramètres, mais pas aux deux simultanément.
EDIST	+ performant que DDM et EDDM	
Page Hinkley Test	Très connu et globalement performant	Dépend de la valeur du seuil λ

II.2.3 *Problématique scientifique spécifique*

Même si de nombreuses méthodes de modélisation des systèmes physiques par apprentissage existent, rares sont celles qui parviennent à combiner l'adaptation réactive aux évolutions de processus avec une robustesse/stabilité définie pour éviter les réapprentissage répétitifs probablement déclenchés par le bruit.

Dans notre cas les distributions sont généralement gaussiennes et nous recherchons bien un instant où la dérive se produit et par conséquent le test de Page Hinkley conviendra. Celui-ci sera utilisé avec une carte de contrôle qui permettra de mettre en évidence le début de la dérive.

II.3 PROPOSITION

II.3.1 Méthodologie générale

Nous proposons une méthodologie de résolution du problème industriel relatif à la maîtrise de la qualité à partir de la conception d'un outil à base de machines d'apprentissage qui surveille l'ensemble des paramètres contrôlables et non contrôlables pour maîtriser la qualité de manière proactive et émettre des alertes pour interdire certains ordonnancements.

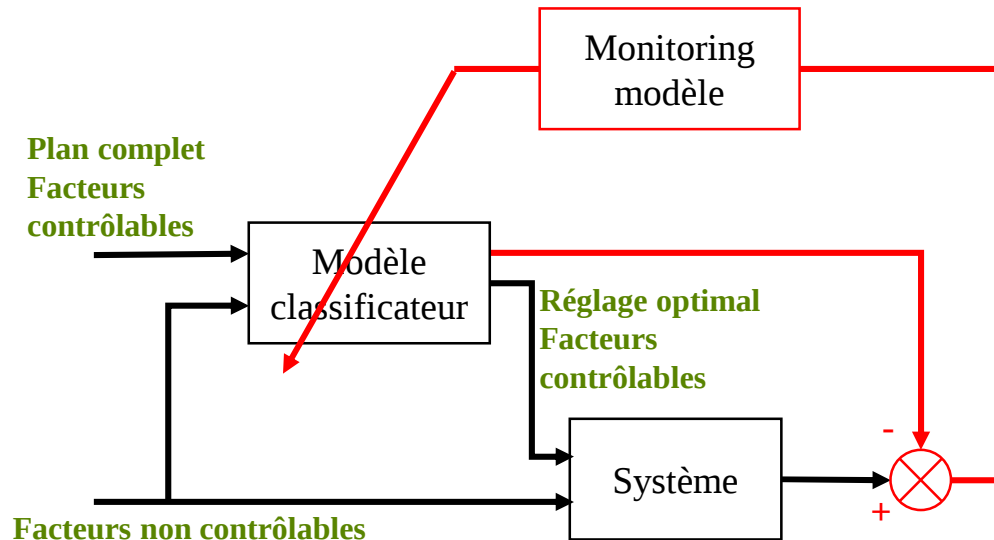


Figure 17 – Proposition : Adaptation régulée du réseau de neurones

Le bloc « système » représente un système physique dont le comportement est humainement difficile à prévoir en raison des nombreux facteurs incontrôlables auxquels il est soumis. Ce système réel est modélisé à l'aide de classificateurs symbolisés par le bloc « modèle classificateur ». Ce modèle est idéalement capable de « singer » le comportement du système réel quel que soient les valeurs prises par les facteurs non contrôlables. Il permet alors d'évaluer l'impact des facteurs contrôlables afin de pouvoir, à l'aide de plans complets par exemple, déterminer les plages de réglage autorisées, où même le réglage optimal de ces facteurs en fonction des facteurs non-contrôlables. Parce que le système physique et son environnement (facteurs non-contrôlables) évoluent, nous proposons un module de surveillance appelé observateur qui comparera les prévisions à la réalité de manière à pouvoir réadapter le modèle dès lors qu'une dérive sera détectée. Il s'agit de la partie rouge de la Figure 17.

La méthodologie de mise en place de ce système de maîtrise de la qualité est présentée Figure 18.

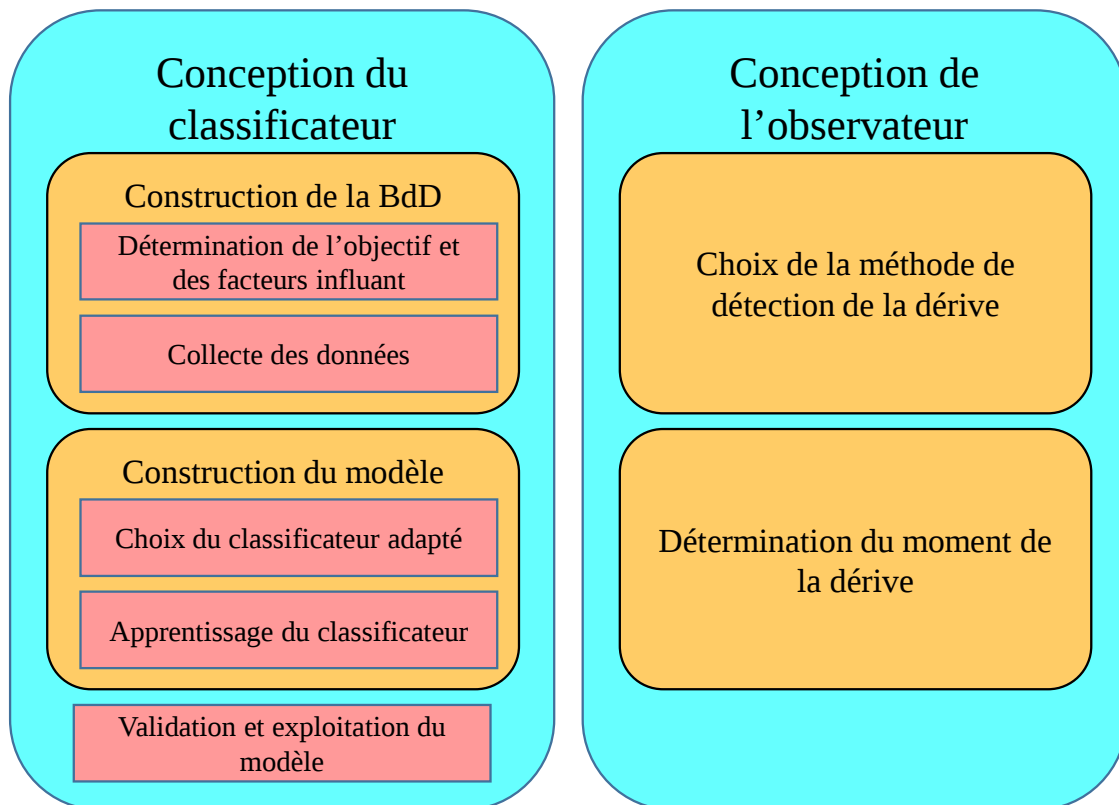


Figure 18 - Méthodologie de mise en place du système de maîtrise de la qualité

II.3.2 Conception du classificateur

La première étape consiste à concevoir le modèle représentatif du système physique. Il s'agira d'un classificateur qui sera capable de classer le résultat d'une combinaison de facteurs influant d'une manière cohérente en fonction du système réel. Suivant le type de problème, plusieurs types de classificateurs peuvent convenir. La première étape réside donc dans la définition des objectifs et des facteurs influents.

a. Détermination de l'objectif et des facteurs influents

Globalement, l'analyse de la situation nécessite deux étapes :

- La détermination de la fonction « objectif » du système à surveiller
- L'identification des paramètres influents

La fonction objectif est la variable à prédire. Il s'agit traditionnellement de la valeur dont le comportement est difficile à comprendre à cause des multiples dépendances à de nombreux facteurs difficiles à évaluer. Elle peut être par exemple, la température d'un four (qui dépend de la température extérieure, de la température des produits qui y entrent, des temps d'arrêt, etc...). Parfois, la fonction objectif peut être une agrégation de plusieurs objectifs unitaires, spécialement dans des contextes de prévision de la non-qualité où le but

est de préserver un niveau d'apparition de défauts bas pour plusieurs causes différentes. Dans tous les cas, nous appellerons N_o le nombre d'objectifs de la fonction objectif.

Pour isoler les paramètres influents, une réflexion avec les experts doit être menée. Le but est de lister les N_f facteurs qui affectent directement la fonction objectif. Cette étape peut être menée à l'aide d'outils tels que le diagramme Ishikawa, par exemple. Le point le plus important étant de ne pas oublier un facteur critique. Pour chaque paramètre F_x , des bornes minimale ($F_{x_{min}}$) et maximale ($F_{x_{max}}$) doivent être estimées dans le but de déterminer une zone de fonctionnement normal du système appelée le « domaine d'apprentissage ». La figure suivante montre un exemple de zones de fonctionnement pour 5 facteurs différents. Ces zones, même si elles semblent continues, peuvent être constituées d'éléments discrets. S'il n'est pas possible de les déterminer avec les experts, il peut être envisagé de les extraire directement des données d'apprentissage. Les unités des ordonnées ne sont volontairement pas mentionnées car chaque paramètre peut avoir sa propre unité.

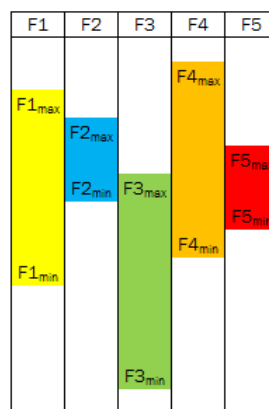


Figure 19 – Zones de fonctionnement normal du système pour 5 facteurs influents.

Ces facteurs peuvent être classés de différentes manières. Par exemple, il est possible de les classer suivant les 5M⁶ (Ishikawa, 1982). Nous proposons ici une classification suivant leur contrôlabilité qui est plus pertinente par rapport à l'application qui en découlera :

- Facteurs d'environnement (liés au milieu) comme la température ou l'humidité par exemple. Ces facteurs sont généralement peu ou non-contrôlables. Même s'ils sont parfois faciles à mesurer, ils peuvent être très difficiles de les asservir.

⁶ Machine (technologie), Méthode (processus), Matériau, Manutention, et Milieu (environnement)

- Facteurs techniques/économiques (liés à la machine et aux méthodes) résultant principalement de l'état de la machine durant l'opération. Ces facteurs sont contrôlables car ils correspondent exactement au réglage de la machine.
- Facteurs humains (liés à la manutention) dans les opérations manuelles. Ils sont difficiles à prendre en compte car ils varient souvent de manière conséquente entre les opérateurs. Les tentatives de contrôler les facteurs humains (poka yoke, standards de travail, etc...) sont donc souvent limitées et présentent de multiples contraintes.

Cette classification, qui peut aussi faire référence aux piliers de l'économie durable, est résumée et illustrée de quelques exemples dans le Tableau 17.

Tableau 17 - Contrôlabilité des facteurs et piliers de l'économie durable.

	Environnement	Technico-économique	Social
Ishikawa /5M équivalence(s)	Environnement	Machine Méthode	Energie humaine
Contrôlabilité	Peu- ou non-contrôlable (facteurs de bruit)	Contrôlable	Beaucoup de contraintes et de limitations
Exemples	Température ambiante	Aspiration	Vitesse de peinture
Dérive associée et défaut	Augmentation en milieu de journée → temps de séchage trop court → "Microbullage"	Diminution due à une fatigue du compresseur → reste de poussières avant laquage → Grains	Lenteur à cause d'un mal de bras. → pré-séchage avant recouvrement → phénomène de bandes
Pour optimiser les réglages initiaux	Plans d'expériences optimaux		Entraînement, Standards
Pour contrôler les dérives	Contrôle en-ligne	Plan de maintenance préventive	Poka-Yoke

b. Collecte de données

La collecte de données doit essentiellement répondre à 3 questions.

Que doit-on collecter ? Au moins les N_f paramètres jugés comme influant. Tous les paramètres incertains que certains jugent inutiles mais que d'autres pensent influant devraient être collectés aussi de manière à ne passer à côté d'aucun élément ayant de l'influence sur la fonction objectif. Les N_o valeurs des objectifs unitaires doivent aussi être collectées à la sortie du système.

Comment doit-on les collecter ? Evidemment, le sens commun voudrait qu'on stocke ces données de manière à pouvoir les traiter informatiquement. La table qui contiendra les données doit pouvoir être adaptable (nombre de paramètres et type de données) dans le cas où un nouveau paramètre devrait être pris en compte ou que l'on décide de suivre un paramètre à la place d'un autre.

Combien de données doit-on collecter ? Dans l'idéal, tous les cas d'utilisation devraient être couverts. Cela signifie que pour chaque F_x , au moins une donnée proche de $F_{x_{max}}$ et une autre proche de $F_{x_{min}}$ doivent se trouver dans la base de données collectées. La période de collecte doit donc être adaptée dans le but de remplir cette condition. Comme cela a été dit précédemment, si les bornes n'ont pas pu être déterminées par les experts (ou s'il est impossible d'attendre le temps de collecte nécessaire pour couvrir tous les cas d'utilisation), les bornes pourront être déterminées à partir des données collectées mais la zone de fonctionnement normal du système sera réduite entre les situations extrêmes rencontrées pendant la collecte des données.

Un autre problème réside dans le fait que la collecte de données est souvent perçue comme une perte de temps. En effet, si cette tâche n'est pas réalisée automatiquement, et si les données ne sont pas utilisées pour autre chose (comme des indicateurs de production par exemple), l'opérateur doit stopper son travail pour collecter les informations et a donc l'impression de « gaspiller » son temps de manière improductive. C'est pourquoi il est très important de travailler à la fois sur 3 objectifs :

- S'interfacer directement avec les machines dès que possible pour récupérer des informations plus fiables en temps masqué.
- Créer des interfaces de saisie simples d'utilisation (interfaces tactiles par exemple) et intuitives pour réduire le temps de saisie au minimum.
- Accompagner les opérateurs dans le changement en insistant sur l'intérêt de cette collecte de données qui ne portera pourtant ses fruits que plus tard, lorsque la base de données sera assez conséquente pour débiter l'apprentissage.

Heureusement, la plupart des entreprises sont intéressées par l'archivage des données de production, principalement pour des raisons de traçabilité, mais elles n'exploitent que très rarement ces informations et, si elles le font, c'est principalement dans le but de suivre des indicateurs pour du management visuel en temps réel.

Il faut ensuite préparer les données en réalisant des étapes de sélection, préprocessing et transformation des données qui correspondent aux premières étapes du processus KDD (Knowledge Discovery in Data) (Patel & Panchal, 2012). Les données doivent être pré-

traitées dans le but, par exemple, de synchroniser les différentes bases de données, de supprimer les données aberrantes ou encore de digitaliser les données qualitatives telles que la couleur par exemple.

c. Choix du classificateur (unique, ensemble ?)

Tandis qu'un classificateur unique peut convenir pour certains problèmes, il faut parfois avoir recours aux ensembles pour obtenir plus de fiabilité dans les résultats. En revanche, utiliser systématiquement un ensemble de classificateur peut représenter une perte de temps pour les cas plus simples.

d. Conception du modèle de prévision par apprentissage

L'exploration de données, aussi connue sous les expressions de « fouille de données », « forage de données », « prospection de données », « datamining », ou encore « extraction de connaissances à partir de données », est une partie du processus KDD qui consiste à analyser une grande quantité de données de manière à en extraire un savoir ou une connaissance. Le processus KDD identifie les nouveaux patterns compréhensibles et utilisables en exploitant les données collectées.

Comme montré dans (Agard & Kusiak, 2005), le volume de données à analyser est souvent conséquent. L'objectif de la classification automatique est donc de trouver une partition dans un jeu de données afin de synthétiser des données complexes ou volumineuses et donc difficilement interprétables autrement. On parle souvent d'organisation d'objets en classes homogènes. L'avantage des démarches probabilistes face aux démarches géométriques est qu'elles proposent un cadre formel permettant de proposer des solutions à des problèmes difficiles à cause du nombre de classes.

e. Exploitation du modèle de prévision

Une fois que le modèle de prévision est construit et validé, nous bénéficions donc d'un outil numérique capable de se comporter de la même manière que le système physique. Si ce dernier a été choisi pour être modélisé, c'est soit parce que son comportement est particulièrement difficile à comprendre et à anticiper, soit parce qu'il est soumis à une utilisation telle qu'il est impossible de prendre du temps (et des matières premières) pour tester son comportement. L'existence d'une représentation numérique de ce système ouvre alors des perspectives intéressantes, la première étant bien évidemment la prévision d'un risque d'occurrences de défauts dans les conditions (paramètres) données. C'est ce premier objectif qui va permettre à l'opérateur de pouvoir tester, et éventuellement adapter, ses réglages de production avant même de produire. Une autre utilisation intéressante est la

détermination de plages d'utilisation des paramètres pour garantir un niveau de qualité correct. En effet, il est possible de déterminer des bornes basses et hautes pour des paramètres clefs et d'interdire simplement la production en dehors de ces bornes. Même si elle ne représente pas les conditions optimales de production, cette restriction qui sera transmise aux opérateurs sous la forme d'un « standard » permettra de limiter l'apparition de la non-qualité.

Ces plages peuvent être construites en mettant en œuvre un ou des plans d'expériences de Taguchi. Dans notre cas, l'exploitation du modèle de prévision pour réaliser ces plans d'expériences nous permet de réaliser autant d'expériences que souhaité en confrontant les facteurs désirés sans devoir solliciter le poste de travail. Il s'agit donc de plans d'expériences simulés (sur un système virtuel) qui présentent donc l'avantage de pouvoir être réalisés en masse et d'être complètement dissociés de la réalité du poste de travail.

II.3.3 Conception de l'observateur

Comme nous l'avons montré précédemment l'observateur va donc être construit sur la base de l'usage de cartes de contrôle pour détecter la dérive puis celui de l'algorithme de Page Hinkley pour définir le moment où celle-ci est significative et donc impose un ré-apprentissage.

a. Détection de la dérive par carte de contrôle

Nous proposons donc d'exploiter une carte de contrôle pour surveiller l'apparition d'une dérive entre le modèle neuronal et le système. Du et al. (Du, Ke, & Su, 2012) ont travaillé sur la combinaison inverse de ces deux outils en proposant un algorithme de reconnaissance de carte de contrôle en utilisant les réseaux de neurones pour donner l'alerte en cas de problème qualité et fournir des indices pour en déterminer les causes.

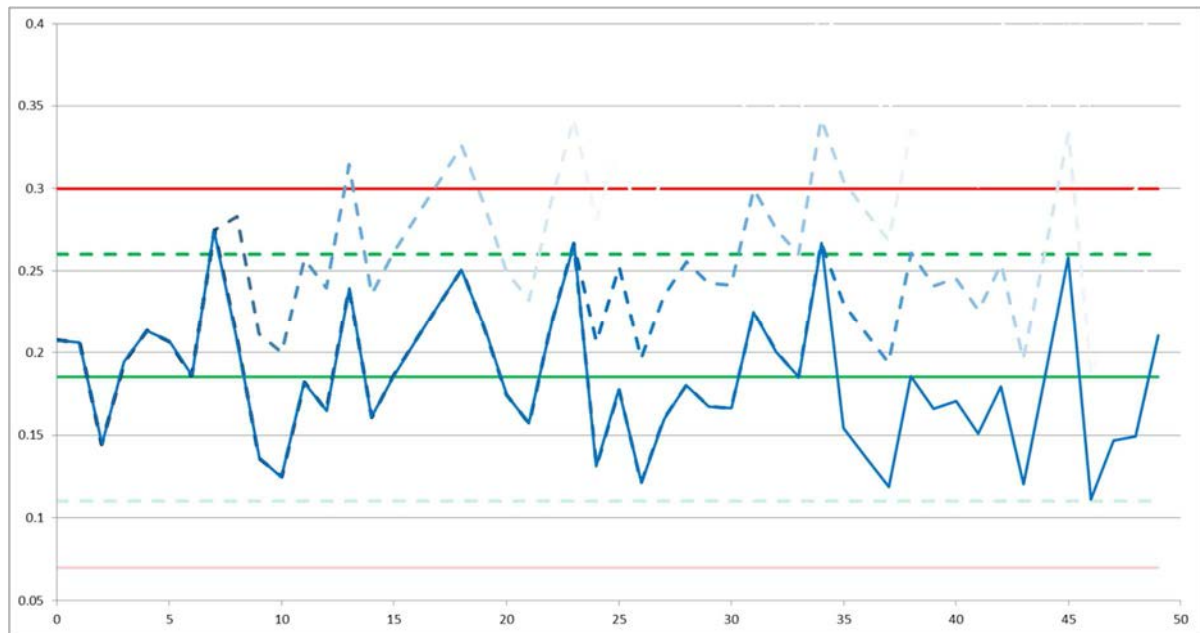


Figure 20 – Une carte de contrôle pour contrôler le modèle de prévision en ligne

Dans notre cas, le but de la carte de contrôle (p-charts) (Sematech, 2012) est de déterminer si le taux de mauvaises classifications de notre système de prévision de la qualité dérive. Le principe est présenté sur la Figure 20. Chaque point de la figure correspond à un taux de mauvaises classifications sur un échantillon de données constitué de n données, n correspondant donc à la taille de l'échantillon (n est choisi en fonction de plusieurs critères tels que la criticité, la fréquence...).

Pour les cartes de contrôle, deux paires de bornes sont normalement nécessaires en fonction du niveau de confiance recherché. Les bornes d'alarmes basse ($LCL_{95\%}$) et haute ($UCL_{95\%}$) pour un niveau de confiance de 95% sont données par :

$$\begin{cases} LCL_{95\%} = p - 1.96\sqrt{\frac{p(1-p)}{n}} \\ UCL_{95\%} = p + 1.96\sqrt{\frac{p(1-p)}{n}} \end{cases} \quad (3)$$

Dans notre utilisation, les bornes d'alarmes sont les seules à avoir un intérêt car le réapprentissage est effectué dès lors qu'un point passe ces bornes. En effet, dans l'utilisation classique des cartes de contrôles, lorsqu'un point passe le seuil d'alarme, un ré-échantillonnage immédiat de pièces est effectué et un nouveau point est créé. Si ce nouveau point est également en dehors des bornes, une dérive est donc effective et une action de correction doit être entreprise. Si au contraire ce nouveau point revient dans les bornes le précédent est considéré comme une anomalie statistique. Cependant, dans notre

carte de contrôle, chaque point correspond à l'ensemble des données de mauvaises classifications obtenues sur une fenêtre temporelle (dans notre cas, une semaine). Vouloir obtenir un deuxième point pour prendre une décision impose donc d'attendre une fenêtre temporelle supplémentaire avant de réagir ce qui peut avoir un impact important sur la qualité de production durant cette semaine. De plus, comme notre taux de mauvaises classifications est déjà une moyenne de plusieurs valeurs, nous supposons que si le point considéré est en dehors des bornes vertes pointillées, cela signifie qu'une multitude de données utilisées pour construire ce taux, donc cette valeur moyenne, étaient en dehors des bornes et que, par conséquent, le système est effectivement en train de dériver. C'est pourquoi, un réapprentissage est nécessaire dès l'apparition d'un point au-delà des bornes vertes pointillées (par exemple échantillon 7 sur la Figure 20). La courbe pointillée bleue sur la Figure 20 représente l'évolution du taux de mauvaises classifications si aucun réapprentissage n'est réalisé.

Les autres bornes, appelées bornes d'interdiction basse ($LCL_{99.73\%}$) et haute ($UCL_{99.73\%}$), normalement définies à un niveau de confiance de 99.73% dont la formule est donnée en (4) n'ont donc pas d'intérêt.

$$\begin{cases} LCL_{99.73\%} = p - 3\sqrt{\frac{p(1-p)}{n}} \\ UCL_{99.73\%} = p + 3\sqrt{\frac{p(1-p)}{n}} \end{cases} \quad (4)$$

Dans ces formules, n correspond à la taille de l'échantillon et p à la ligne centrale (CL) qui correspond au taux de mauvaises classifications obtenu sur les données de validation lors de l'apprentissage initial (Noyel, Thomas, Charpentier, Thomas, & Brault, 2013).

Parce que la courbe représente un taux de mauvaises classifications, les bornes inférieures n'ont pas, non plus, d'importance car une dérive dans ce sens signifierait que notre système de prédiction réalise des prévisions de plus en plus justes. Il n'y a donc au final qu'une seule borne intéressante dans notre cas, $UCL_{95\%}$. La décision de réapprentissage est prise dès lors qu'un point est situé au-dessus de cette borne. L'étape suivante consiste à déterminer sur combien de données le réapprentissage doit être effectué.

b. Détermination du moment de la dérive avec PHT

Les cartes de contrôle permettent de déclencher le réapprentissage, mais, pour corriger notre système de prévision, nous avons besoin de déterminer le nombre de données sur

lesquelles réapprendre. Même si la carte de contrôle met en relief « le point de données » où la dérive est prouvée, il ne peut pas déterminer depuis quand cette dérive a lieu, quelle est la date de l'événement qui a généré le départ de cette dérive. Il ne peut donc pas fournir d'indication sur le nombre de données à utiliser pour le réapprentissage. Deux possibilités sont envisageables :

- Le temps écoulé depuis le dernier réapprentissage n'est pas trop important : il est possible de réaliser le nouveau réapprentissage sur toutes les nouvelles données disponibles depuis le dernier réapprentissage. De cette façon, la tâche est simple et rapide (et similaire à l'apprentissage classique)
- Le temps écoulé depuis le dernier réapprentissage est trop important : pour économiser du temps, il est possible de recommencer un apprentissage complet sur un nombre défini de données. On peut parler de « fenêtre glissante ». Cette solution autorise le système à « oublier » les comportements anciens qui ne sont probablement plus adaptés pour obtenir un comportement du système plus flexible. Le bon nombre de données doit être défini correctement. S'il est trop faible, le système risque d'apprendre le bruit. S'il est trop long, le système ne sera pas assez flexible. L'idée est de déterminer le point d'inflexion de la courbe représentant le taux de mauvaises classifications.

Comme nous l'avons vu dans la partie II.2.2.4, pour déterminer le point d'inflexion, différentes méthodes peuvent être utilisées. Dans cette comparaison, il apparaît que PHT et SPC sont moins chronophages qu'ADWIN et FCWM. Ce point est crucial car l'un des objectifs principaux est de réduire le temps de calcul en optimisant la taille des lots de données.

Le but du PHT est de détecter un saut de la moyenne dans un signal constant pollué par un bruit blanc (Page, 1954); (Hinkley, 1971); (Basseville, 1986). Il est capable de déterminer si un saut a eu lieu et quand il s'est produit. Dans notre cas, le signal considéré est le taux de mauvaises classifications obtenu sur différentes données et nous cherchons à déterminer le moment de la désynchronisation entre le système physique et le modèle. Ce signal peut être représenté par une séquence de variables aléatoires gaussiennes $E=[e_i]$, $i=1,\dots,n$ de variance σ^2 et de moyenne m_i . Avec l'hypothèse que seulement un seul saut a eu lieu à un instant inconnu noté r avec $1 \leq r \leq n$, détecter cette dérive revient à accepter l'hypothèse H_1 d'un changement contre l'hypothèse H_0 de non changement :

$$\begin{cases} H_0 : & e_i \sim N(m_0, \sigma^2), & P(e_i) = P_0(e_i) & i = 1, \dots, n \\ H_1 : & e_i \sim N(m_0, \sigma^2), & P(e_i) = P_0(e_i) & i = 1, \dots, r-1 \\ & e_i \sim N(m_1, \sigma^2), & P(e_i) = P_1(e_i) & i = r, \dots, n \end{cases} \quad (5)$$

L'utilisation de ce test implique que les deux valeurs moyennes m_0 et m_1 sont connues *a priori*. Dans notre cas, la moyenne m_0 peut être estimée avec la moyenne du taux de mauvaises classifications obtenu lors de la conception du modèle, pendant la phase de validation sur le lot de données de validation. Cependant, la moyenne m_1 est inconnue et la valeur absolue minimale de l'amplitude de la dérive δ_m à détecter doit être fixée. Ces deux tests sont réalisés en parallèle dans le but de détecter respectivement une augmentation ou une diminution de cette moyenne. Ces tests peuvent être calculés récursivement pour détecter une augmentation de moyenne par :

$$\begin{cases} U_0 = 0, & U_i = U_{i-1} + e_i - m_0 - \frac{\delta_m}{2}, & i \geq 1 \\ \gamma_0 = 0, & \gamma_i = \min(\gamma_{i-1}, U_i) & i \geq 1 \end{cases} \quad (6)$$

et le moment r_i de la dernière augmentation de moyenne est donnée par :

$$r_i = \max(i | U_i = \gamma_i) \quad (7)$$

Pour la diminution de moyenne, le test est donné par :

$$\begin{cases} T_0 = 0, & T_i = T_{i-1} + e_i - m_0 + \frac{\delta_m}{2}, & i \geq 1 \\ \eta_0 = 0, & \eta_i = \min(\eta_{i-1}, T_i) & i \geq 1 \end{cases} \quad (8)$$

et le moment r_d de la dernière diminution de moyenne est donnée par :

$$r_d = \max(i | T_i = \eta_i) \quad (9)$$

Donc, si la carte de contrôle détecte qu'un réapprentissage est nécessaire au moment n , le réapprentissage doit être réalisé en utilisant les données collectées entre les moments r et n avec r donné par :

$$r = \min(r_i, r_d) \quad (10)$$

Les tests (6) et (8) utilisent la valeur absolue minimale de l'amplitude de la dérive δ_m qui peut être fixée comme un multiple de l'écart-type σ de l'erreur obtenue sur les données de validation. Dans notre cas, δ_m est fixé à :

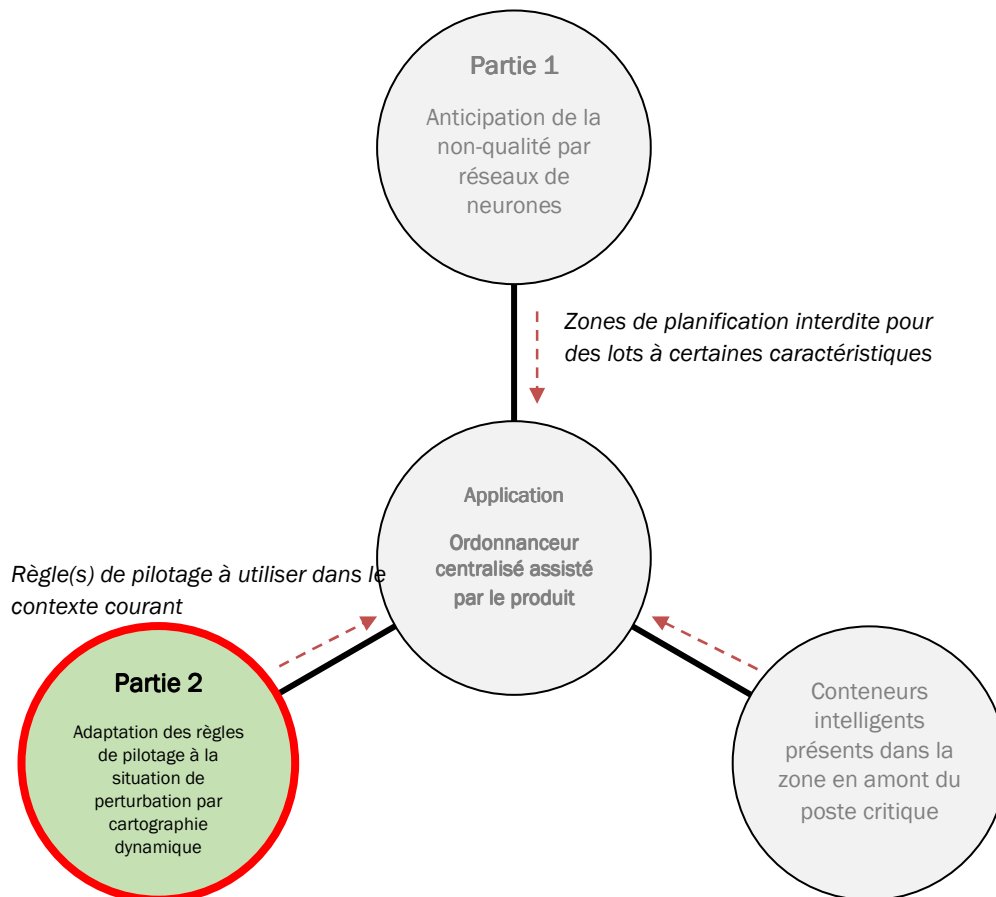
$$\delta_m = \frac{\sigma}{3} \quad (11)$$

II.4 CONCLUSION

Nous avons présenté dans cette partie notre méthodologie pour assurer dynamiquement le suivi et le contrôle des paramètres influant sur la qualité des produits. Celle-ci impose de concevoir un classificateur et un observateur adaptés. Le système de production ainsi contrôlé pourra voir ses réglages évoluer en fonction de l'évolution du contexte et de ces paramètres. Dans la partie IV.1, nous présenterons une application de cette démarche dont le pilote est déjà mis en place dans l'entreprise Acta Mobilier.

Ce faisant, ces réglages vont donc impacter la planification des lots dans l'atelier. Aussi pour ajuster l'équilibre charge/capacité des postes critiques de celui-ci il faudra faire évoluer les règles de pilotage. La partie suivante analysera ce problème et présentera notre proposition.

III EVOLUTION REACTIVE DES REGLES DE PILOTAGE EN FONCTION DE LA SATURATION DE L'OUTIL DE FABRICATION



III.1 INTRODUCTION

Cette partie traite de l'évolution réactive des règles de pilotage en fonction de l'état de saturation de l'outil de fabrication. En effet, à cause du taux de non-qualité fluctuant ajoutant de la perturbation aux flux de produits déjà complexes, il existe une forme de saturation de l'atelier qui paralyse la fabrication et engendre des retards considérables. Les règles de pilotage permettant de rendre prioritaire un lot par rapport à un autre sont un vecteur d'amélioration des flux. Cependant, les évolutions de situation conséquentes qui surviennent dans le cas de forte présence de non-qualité obligent à considérer ces évolutions des règles de pilotage de manière dynamique.

III.2 ETAT DE L'ART

III.2.1 Introduction : impact de la non-qualité sur la perturbation des flux

Il existe 4 manières de réagir suite à la détection d'une non-qualité :

- Soit le produit est jeté. C'est le cas des productions en série de produits relativement peu coûteux.
- Soit le produit est envoyé sur un poste de travail dédié aux réparations avant de reprendre le cours de sa fabrication normale,
- Soit le produit est renvoyé à une étape précédente de sa fabrication.
- Soit le produit est entièrement refabriquée. Ce dernier point est en fait un cas particulier du 3ème où le produit est renvoyé au tout premier poste de travail.

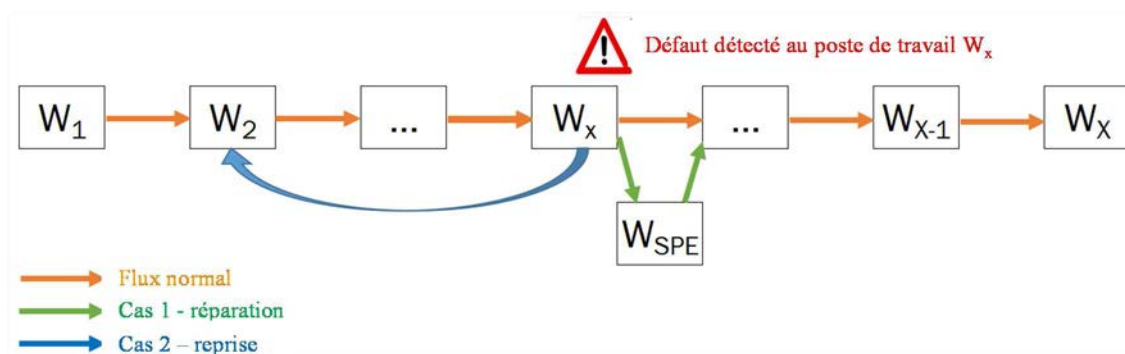


Figure 21 - Réactions possibles suite à la détection d'un défaut

Si nous nous plaçons dans le cas de produits coûteux, deux réactions seulement sont possibles pour réparer un produit défectueux détecté au $x^{\text{ème}}$ poste sur les X existants. Le premier cas est communément appelé réparation :

- 2 nouveaux flux créés : $W_x \rightarrow W_{spe}$ et $W_{spe} \rightarrow W_{x+1}$
- Nombre maximal de flux additionnels : $M_{apf} = 2 \cdot N_{spe} \cdot X$

avec N_{spe} le nombre de postes de travail dédiés aux réparations.

Le deuxième cas est communément appelé reprise :

- Seulement un nouveau flux créé : $W_x \rightarrow W_{x-y}$
- Nombre maximal de flux additionnels : $M_{afp} = \sum_{i=2}^X (i - 1) = \frac{X}{2} (X - 1)$

Cette reprise est finalement très similaire aux boucles de production définies précédemment puisqu'il s'agit d'une répétition d'opérations déjà effectuées dans le but d'aboutir à un produit acceptable qualitativement. Pour la suite de nos travaux, nous définissons donc comme « reprise » toute perturbation de la gamme de fabrication linéaire d'un produit découlant soit d'un problème qualité, soit d'une boucle de production. Conformément à cette définition, une reprise (qu'elle soit prévue par la gamme de fabrication, ou non-prévue suite à la détection d'un défaut) engendre donc plusieurs types de complications au niveau du pilotage de flux :

- Augmentation du nombre de flux (principalement création de flux en « contre sens »).
- Augmentation du volume des flux prévus.
- Perturbation des ordres des files d'attente ou le FIFO n'est plus la règle.
- Augmentation du temps de production, des temps de manutention, des temps d'attente dans les files d'attente, et donc des risques de chocs, casses ou pertes.
- Augmentation des retards de production et/ou des retards de livraison.

(Colledani & Tolio, 2006) ont montré que les reprises et les boucles de production influencent de manière conséquente le pilotage de la production. Les perturbations dues aux reprises présentent de surcroît la particularité d'être fluctuantes, souvent imprévisibles et particulièrement importantes lorsque les reprises ont lieu en aval du processus de production (Love, Li, & Mandal, 1999).

Cependant, même en diminuant et en stabilisant les taux de reprises, la perturbation résultante sur les flux de produits reste conséquente et son caractère cumulatif (que l'on pourrait assimiler au bullwhip effect ou à l'effet Forrester) rend la planification très difficile, voire impossible. En effet, les produits défectueux ont tendance à désorganiser les lots de fabrication (division, regroupage... problématique du Burbidge Effect) à modifier les gammes (étapes de réparation supplémentaires qui surchargent les plannings de travail... problématique du Houlihan Effect) ou à changer les priorités d'ordonnancement (Disney & Towill, 2003). La non-qualité semble donc aussi être une des causes primaires

du bullwhip effect. Dans une analyse statistique par rapport à un problème de poste de travail (Skidmann, Schweitzer, & Nof, 1985) soulignent le phénomène oscillatoire qui est également mentionné dans d'autres travaux dont l'objectif est d'évaluer la performance d'un atelier de type multi-produits (Skidmann, Schweitzer, & Nof, 1985).

Dans tous ces travaux, les reprises sont prises en compte de manière statique ou sous forme de ratios constants mais pas de manière dynamique. Or, pour rester en mode réactif par rapport à l'information portée par le produit, cette information directement corrélée à l'état de saturation de l'atelier doit impérativement pouvoir être prise en compte dynamiquement et avec sa valeur à « l'instant t ». Il faut donc « amener de l'intelligence » et surtout de l'observation dans le processus de fabrication, c'est le rôle des indicateurs de performance portant sur les flux.

III.2.2 Les indicateurs de contrôle de flux de la littérature

Afin d'évaluer la perturbation des flux engendrée par la non-qualité, il convient de lister et de définir correctement les différents indicateurs possibles pour évaluer soit la non-qualité, soit la perturbation des flux directement. Comme tout indicateur de performance, ces indicateurs correspondent à des triplets (objectif, mesure, variable) qui expriment l'efficacité d'un système par rapport à une norme préétablie (AFGI, 1992). Ils peuvent être globaux (à l'échelle de l'entreprise) ou locaux (à l'échelle du poste goulot par exemple). Chacun peut avoir son propre comportement en fonction de l'augmentation progressive du taux de reprises et offre donc une seule vue du problème. Dans un premier temps, les indicateurs que nous avons analysés sont classés par famille de comportement dont chacune offre une vue différente du problème. Il convient aussi de vérifier la « disponibilité » de ces indicateurs. En effet, certains sont facilement disponibles en temps réel tandis que d'autres sont parfois mesurables uniquement *a posteriori*, c'est-à-dire une fois que la fabrication de tous les lots est terminée. Dans ce cas, leur utilisation pour des fonctions de prédiction n'est pas possible. Le Tableau 18 présente quelques indicateurs intéressants relatifs à la non-qualité tandis que le Tableau 19 présente d'autres indicateurs intéressants relatifs à la perturbation des flux.

Tableau 18 - Indicateurs de non-qualité

Indicateur	Vue et disponibilité
N_{config} ou N_{defaut} (équivalent dans le cas de lots unitaires) : Pourcentage de pièces qui ont présenté au moins un défaut	Capacité à réaliser le produit juste du premier coup. Disponible sur une période donnée si les défauts sont relevés dès leur création à chaque poste susceptible de générer des défauts.
$N_{travail}$: Pourcentage de pièces effectivement travaillées sur la machine	Densité du flux ou charge de travail réelle sur l'atelier. Facilement disponible par poste pour $N_{travail}$

N_{defaut} permet de se faire une idée de la capacité de l'atelier à réaliser une pièce correctement du premier coup. Couplé à $N_{travail}$ qui représente le volume de travail total sur l'atelier, on peut se faire une idée du nombre de pièces qui passent en réparation alors qu'elles y étaient déjà passées. En revanche, il est impossible d'associer cette information à la pièce avec ces deux seuls indicateurs. Il nous manque donc une vue intéressante qui correspond à la capacité à réparer correctement du premier coup. Nous proposons un nouvel indicateur, le Ratio d'Opération présenté dans la partie III.3.4, pour pouvoir utiliser cette vue.

Tableau 19 - Indicateurs de flux

Indicateur	Vue et disponibilité
N_{EC} et N_{stock} : Nb de pièces en en-cours global et niveau de stock en amont d'un poste de charge	Densité du flux ou charge de travail réelle sur l'atelier N_{EC} disponible en temps réel en tant que valeur offerte traditionnellement par les ERP. N_{stock} disponible de la même manière ou mesurable directement en amont des postes de travail critiques.
C_{max} et T_A : Temps pour terminer la fabrication de toutes les pièces à produire et temps moyen d'attente des pièces	Vitesse d'écoulement du flux C_{max} disponible a posteriori, une fois que toute la production a été réalisée ou estimable mais avec beaucoup de dispersion. T_A disponible plus facilement sur le terrain (ex : mise en place de minuteurs).
R_{max} : Retard maxi	Ponctualité Disponible a posteriori puisque le lot doit être terminé pour pouvoir évaluer son retard par rapport à une date d'expédition prévue.
N_{retard} : Pourcentage de retards.	Satisfaction client Disponible a posteriori pour la même raison que R_{max} .

III.2.3 Les méthodes pour le basculement des règles de pilotage

La volonté de basculer d'une règle de pilotage à l'autre afin d'optimiser l'ordonnancement dans des ateliers dont l'état change au cours du temps est un problème que l'on retrouve dans la littérature.

Mebarki (Mebarki, 1995) assume que, bien que l'usage des règles de priorité soit très répandu pour l'ordonnancement dynamique, le principal inconvénient vient du fait qu'aucune règle n'est globalement meilleure que les autres car son efficacité dépend de différents facteurs tels que les critères de performance évalués, les conditions opératoires et l'état de l'atelier. Il propose dans sa thèse un système basé sur la simulation qui établit un diagnostic de l'atelier à partir de symptômes tels que des goulots d'étranglement, le nombre de retards... Même si les résultats montrent de fortes améliorations, notamment sur le retard conditionnel, la mise en place sur le terrain pour impliquer les opérateurs est particulièrement délicate.

Hoc (Hoc, Mebarki, & Cegarra, 2004) insiste plus particulièrement sur cet aspect terrain et les facteurs humains dans l'ordonnancement. La méthode la plus formelle consiste à fournir une solution optimale ou générée grâce à des règles de priorité, mais l'opérateur humain ne sera plus capable d'améliorer cette solution faute de lisibilité. On parle d'un phénomène de contentement que l'on cherche à éviter en fournissant à l'opérateur plusieurs solutions parmi lesquelles choisir. Il aborde l'ordonnancement dynamique et l'ordonnancement réactif avec une vision « homme-machine » ou soutien à la coopération « homme-homme » en citant des travaux qui mettent en évidence que les algorithmes intègrent difficilement tous les paramètres sur lesquels les « ordonnanceurs humains » jouent habituellement tels que des changements de priorité, des changements de taille de lots, des divisions entre plusieurs machines, des recouvrements d'opérations, des interruptions au profit de tâches plus urgentes, des renégociations de dates avec les clients, des utilisations non standard de machines, etc... L'enjeu est de concevoir des interfaces permettant de mettre à disposition l'information nécessaire à l'opérateur pour lui permettre d'utiliser son intelligence à des fins de réordonnancement réactif sur le terrain.

Dans notre contexte où l'état de saturation de l'atelier évolue de manière conséquente à cause de la complexité des gammes perturbées par les non-qualités, il est clair que la règle de pilotage à utiliser doit, elle-aussi, évoluer en fonction des situations.

III.2.4 Problématique spécifique

Un système d'indicateurs adapté est donc nécessaire pour suivre avec pertinence l'évolution des flux physiques, mais dans notre contexte les indicateurs classiquement utilisés ne sont pas suffisants. Pour que les opérateurs (comme nous l'avons montré au – paragraphe I.2.4-) puissent gérer ces indicateurs dans le temps, un système adapté des principes du management visuel sur le terrain et offrant une meilleure compréhension du problème devra être proposé.

III.3 PROPOSITION

Notre apport réside dans la mise en évidence des seuils de saturation dans la méthodologie de détermination des zones de comportement et dans la technique de détermination de la meilleure règle de pilotage via une cartographie dynamique visuelle.

Il est facile de déterminer, par simple application de la théorie des files d'attente, la saturation d'un atelier par l'augmentation des réparations dues à la non-qualité à cause de l'effet Forester. Le phénomène peut être comparé au mécanisme de création d'un bouchon autoroutier où les voitures finissent par s'immobiliser alors que l'élément déclencheur n'était qu'un freinage de la part d'un seul automobiliste. Pour comprendre ce phénomène complexe, il est théoriquement possible d'analyser des données réelles issues du terrain qui le représentent dans son intégralité. Cependant, cette analyse n'est pas toujours représentative pour deux raisons :

- Les données doivent représenter le phénomène dans son intégralité. Elles doivent donc contenir les valeurs de tous les facteurs influents dans tous les états extrêmes (état maîtrisé et état saturé) et intermédiaires ce qui n'est pas toujours facile à obtenir.
- Certains facteurs influents peuvent avoir évolué au cours de l'acquisition des données (changement de règle de priorité par exemple) et peuvent engendrer une mauvaise interprétation des données si l'on n'en tient pas compte dans l'analyse.

Un autre inconvénient majeur est qu'elles ne permettent pas de tester d'autres « solutions » que celles utilisées lors de l'acquisition des données.

La simulation est un outil courant pour les problèmes complexes de production, le but étant de trouver les conditions optimales en fonction des contraintes du terrain. Il s'agit généralement de problèmes d'optimisation qui doivent être résolus de manière itérative. Le principal inconvénient n'est autre que l'étape de fabrication du modèle qui requiert une

connaissance experte du phénomène à modéliser. Dans le domaine de la chimie par exemple, les modèles sont partagés par plusieurs chercheurs de leur communauté et chacun cherche à améliorer le modèle d'origine pour le rendre de plus en plus fidèle au comportement du phénomène ou composant réel. Concernant des travaux comme les nôtres, la qualité de notre modèle est soumise principalement à notre jugement. L'explication détaillée du modèle est donc un gage important de validation possible. Cependant les avantages de l'usage de la simulation pour ce type de problème sont nombreux. La simulation permet de prendre en compte une multitude de variabilités par l'intermédiaire de lois de distribution telles que les Gaussiennes ou les lois de Poisson. Elle permet aussi, dans une certaine limite, l'interpolation et l'extrapolation et autorise donc la recherche de solutions en dehors des conditions connues traditionnellement. On peut ainsi envisager l'application d'une nouvelle règle de pilotage qui n'a jamais été testée en réel. Afin de pouvoir mettre en évidence ces comportements faciles à imaginer mais difficiles à mesurer, nous proposons donc d'utiliser la simulation.

Parmi les deux grandes familles de simulation (continue ou discrète), nous utiliserons les simulations à événements discrets où chaque événement arrive à un instant donné et modifie l'état du système. Ce type de simulation est particulièrement adapté au suivi de production où chaque événement est déterminé par le début ou la fin de travail d'un lot sur un poste de travail.

III.3.1 Modèle de simulation

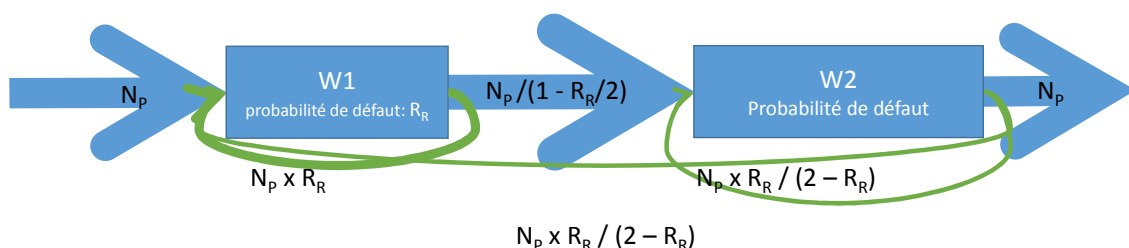


Figure 22 - Modèle réduit utilisé pour les simulations

Le modèle présenté sur la Figure 22 a été implémenté dans le logiciel ARENA qui permet d'analyser l'impact d'une augmentation du nombre de reprises sur un procédé de fabrication. L'apparente simplicité de ce modèle est justifiée par les travaux sur la réduction de modèle qui montrent qu'on parvient à réduire le problème par l'usage d'un modèle très

simple (Thomas & Charpentier, 2005), et par d'autres travaux du même ordre qui ont conduit eux-aussi à des systèmes à une ou deux machines (Rabiee, Zandieh, & Jafarian, 2012). Le modèle considère un flux de pièces à l'unité et chaque produit est transféré directement à la station de travail sans temps supplémentaire correspondant à l'attente de la fin de lot avant transfert. Cela permet ainsi d'éviter le problème de divisions des lots entre les pièces à réparer et les pièces de bonne qualité (Zargar, 1995).

Nous avons testé en termes de scénarios, une des règles de priorité parmi les plus couramment utilisées : EDD (au plus tôt à la date d'échéance). La date d'échéance par produit requise pour la mise en œuvre de la règle de priorité EDD a été déterminée suivant l'hypothèse du Total Work Content en prenant une marge proportionnelle au temps de réalisation du produit (Cheng, 1987). Les probabilités d'occurrence des reprises ont été définies arbitrairement. Dans notre modèle de simulation, chaque machine a la même probabilité de générer des défauts (12) et chaque pièce défectueuse peut être réparée sur la même machine ou une machine précédente avec une probabilité équivalente (12).

$$\begin{cases} T_{\text{reprises}}(M_1 \rightarrow M_1) = T_{\text{reprises}}(M_2 \rightarrow M_1) + T_{\text{reprises}}(M_2 \rightarrow M_1) \\ T_{\text{reprises}}(M_2 \rightarrow M_1) = T_{\text{reprises}}(M_2 \rightarrow M_2) \end{cases} \quad (12)$$

Ces équations peuvent être simplifiées en ne considérant qu'une seule probabilité de reprises T_r sur les machines M_1 et M_2 . Le passage sur la machine 1 est toujours suivi d'un passage sur la machine 2, même dans le cas d'une reprise car on considère que si la pièce nécessite à nouveau la tâche 1, la(les) tâche(s) suivante(s) a(ont) été(s) compromise(s). Les deux machines sont paramétrées avec la même cadence de production. Le 100% en entrée de la machine 1 matérialise un nombre de pièces total qui va entrer sur la chaîne de production $M_1 - M_2$. Cette quantité de pièces va augmenter entre les 2 machines à cause des reprises alors que la quantité de sortie sera à nouveau égale à la quantité d'entrée vu qu'il n'y a ni création, ni perte de pièces sur la chaîne.

III.3.2 Intégration des variabilités dans le modèle

La variabilité est une notion très importante car elle est une des composantes principales des phénomènes réels menant à leur incompréhension. L'avantage d'une simulation à événements discrets est qu'elle permet d'implémenter cette variabilité sur des systèmes virtuels et ainsi de reproduire ces phénomènes.

La variabilité implémentée dans le modèle pour les temps de génération des commandes, les temps de traitement sur les postes de travail et la probabilité d'occurrence de défauts impliquent des perturbations (ex. Bullwhip Effect). Cette variabilité a été prise selon différentes lois de distribution explicitée dans le Tableau 20, la plus importante étant la loi normale attribuée au temps de travail sur les machines qui est connue en théorie des probabilités et en statistique comme l'une des lois les plus adaptées pour modéliser des phénomènes issus de plusieurs événements aléatoires. D'autres lois auraient pu être utilisées pour parvenir à des résultats équivalents.

Tableau 20 - Type de variabilité en fonction du paramètre

Paramètres variable	Variabilité
Création des entités dans le modèle	Loi exponentielle : EXPO(10)
Nb de pièces dans le lot (NbPcs)	Loi uniforme : UNIF(5, 70)
Temps de travail sur la machine 1	Loi normale : NORM(NbPcs, 2)/5
Taux bouclage après machine 1	Fonction du temps de la simulation
Temps de travail sur la machine 2	Loi normale : NORM(NbPcs, 2)/5
Taux bouclage après machine 2	Fonction du temps de la simulation

III.3.3 Validation du modèle

La validation de la structure du modèle est évidente dès lors que l'on réduit l'atelier (en particulier celui de Acta Mobilier) à ses postes critiques et aux flux bouclés de produits qui le traversent. Dans notre contexte d'utilisation, le premier poste correspond à l'atelier de laquage tandis que le second correspond à l'atelier de polissage qui évalue la qualité de la pièce pour un renvoi au premier poste ou au second.

En confrontant le comportement de notre modèle avec la réalité, nous constatons bien l'apparition du phénomène de saturation lorsque nous faisons augmenter le taux de réparation. Ce phénomène qui paralyse l'atelier dans la réalité est évoqué en partie III.3.5. De même nous parvenons à mettre en évidence le fait qu'il faille changer de règle de priorisation pour le travail des lots en fonction de cet état de saturation, un fait qui avait été constaté et affirmé sur le terrain par des standards pas forcément justifiés. Ce phénomène est présenté en partie III.3.6.

L'objectif de notre modèle n'étant pas tant d'associer des valeurs d'indicateurs à ces phénomènes que de les mettre en évidence pour pouvoir les transformer en informations claires pour les utilisateurs, le fait que son comportement reflète celui du système réel sur ces deux sujets au moins est une validation suffisante pour la suite de nos travaux.

III.3.4 Un nouvel indicateur de non-qualité : le ratio d'opérations

Dans notre cas particulier où le nombre de reprises peut être conséquent même pour une même pièce, nous proposons un nouvel indicateur que nous appellerons le Ratio d'Opérations (RO). Nous le définissons comme correspondant au nombre d'opérations réalisées en réalité par rapport à ce qui était prévu dans la gamme. Comme il y aura donc une valeur par pièce/lot, nous proposons, pour une utilisation à une échelle globale, de prendre la moyenne (RO_{moy}) ou l'écart-type (RO_{ET}). Ce nouvel indicateur témoigne de la capacité à réparer correctement le produit défectueux et donc, donne une indication sur la prévisibilité du respect des délais. En effet, si le RO_{moy} est proche de 1, il y a peu de réparations donc peu de non-qualité. S'il commence à s'en éloigner, les produits sont beaucoup réparés mais de manière efficace. S'il s'approche de 2 et au-delà, on peut supposer qu'on commence à réparer des pièces qui l'ont déjà été et que cette perte de temps aura forcément une conséquence néfaste sur la tenue des délais.

Cet indicateur est disponible *a posteriori* par comparaison avec la gamme de fabrication prévue, uniquement si la traçabilité des produits est correctement sauvegardée.

III.3.5 Mise en évidence de la saturation rapide de l'atelier

Le modèle présenté précédemment peut être exploité de manière à mettre en évidence la problématique de saturation évoquée précédemment. Ainsi, en augmentant progressivement le taux de non-qualité R_R et en appliquant la règle EDD pour déterminer la priorité des lots dans les files d'attente, on obtient une évolution du nombre de retard qui suit la courbe présentée sur la Figure 23.

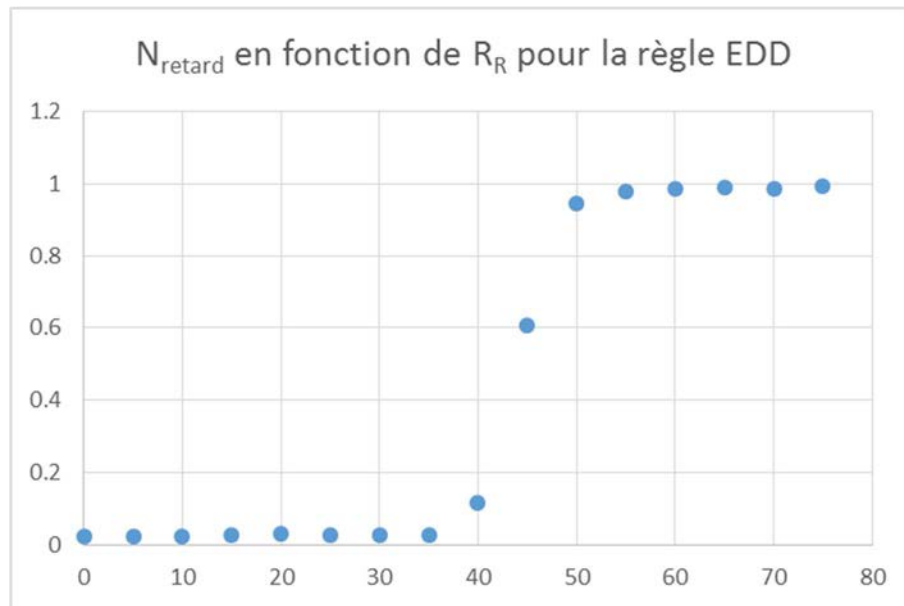


Figure 23 - Nombre de retards en fonction du taux de non-qualité pour la règle de pilotage EDD.

On peut voir que le nombre de retards reste relativement faible jusqu'aux alentours de 35% de réparations. En revanche, dès cette borne atteinte, le nombre de retards augmente de manière exponentielle pour atteindre quasiment 100% avant le seuil de 50% de réparations. Il existe donc bien un phénomène de saturation rapide de l'atelier dans un contexte de perturbations par les boucles de réparations.

Cette vitesse de saturation dépend de la structure du système de production, de la charge initiale du système mais aussi de la règle de pilotage utilisée. Par exemple, la règle FIFO peut mettre en évidence le même phénomène de saturation mais de manière moins abrupte que la règle EDD, comme présenté sur la Figure 24.

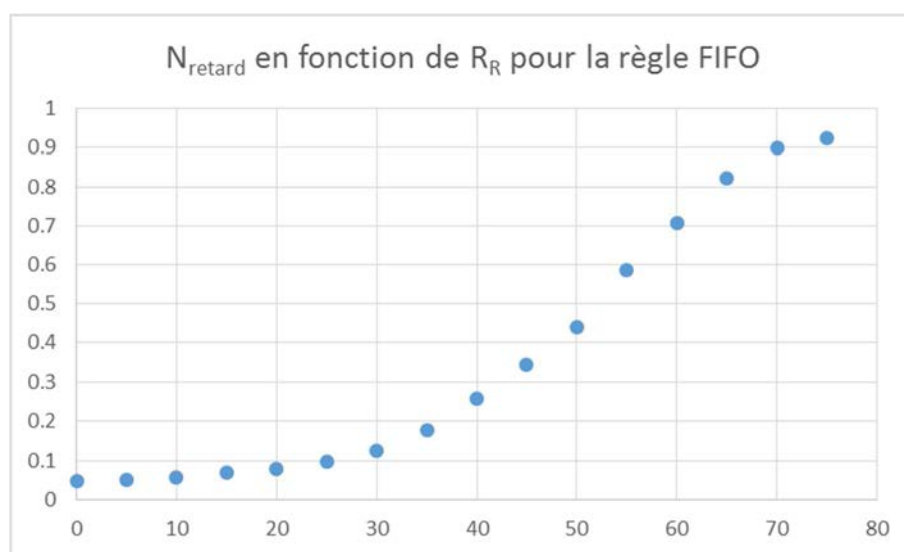


Figure 24 - Nombre de retards en fonction du taux de non-qualité pour la règle de pilotage FIFO.

Avec la règle FIFO, on voit que le nombre de retards augmente pour un taux de réparations plus faible mais moins rapidement qu'avec la règle EDD pour finalement atteindre la saturation plus tard que la règle EDD. En comparant ces deux graphiques, il peut être possible de déterminer des plages de valeurs sur lesquelles une règle de pilotage sera plus intéressante qu'une autre.

III.3.6 Mise en évidence de la supériorité d'une règle par rapport à une autre dans un contexte donné

Les règles de pilotage proposées dans la littérature (Lopez, 1991) trouvent chacune leur intérêt dans des conditions de fabrication spécifiques, ce qui justifie l'idée de basculer de l'une à l'autre de manière dynamique en fonction de l'évolution de ces conditions de fabrication. Afin de mettre en évidence le concept de seuil de basculement, nous avons réalisé des simulations à partir du modèle présenté en III.3.1 pour tester l'impact d'une règle de pilotage FIFO ou EDD lors d'une croissance du taux de non-qualité allant de 0 à 65%. Nous avons choisi les règles à tester en s'appuyant sur les réflexions menées en I.3.2. L'objectif n'est pas d'évaluer mathématiquement les seuils, ni de justifier la supériorité d'une règle en fonction de conditions de fabrication mais simplement de mettre en évidence l'existence d'un seuil de basculement au-delà duquel il est judicieux de changer de règle de pilotage. Dans le cadre d'une application industrielle, il faut bien évidemment régler le modèle de manière à représenter fidèlement la réalité et choisir des règles pertinentes, à la fois du point de vue scientifique par une justification dans la littérature, et du point de vue industriel en vérifiant la faisabilité de l'utilisation de ces règles. Le cas présenté ici est donc purement illustratif.

Les résultats évalués sur le nombre de retards sont présentés sur la Figure 25. Dès l'apparition de quelques défauts, la règle FIFO entraîne des retards tandis que la règle EDD permet d'atteindre plus de 20% de défauts sans engendrer le moindre retard. Il est donc plus intéressant d'utiliser la règle EDD pour de faibles taux de non-qualité. En revanche, lorsqu'on atteint environ 30% de réparations, la tendance s'inverse et c'est alors la règle FIFO qui engendre moins de retard que la règle EDD.

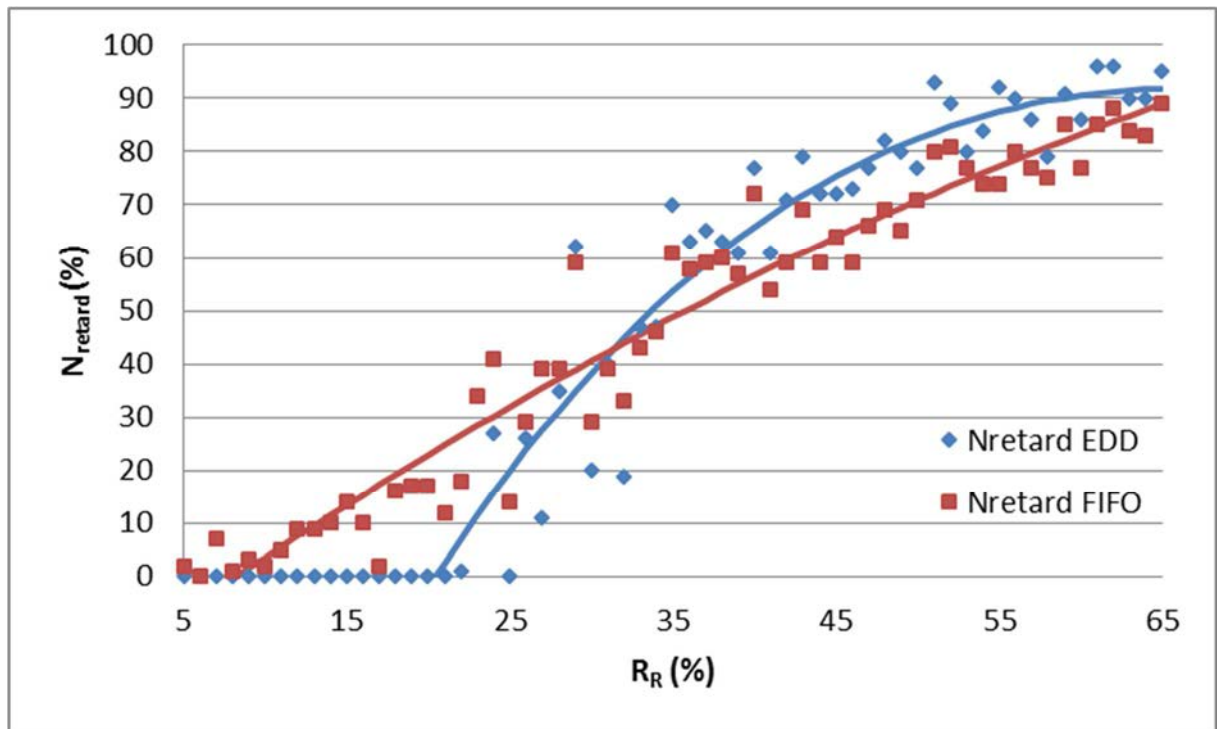


Figure 25 - Changement de règle au cours d'une évolution du taux de non-qualité

Il existe donc un seuil de basculement entre EDD et FIFO qui est ici fonction des indicateurs N_{retard} et R_R et que l'on pourrait noter $S_{FIFO \rightarrow EDD}(N_{retard}, R_R) = 32\%$. Il serait possible de déterminer de la même manière des éventuels seuils de basculement pour chaque couple de règles et pour chaque couple d'indicateurs sous la forme $S_{Règle\ 1 \rightarrow Règle\ 2}(Indicateur_1, Indicateur_2)$.

III.3.7 La détermination des seuils d'absorption et de saturation

Pour avoir un aperçu plus global et donc plus réaliste du problème de perturbation des flux par la non-qualité, suivre un seul indicateur en fonction du taux de réparation n'est pas suffisant. Il devient important de travailler sur une combinaison de ces indicateurs.

La combinaison de deux, ou plus, des indicateurs présentés précédemment en fonction du taux de réparations permet de combiner plusieurs vues différentes afin de porter un jugement plus complet sur l'état d'un atelier de production. Ce sont ces combinaisons qui permettent de matérialiser les seuils nécessaires à la compréhension des phénomènes de perturbations des flux et à la création de cartographie. L'indicateur N_{defaut} qui évalue le nombre de pièces ayant présenté au moins un défaut, est couramment utilisé dans les entreprises car il résulte directement des méthodes de contrôle de la qualité. Il est donc directement lié avec la capacité de l'atelier à réaliser le produit juste du premier coup. Seul,

il ne permet pas de comprendre et prévoir l'impact sur la perturbation des flux alors que des combinaisons réfléchies peuvent y parvenir. La Figure 26 présente une double vue du même atelier de production géré avec la même règle de pilotage. 100 simulations ont été lancées avec chacune un taux de reprises différent.

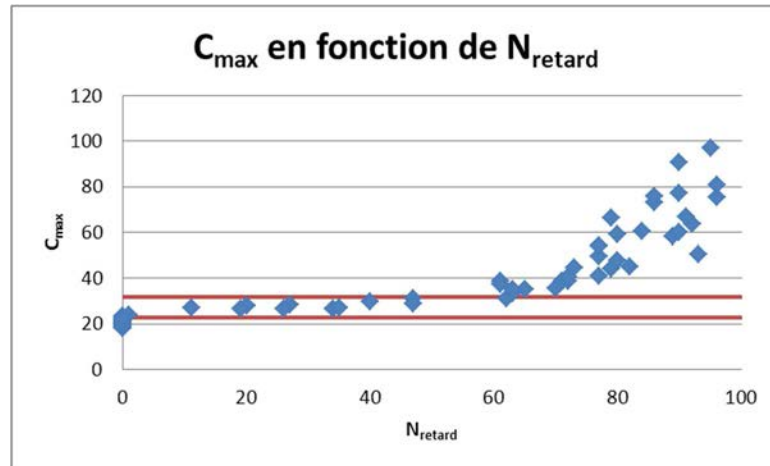


Figure 26 - Combinaison des indicateurs C_{max} et N_{retard}

Chaque point de la Figure 26 correspond à une simulation et donc à un taux de reprises. Pour chaque simulation, nous plaçons le point sur le graphique en fonction de la valeur atteinte par les deux indicateurs que nous souhaitons mettre en vis-à-vis, nous offrant ainsi une vue particulière de l'évolution du comportement de ces deux indicateurs en fonction du taux de reprises. La position des indicateurs entre abscisse et ordonnée n'a pas d'importance dans notre cas car nous regardons leurs évolutions simultanées et non l'un en fonction de l'autre. Ces combinaisons permettent de mettre en évidence un ou plusieurs seuils déterminant des zones à différents niveaux de perturbation des flux. Nous nous intéressons ici uniquement à l'existence de ces seuils et à la manière de les rendre visibles en combinant des indicateurs judicieusement choisis. **Les valeurs des seuils** présentés dans la figure **n'ont donc pas d'intérêt** particulier car ils dépendent essentiellement des paramètres de simulation.

Sur la Figure 26, on peut diviser le graphique en trois zones de comportement différentes. La première zone (sous le premier seuil d'absorption) correspond à la zone d'absorption. Il s'agit des points situés sur l'axe des ordonnées (lorsque N_{retard} vaut 0). Ceci signifie que la durée d'écoulement de production de toutes les pièces peut s'allonger quelques peu sans générer de retard. Cette capacité d'absorption est par exemple possible pour un atelier travaillant généralement en sous-capacité. La deuxième zone correspond à l'augmentation rapide du nombre de retards (jusqu'à 50% environ sur le graphique) alors

que le C_{max} n'augmente pas de manière conséquente. On peut expliquer ce comportement par le fait qu'une augmentation du C_{max} (une grosse réparation par exemple) mettrait en retard toutes les pièces qui devraient être produites après. Ce comportement est renforcé par l'utilisation de la règle de priorité EDD entre les pièces qui ne cherche pas à « sauver » un maximum de pièces du retard en « en condamnant » quelques-unes, mais qui fait toujours passer la plus urgente en priorité, même si cette dernière est déjà en retard. La dernière zone correspond à une augmentation très marquée du C_{max} avec une variabilité croissante. Cette zone témoigne bien du fait que lorsque le taux de reprises est trop important il devient quasiment impossible de prévoir correctement le temps nécessaire à la fabrication d'une quantité de lots donnés. Cette augmentation correspond à l'incapacité du système à absorber les reprises et se traduit sur le terrain par une augmentation du nombre de produits dans les files d'attente devant les postes de travail.

Le graphique présenté ainsi que les autres détaillés dans (Noyel, Thomas, Thomas, Charpentier, & Brault, 2014) ont permis de mettre en relief l'existence de différents seuils de fonctionnement liés à la variation du taux de reprises. Ces seuils semblent délimiter trois zones dont les caractéristiques sont les suivantes :

- Zone 1 (taux de reprises faible) : Absorption des reprises par le système, pas d'impact (ou très peu) sur la majorité des indicateurs considérés.
- Zone 2 (taux de reprises proche du seuil) : Zone d'alerte ; dégradation très rapide de la plupart des indicateurs. Le système est en cours de saturation.
- Zone 3 (taux de reprises supérieur au seuil) : Le système est saturé, les files d'attente accumulent de plus en plus de produits, les indicateurs présentent de plus en plus de variabilité.

Il est évident que la position de ces seuils et donc, la grandeur de ces zones est fonction de différents paramètres définis de manière figée dans notre modèle de simulation. Parmi ces paramètres, il est possible de citer :

- La marge totale qui correspond au temps disponible pour la réalisation du produit, c'est-à-dire le temps entre le lancement en fabrication et la date souhaitée d'expédition ou la marge restante qui correspond à la différence entre le temps avant livraison et le temps de production prévu par la gamme.
- La capacité des machines considérées comme goulots et qui sont donc les principales génératrices de files d'attente.
- La charge globale de l'atelier qui a un impact direct sur les files d'attente.

Lors d'une application industrielle, deux possibilités s'offrent à nous :

- Déterminer les formules permettant de calculer systématiquement la valeur de ces seuils. Cette solution est particulièrement compliquée car le phénomène de saturation dépend de plusieurs indicateurs et certainement de leurs interactions.
- Utiliser l'analyse de données réelles.

A chacune de ces zones correspondra des règles de pilotage différentes afin de maintenir le système « sous contrôle ».

III.3.8 Une cartographie visuelle

Comme présenté dans la partie III.2.3, les systèmes de changements automatiques de règles de pilotage ont déjà été abordés dans la littérature. Cependant, afin de laisser à l'opérateur humain le choix du prochain lot sur lequel travailler pour bénéficier de son intelligence et de sa réactivité, nous proposons ici une approche graphique avec une cartographie visuelle permettant de situer l'état d'encombrement de l'atelier et donc de spécifier aux utilisateurs le comportement à adopter en matière de priorité des lots comme présenté sur la Figure 27.

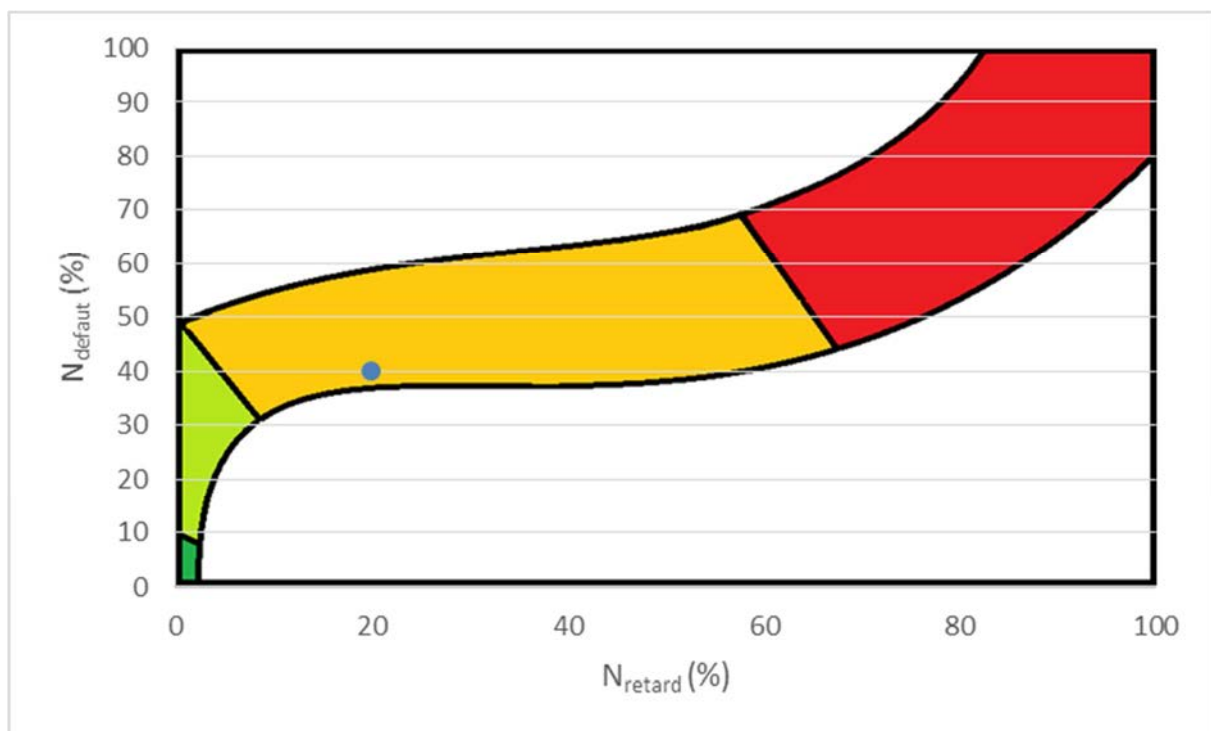


Figure 27 - Cartographie des états de saturation de l'atelier construite à partir des seuils.

Cette cartographie est donc volontairement présentée en deux dimensions afin qu'elle soit facilement compréhensible par tous, chacun des axes représentant alors de manière claire un indicateur clef du taux de non-qualité et/ou de la perturbation des flux. Le point bleu représente l'état de saturation actuel de l'atelier. Suivant la zone dans laquelle il se

situé (ici zone orange), un type de règle ou certaines consignes peuvent être transmises aux utilisateurs.

Les avantages de cette méthode sont donc clairement tournés vers l'application facile et rapide sur le terrain ainsi que vers la communication et la responsabilisation des employés. Même si, contrairement aux algorithmes, la règle de pilotage la mieux adaptée dans l'absolu ne sera pas toujours mise en évidence, cette cartographie présente d'autres avantages présentés dans le Tableau 21.

Tableau 21 - Comparaison des concepts de cartographie visuelle et d'algorithme traditionnel.

	Cartographie visuelle	Algorithme
Complexité	Obligation de simplifier pour ne représenter que 2 ou 3 règles de pilotage différentes correspondant à 2 ou 3 états d'encombrement de l'atelier uniquement.	Possibilité de complexifier les règles de pilotage (en les hybridant) de manière à s'adapter au mieux à la situation.
Visibilité / Communication	Possibilité pour tout le monde d'avoir l'information importante sur l'état de l'atelier en un coup d'œil.	Souvent une « boîte noire » qui ne sera pas à la portée de n'importe quel utilisateur.
Standardisation / Globalisation	Nécessité d'homogénéisation sur la globalité de l'unité de production.	Possibilité de faire du « cas par cas ».

III.3.8.1 Construction à partir des seuils

La création de ce type de cartographie peut se faire « manuellement » en positionnant chaque seuil à partir de comparaisons de comportements de règles de priorisation dans des états de saturation donnés. Cette tâche chronophage nécessite une bonne visualisation de ce que l'on souhaite créer et de l'intuition que l'expert peut avoir quant aux déterminations des limites entre classes.

III.3.8.2 Construction par clustering

Une approche plus performante, appelée clustering, consiste à déterminer automatiquement, et par apprentissage, les regroupements de données en différentes classes ou zones de l'espace et de délimiter les limites de chacune de ces zones. Nous avons présenté succinctement au paragraphe II.2.1 cette approche. De nombreux algorithmes de

clustering ont été proposés (Jain, Murty, & Flynn, 1999) dont les principaux sont répertoriés et comparés dans le Tableau 22.

Tableau 22 - Quelques algorithmes de clustering et leurs particularités.

Algorithmes	Référence	Avantages	Faiblesses
K-means / K-medoids	(Singh & Chauhan, 2011)	Très connu Trace les contours des clusters Déjà implémenter dans les programmes de calculs Fournit de bons résultats	Le nombre de clusters doit être défini à l'avance Changer le nombre de clusters peut arbitrairement changer l'appartenance des éléments à classer Risque de tomber dans un minimum local Dépendant de l'initialisation et de la distance choisie
Agglomération hiérarchique	(Day & Edelsbrunner, 1984)	Propriété de monotonie	Pas approprié pour de gros volumes de données
Division hiérarchique	(Guénoche, Hansen, & Jaumard, 1991)	Peut donner de meilleurs résultats que l'algorithme d'agglomérations s'il y a peu de clusters	Rarement utilisé
Vector quantization	(Equitz, 1989)		Réservé pour la compression d'images
Maps auto-organisatrice	(Villa-vialaneix, Olteanu, & Cierco-Ayrolles, 2013)	Rapide Maintient les bornes topologiques	Voisinage non modifiable

Notre objectif est d'obtenir simplement et rapidement un résultat fiable. De plus, comme on peut considérer le nombre de clusters comme connu, nous proposons d'utiliser l'algorithme K-means. Le but de cet algorithme est de rassembler les N points observés dans K clusters avec une contrainte : chacun des N points doit être associé au cluster pour lequel la distance à la moyenne est la plus faible.

III.3.9 Les règles de pilotage associées

Afin d'associer les règles de pilotage (ou des consignes) à la cartographie établie précédemment, nous proposons de réaliser un plan d'expériences de Taguchi dont les facteurs correspondront à :

- Indicateur de l'axe des abscisses de la cartographie
- Indicateur de l'axe des ordonnées de la cartographie
- Règles de pilotage à tester

Pour une cartographie en deux dimensions, le plan d'expériences comprendra donc 3 facteurs. Le nombre de niveaux à envisager sera déterminé en fonction du nombre de règles

de pilotage à tester et des matrices de Taguchi disponibles. Par exemple, si nous envisageons au moins un facteur à 3 niveaux, une seule matrice est possible parmi les matrices d'expériences fractionnaires de base, la « L_8 » (4 facteurs à 2 niveaux et 1 facteur à 4 niveaux). Il est en revanche possible d'utiliser les matrices d'expériences fractionnaires issues des matrices de base (lesquelles présentent cependant l'inconvénient de nécessiter beaucoup plus d'expériences).

Il peut aussi être fortement intéressant (voire nécessaire suivant la complexité des cas) d'étudier les interactions. Les matrices d'expériences complètes permettent à la fois de tester ces interactions en autorisant un nombre de niveaux plus élevé que les matrices de base. La matrice L_{16} (4 facteurs à 2 niveaux + 11 interactions) permet notamment d'évaluer les interactions de premier, de deuxième et éventuellement de troisième ordre.

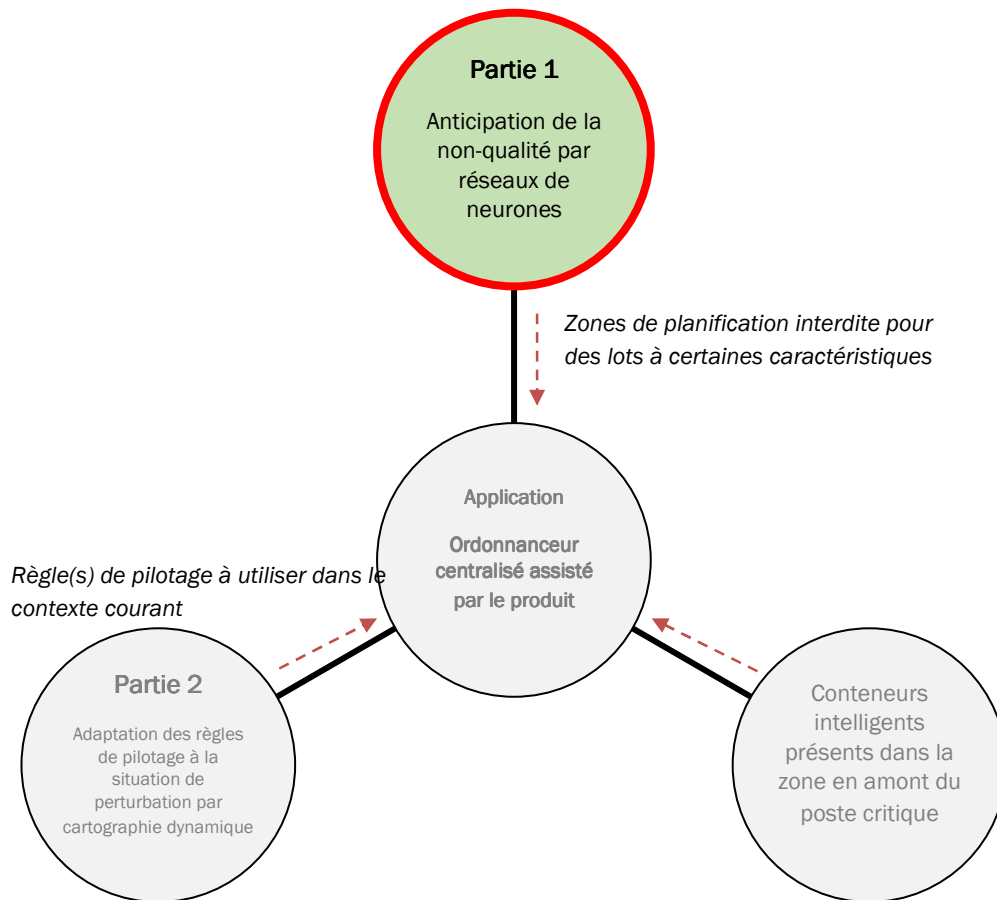
IV APPLICATIONS INDUSTRIELLES

IV.1 INTRODUCTION

Dans cette partie applicative nous abordons l'implémentation dans l'entreprise par 4 chapitres. Les chapitres 1 et 2 reprennent respectivement les propositions des parties II (anticipation de la non-qualité sur des processus de fabrication instables) et III (analyse visuelle de la perturbation des flux). Le chapitre 3 concerne l'apport de « containers intelligents » dans le cadre d'un système contrôlé par le produit sur un poste de travail stratégique (goulot) fortement contraint. Et le chapitre 4 se concentre sur la mise en œuvre du ré-ordonnancement dynamique rapide tout en maintenant le niveau de qualité requis sans devoir faire appel à des approches d'ordonnancement analytiques optimales trop chronophages pour une utilisation en mode réactif.

Ces 4 chapitres sont présentés comme 4 outils qui ont été implémentés dans l'entreprise. Ils peuvent éventuellement être exploités indépendamment les uns des autres suivant les problématiques rencontrées dans les entreprises mais sont interconnectés dans la solution finale adaptée à l'entreprise Acta-Mobilier. Chacun des outils est donc apporté comme une brique d'un système global de supervision et d'aide à la décision.

IV.2 MAITRISE DE LA QUALITE ON-LINE SUR DES PROCESSUS INSTABLES



Dans ce premier chapitre, nous allons présenter le déploiement de la stratégie de contrôle de la qualité décrit dans la partie scientifique II sur un atelier pilote de la société ACTA mobilier.

IV.2.1 Justification du cas d'implémentation

La majorité des défauts qualité est produite lors de l'opération de laquage. Cette étape incontournable de la fabrication de toutes les pièces peut être réalisée soit sur la ligne de laquage robotisée soit dans les cabines de laquage manuel. La méthodologie décrite précédemment a été implantée sur le robot de laquage pour deux raisons. Premièrement, il est considéré comme l'un des goulots structurels de l'entreprise car, destiné aux moyennes

et grandes séries, la grande majorité des pièces y passe plusieurs fois pendant des temps de fabrication très longs à cause du séchage intégré. Et deuxièmement, son caractère plus automatisé permet des analyses mathématiques inenvisageables sur les opérations manuelles et donc trop marquées par la variabilité des cabines de laquage manuel destinées aux petites séries et réparations isolées.

IV.2.2 Description du poste de travail concerné

Le robot de laquage est constitué d'une zone de chargement manuel, d'une zone de laquage supposée fonctionner en continu en fonction d'un programme choisi par l'opérateur, d'un tunnel de séchage et d'une zone de contrôle/déchargement présentée sur la Figure 28.

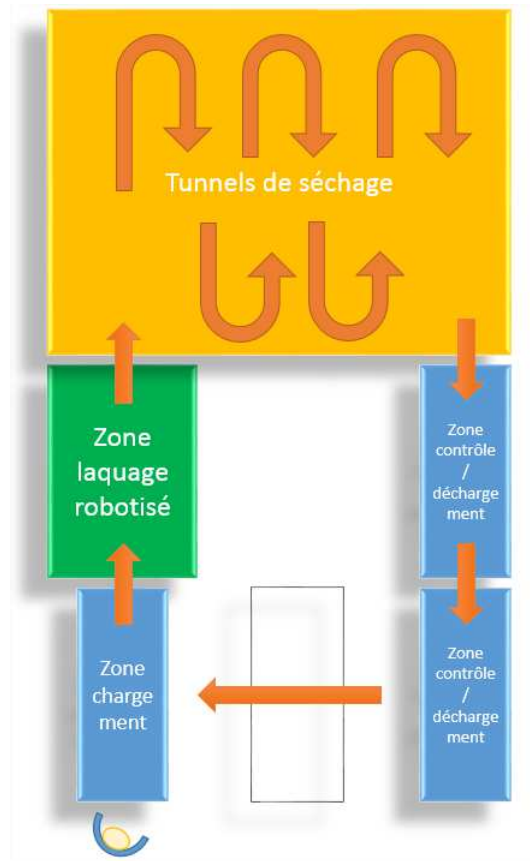


Figure 28 - Schéma structurel du robot de laquage

Les unités mobiles qui transportent les pièces à travers les différentes zones et d'une zone à l'autre sont appelées « tables ». L'opérateur charge les pièces sur une table dans la zone de chargement. Dès que la table est pleine, elle se déplace vers la zone de laquage robotisée. Une fois la table laquée, elle rejoint les différents tunnels de séchage qui réalisent les différentes étapes de séchage (de-solvatation à température ambiante, séchage,

refroidissement, etc...) avant de sortir du côté contrôle et déchargement. C'est la zone de laquage robotisée qui cadence la production car toute la chaîne fonctionne en flux poussé. Lorsqu'une table quitte la zone de laquage robotisée, elle pousse toutes les autres tables présentes dans le tunnel de séchage, faisant ainsi sortir la plus ancienne. A cause de cette particularité, le temps passé par un lot dans la zone de séchage n'est pas constant et dépend non pas des caractéristiques du lot mais de la somme des temps opératoires des types de lots laqués dans la zone de laquage robotisé après lui, en y ajoutant les aléas de production tels que les pannes ou les attentes approvisionnement, par exemple. La chaîne complète permet d'accueillir environ 80 tables en même temps.

Comme le contrôle qualité ne peut s'effectuer qu'en zone de contrôle, il a donc lieu au minimum six heures après le laquage. Si le lot se répartit sur moins de 80 tables, ce qui correspond à la grande majorité des cas avec une moyenne de 3 ou 4 tables, il sera alors complètement réalisé bien avant qu'on puisse débiter le contrôle qualité. Il est donc impératif d'avoir de la visibilité sur le niveau de qualité qui résultera des paramètres de fabrication bien avant de pouvoir contrôler la première pièce réalisée.

IV.2.3 Conception du classificateur

Comme nous l'avons vu au paragraphe II.3.2 et tout particulièrement à la Figure 18, la conception du classificateur passe par un certain nombre d'étapes dont les deux premières concernent la création de la base de données (BdD).

IV.2.3.1 Création de la Base de Données

Afin de créer la BdD, il est nécessaire de déterminer quel est l'objectif et quelles données doivent être collectées.

a. Détermination de l'objectif et des facteurs influents

Dans notre cas, nous avons un objectif multiple : prévoir le risque d'apparition (%) d'une liste de défauts. Chaque défaut correspondra donc à un objectif. Parmi la liste exhaustive de défauts détectés à la sortie du robot de laque, certains résultent directement de facteurs contrôlables tandis que d'autres apparaissent comme purement aléatoires tels que les coups (chocs), par exemple. Quelques-uns de ces défauts et leurs causes éventuelles sont présentées dans le Tableau 23 dont certaines descriptions sont tirées de (<http://www.glasurit.com>, s.d.).

Tableau 23 - Liste des défauts et de leurs causes possibles

Défaut	Description	Causes éventuelles
Coulures	Epaississement de la laque en forme de gouttes ou de vague sur les surfaces verticales.	Dilution trop lente. Viscosité trop basse. Support à peindre trop froid. Epaisseur de couche trop élevée. Temps d'évaporation trop court. Distance trop courte entre le pistolage et l'objet. Buse de pistolage trop grande. Pistolage irrégulier.
Grain sur face / dos / chant	Petites protubérances souvent irrégulières dans le film peinture, dues à des particules étrangères (ex : poussières) qui apparaissent quel que soit la taille, le type ou la forme et avec répartition aléatoire.	Nettoyage insuffisant des surfaces avant peinture. Problème de poussière dans la cabine dus à des filtres encrassés ou à une mauvaise circulation d'air et une surpression non conforme aux préconisations. Aspiration d'air souillé (produit de polissage, poussières fines,...) provenant d'autres secteurs de l'entreprise.
Coup	Marque survenant suite à un choc entre pièce ou un autre objet.	Manutention des pièces.
Eclat	Coup important ayant engendré la perte d'un morceau de laque ou de pièce.	
Manque laque sur face / dos / chant Perce	Epaisseur de laque se réduisant à certains endroits engendrant la réapparition du support direct ou par transparence.	Quantité de laque insuffisante.
Pièces collées	Deux pièces liées en très elles par la laque qui a séché.	Déplacement brusque de deux tables ayant engendré le glissement / rapprochement de deux pièces non sèches.
Microbullage	Défaut de surface sous forme de petites cloques, provoqué par des solvants emprisonnés dans le film de peinture.	Epaisseur de couche trop élevée. Durcisseur ou diluant trop rapides. Temps d'évaporation trop court entre les différentes passes. Temps d'évaporation trop long avant séchage au four ou aux infrarouges (IR). Distance trop faible en cas de séchage IR d'où à des températures trop élevées. Temps d'évaporation trop court entre les différentes couches en cas d'application mouillé sur mouillé.
Laque pas sèche	Laque entre marquée au touché.	Temps de séchage insuffisant lors dû à un passage trop rapide dans le four.
Cratère / cuvette	Cavité circulaire de 0.5 à 3 mm de diamètre où la laque n'a pas adhéree.	Imprégnation des vêtements de travail, gants en caoutchouc. Lubrifiants des pièces mobiles, nettoyage insuffisant des épurateurs d'huile et d'eau, filtres encrassés au plafond et au sol. Utilisation non conforme des produits annexes (additifs), diluants (ou durcisseurs) non appropriés, impureté dans la laque. Produit de ponçage non adaptés, colles des bandes adhésives. Aspiration d'air chargé d'impuretés, étanchéifiassions et isolation du bâtiment.
Refus	Les refus sont des cratères de surface beaucoup plus grande.	
Goutte d'eau	Petite auréoles ou boursoufflures claires et blanchâtres apparaissant sur la surface du film lors de l'évaporation de l'eau mélangée au calcaire et au sel. Les surfaces intérieures sont généralement saines, les bords présentes de légères proéminences.	Couches trop épaisses. Séchage insuffisant. Mauvais dosage de durcisseur.

Si nous prenons par exemple l'objectif d'anticiper la création de coulure, nous devons, *a priori*, nous intéresser à la vitesse de dilution, la valeur de la viscosité, la température du support, l'épaisseur de la couche de laque, le temps d'évaporation, la distance entre le pistolet et la pièce, la taille de la buse de pistolage, et la manière de « pistoler ». Nous devons donc en déduire des facteurs mesurables présentés dans le Tableau 24.

Tableau 24 - Facteurs à suivre en fonction des causes possibles de l'apparition d'une coulure.

Causes possibles	Facteur
Dilution trop lente	Temps de dilution
Viscosité trop basse	Valeur viscosité
Support à peindre trop froid	Température de la pièce
Epaisseur de couche trop élevée	Grammage
Temps d'évaporation trop court	Temps passé dans la première colonne du four
Distance trop courte entre pistolet et objet	Hauteur de la tête de laquage
Buse de pistolage trop grande	Diamètre de la buse
Pistolage irrégulier	?

Une étude (que nous ne détaillerons pas ici) a donc eu lieu et a été validée avec les experts de manière à statuer sur les points suivants :

- Quels défauts sont cruciaux et méritent que l'on cherche à prévoir leur apparition ?
- Quels facteurs sont responsables de l'apparition de ces défauts ?
- Parmi ces facteurs, quels sont ceux qui sont contrôlables et quels sont ceux que l'on doit subir ?

La liste de facteurs à analyser est présentée dans le Tableau 25.

Tableau 25 - Facteurs à analyser

Facteurs contrôlables	Le taux de chargement de la table : ce taux a une influence sur l'écartement des pièces et donc sur la façon dont la laque atteindra l'intégralité du chant. Ce facteur est donc supposé avoir un lien avec des défauts atteignant les chants principalement.
	Le nombre de passes : les produits peuvent subir un deuxième laquage (2 ^{ème} passe) après le séchage de leur 1 ^{ère} passe, voire éventuellement une 3 ^{ème} dans le cas de laque très particulières. <i>Facteur discret à binariser</i>
	Le temps par table : lorsqu'on connaît le nombre de passes, ce facteur traduit directement la vitesse de la tête et peut donner une indication sur le potentiel soulèvement de poussière par exemple, en plus de la quantité de laque déposée effectivement sur les pièces.
	La quantité de laque.

	Le grammage : il s'agit de la donnée la plus précise concernant le volume de laque déposé puisqu'il est mesuré directement sur une pièce d'essai laquée en même temps que le lot à travailler.
	Le nombre de couches : lors du laquage, la tête robotisée va se déplacer pour couvrir intégralement la table au moins une fois (1 couche) mais, suivant les programmes utilisés, elle peut venir recouvrir la première couche d'une seconde, puis éventuellement d'une troisième. <i>Facteur discret à binariser</i>
	Le nombre de produits : au même titre que le taux de chargement mais sous un nouvel angle, ce facteur permet de se faire une idée sur l'écartement des produits.
	Le temps de séchage.
Facteurs non contrôlables	Température au plus proche du produit.
	Pression atmosphérique.
	Humidité au plus proche du produit.

Le nombre de couches et le nombre de passes sont tous deux des facteurs discrets qui peuvent prendre 3 états et qui doivent donc être binarisés (Thomas & Thomas, 2009).

Certains des facteurs étant intercorrélés et quelques peu soumis au planning (nombre de pièces et taux de chargement), on peut dans un premier temps les considérer comme non contrôlables et comme des contraintes issues du planning. Le système de prévision de la qualité pourra donc se schématiser par la Figure 29.

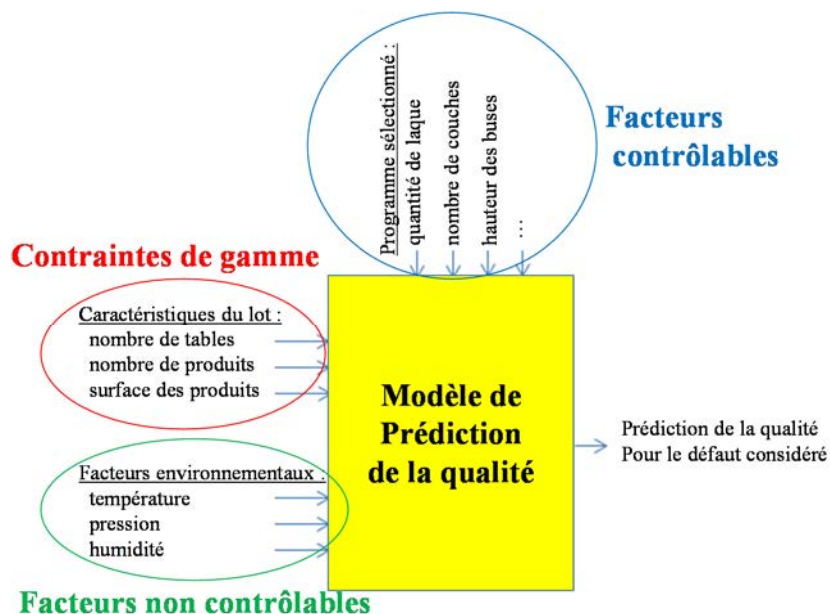


Figure 29 - Représentation des entrées et sortie du système de prévision des défauts.

b. La collecte des données

Sur ce poste de travail relativement récent dans l'entreprise, la volonté de collecter des informations pour pouvoir analyser le fonctionnement s'était faite sentir mais n'avait abouti qu'à la création d'un fichier Excel rempli manuellement par les opérateurs. Cette méthode présentait de multiples inconvénients tels que :

- le temps d'ouverture du fichier Excel de plus en plus long à cause du volume de données grandissant
- les erreurs de saisie
- la caractérisation du lot travaillé pas toujours exploitable
- le traitement de ces données qui n'était fait qu'*a posteriori* et dans des délais dépendant de la charge de travail

Nos travaux sur la qualité supposaient de faire saisir toujours plus d'informations aux opérateurs pendant plusieurs mois sans qu'ils puissent en voir l'intérêt direct. L'idée était donc d'obtenir les informations nécessaires à nos travaux tout en offrant un nouveau service aux opérateurs de ce poste de travail par l'intermédiaire d'un outil de travail qui saurait se rendre indispensable : un outil de TPM⁷ adapté développé en interne que nous avons nommé POTER.



Figure 30- Menu général du logiciel POTER

Les objectifs initiaux du développement de ce logiciel étaient de fournir aux opérateurs des informations sur leur productivité, aux responsables de production des outils/indicateurs de management simultané de la qualité et de la productivité au jour le jour, aux commerciaux et au service industrialisation des retours systématiques sur les coûts de fabrication, au service maintenance un suivi des pannes et des dérives machines et à tous une traçabilité partielle des pièces.

⁷ Total Productive Maintenance

Le Tableau 26 présente un extrait de son cahier des charges fonctionnel.

Tableau 26 - Analyse fonctionnelle de la solution proposée

F1	Ne pas faire perdre trop de temps à l'opérateur
F11	Limiter les chargements à moins d'une minute
F12	Réduire l'utilisation du clavier ou de la souris (Ergonomie)
F2	Limiter les erreurs de saisie
F21	Identifier automatiquement le lot travaillé (badgeage)
F22	Automatiser la récupération d'un maximum de données
F3	Historiser la fabrication
F31	Horodater le passage d'un lot sur une machine
F32	Sauvegarder le réglage machine correspondant au travail effectué
F33	Sauvegarder les paramètres environnements correspondant au moment du travail effectué
F34	Sauvegarder les arrêts et pannes par catégories
F4	Fournir des indicateurs de travail
F41	Afficher en continu la productivité d'une équipe sur un poste de travail
F42	Afficher le pareto des problèmes qualité
F43	Synthétiser les coûts de fabrication dans des bilans de production
F44	Suivre l'évolution des pannes et dérives machines
F5	Assurer la traçabilité des pièces
F51	Fournir pour un lot donné le dernier poste de travail visité
F52	Identifier comme lot fils toute partie de lot devant être séparée d'un lot père
F6	Fournir contextuellement des informations spécifiques de production à l'opération
F61	Fournir des standards de travail de divers formats (texte, image, vidéo...)
F7	S'adapter à n'importe quel poste de travail

Aujourd'hui, le logiciel a été déployé en réalisant les fonctions F1, F21, F3, F4, F5 et F7. La fonction F22 est pour l'instant mise de côté faute de moyens. La fonction F6 correspondant à de la communication descendante est en cours de déploiement. Il sera alors capable de transmettre une information dynamique (alertes, standards de travail sous divers formats (image, document, vidéo...)) qui est directement fonction du lot travaillé. Cette fonctionnalité achèvera de le rendre indispensable aux opérateurs.

Les données de production sont les données d'entrée de notre réseau de neurones. Elles doivent donc être impérativement collectées plusieurs mois avant de réaliser l'apprentissage même si elles n'ont en apparence pas d'intérêt pour les indicateurs de productivité opérationnelle offerts aux opérateurs.

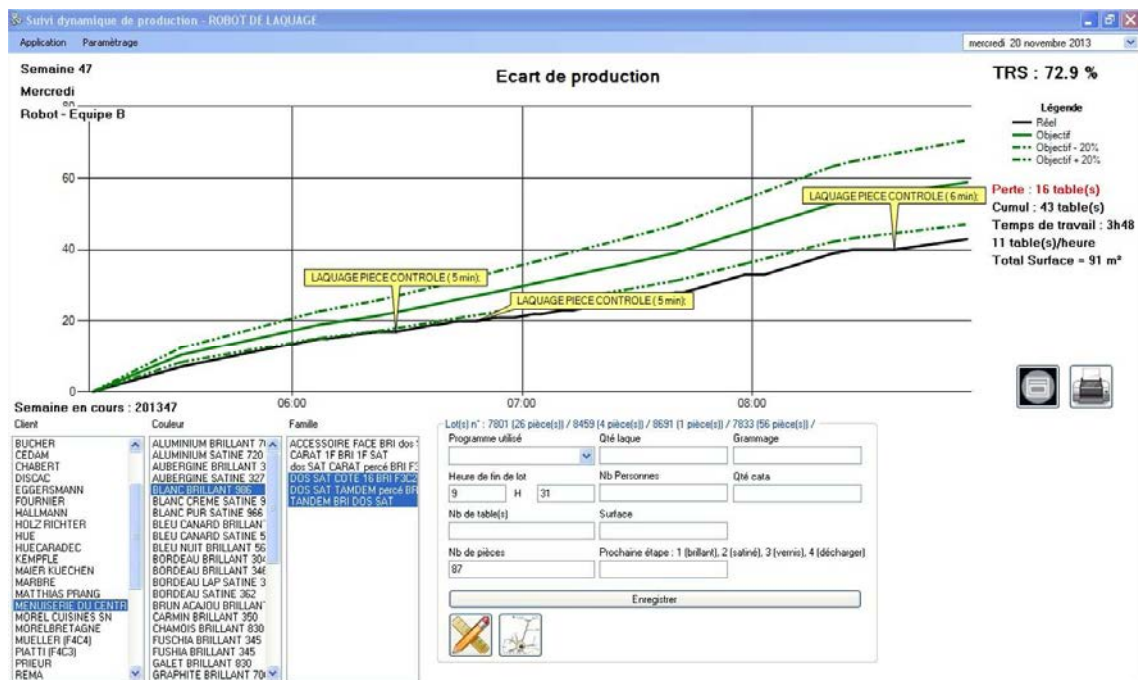


Figure 31 - Suivi dynamique de production

L'opérateur sélectionne son lot par l'intermédiaire des trois listes déroulantes en bas à gauche qui ne contiennent que les lots à fabriquer de la semaine. Cette étape est en-cours de remplacement par l'identification automatique des lots par code-barres 2D. Il est aussi envisagé d'y faire figurer uniquement le planning des prochains lots à fabriquer. L'encart en bas à droite permet de saisir les données de production nécessaires, et ce avec une volonté de récupérer un maximum d'informations automatiquement comme par exemple la surface de pièces à travailler transmise par l'ERP ou encore la quantité de laque utilisée transmise directement par le robot de laquage. Les données à compléter sont paramétrables de manière à pouvoir collecter temporairement si besoin une information supplémentaire. Enfin, la partie supérieure de la fenêtre permet à l'opérateur d'avoir une idée plus précise de son travail.

Les données de qualité sont les données de sortie de notre réseau de neurones. Elles doivent donc être impérativement collectées plusieurs mois avant de réaliser l'apprentissage. L'intérêt de collecter ces données est beaucoup plus facilement compréhensible par l'opérateur qui sait qu'il s'agit en premier lieu de la satisfaction client. L'opération était d'ailleurs déjà réalisée mais de manière manuelle et peu cadrée, ce qui ne facilitait pas l'analyse. L'interface tactile proposée permet le « batonnage ». En effet, l'opérateur, lors du déchargement des pièces, inspecte chaque pièce et touche le bouton correspondant à l'éventuel défaut détecté pour incrémenter le compteur.

Saisie des défauts : prod 47762

MEL
11 table(s)
Nb de pièces : 104
Temps au four : 7h15

200114
1c SAT face
Taux de rebus : 31.73 %

BLANC BRILLANT 986
Grammage : 85
Nb pcs défaut : 33

	BOURRELET 0	COULURES 0	GRAIN SUR FACE 32	GRAIN SUR DOS 0	GRAIN SUR CHANTS 0	TACHE SUR FACE 0	TACHE SUR DOS 0	TACHE SOUS LAQUE 0
	GOUTTE DE LAQUE 0	COUP 0	ECLAT 0	RAYURE 0	MANQUE LAQUE SUR FACE 0	MANQUE LAQUE SUR DOS 0	MANQUE LAQUE SUR CHANTS 0	REFUS 1
	PERCE 0	DEFAUT APPRET 0	DEFAUT PONCAGE 0	EMPREINTE 0	SILICONE 0	TROU 0	PIECES COLLEES 0	TRACES DE MOUSSE FOUR 0
	MICROBULLAGE 0	LAQUE PAS SECHE 0	CRATERE / CUVETTE 0	DIVERS 0	RACCORD 0	GOUTTE D'EAU 0	LAQUE PAS SECHE 0	

Figure 32 - Interface tactile de saisie des défauts

Pour des raisons de précision, nous avons préféré réaliser un système de prévision par défaut possible.

IV.2.3.2 Choix et réglage du classificateur pertinent

A ce stade, une BdD est disponible pour mettre en œuvre des stratégies d'apprentissage pour construire le classificateur. Classiquement, cette BdD est subdivisée en deux parties l'une pour l'apprentissage proprement dit, et l'autre pour la validation du modèle appris.

Au paragraphe II.2.1, une étude bibliographique nous a permis de restreindre les types de classificateurs susceptibles d'être utilisés pour notre application à quatre familles (Figure 12) parmi lesquels quatre outils (un par famille) ont été plus particulièrement sélectionnés. Même si certains critères (données bruitées, besoin ou non d'un modèle de connaissance...) nous permettraient *a priori* de privilégier certains outils par rapport à d'autres, il est difficile d'assurer avec certitude que tel ou tel outil sera le plus efficace pour l'application considérée. Aussi, les quatre outils seront tous testés et comparés. Pour chacun de ces outils divers réglages et/ou choix d'algorithmes doivent être effectués afin d'assurer l'efficacité de l'apprentissage.

a. Réglage de l'algorithme kpp

Comme vu au paragraphe II.2.1.1, le premier réglage qui doit être fait concernant l'algorithme kpp est le choix de la métrique pour calculer la distance entre les échantillons.

Or, nous sommes en présence de données réelles industrielles, et de ce fait, susceptibles d'être polluées par du bruit et des valeurs aberrantes. Nous avons donc choisi d'utiliser la distance Mahalanobis qui est moins sensible à ces problèmes.

Le deuxième réglage concerne le choix du nombre de voisins k . Pour chacun des modèles construits, ce nombre a été déterminé par essais-erreurs en le faisant varier et en sélectionnant la valeur de k qui minimise le taux de mauvaises classifications sur le jeu de données de validation.

b. Réglage de l'arbre de décision

Comme vu au paragraphe II.2.1.2, divers algorithmes ont été proposés pour apprendre des arbres de décisions. Nous avons sélectionné l'algorithme CART qui, selon la bibliographie semble être l'un des plus performants. Cet algorithme est capable de déterminer automatiquement la dimension de l'arbre résultant, même si une stratégie de « pruning » par essais-erreurs sur un jeu de validation permet d'améliorer les performances du modèle. Une telle stratégie est cependant assez gourmande en temps de calcul et nous avons choisi de ne pas la mettre en œuvre.

c. Réglage des SVM

Comme expliqué au paragraphe II.2.1.3, la performance des SVM dépend en grande partie du réglage de deux types de paramètres que sont le poids des variables ressorts dans le critère à optimiser et la largeur du noyau RBF. Même si l'algorithme utilisé (Grandvalet & Canu, 2008) est sensé adapter automatiquement ces paramètres, il est tout de même très sensible à la valeur initiale choisie pour la largeur du noyau, et là encore, une stratégie par essais-erreurs sur un jeu de validation a été mise en place pour déterminer la largeur initiale du noyau RBF.

d. Réglage des MLP

Durant le paragraphe II.2.1.4, nous avons vu qu'il existe un grand nombre de réseaux de neurones différents. Celui que nous avons privilégié pour notre application est le perceptron multicouches (MLP) dont la mise en œuvre nécessite de faire un certain nombre de choix, à commencer par les fonctions d'activations des neurones cachés et de sortie, ainsi que le nombre de neurones à introduire sur chaque couche. Le nombre de neurones sur la couche d'entrée correspond au nombre de variables d'entrée du modèle. Pour ce qui est de la sortie, nous avons choisi de construire un modèle par type de défaut ce qui fait que chaque modèle va inclure un et un seul neurone de sortie. Comme le problème est d'obtenir une probabilité d'apparition de défaut, la fonction d'activation du neurone de sortie est choisie sigmoïdale.

Pour les neurones cachés la fonction d'activation choisie est la tangente hyperbolique qui permet d'assurer les capacités d'approximateur universel parcimonieux du MLP.

Aucune technique à ce jour ne permet de définir *a priori* le nombre de neurones à inclure dans la couche cachée, et les algorithmes d'apprentissage effectuant une recherche locale de minimum, l'efficacité de l'apprentissage dépend du jeu de paramètres initiaux. Les trois étapes (initialisation, apprentissage et sélection de la structure optimale) seront donc effectuées séquentiellement en utilisant les algorithmes suivant :

- Initialisation semi aléatoire des poids des synapses permettant de recommencer l'apprentissage à partir de divers points dans l'espace de recherche et ce afin de palier au problème d'obtention de minimum locaux (Nguyen & Widrow, 1990)
- Apprentissage des paramètres en utilisant l'algorithme de Levenberg-Marquard exploitant un critère robuste permettant de réduire l'impact de la présence de valeurs aberrantes dans les données (Thomas, Bloch, Sirou, & Eustache, 1999)
- Elimination des poids pour supprimer les paramètres inutiles afin de trouver la structure optimale et ainsi éviter les problèmes de surapprentissage (Thomas and Suhner, 2014)

e. Constitution d'un ensemble classificateur

Chacun des outils précédemment décrits (kpp, DT, SVM, MLP) est capable de modéliser notre système avec plus ou moins de précision. Au paragraphe II.2.1.5, nous avons vu que l'utilisation d'un ensemble classificateur est susceptible d'améliorer ces résultats. Pour évaluer la pertinence d'utiliser un ensemble classificateur, nous avons comparé les performances des divers classificateurs individuels obtenus avec divers ensembles classificateurs dont nous avons décrit la construction. Tout d'abord, pour construire un ensemble classificateur, il est nécessaire de déterminer le type de classificateur individuel à intégrer. Nous avons présumé que l'utilisation de classificateurs de types différents (kpp, DT, SVM ou MLP) va permettre d'améliorer les performances de l'ensemble. Pour évaluer ce fait, nous avons construit des ensembles utilisant uniquement un type de classificateur (kpp, DT, SVM ou MLP) ainsi que des ensembles utilisant les quatre types de classificateurs simultanément.

Le deuxième point qu'il faut déterminer est la méthode de sélection des classificateurs individuels. Au paragraphe II.2.1.5, deux stratégies ont été proposées reposant l'une sur la précision et l'autre sur la diversité des classificateurs. Afin d'évaluer l'impact de la diversité sur le choix des classificateurs, ces deux approches ont été utilisées concurremment.

Le troisième point concerne la fusion des classificateurs individuels. Au paragraphe II.2.1.5, nous avons vu que cette fusion était généralement réalisée par vote majoritaire, mais nous avons également proposé, dans le cas où les classificateurs individuels fournissent une sortie numérique, d'utiliser la moyenne des sorties des classificateurs individuels. Enfin une troisième approche proposant d'exploiter un MLP pour simultanément effectuer la sélection et la fusion des classificateurs a également été proposée. Là encore, afin d'évaluer l'impact de ce choix sur les performances de l'ensemble, ces trois stratégies ont été testées et comparées.

Ces différents choix nous ont conduit à construire 20 types d'ensembles classificateurs différents dont il nous faudra comparer les performances entre eux mais aussi, avec celles des meilleurs classificateurs individuels.

f. Critère de comparaison des classificateurs

Dans les problèmes de classification, l'objectif est de réduire le nombre de données mal classées. Dans ce sens, le critère de comparaison classique est le taux de mauvaises classifications appelé aussi taux d'erreurs ou « zero-one score function » (Hand et al., 2001).

$$S_{01} = \frac{1}{N} \sum_{n=1}^N I(y_n, \hat{y}_n) \quad (13)$$

où $I(a, b) = 1$ quand $a \neq b$ et 0 sinon, et N le nombre de patterns.

Si l'on souhaite associer différents coûts aux différentes erreurs de classification, une matrice de perte générale peut être construite (Bishop, 1995) mais ce problème n'est pas abordé ici.

Le meilleur taux de mauvaises classifications $S_{01_{\min}}$ est obtenu sur l'ensemble de données de validation avec la meilleure approche. Un test d'hypothèse statistique de McNemar est utilisé pour déterminer dans quels cas le taux de mauvaises classifications des autres approches est statistiquement différent de celui obtenu avec la meilleure approche. L'hypothèse nulle (où l'approche testée est statistiquement égale à la meilleure) est testée, quand H_0 et son alternative H_1 sont :

$$\begin{cases} H_0 : S_{01} = S_{01_{\min}} \\ H_1 : S_{01} \neq S_{01_{\min}} \end{cases} \quad (14)$$

L'hypothèse nulle H_0 est rejetée avec un risque de 5% si :

$$U = \frac{|N_{10} - N_{01}|}{\sqrt{N_{10} + N_{01}}} > 1.96 \quad (15)$$

où N_{10} est le nombre de cas où le meilleur classificateur donne un résultat correct tandis que le classificateur auquel il est comparé donne un résultat faux, et N_{01} est le nombre de cas strictement inverses.

IV.2.3.3 Création et évaluation des classificateurs

a. Les bases de données

Nous allons focaliser notre présentation sur un des trente types de défauts collectés, le défaut « grain ». Sur la période considérée nous avons pu collecter 2270 données dont 273 appartiennent à la classe 1 (défaut présent). Ce jeu de données est subdivisé aléatoirement en deux, un pour l'apprentissage (1202 données dont 146 positifs) et un pour la validation (1068 données dont 127 positifs).

L'apprentissage de certains types de classificateurs, comme les arbres de décision, est un processus déterministe. Pour introduire de la diversité entre les classificateurs construits, un algorithme de « bagging » est utilisé sur le jeu de données d'apprentissage afin de créer 100 jeux de données d'apprentissage différents. Cette étape n'est pas nécessaire pour les MLP qui donnent des résultats similaires avec ou sans « bagging ». Nous pouvons donc ainsi créer 100 arbres de décision (DT), 100 kpp, 100 SVM et 100 MLP que nous pourrions comparer comme classificateurs individuels mais aussi intégrer dans les différents ensembles construits.

b. Création d'un classificateur

Afin de simplifier la présentation, nous allons présenter et discuter uniquement le cas d'un modèle MLP, le même type de travail pouvant être mené pour chacun des 400 modèles construits.

Le réseau de neurones de départ contient 15 neurones d'entrée dont 6 sont binaires. L'apprentissage initial est effectué avec 25 neurones dans la couche cachée. La phase de pruning permet d'éliminer les entrées non significatives et les neurones cachés inutiles. Après cette étape de pruning, sur le modèle considéré 6 neurones cachés et 1 entrée (le nombre de passes) ont été éliminés. Dans le cas où cette entrée serait majoritairement éliminée lors de la phase de pruning pour les 100 modèles différents, cela signifierait que cette variable (nombre de passes) n'a pas un impact significatif sur la création du défaut considéré. Ce n'est pas le cas ici.

Le modèle ainsi construit est testé sur le jeu de validation constitué de 1068 données dont 127 positifs (Figure 33). Sur ces 127 positifs, 112 sont prédits par le réseau de neurones conduisant à un taux de non-détections de 11,8%. La proportion de faux positifs est de 19,2%. Ce taux important peut être expliqué partiellement par le fait qu'une partie des défauts n'a pas été enregistrée à la sortie de la machine sur une période correspondant à la formation d'un nouvel opérateur(entre les données 600 et 800 Figure 33).

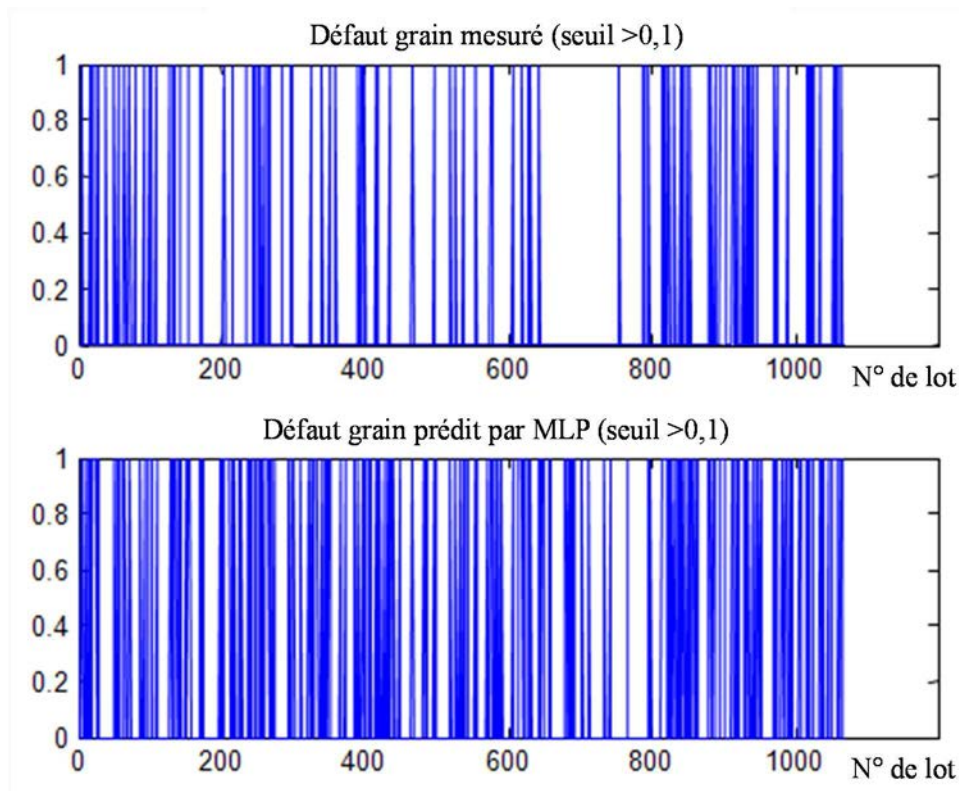


Figure 33 - Comparaison des défauts de grains détectés à la sortie du robot de laquage (graphique supérieur) avec les défauts prévus par le réseau de neurones (graphique inférieur)

Il est évident que ce défaut est largement expliqué par les conditions de production archivées et il est donc possible d'utiliser un classificateur en amont du poste de travail pour limiter le risque de « grains ». L'expérience pourra être répétée avec les autres défauts évoqués précédemment. En suivant la logique de Pareto, l'idéal serait de travailler en priorité sur les 20% des défauts qui sont à l'origine des 80% de non-qualité. Pourtant, certains défauts ne peuvent être prédits par les classificateurs avec les facteurs décrits précédemment. C'est le cas par exemple du défaut "coups" pour lequel la phase de validation après l'apprentissage conduit à 73% de non-détections pour 11% de faux-positifs. Ce facteur non-prévisible dépend soit d'autres facteurs qu'il faudrait savoir identifier, quantifier et collecter afin de pouvoir le prédire correctement, soit de la

complexité de prévision qui nécessiterait un ensemble de classificateurs plutôt qu'un seul. Au total, sur les 30 défauts identifiés, 7 peuvent être expliqués en utilisant les valeurs des facteurs collectés avec un seul classificateur.

c. Précision des classificateurs individuels

Le paragraphe précédent a montré l'étude menée sur un classificateur. 400 classificateurs de 4 types différents ont été construits aussi est-il nécessaire de les comparer pour sélectionner le ou les plus pertinents.

La diversité entre les classificateurs est un point clef mais moins important que la précision des différents classificateurs. La Figure 34 présente la distribution des taux de mauvaises classifications des différents types de classificateurs. Les rectangles représentent le deuxième et le troisième quartile de la distribution.

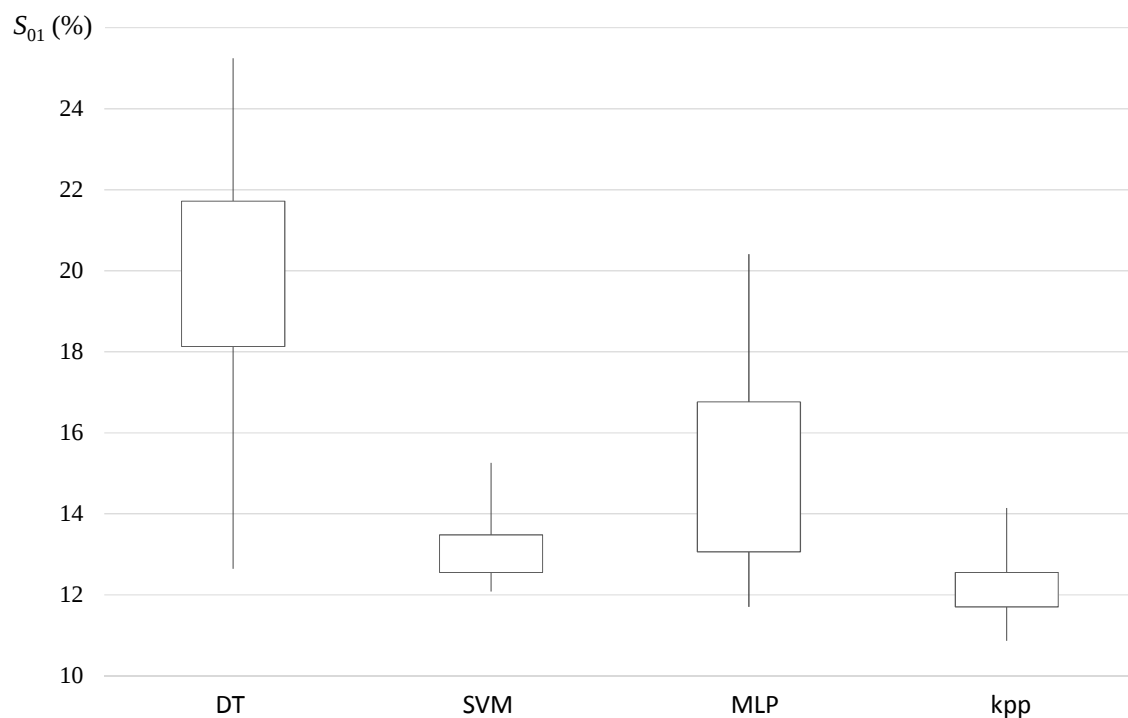


Figure 34 - Distribution des taux de mauvaises classifications

La Figure 34 montre que les kpp donnent généralement les meilleurs résultats tandis que les arbres de décision donnent les plus faibles. Les différents modèles SVM donnent des résultats très proches les uns des autres tandis que les réseaux de neurones présentent une distribution plus forte.

Le Tableau 27 montre les résultats obtenus avec le meilleur classificateur individuel de chaque type de classificateur. Les modèles kpp donnent les meilleurs résultats pour S_{01} mais avec un très mauvais taux de non-détections (comme pour les modèles SVM).

Les réseaux de neurones et les arbres de décision donnent des résultats similaires avec un taux de non-détections meilleur mais au détriment du taux de fausses alarmes.

Tableau 27 - Résultats obtenus par le meilleur réseau de neurones, arbres, kpp et SVM.

	S_{01}	% fausse alarme	% non-détection	U
meilleur MLP	11.8%	6.7%	49.6%	4.16
meilleur DT	12.6%	6.9%	55.1%	5.12
meilleur kpp	10.9%	1.9%	77.2%	3.40
meilleur SVM	12.1%	1.9%	87.4%	4.54

Le Tableau 28 présente la matrice de confusion pour les deux meilleurs classificateurs individuels (NN et kpp) qui se comportent différemment. En effet, les modèles kpp favorisent le taux de fausses alarmes tandis que les réseaux de neurones sont plus équilibrés, ce qui a tendance à améliorer le taux de non-détections.

Tableau 28 - Matrice de confusion des deux meilleurs classificateurs (NN et kpp)

	meilleur classificateur MLP			meilleur classificateur kpp		
	Défaut (predit)	non défaut (predit)	précision	Défaut (predit)	non défaut (predit)	précision
Défaut	85	32	72.65%	29	98	22.83%
Non défaut	93	848	90.12%	18	923	98.09%

d. Diversité entre les classificateurs individuels.

Nous avons vu que la diversité entre classificateurs peut être déterminante pour la construction d'un ensemble classificateur. Nous allons donc étudier cette diversité entre les différents classificateurs en utilisant la mesure de diversité DF donnée en équation (2). Il s'agit d'une équation de comparaison qui donne une valeur inversée. Plus la mesure est faible, plus la diversité est grande. Cette mesure indique le percentile de données pour lequel les deux classificateurs considérés sont incorrects.

La Figure 35 présente la distribution de la mesure de diversité pour les 4 classificateurs étudiés individuellement puis ensemble. Le rectangle représente le deuxième et le troisième quartile de la distribution. La figure montre que, même si les approches SVM et kpp conduisent à différents modèles, ces derniers présentent moins de diversité que ceux

obtenus en utilisant d'autres approches. L'approche avec les arbres de décision est celle qui conduit à un maximum de diversité entre les modèles même si la plus petite valeur de diversité a été obtenue avec les réseaux de neurones. Cependant, cette dernière approche est celle qui présente le plus de dispersion. Parfois, les réseaux de neurones sont très similaires ($DF = 0.17$) et parfois ils le sont beaucoup moins ($DF=0.011$).

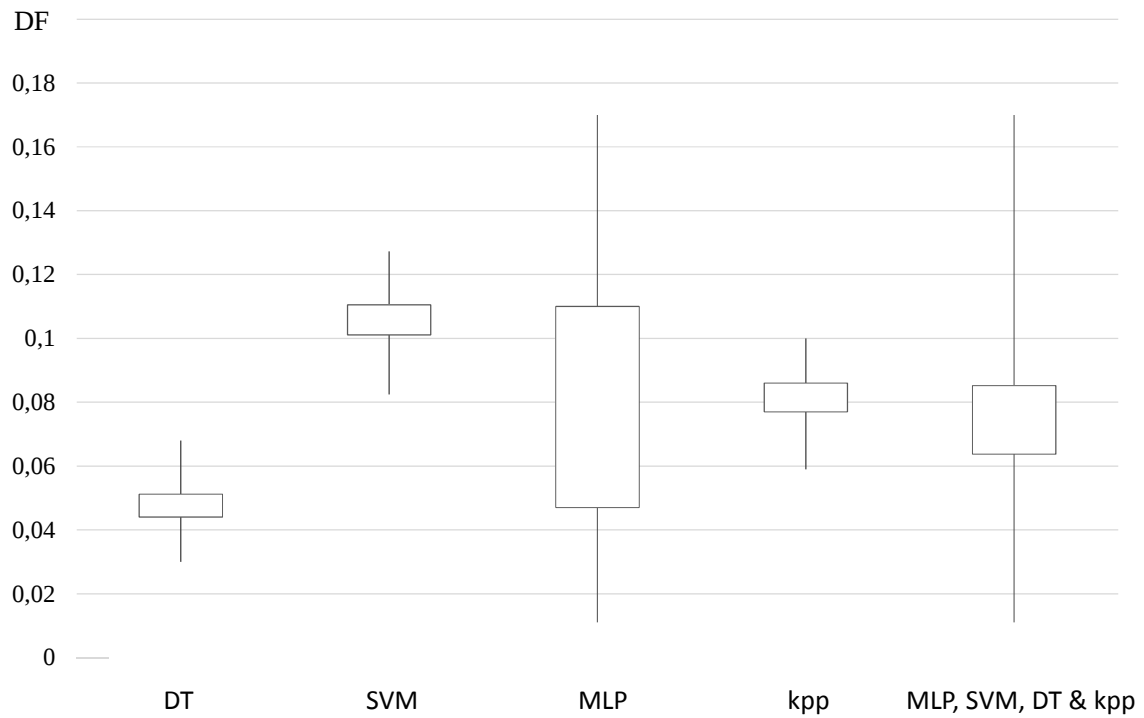


Figure 35 - Distribution des mesures de diversité (DF)

e. Précision des ensembles classificateurs

Indépendamment, les kpp classificateurs sont performants en termes de mauvaises classifications et les réseaux de neurones et les arbres de décision sont tous deux performants en termes de non-détections. Ces résultats peuvent donc vraisemblablement être améliorés en utilisant des ensembles.

Le Tableau 29 présente les taux de mauvaises classifications, de fausses alarmes et de non-détections pour le meilleur ensemble de classificateurs dans lequel les classificateurs individuels ont été sélectionnés en fonction de leur précision et de leur diversité. La fusion de leur avis a été réalisée en utilisant le vote, la moyenne ou des réseaux de neurones. Le meilleur résultat ($S_{01_{\min}} = 7.7\%$) est obtenu avec l'ensemble de classificateurs qui regroupe les 4 types de classificateurs individuels choisis sur leur diversité. Le test de McNemar est réalisé pour déterminer si les autres approches sont statistiquement différentes ou non de la

meilleure. Ce test montre que les 4 meilleurs classificateurs individuels présentés dans le Tableau 27 sont moins performants que le meilleur ensemble de classificateurs ($U > 1.96$). Donc, l'utilisation d'un ensemble de classificateurs améliore bien la précision du système de prévision.

Tableau 29 - Résultats obtenus en utilisant le meilleur ensemble de classificateurs.

	sélection	fusion	taille	S_{01}	% fausse alarme	% non-détection	U
MLP ensemble	accuracy	vote	12	8.8%	5.2%	35.4%	1.39
		moyenne	13	8.6%	6.4%	25.2%	1.13
DT ensemble		vote	12	10.3%	2.1%	70.9%	2.95
kpp ensemble		vote	23	10.3%	0.5%	82.7%	2.77
SVM ensemble		vote	1	12.1%	1.9%	87.4%	4.54
		moyenne	9	11.4%	2.0%	81.1%	4.00
classificateur ensemble		vote	6	9.0%	3.1%	52.8%	1.65
MLP ensemble	diversity	vote	17	10.7%	8.1%	29.9%	3.57
		moyenne	22	9.0%	5.4%	35.4%	1.94
DT ensemble		vote	10	10.1%	1.4%	74.8%	2.71
kpp ensemble		vote	28	10.2%	1.6%	74.0%	2.71
SVM ensemble		vote	23	11.7%	1.6%	86.6%	4.37
		moyenne	1	12.1%	1.9%	87.4%	4.54
classificateur ensemble		vote	24	7.7%	4.0%	34.7%	-
MLP ensemble	MLP		92	10.6%	6.2%	43.3%	2.94
DT ensemble			97	13.1%	6.4%	62.2%	5.91
kpp ensemble			100	13.1%	4.5%	77.2%	5.88
SVM ensemble			98	12.1%	2.6%	82.7%	4.27
classificateur ensemble			394	12.0%	1.1%	92.9%	4.20

L'utilisation de la stratégie de réseaux de neurones pour fusionner les classificateurs n'est pas apte à trouver une structure satisfaisante ($U > 1.96$). Ce fait n'est pas surprenant pour les ensembles de kpp et SVM à cause de leur faiblesse en matière de diversité. Ce n'est pas surprenant non plus pour les ensembles d'arbres de décision à cause de leur faiblesse en matière de précision.

Quatre ensembles de classificateurs donnent des résultats statistiquement équivalent au meilleur ($U < 1.96$) :

- L'ensemble de réseaux de neurones sélectionnés sur la précision avec une fusion basée sur le vote
- L'ensemble de réseaux de neurones sélectionnés sur la précision avec une fusion basée sur la moyenne
- L'ensemble des réseaux de neurones sélectionnés sur la diversité avec une fusion basée sur la moyenne
- L'ensemble de classificateurs sélectionnés sur la précision avec une fusion basée sur le vote.

La fusion basée sur la moyenne améliore sensiblement les résultats pour les ensembles de réseaux de neurones (sélectionnés sur la précision comme sur la diversité). Si les meilleurs résultats sont obtenus avec les ensembles de classificateurs sélectionnés sur la

diversité, les ensembles de classificateurs sélectionnés sur la précision sont les plus parcimonieux (seulement 6 classificateurs individuels).

Un autre avantage des ensembles de classificateurs est que le vote peut être utilisé comme un intervalle de confiance sur le résultat. Par exemple, si 40% des classificateurs votent pour un défaut et 60% votent contre, on va prédire qu'il n'y aura pas de défaut, mais il sera également possible d'indiquer un niveau de confiance faible en la prédiction comparativement au cas où 100% des classificateurs sont d'accord pour prédire l'absence de défaut, par exemple.

Tableau 30 - Matrice de confusion des deux meilleurs ensembles de classificateurs

	Classificateur ensemble (diversité)			MLP ensemble (précision & moyenne)		
	Défaut (predit)	non défaut (predit)	précision	Défaut (predit)	non défaut (predit)	précision
Defaut	39	88	30.71%	95	32	74.80%
Non défaut	23	918	97.56%	60	881	93.62%

Le Tableau 30 présente la matrice de confusion pour les deux meilleurs ensembles de classificateurs suivants :

- Ensemble de classificateurs sélectionnés sur la diversité
- Ensemble de réseaux de neurones sélectionnés sur la précision avec fusion basée sur la moyenne.

Cette matrice montre que les ensembles de réseaux de neurones tendent à avoir un comportement équilibré concernant les taux de non-détections et de fausses alarmes tandis que les ensembles de classificateurs favorisent le taux de fausses alarmes. En comparant ces résultats avec ceux du Tableau 27, l'utilisation d'ensembles de réseaux de neurones améliore à la fois les taux de non-détections et de fausses alarmes.

f. Exploitation du classificateur sélectionné

L'exploitation du classificateur sélectionné peut se faire selon deux approches :

- Alarme. En analysant les données d'entrée à travers le réseau de neurones, il devient possible de prédire l'occurrence de défauts et d'avertir l'opérateur que les conditions sont remplies pour créer le risque.
- Limitation. En utilisant le réseau de neurones pour limiter la valeur des facteurs d'entrée par une limite basse et une limite haute et en interdisant la production dès lors qu'un des facteurs est hors limites. S'il s'agit d'un facteur contrôlable, l'opérateur peut le modifier pour débloquer la production du lot. Sinon, la production du lot sera planifiée plus tard, lorsque les conditions de production (notamment météorologique) auront évolué pour redevenir acceptables. Le réseau de neurones peut alors être utilisé pour réaliser des plans d'expériences complets à la place du système de production physique dans le but de réduire

les coûts. Pour plus de sécurité, le résultat des plans d'expériences doivent être validés sur le système réel.

Dans ce dernier cas, suivant la philosophie de Taguchi, l'idée est d'exploiter le modèle ainsi construit pour simuler un plan d'expériences complet sur les paramètres contrôlables en fonction de l'état courant des paramètres non contrôlables (ou imposés par la gamme du lot de fabrication considéré) afin de déterminer les bornes de réglage optimal des paramètres contrôlables. La Figure 36 présente un exemple de réalisation d'une telle simulation de plan d'expériences en exploitant comme classificateur, le meilleur MLP. Le nombre de modalités est fixé à 10 pour l'ensemble des facteurs contrôlables.

Les facteurs non contrôlables sont eux fixés par les conditions courantes dans lesquelles évolue le système (température actuelle, pression actuelle ...) ou par la gamme de fabrication du lot considéré (nombre de passes, nombre de couches ...). Dans l'exemple considéré ces paramètres imposés par la gamme du lot sont fixés à :

- 1 pour le nombre de passes et de couches,
- Leurs valeurs médianes pour les temps par table, quantité de laque et nombre de pièces.

Les facteurs non-contrôlables imposés par l'environnement du système sont pour l'exemple fixés à leurs valeurs médianes.

L'effet de chaque facteur x_i à un niveau A_i est donné classiquement par la moyenne des résultats (occurrence de défaut) obtenu quand $x_i = A_i$ moins la moyenne des résultats obtenus avec toutes les expériences. Quand un effet est positif, cela implique que le niveau considéré augmente la probabilité d'apparition d'un défaut. Quand il est négatif, le niveau considéré réduit donc la probabilité d'apparition d'un défaut. L'interaction entre les facteurs peut être étudiée de la même manière.

La Figure 36 présente les résultats obtenus avec le meilleur réseau de neurones. Elle montre que le taux de chargement de la table a un impact relativement faible sur l'apparition de défauts. Par contre, l'augmentation du grammage a un effet très important. Un fort grammage aura tendance à générer beaucoup plus de défauts. Le temps de séchage a aussi un impact. C'est cette fois un temps de séchage trop court qui aura tendance à générer plus de défauts. Ce plan d'expériences nous permet donc de fixer, pour un lot considéré, et dans les conditions courantes d'utilisation, des valeurs limites de réglages pour les trois facteurs contrôlables afin de limiter le risque de défaut « grain ».

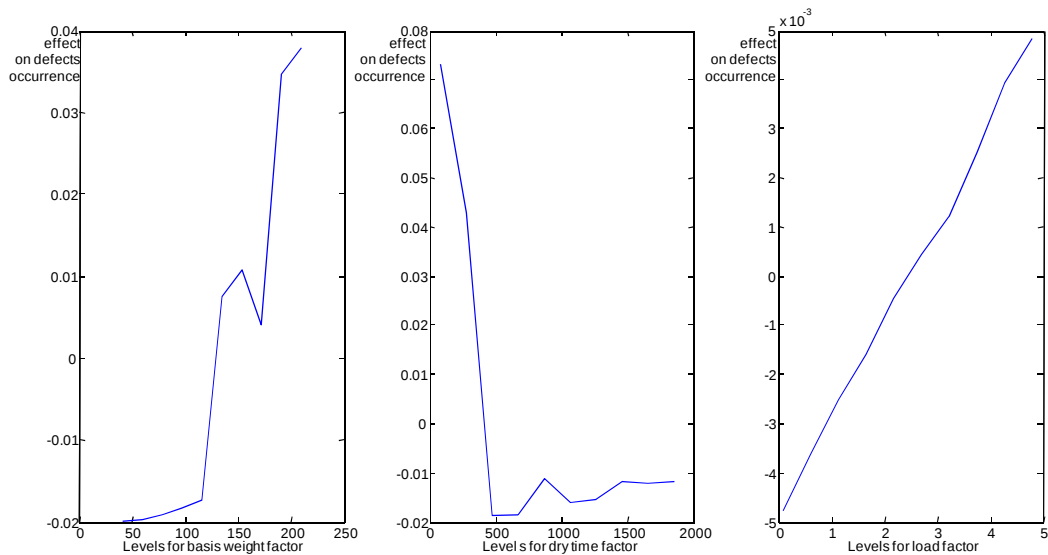


Figure 36 - Résultats du plan d'expériences en utilisant le meilleur réseau de neurones.

La Figure 37 présente un travail similaire réalisé avec le meilleur ensemble classificateur. A cause de la diversité des différents classificateurs constituant l'ensemble, l'impact de chacun des effets est présenté sous la forme d'une enveloppe qui inclut donc les résultats de tous les classificateurs. L'information globale fournie par l'ensemble est donc plus complète que celle fournie par un seul classificateur. Si l'ensemble de réseaux de neurones confirme bien que le taux de chargement a un impact très faible sur l'apparition de défauts, il affine cependant les résultats pour le grammage et le temps de séchage en montrant qu'il existe des bornes hautes et basses au-delà desquelles l'occurrence de défauts augmente. Cette information ne pouvait pas être obtenue par l'utilisation seule du meilleur réseau de neurones. Ces résultats qui se traduisent comme de la connaissance du processus se révèlent très utiles pour le réglage des paramètres à leur valeur optimale sous la contrainte des facteurs incontrôlables comme les facteurs météorologiques par exemple.

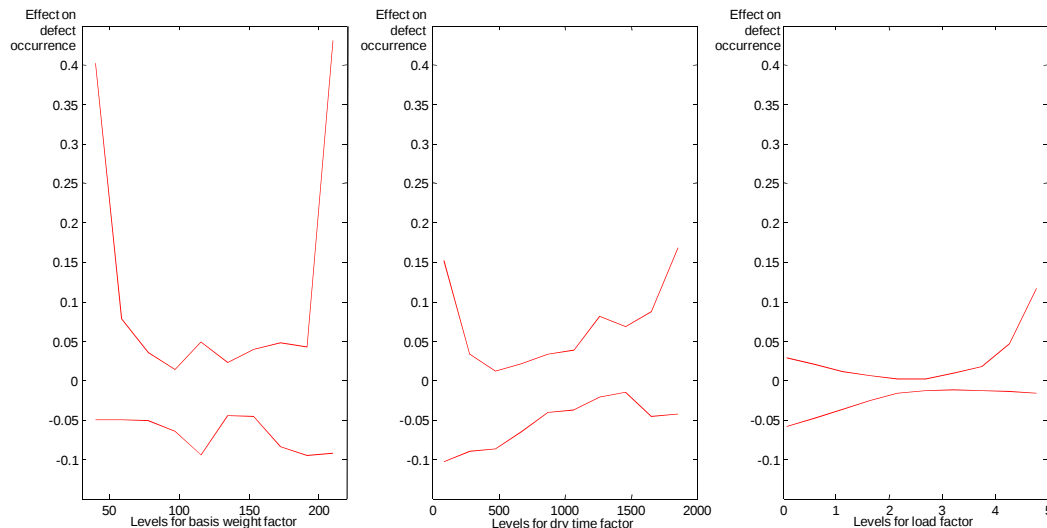


Figure 37 - Résultats du plan d'expériences en utilisant un ensemble classificateur.

Ces travaux montrent que la méthode proposée est bien crédible et qu'il est possible de construire un classificateur (unique ou ensemble) capable de prédire le risque d'apparition de défauts et exploitable pour déterminer des plages de réglages pour les paramètres contrôlables.

Pendant, ce modèle est statique alors que le système réel (ou son environnement) est susceptible d'évoluer au cours du temps. La deuxième étape de notre méthode propose alors une procédure d'adaptation requise dans le but d'adapter le modèle au système réel dès qu'un changement est détecté.

IV.2.4 Conception de l'observateur

Au paragraphe II.3.3, nous avons proposé une procédure permettant d'adapter notre classificateur aux évolutions du système ou de l'environnement.

IV.2.4.1 Besoin d'adaptation du classificateur

Parmi les raisons de dérive énoncées au paragraphe II.2.2 qui sont à l'origine des désynchronisations entre le système physique et le modèle construit, la première rencontrée pour notre application concerne l'évolution des paramètres d'entrée. En effet, le modèle ayant réalisé son apprentissage sur des plages données de paramètres d'entrée, les résultats qu'il fournira ne peuvent être valides que si ces entrées restent dans les plages qui ont permis l'apprentissage. Cette partie (Figure 38) présente l'évolution d'un des paramètres environnementaux (la température), dans la base de données utilisée pour les phases d'apprentissage et de validation (en bleu) et dans la base de données correspondant à la mise en exploitation du modèle classificateur. Cette évolution fait apparaître des domaines

de fonctionnement dans la base de données d'exploitation qui n'existaient pas dans la base de données d'apprentissage (températures négatives).

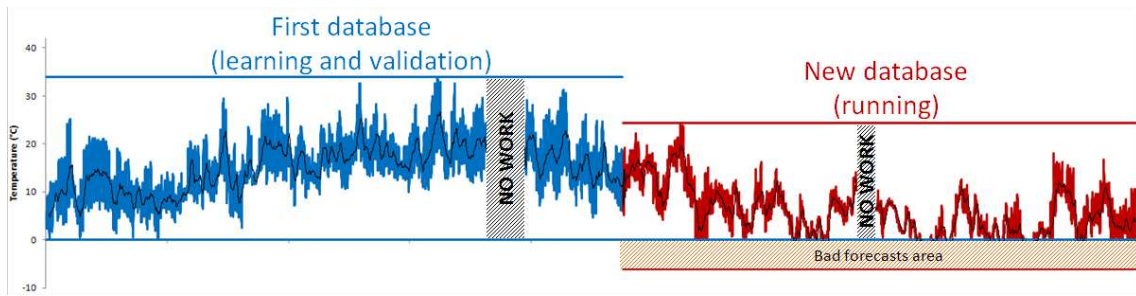


Figure 38 - Différence entre le domaine d'apprentissage et le domaine d'utilisation.

Après mise en service du classificateur, 446 données furent collectées et l'exploitation de ces données a conduit à 73% de non-détections et 32% de faux positifs. Ces résultats très décevants peuvent s'expliquer par la différence de conditions de fabrication entre les 2 périodes (apprentissage + validation et utilisation) constatée sur la température mais aussi sur d'autres paramètres d'entrée.

Ce type de changements de comportement est considéré comme une dérive lente. Il peut exister des changements de comportement beaucoup plus brusques qui résultent généralement d'événements ayant un impact sur le processus de fabrication.

Jusqu'à ce changement qui peut ne pas être connu des opérateurs ou experts, le modèle fournira des résultats conformes à la réalité. Un exemple d'une telle évolution brusque est présenté sur la Figure 39 où l'on peut voir une dérive franche sur les statistiques d'apparition du défaut « grain sur chants » aux alentours du 22 juin. Les données collectées dans l'environnement du poste de travail ne permettent pas de savoir s'il y a eu un changement sur un paramètre d'entrée influent oublié ou si le domaine d'utilisation était trop différent du domaine d'apprentissage. Dans le cas présenté, il pourrait s'agir par exemple d'un défaut du compresseur qui ne permettait plus une aspiration suffisante pour le dépoussiérage total des pièces. Quelle qu'en soit la raison, le modèle du système construit avant le changement n'est plus représentatif du comportement de ce même système après changement.

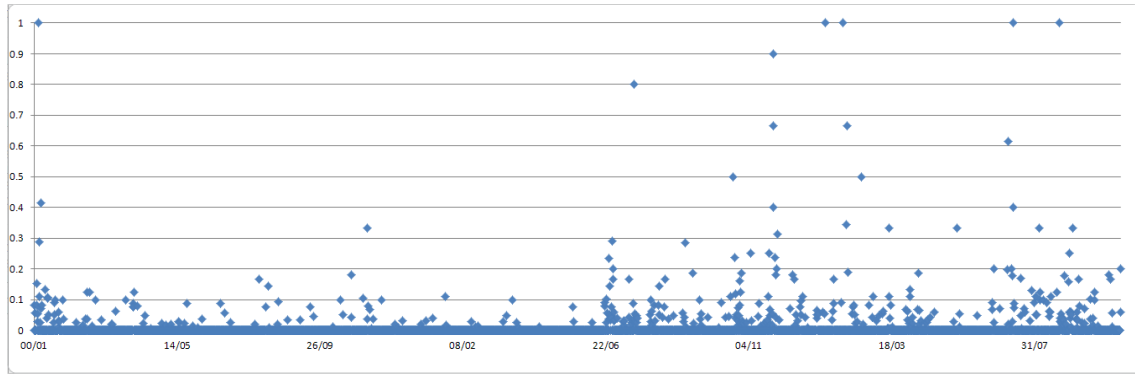


Figure 39 - Historique d'apparition de grains sur les chants.

Il est donc nécessaire d'être capable de détecter de tels changements et d'adapter le modèle si besoin.

IV.2.4.2 Déclencher le réapprentissage avec les cartes de contrôle

En donnant au système de prévision de la qualité la capacité de vérifier ses hypothèses en restant informé de la réalité, il devient capable de reconnaître ses erreurs et de réagir en fonction. Cette capacité lui est conférée d'une part par le biais d'une carte de contrôle permettant de déterminer quand une adaptation du modèle est rendue nécessaire, et d'autre part, par un test de Page Hinkley permettant de déterminer le volume de données sur lequel doit avoir lieu le réapprentissage comme présenté dans la partie II.3.3. La Figure 40 présente une resynchronisation du système suite à la détection d'une dérive dans ses prévisions par l'utilisation d'une carte de contrôle. La courbe en rose représente la prévision s'il n'y avait pas eu de réapprentissage qui aurait été hors limites dès le troisième lot de données. Grâce au réapprentissage qui a lieu après le deuxième lot de données, les prévisions sur le troisième lot de données restent fiables et ne nécessitent pas un deuxième réapprentissage.

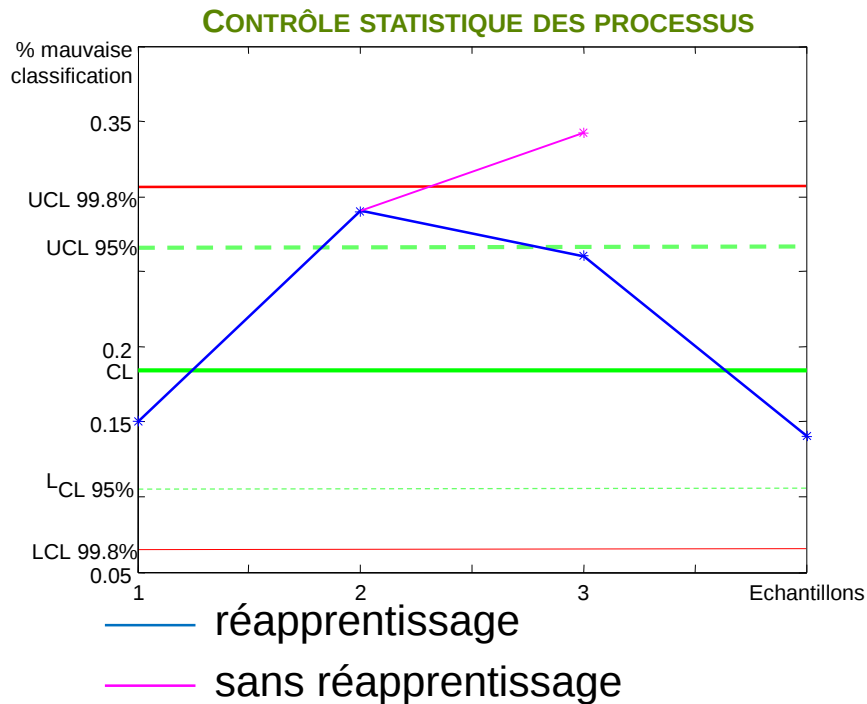
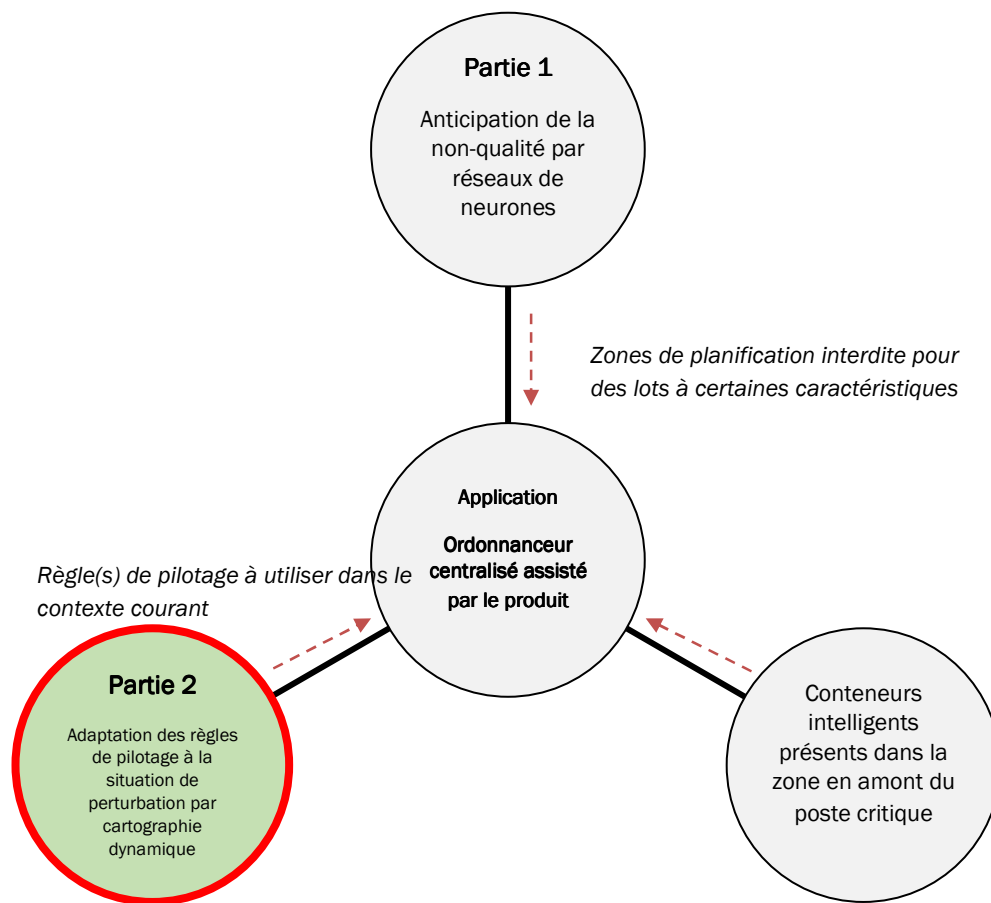


Figure 40 - Réaction face à la détection d'une dérive par le système lui-même.

IV.2.5 Résultats et discussions

L'implémentation de ce système de prévision de la non-qualité sur le robot de laquage de l'entreprise permet de limiter la production de non-qualité dont l'importance a un impact démontré (cf partie III.2.1) sur les flux. Cependant, suivant le type de défaut considéré, un classificateur seul, dont l'apprentissage sera relativement simple, peut parfois suffire, mais il sera parfois nécessaire de mettre en place un ensemble de classificateurs, ce qui sera une tâche plus conséquente. Dans certains autres cas (concernant des facteurs majoritairement liés à l'humain), même un ensemble de classificateurs ne pourra pas prévoir avec fiabilité l'apparition de défauts. Le système proposé aura donc bien pour vocation de diminuer l'apparition de la non-qualité et non de la supprimer totalement et nécessitera une analyse poussée pour chacun des défauts envisagés. Il y aura donc autant de systèmes « Classificateur-Observateur » que de défauts à surveiller, chacun étant capable de surveiller sa propre dérive et de déclencher un réapprentissage si besoin. Limiter l'apparition de la non-qualité a une conséquence directe sur la simplification et la rationalisation des flux. Ces derniers vont être réduits en nombre et en volume et leur pilotage va devenir alors moins compliqué.

IV.3 EVOLUTION REACTIVE DES REGLES DE PILOTAGE EN FONCTION DE L'ETAT DE SATURATION DE L'ATELIER



IV.3.1 Application et perspective d'implémentation

Cette partie n'a pas pu à l'heure actuelle être implémentée dans l'entreprise. De nombreuses règles de priorité sont mises en œuvre en fonction de critères principalement visuels tels que le nombre de lots dans une file d'attente par exemple. A cause de la multiplicité de ces règles et de la visibilité uniquement locale des opérateurs du poste, les décisions de priorisation sont parfois difficiles et souvent non adaptées réellement à la situation. L'apport du système de cartographie visuelle serait non négligeable en tant qu'aide à la décision offrant une vue dynamique plus complète de la situation de saturation de la zone en question.

En revanche, une application a été menée à partir de données simulées à l'aide du modèle présenté dans la partie III.3.1. C'est cette application que nous détaillerons ici. Les données ont été collectées en lançant différents scénarios de simulation pour trois valeurs différents de N_{default} , N_{wip} et de règles de pilotage (EDD, FIFO et SPT).

IV.3.2 Construction de la cartographie des états de saturation par clustering

Comme évoqué dans la partie III.3.8.2, l'algorithme K-means est utilisé dans le but de définir correctement les zones de comportement de la cartographie. Le traitement des données a été réalisé avec SCILAB qui contient nativement une version standard de l'algorithme K-means dans la librairie de fonctions du logiciel. Les paramètres choisis sont les suivants :

- Nombre de clusters : 4
- Distance : Euclidienne

La Figure 41 montre les résultats obtenus sur les données en utilisant l'algorithme K-means avec les réglages décrit ci-dessus.

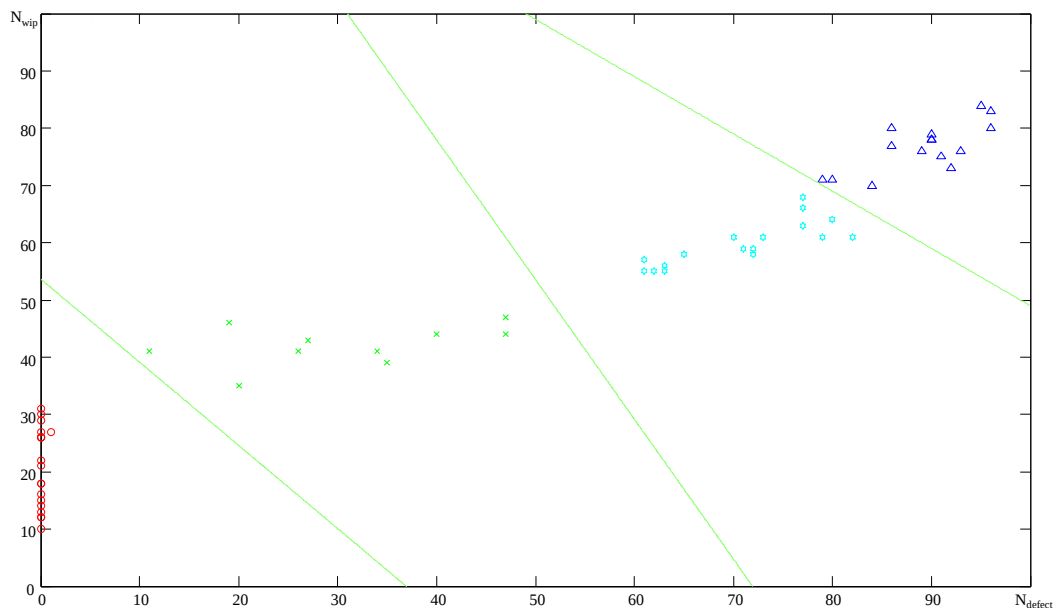


Figure 41 - Résultats obtenus avec l'algorithme K-means avec la distance Euclidienne

Les bornes entre les différentes zones sont tracées et semblent en accord avec la cartographie construite plus intuitivement dans la partie III.3.8.1.

L'étape suivante consiste à associer à chacune de ces zones la règle de pilotage (ou la consigne) la plus adaptée à fournir à l'opérateur.

IV.3.3 Association des règles de pilotage

Afin d'associer les règles de pilotage (ou des consignes) à la cartographie établie précédemment, nous avons réalisé un plan d'expériences de Taguchi dont l'enjeu a été d'identifier l'influence des variables $N_{\text{défaut}}$, N_{wip} et de la règle de pilotage sur le système. La définition des niveaux a conduit à choisir les modalités décrites dans le Tableau 31. Celles-ci l'ont été de manière à correspondre à des états strictement différents.

Tableau 31 - Niveaux des facteurs du plan d'expérience.

Facteurs	Niveaux
$N_{\text{Défaut}}$	Faible (30%), moyen (45%), fort (60%)
N_{WIP}	Faible, moyen, fort
Règle de pilotage	EDD, FIFO, SPT

D'un point de vue opérateur, la règle la plus adaptée au fonctionnement de l'atelier est la règle FIFO avec un traitement des lots dans l'ordre de leur arrivée sans besoin d'information supplémentaire. La règle EDD est aussi applicable aisément dès lors que les dates de départ sont clairement indiquées sur des palettes organisées de manière à pouvoir choisir n'importe laquelle. En revanche la règle SPT est plus difficile à mettre en place car, dépendant du temps de travail, il faut alors connaître, au moment du choix du prochain travail à réaliser, les temps gamme de tous les lots disponibles. Cette information disponible dans l'ERP n'est pas forcément affichée directement sur la fiche suiveuse des lots étant donné qu'il y en a une par poste de travail/ phase de la gamme. La mise en place de cette dernière règle nécessite donc la présence d'un système d'information performant. On peut imaginer éventuellement la présence de cette information directement sur le produit/lot lui-même, grâce à des technologies Auto-ID, qui permettront de la restituer au bon moment ou par consensus avec les lots/produits voisins.

L'objectif mesuré du plan d'expériences est le nombre de retards, conformément aux autres analyses menées précédemment. La Figure 42 montre les résultats sous la forme des graphiques d'effets de Taguchi.

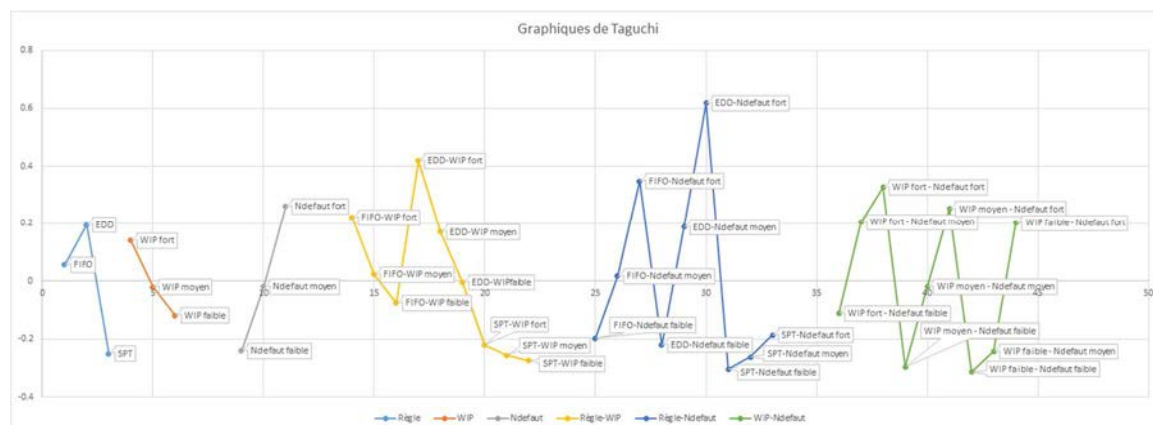


Figure 42 - Plan d'expériences de Taguchi présentant l'influence du WIP, de Ndefaut, des règles de pilotage et de leurs interactions.

Les 3 premières courbes montrent l'impact des 3 facteurs individuellement. Analysées seules, elles confirment que le cas le plus favorable pour ne pas avoir de retards et d'utiliser une règle SPT dans des conditions de Work In Process faible et avec peu de défauts. Cependant lorsqu'on s'intéresse aux interactions (les 3 courbes de droite), on se rend compte que, dans des conditions de Work In Process et de taux de réparations imposées, la règle de pilotage la plus adaptée n'est pas toujours la même. On peut construire le Tableau 32 qui présente la règle la plus adaptée à une situation donnée par la valeur des deux indicateurs considérés.

Tableau 32 - Meilleure règle de pilotage en fonction des valeurs d'indicateurs imposés.

		N _{Défaut}		
		30%	45%	60%
WIP	Fort	SPT	SPT	SPT
	Moyen	EDD	SPT	SPT
	Faible	EDD	EDD	SPT

La limite correspondant au basculement entre EDD et SPT est donc la courbe verte centrale sur la Figure 41. Les deux autres courbes sont donc des limites d'état de saturation mais qui ne justifient pas le passage d'une règle de pilotage à l'autre parmi les 3 règles envisagées. On peut en revanche associer à chacune des zones des alertes et des recommandations telles que présentées dans le Tableau 33.

Tableau 33 - Préconisation de travail en fonction des zones de saturation de l'atelier.

	Zone 1	Zone 2	Zone 3	Zone 4
Règle de pilotage préconisée	EDD	EDD	SPT	SPT
Recommandations		Augmenter la capacité du poste de réparation.	Attribuer une personne au suivi des pièces potentiellement en retard.	Définir les lots à prioriser au détriment des autres.

La recommandation de la zone 2 a pour objectif de réduire rapidement le nombre de réparations en cours et donc de déplacer le point représentant l'état de saturation de l'atelier vers la gauche afin de revenir en zone 1. Dès lors que le point repasse en zone 1, on peut revenir à une capacité normale pour le poste de réparation.

La recommandation de la zone 3 a pour objectif d'agir en plus sur le nombre de retards. L'action combinée sur les retards et sur les réparations aura tendance à déplacer le point simultanément vers le bas et vers la gauche avec pour objectif de rejoindre la zone 2 rapidement.

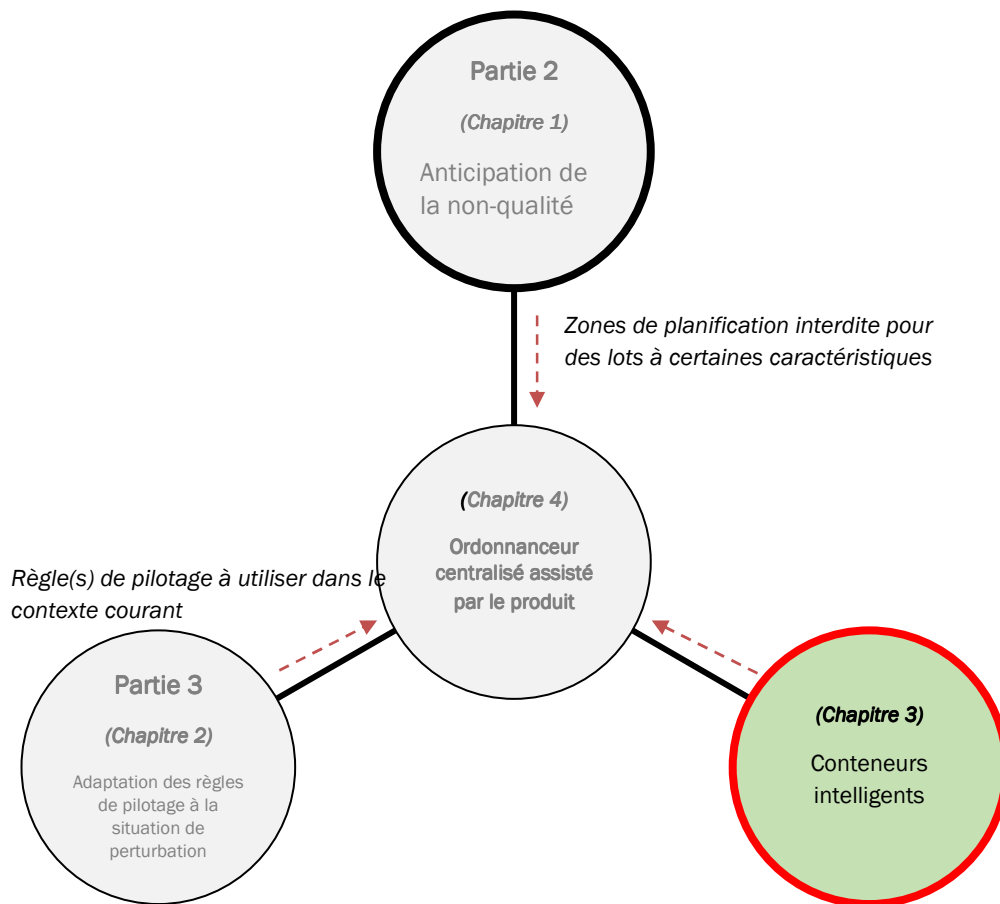
La recommandation de la zone 4 est plutôt palliative tant la situation est grave. Elle tente de « sauver » les lots d'une importance haute en abandonnant temporairement le traitement des autres. Cette zone, matérialisée en rouge sur la cartographie, est idéalement à ne jamais atteindre !

IV.3.4 Résultats et discussions

Dans ce chapitre, nous avons travaillé sur des données simulées avec ARENA sur le modèle présenté dans la partie III.3.1 afin de construire une cartographie. Nous avons déterminé 4 clusters de données qui correspondent à 4 zones de saturation de l'atelier. Grâce à un plan d'expériences réalisé sur les mêmes données, nous avons pu identifier les 2 règles les mieux adaptées aux différentes situations et les répartir sur les 4 zones représentées dans le Tableau 33. Comme évoqué dans la partie III.3.7, l'avantage principal

de cette méthode est donc son aspect simple et visuel adapté à l'utilisation intuitive sur le terrain. Elle est clairement tournée vers l'application facile et rapide ainsi que vers la communication et la responsabilisation des ouvriers.

IV.4 CONTAINER INTELLIGENT



IV.4.1 Choix de la stratégie d'identification

IV.4.1.1 Avantages d'une identification à la pièce

Dans les entreprises fabriquant du sur-mesure, notamment sur des produits à forte valeur ajoutée, l'identification à la pièce est souvent perçue comme incontournable. En effet, parce que chaque pièce fabriquée est différente, et parce que la diversité des produits a tendance

à augmenter, l'identification à la pièce permet la traçabilité unitaire en matière d'historique et de localisation.

Avant l'arrivée de la technologie RFID, de nombreuses autres technologies d'identification avaient été testées sans succès et de ce fait l'entreprise en restait toujours à l'utilisation de fiches suiveuses. Avec les promesses de la RFID, cette technologie attractive méritait d'être étudiée en profondeur pour tirer parti de ses nombreux avantages. Une étude d'adaptation technologique a donc été menée.

IV.4.1.2 Adaptation technologique

Les supports RFID sont de natures diverses et variées. L'utilisation de la technologie ainsi que les contraintes liées au produit à identifier vont conditionner le choix du support. De par le fait que le système devra fonctionner en boucle ouverte, c'est-à-dire que chaque panneau aura un tag non récupérable et qu'il faudra donc acheter autant de tags que de pièces à fabriquer, l'idée initiale était de se tourner vers le tag le moins cher possible : l'inlay, initialement conçu pour être directement collé sur la pièce à suivre.

Pendant, même si cette façon de taguer la pièce est des plus simples, elle pose le même problème de décollement/recollement qui a pu exister avec les codes à barres et qui avait finalement conduit à l'abandon de ces derniers au profit des fiches suiveuses.

Si l'on décide d'insérer le tag à l'intérieur du panneau pour éviter ce problème, il est nécessaire de le dissimuler en garantissant une qualité de surface irréprochable dans le temps, et ce de manière facilement industrialisable.

De nombreux essais ont été menés, mais toujours, après séchage et/ou tests de vieillissement, l'emplacement du tag était ou redevenait visible à l'œil alors qu'aucune déformation n'était perceptible au toucher. En vue de mieux apprécier la sensibilité de l'œil et de nous faire une idée de la déformation physique acceptable pour rester invisible, nous avons fait établir un profil de surface par le Centre Interrégional de Métrologie de la maison de l'entreprise de l'Yonne présenté en Figure 43.

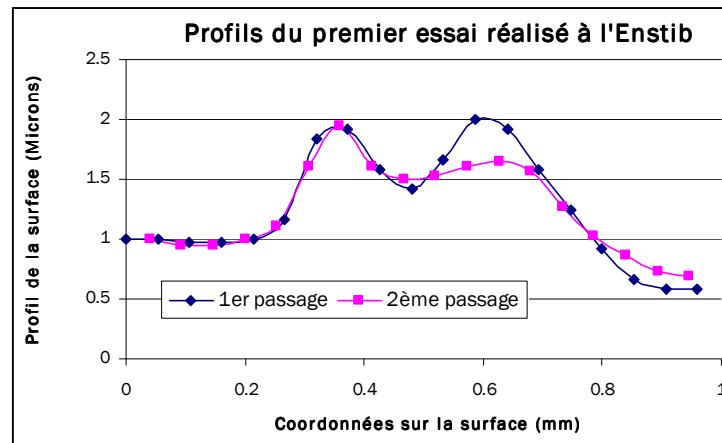


Figure 43 : Profil réel de l'état de surface.

Le résultat obtenu montre qu'il s'agit d'une bosse de laque qui pourrait être due à la colle qui gonflerait légèrement ou à un effet de bord dans le cas où un micro-creux existerait entre l'insert et le panneau. La largeur du liseré visible a pu être mesurée à environ 0,7 mm et la différence de planéité maximale à environ 1,3 microns. Cette différence extrêmement faible et pourtant visible à l'œil nu témoigne de l'ampleur du problème car elle signifie qu'il faut réaliser une planéité de surface quasiment parfaite pour que l'insertion ne soit pas visible.

Un plan d'expériences a donc été mené de manière à analyser l'influence des facteurs jugés critiques sur la qualité de l'insertion (voir Annexe en partie VII). L'enjeu de ce plan d'expériences était donc de pouvoir trouver une solution d'implantation de la technologie RFID afin de l'utiliser sans décollements/recollements successifs des inlays. Utiliser des tourillons comme « rebouchage » est une solution couramment utilisée lorsqu'une erreur de perçage est faite et qu'il existe un tourillon du diamètre correspondant. Toutefois, de par son côté aléatoire, cette solution n'est utilisée que dans des zones qui ne sont pas destinées à être visibles (derrière les charnières par exemple). Cependant, la faisabilité des expériences est bien démontrée par les « réparations » qui ont déjà été faites dans l'entreprise.

Etant donné que le moindre grain de poussière est mis en relief par la laque, l'objectif n'est donc pas de n'avoir aucune marque à l'endroit d'insertion, mais d'en obtenir une suffisamment discrète pour ne pas engendrer de refus client. La mesure des résultats passe par l'intermédiaire d'une échelle qualitative visuelle allant de 0 à 5 dont seuls les niveaux 0 et 1 seront acceptables pour le client. Pour être plus précis, il aurait été possible d'envisager le même type de mesure que celui présenté en Figure 43, mais cette solution

onéreuse et chronophage ne se justifie pas à cause du l'appréciation principalement visuelle de la non-qualité.

Suite à l'échec de ce nouveau plan d'expériences, qui ne fournira pas non plus de solution acceptable pour le client et face à l'incapacité des experts de fournir d'autres potentielles explications à ces incontournables défauts de surface, la technologie RFID finit par être mise de côté. D'autres solutions d'identification sont alors envisagées telles que le marquage laser, le marquage par micro-percussion, l'encre invisible mais chacune présente des contraintes rédhibitoires.

IV.4.1.3 Inconvénients d'une identification à la pièce

Au vu des essais décrits précédemment et en annexe (partie VII), il est évident que l'inconvénient principal est donc la difficulté de réaliser techniquement cette identification à la pièce.

IV.4.2 L'identification intuitive du « noyau élémentaire »

Il a donc fallu rechercher une granularité d'identification différente en s'intéressant au « lot » plutôt qu'à la « pièce ». Des recherches sur le terrain ont permis de valider le fait que des groupes de pièces, pourtant différentes, suivaient ensemble la même gamme de production. L'idée était d'isoler les caractéristiques similaires de ces pièces de manière à prévoir à l'avance ces regroupements qu'on nommera par la suite « noyaux élémentaires ». Les fiches suiveuses étaient d'ailleurs naturellement éditées en fonction de ces groupes de pièces à caractéristiques similaires. Les caractéristiques retenues ont été les suivantes :

- Semaine
- Client
- Famille de pièce (incluant le modèle et la finition)
- Couleur
- Circuit court

L'ERP de l'entreprise ne permettant pas nativement ce type de regroupement, un programme « lotisseur » a été développé sur la base des techniques classiques de « groupements analogiques » issues de la technologie de groupes (TG) apparue dès 1959 (McAuley, 1972). Il s'agit d'une méthode qui consiste à regrouper les pièces pour les concevoir et les fabriquer en tirant profit de leurs analogies/morphologie. Ces analogies peuvent être définies à partir d'une pièce imaginaire qui présente tous les détails géométriques existant dans les pièces qui composent une famille. Il est donc nécessaire de faire une analyse géométrique des pièces en tenant compte de leurs gammes d'usinage. La

TG est un bon outil pour rationaliser et ordonner la production, des pièces aux différents stades de leur réalisation. Nous l'avons donc utilisé pour formaliser des « noyaux élémentaires » par regroupement des pièces à caractéristiques similaires sous une référence de lot unique.

Au niveau du poste de travail ces noyaux élémentaires peuvent être regroupés suivant des pré-requis machine (par exemple, au laquage, mieux vaut regrouper les lots de même couleur tandis qu'à l'usinage, mieux vaut regrouper les lots de la même épaisseur qui sont donc à usiner dans le même panneau) en lot de fabrication. Ce sont donc ces lots de fabrication qu'il faudra par la suite planifier sur les postes de travail.

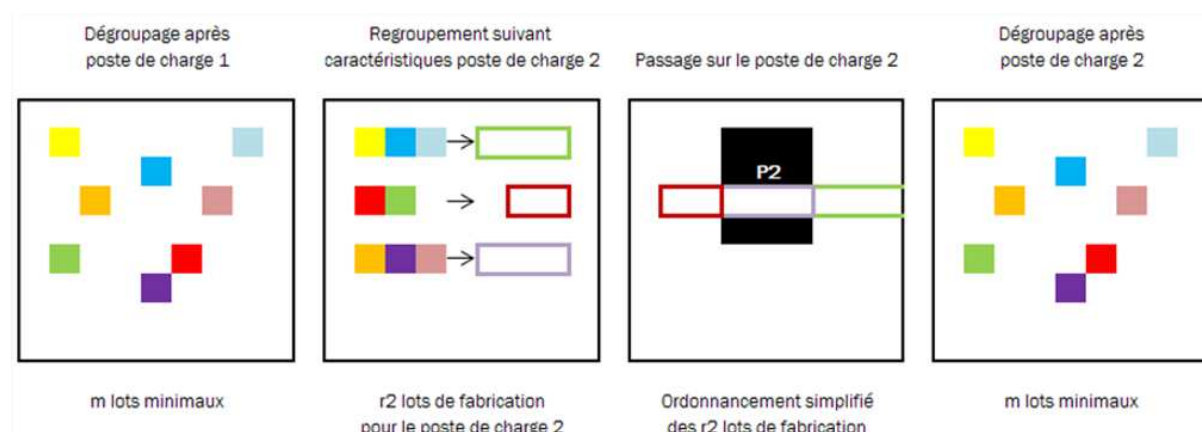


Figure 44 : Regroupement des noyaux élémentaires en lots de fabrication.

Ce nouveau niveau de granularité oblige à gérer les éventuelles « sorties » de lots. En effet, principalement pour des raisons qualité, certaines pièces sont amenées à quitter leur lot d'origine pour emprunter une gamme différente ou réaliser une reprise dans leur propre gamme. Cette division physique du lot est informatisée par le formulaire de « déclaration d'une réparation » qui confère un nouveau numéro de lot à la(les) pièce(s) qui vont être séparées du lot initial et qui permet l'impression du nouvel identifiant unique. Ce processus est appelé « généalogie » en raison de sa représentation possible sous forme d'arbre (présenté sur la Figure 45) rappelant l'arbre de généalogie d'une famille. Les lots prennent d'ailleurs le nom de lot parent et de lot fils.

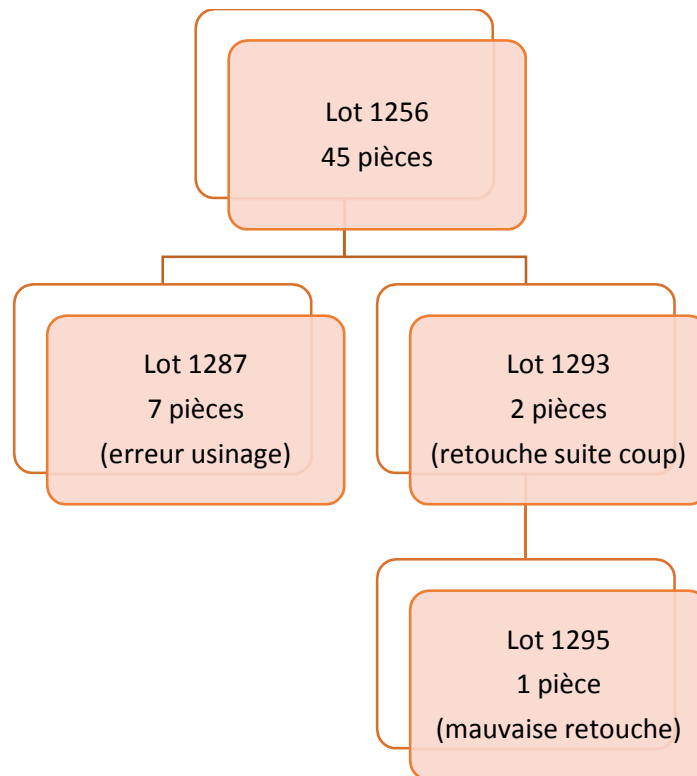


Figure 45 - Exemple de généalogie

Sur cet exemple, le lot initial de 45 pièces (Id 1256) a d'abord été divisé en 2 lots immédiatement après l'usinage à cause d'une erreur de découpe. Il a alors donné naissance à un nouveau lot de 7 pièces (Id 1287). Les 38 pièces restantes ont continué leur chemin dans le lot 1256. Lors de son passage dans l'atelier polissage, 2 pièces ont dû être réparées et ont donc formé ensemble un nouveau lot fils (Id 1293). Les 36 autres pièces ont pu continuer leur chemin. On a donc à cet instant 3 lots différents enfants du lot initial 1256 qui circulent indépendamment dans l'atelier. Si une des deux réparations n'est pas correcte et doit être reprise, ce sera alors 4 lots qui circuleront dans l'atelier dont 2 lots unitaires (les lots 1293 (qui a donc perdu une pièce) et 1295).

Ce système de généalogie permet de conserver l'information de traçabilité au sens localisation comme identification tout au long de la chaîne de production.

IV.4.2.1 Avantages du suivi au lot

La traçabilité au lot permet aux opérateurs de n'interagir avec le MES que deux fois (lors du début de travail et fin de travail) par lot, ce qui engendre un gain de temps (et de mouvement) non négligeable par rapport à une déclaration de production à la pièce.

Le lotissement n'intervient pas après le calcul des besoins mais dès la saisie de commande par la logistique ce qui permet encore plus de réactivité pour le déclenchement de l'usinage.

IV.4.2.2 Inconvénients du suivi par lot

En identifiant les lots et non plus les pièces, il n'est plus possible de connaître exactement le détail du passage d'une pièce sur une machine. Par exemple, concernant la remontée d'informations, nous aurons un temps ou une consommation matière pour un lot et non par pièce. Il faudra alors utiliser des proratas en fonction du nombre de pièces ou de la surface travaillée afin de revenir à des temps ou consommations unitaires. Cette dégradation de la précision n'est cependant pas limitante face au gain d'information que le système représente.

IV.4.2.3 Perspectives

Une constatation sur le terrain concerne les lots conséquents qui vont être répartis sur plusieurs palettes. Ainsi, les palettes peuvent se retrouver séparées sur plusieurs postes de travail et donc à des étapes/phases différentes de la gamme de fabrication. Ce point n'est pas forcément bloquant d'un point de vue traçabilité purement informatique car il suffit d'identifier les deux palettes avec le même « id » de lot pour que chacune des palettes puissent être « badgées » lors du début et de la fin de travail. Cependant, le contenu exact de chacune des palettes n'est pas connu et il n'est, par conséquent, pas possible de retrouver une pièce sans avoir à « dépiler » les deux palettes. Une solution serait d'intégrer à ce lotisseur la notion de volume physique en limitant le nombre maximum de pièces (ou la surface ou le volume) pour s'adapter mieux au contenant et créer un seul lot par contenant. Cette solution permettrait aussi de catégoriser les types de pièces par contenant (par exemple les petites pièces sur une palette et les grandes sur une autre) de manière à faciliter le travail de certains postes de travail.

IV.4.3 La notion de container intelligent

Au-delà de la simple identification, la volonté de rendre le container intelligent s'est traduite par de multiples raisons :

- La traçabilité avec le suivi dans l'atelier d'un groupe de produits ayant la même gamme de fabrication
- Le partage d'informations de priorisation telles que sa date de livraison ou son degré d'urgence

- L'aiguillage automatique par l'embarquement de connaissances sur les chemins (gammes de fabrication) que les containers doivent emprunter dans l'atelier et les différentes alternatives possibles
- La communication avec les machines et autres systèmes tels que le système de prévision de la qualité présenté précédemment.
- La collecte d'informations « terrains » au fil de la fabrication telles que le temps passé dans un four de séchage par exemple.

Dans notre vision du système, c'est ce container même qui est au cœur de la décision. Il est possible d'envisager différentes manières d'organiser la décision. Nous en proposons un exemple sur la Figure 46 où le container ne signale sa présence que si les « conditions qualité » sont réunies, permettant ainsi de ne pas surcharger le travail de l'ordonnanceur qui ne travaille alors qu'avec les lots présents dans la zone et déjà validés par le système qualité. Ces conditions lui permettent donc d'économiser du temps de traitement et d'envisager qu'un simple algorithme traditionnel de programmation linéaire puisse réaliser cette fonction qui semblait particulièrement compliquée sans cette décomposition des tâches entre les différents systèmes acteurs. C'est en revanche à la charge de l'ordonnanceur de s'assurer de la règle de pilotage à utiliser pour son ordonnancement et de demander les informations nécessaires aux lots qui se sont signalés comme disponibles à l'ordonnancement. La nature des informations demandées varient donc en fonction de la règle de pilotage en cours préconisée par la cartographie visuelle. Une règle FIFO ne requerra que la date d'entrée dans la zone à ordonnancement dynamique tandis qu'une règle EDD requerra la date d'expédition souhaitée par exemple.

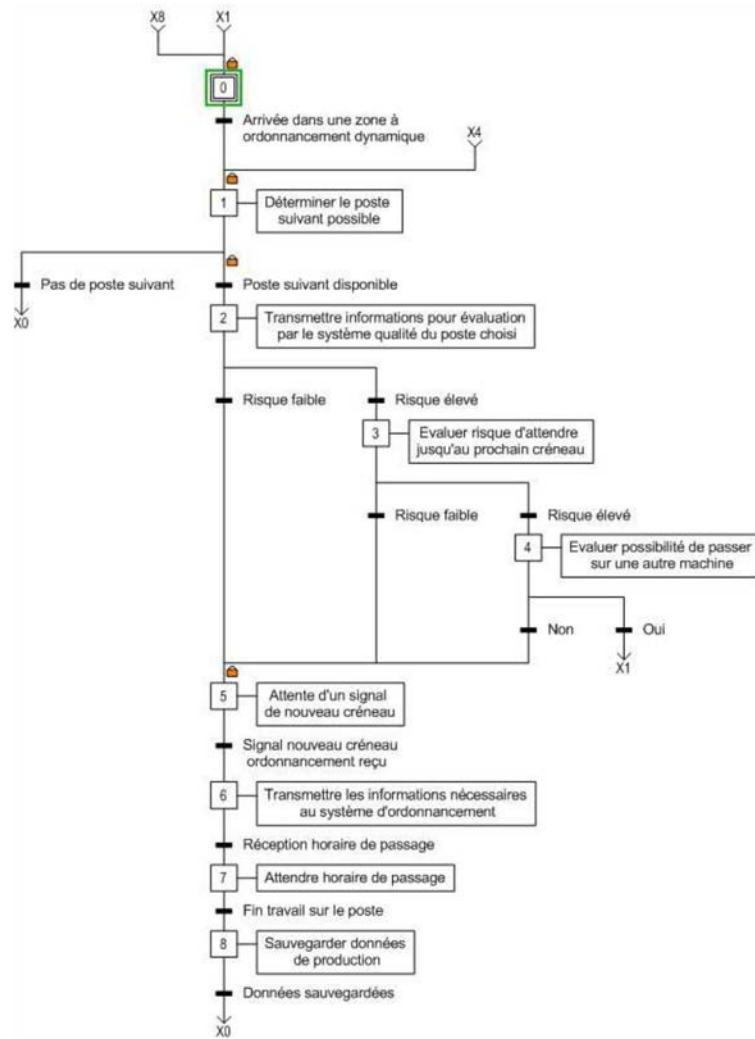


Figure 46 - Exemple de logigramme d'un container intelligent

IV.4.4 Résultats et discussions

Même si la volonté première de l'entreprise était tournée vers l'identification à la pièce, les contraintes techniques de lien entre le produit et son identifiant infotronique ont conduit à reconsidérer l'unité de gestion. L'observation sur le terrain a permis de définir le « noyau élémentaire » à considérer par groupage des pièces de caractéristiques équivalentes que nous pouvons considérer comme un « container » puisque les pièces suivront la même gamme et donc le même « chemin » dans l'atelier. C'est la volonté de donner à ces containers une capacité d'interaction avec les systèmes voisins qui permet d'envisager la simplification du problème d'ordonnancement au robot par application du concept de « contrôle par le produit ». Dans l'utilisation actuelle du système, le container a déjà atteint un certain niveau d'intelligence illustré sur la Figure 47 (Meyer, Främling, & Holmström, 2009) car c'est lui qui fournit les informations de production à l'utilisateur (standards,

consignes particulières, etc...) et est capable de mémoriser une information saisie par l'opérateur pour la diffuser plus tard.

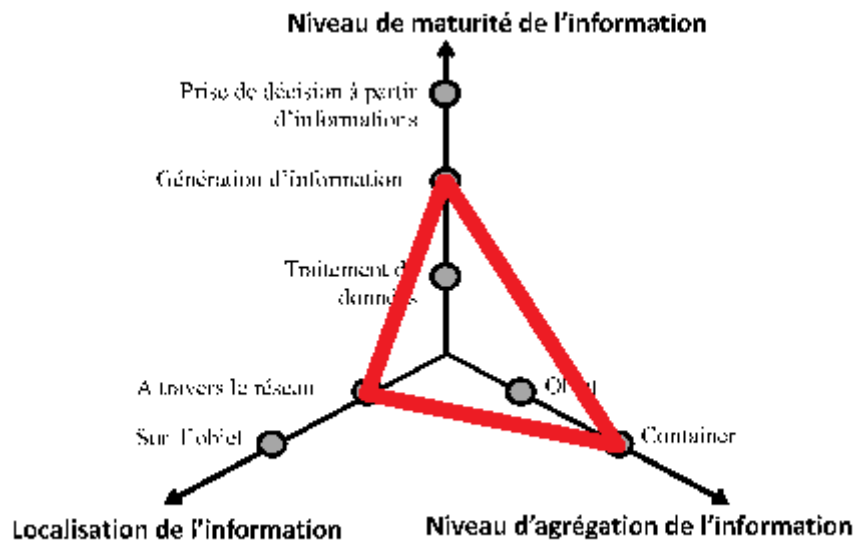


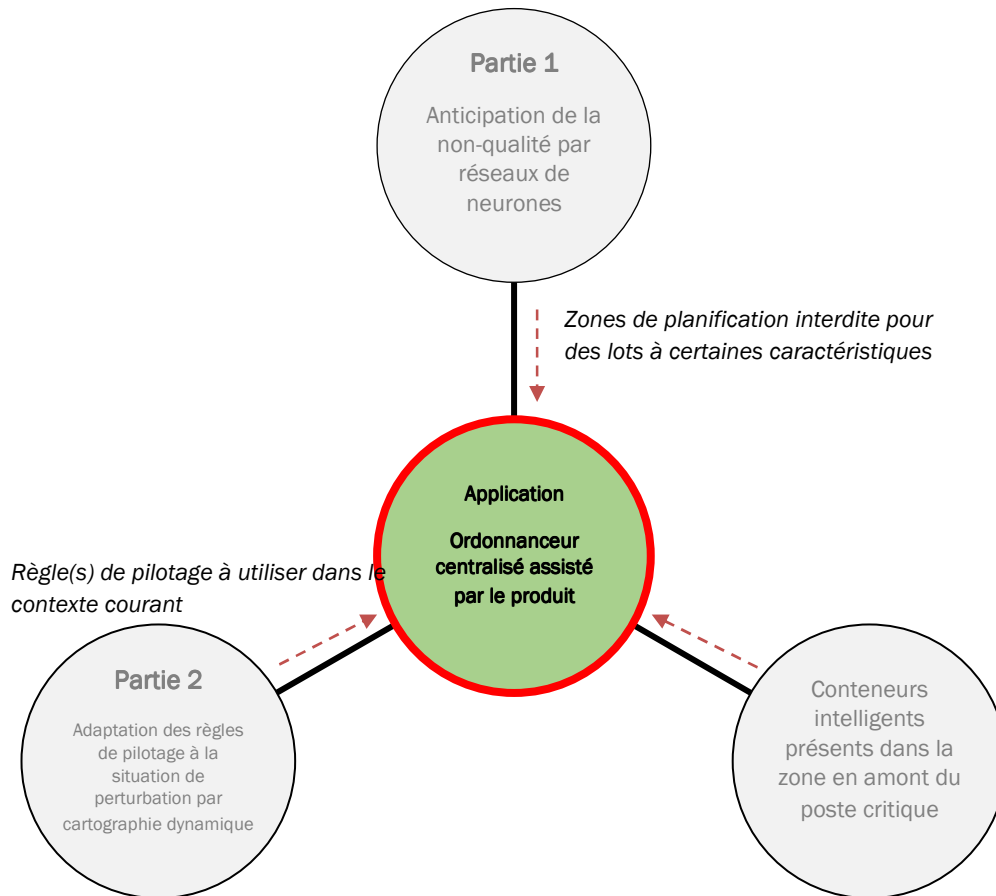
Figure 47 – Description du niveau d'intelligence du système en place

Dans le futur, il est envisagé d'utiliser les technologies infotroniques pour identifier ces containers afin de pouvoir changer la localisation de l'information.

Le logigramme présenté en Figure 46 (et en référence à la Figure 47) nous permet déjà de situer le *niveau de maturité de l'information* au niveau « prise de décision à partir d'information » bien que le système lui-même ne soit capable que de proposer des alternatives et laisse encore l'opérateur choisir. Le niveau « génération d'informations » est quant à lui déjà mis en œuvre lors de l'interactivité du container avec le système d'ordonnancement puisqu'il calcule et transmet sa disponibilité en fonction d'une demande auprès du système qualité et des informations qu'il porte lui-même.

On peut donc par extension qualifier notre noyau élémentaire d' « intelligent ».

IV.5 ORDONNANCEUR CENTRALISE ASSISTE PAR LE PRODUIT



IV.5.1 Descriptif du problème

Nous considérons toujours ici que le poste goulot de l'atelier est le robot de laquage et que l'optimisation de son travail est une piste de progrès indéniable pour l'entreprise. Un stock tampon en amont de celui-ci permet de ne jamais l'affamer et les postes de travail en aval vont être asservis à son rythme. De plus, le robot est l'un des principaux responsables de la non-qualité comme présenté précédemment tout en étant aussi lui-même un poste de réparation de ces non-qualités. Son fonctionnement en nombre de tables figé oblige à penser un ordonnancement non plus à l'unité « temps » (heures, minutes, ...) mais à l'unité « table ». Le principe de retournements successifs (laquage du dos, puis de la face, puis encore éventuellement un vernis) oblige à imaginer des algorithmes de type fenêtres

interdites et à placer des lots de manière spatiale dans les fenêtres restantes autorisées comme présenté sur la Figure 48.

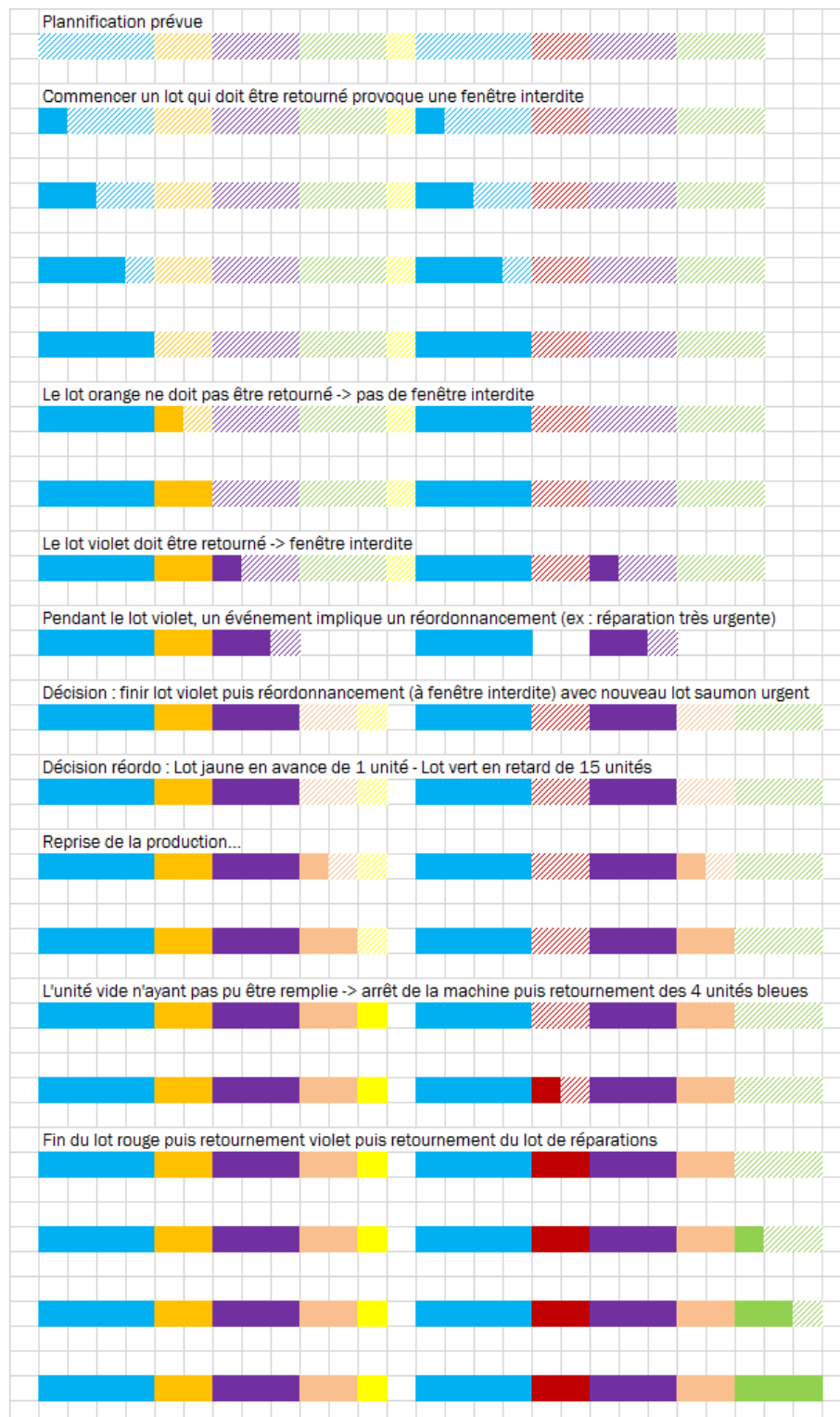


Figure 48 - Illustration de la problématique d'ordonnancement sur le robot de laquage.

Si on ajoute à cela la nécessité de réactivité pour pouvoir rajouter en dynamique la réparation d'un lot de pièces, ou s'adapter à un manque de matière première, le problème

devient particulièrement complexe. Pour autant, si le nombre de lots arrivant dans la zone de stockage n'est pas trop important, une solution analytique pourrait être utilisée. Dans la suite nous proposons une première solution mathématique puis une deuxième qui mettrait en œuvre une application du contrôle par le produit.

IV.5.2 Solution par programmation linéaire

Les problèmes d'ordonnancement de complexité raisonnable sont souvent abordés avec des programmes linéaires. Nous proposons ici un programme linéaire qui a été implémenté avec Xpress IVE et qui a pu permettre de résoudre de manière optimale les problèmes d'ordonnancement à faible échelle sur le robot de laquage.

IV.5.2.1 Hypothèses

- Uniquement sur le robot de laquage.
- Contraintes spatiales plutôt que temporelles.
- Déchargements puis rechargements de lots non autorisés.

IV.5.2.2 Paramètres

- N_l : nombre de lots à produire
- B_i avec $i \in [1, N_l]$: nombre de passages du lot i sur la machine
- N_t : nombre de tables avant planification du passage suivant sur la machine
- $M=1000$: paramètre de grande valeur nécessaire à la détermination des contraintes alternatives
- e_i avec $i \in [1, N_l]$: occupation spatiale (nombre de tables) du lot i
- d_i avec $i \in [1, N_l]$: la date d'exigibilité du lot i

IV.5.2.3 Variables

- Soit $s_{i,j}$ le numéro de table pour le début du $j^{\text{ème}}$ passage du lot i sur la machine
- Soit $z_{i,j,k,l}$ les variables binaires de décision pour assurer l'unicité de la ressource. Elles correspondent aussi à la notion de précédence (0 si le $j^{\text{ème}}$ passage du lot i précède le $l^{\text{ème}}$ passage du lot k)
- Soit p la variable de priorité maximale assurant la linéarisation de notre objectif de minimisation du s_{i,B_i} maximal
- Soit r_i les variables binaires correspondant au retard (1) ou non (0) du lot i

IV.5.2.4 Contraintes

Les passages doivent être faits dans l'ordre et être espacés de N_t tables pour éviter les déchargements/rechargements :

$$\forall i \in [1, N_l] \quad \forall j \in [1, B_i - 1] \quad s_{i,j+1} = s_{i,j} + N_t \quad (16)$$

La machine ne peut travailler que sur un lot à la fois (contraintes alternatives d'unicité de la ressource) :

$$\forall i \in [1, N_l] \quad \forall j \in [1, B_i] \quad \begin{cases} \forall k \in [1, N_l], k \neq i \\ \forall l \in [1, B_k], l \neq j \end{cases} \quad (17)$$

$$\begin{cases} M \times z_{i,j,k,l} + s_{k,l} \geq s_{i,j} + e_i \\ s_{k,l} \leq s_{i,j} - e_k + M \times (1 - z_{i,j,k,l}) \end{cases}$$

Contraintes de linéarisation de l'objectif. Dans cette formule, le e_i est un paramètre de réglage que l'on a fixé à une valeur importante pour bien privilégier la date de sortie de la dernière table, sans quoi, l'équation va privilégier la date de placement du dernier lot.

$$\forall i \in [1, N_l] \quad s_{i,B_i} \leq p + e_i \quad (18)$$

Le lot est soit en retard, soit à l'heure (contraintes alternatives de détermination du retard) :

$$\forall i \in [1, N_l] \quad \begin{cases} M \times (1 - r_i) + s_{i,B_i} + N_t + e_i \geq d_i \\ s_{i,B_i} + N_t + e_i \leq d_i + M \times r_i \end{cases} \quad (19)$$

IV.5.2.5 Objectifs possibles

La fonction objectif de notre programme linéaire peut être de différentes formes. On peut envisager que ce soit un paramètre directement fourni par la cartographie et donc adapté à l'état de saturation de l'atelier. Nous présentons ici 3 possibilités correspondant aux paramètres, variables et contraintes décrits dans les paragraphes précédents.

- Minimiser la somme de toutes les variables de début des derniers passages :

$$\text{Minimiser } \sum_{i \in [1, N_l]} s_{i,B_i} \quad (20)$$

- Minimiser le s_{i,B_i} maxi (équivalent au cmax) :

$$\text{Minimiser } p \quad (21)$$

- Minimiser le nombre de retards :

$$\text{Minimiser } \sum_{i \in [1, N_l]} r_i \quad (22)$$

IV.5.2.6 Temps de calcul

Lorsqu'on choisit l'objectif portant sur le nombre de retards (23), la résolution du problème pour 8 lots prend environ 7 secondes. La même résolution pour un lot de plus prend environ 22 minutes. Par régression polynomiale d'ordre 2, on peut donc s'attendre à ce que la résolution d'un problème à 10 lots prenne environ 36h ce qui n'est évidemment pas envisageable dans les conditions recherchées de réactivité.

Lorsqu'on cherche à ne travailler que le C_{max} (22), c'est-à-dire que l'on cherche uniquement à réduire le temps de passage des lots disponibles sur la machine en question en estimant que leur ordre d'arrivée est lié à leur ordre de départ, résoudre un problème à 9 lots prend environ 10 secondes, il faut 3 minutes pour résoudre un problème à 10 lots.

Si l'on utilise uniquement la minimisation de la somme des variables de début des derniers passages (21), les temps de calculs sont encore améliorés et la solution fournie, même si ce n'est pas la meilleure à cause de l'approximation d'optimisation, est exploitable.

Ce programme linéaire montre bien la complexité du problème et interdit la recherche simple d'une solution exacte dès lors que le nombre de lots à ordonnancer excède la dizaine sans passer par des heuristiques plus compliquées.

L'outil qualité abordé dans la partie II avec son application détaillée dans la partie IV.1 va permettre de limiter le nombre de lots à ordonnancer pour simplifier la dimension du problème conformément au logigramme de comportement d'un container intelligent présenté en Figure 46. Cette solution simplifiée est présentée dans la partie suivante.

IV.5.3 Une application du contrôle par le produit

En utilisant la cartographie dynamique proposée dans la partie III, le système global est capable de connaître la règle de pilotage la plus pertinente en fonction de l'état de saturation de l'atelier. Une solution simpliste serait donc de déployer une batterie d'algorithmes utilisant ces règles pour réaliser un ordonnancement en exploitant toujours les informations portées par les containers intelligents (application du concept de « système contrôlé par le produit »). Ainsi un ordonnancement FIFO à partir des dates d'arrivée dans la zone à ordonnancement dynamique ne prendra que quelques secondes de calcul même si le nombre de lots est conséquent. De la même manière, un ordonnancement EDD ne nécessitera que le calcul des dates de fin de travail souhaitée pour le poste concerné, tout comme un ordonnancement SPT ne nécessitera que l'évaluation des temps gamme des lots concernés. On peut donc imaginer un système de basculement présenté sur la Figure 49 entre le programme linéaire qui va pouvoir déterminer la solution optimale tant que les conditions le permettent (peu de lots validés en terme de qualité dans la zone à ordonnancement dynamique) et le panel d'algorithmes simples parmi lesquels la cartographie choisira le plus adapté à la situation de saturation de l'atelier.

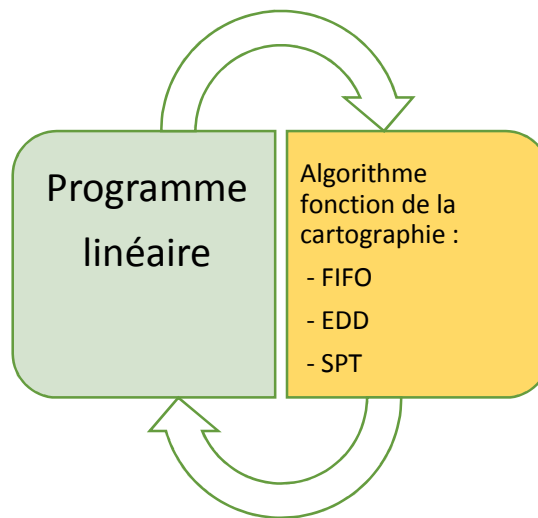


Figure 49 - "Basculement" entre le programme linéaire et les algorithmes simples.

IV.5.4 Résultats et discussions

L'ordonnanceur présenté dans cette partie n'est donc, à ce jour, qu'un système de « basculement » entre un programme linéaire et un panel d'algorithmes simples. Le programme linéaire permet de trouver la solution optimale à un problème d'ordonnancement mais il n'est capable de fournir des solutions dans un temps acceptable uniquement lorsque le nombre de lots à ordonnancer est suffisamment petit. Dès lors que le nombre de lots présents dans la zone à ordonnancement dynamique et validés en niveau de qualité excède une certaine valeur (une dizaine), le système permet de basculer sur le panel d'algorithmes parmi lesquels il faut donc choisir le plus pertinent en fonction de l'état de saturation de l'atelier. Cette information est fournie par la cartographie présentée dans la partie III.

Partie

V

V CONCLUSION ET PERSPECTIVES

Depuis plus de 40 ans, les méthodes de « juste à temps » et de la « qualité totale » sont perçues comme fortement liées. Cette expérience dans l'entreprise Acta-Mobilier, fabricant de panneaux laqués haut de gamme, nous en a fait vivre une application réelle et nous a conduits à travailler conjointement sur la maîtrise de la qualité et sur la maîtrise des flux de production.

Le premier point a été particulièrement intéressant à aborder du fait de la complexité de l'analyse de l'apparition des non-qualités sur le robot de laquage. Là où les experts échouaient, des méthodes « intelligentes » telles que les réseaux de neurones ont pu mettre en évidence des facteurs influents et des interactions insoupçonnées. Ce système d'anticipation de la non-qualité a pu permettre le réglage des processus de fabrication en vue de limiter fortement l'apparition de certains défauts. Toutefois, même des systèmes d'apprentissage élaborés ne parviennent pas à anticiper l'apparition d'autres défauts trop liés au comportement des acteurs, tels que les coups par exemple. Ces derniers devraient être abordés d'une toute autre manière (envisager des transferts mécaniques par exemple) et constituent donc de la « non-qualité résiduelle ».

La simplification des flux qui a résulté de ce premier projet, en termes de nombre ou de densité, a permis de simplifier le problème de pilotage des flux. L'outil POTER a aussi été développé dans ce but même si les premières fonctionnalités déployées sur le terrain ont d'abord été tournées vers la traçabilité et la mise à disposition de la bonne information au bon moment. Le projet en lui-même comprend une multitude d'améliorations telles que les OF uniques par lot équipés de code à barres 2D, le déploiement des douchettes pour identifier les lots rapidement, la déclaration automatique d'anomalies, etc... Ces concepts ont eu pour conséquence une réelle amélioration organisationnelle qui permet de réduire fortement les temps habituellement alloués à l'identification ou à la recherche de pièces. Le concept de cartographie n'a pu être implémenté que sur des données simulées pour le moment mais est en cours de mise en place. Elle permettra de fournir aux opérateurs et au système d'ordonnancement des informations utiles respectivement à la compréhension de l'état de saturation de l'atelier et à l'adaptation des règles de pilotage. Son principal avantage est lié au management visuel puisqu'elle est particulièrement adaptée à l'utilisation intuitive et responsabilise les opérateurs sur le terrain.

Les derniers concepts abordés dans les parties IV.4 et IV.5 de cette thèse sont toujours en cours de développement dans l'entreprise et mon embauche en tant que responsable du système d'information ainsi que mon contrat de chercheur associée au CRAN vont me permettre de les poursuivre, tant de manière industrielle que scientifique, dans le cadre d'une prochaine thèse. Du côté des containers intelligents qui portent l'information dynamiquement, nous luttons contre la myopie qui existait autour du robot de laquage mais aussi sur les autres postes de l'entreprise. En continuant d'augmenter le niveau d'intelligence du système (maturité, localisation et niveau d'agrégation) notamment grâce à l'utilisation de technologies infotroniques, nous envisageons de mettre en place une réelle vision dynamique des flux de la même manière que nous avons mis en place une vision dynamique de la qualité. En effet, à la manière de notre système qualité qui est composé d'un observateur (carte de contrôle) et d'un régulateur (algorithme de réapprentissage du réseau de neurones concerné), notre système de pilotage des flux sera lui aussi composé d'un observateur (cartographie visuelle) et d'un régulateur (réordonnancement dynamique). La mise en service de ces deux systèmes dont les fonctionnements seront intimement liés peut permettre l'automatisation des processus de Lean Management dans l'entreprise, ce qui est une problématique intéressante pour une prochaine thèse CIFRE.

Partie

VI

VI FIGURES, TABLEAUX ET REFERENCES

VI.1 FIGURES

Figure 1 - Stand Peugeot fabriqué par Acta pour un mondial de l'auto	6
Figure 2 - Le service "couleur à la demande" permet de maintenir la compétitivité de l'entreprise.....	6
Figure 3 - Chaîne de production de l'entreprise Acta-Mobilier.....	7
Figure 4 - Méthodologie d'analyse du problème industriel.....	9
Figure 5 - Dérivation des exigences et contraintes pour définir le problème industriel.....	10
Figure 6 : Multiplicité des « boucles »	18
Figure 7 - Résumé des pistes de solutions groupées par axe de travail.....	21
Figure 8 – Mécanismes et architectures de décision	29
Figure 9 – Classification de l'intelligence en trois dimensions.....	31
Figure 10 – Interconnexion des outils et corrélation avec l'organisation structurelle du document.	34
Figure 11 - Relation entre le système de prévision et le système physique.....	39
Figure 12 - Taxonomie des approches de classification (Jain et al. 1999).....	43
Figure 13 - Exemple d'arbre de décision.....	44
Figure 14 - Exemple de réseau de neurones.....	47
Figure 15 - Exemple d'ensemble de classificateurs	50
Figure 16 – Situations où les prévisions sont possibles ou impossibles.....	56
Figure 17 – Proposition : Adaptation régulée du réseau de neurones	59
Figure 18 - Méthodologie de mise en place du système de maîtrise de la qualité	60
Figure 19 – Zones de fonctionnement normal du système pour 5 facteurs influents.....	61
Figure 20 – Une carte de contrôle pour contrôler le modèle de prévision en ligne	66
Figure 21 - Réactions possibles suite à la détection d'un défaut.....	75
Figure 22 - Modèle réduit utilisé pour les simulations.....	81
Figure 23 - Nombre de retards en fonction du taux de non-qualité pour la règle de pilotage EDD.	85
Figure 24 - Nombre de retards en fonction du taux de non-qualité pour la règle de pilotage FIFO.	85
Figure 25 - Changement de règle au cours d'une évolution du taux de non-qualité.....	87
Figure 26 - Combinaison des indicateurs C_{max} et N_{retard}	88
Figure 27 - Cartographie des états de saturation de l'atelier construite à partir des seuils.....	90
Figure 28 - Schéma structurel du robot de laquage.....	99
Figure 29 - Représentation des entrées et sortie du système de prévision des défauts.....	103
Figure 30- Menu général du logiciel POTER.....	104
Figure 31 - Suivi dynamique de production	106
Figure 32 - Interface tactile de saisie des défauts	107
Figure 33 - Comparaison des défauts de grains détectés à la sortie du robot de laquage (graphique supérieur) avec les défauts prévus par le réseau de neurones (graphique inférieur).....	112
Figure 34 - Distribution des taux de mauvaises classifications	113
Figure 35 - Distribution des mesures de diversité (DF).....	115
Figure 36 - Résultats du plan d'expériences en utilisant le meilleur réseau de neurones.	119
Figure 37 - Résultats du plan d'expériences en utilisant un ensemble classificateur.	120
Figure 38 - Différence entre le domaine d'apprentissage et le domaine d'utilisation.	121
Figure 39 - Historique d'apparition de grains sur les chants.....	122
Figure 40 - Réaction face à la détection d'une dérive par le système lui-même.....	123
Figure 41 - Résultats obtenus avec l'algorithme K-means avec la distance Euclidienne.....	125
Figure 42 - Plan d'expériences de Taguchi présentant l'influence du WIP, de Ndefaut, des règles de pilotage et de leurs interactions.	127
Figure 43 : Profil réel de l'état de surface.....	131
Figure 44 : Regroupement des noyaux élémentaires en lots de fabrication.	133
Figure 45 - Exemple de généalogie	134
Figure 46 - Exemple de logigramme d'un container intelligent	137
Figure 47 – Description du niveau d'intelligence du système en place	138
Figure 48 - Illustration de la problématique d'ordonnancement sur le robot de laquage.	140
Figure 50 - "Basculement" entre le programme linéaire et les algorithmes simples.	144
Figure 44 : Diagramme Ishikawa pour le problème de visibilité de l'insertion.....	171

Figure 45 : Graphique représentant l'influence des 3 facteurs retenus.....	172
--	-----

VI.2 TABLEAUX

Tableau 1 - Résumé des actions possibles vis-à-vis de la minimisation du coût de revient	11
Tableau 2 - Résumé des actions possibles vis-à-vis de l'anticipation de la charge	12
Tableau 3 - Résumé des actions possibles vis-à-vis du manque de réactivité.....	13
Tableau 4 - Exemples de contraintes liées aux produits.....	16
Tableau 5 - Résumé des actions possibles vis-à-vis de la simplification de la planification	16
Tableau 6 - Analyse comparative de différentes technologies de traçabilité	17
Tableau 7 - Résumé des actions possibles vis-à-vis de la fiabilisation du suivi terrain	19
Tableau 8 - Résumé des actions possibles vis-à-vis de la limitation du nombre de retard	19
Tableau 9 - Résumé des actions possibles vis-à-vis de l'engorgement de l'atelier	20
Tableau 10 - Résumé des axes de travail.....	33
Tableau 11 - Avantages et limites de la méthode Kpp	44
Tableau 12 - Avantages et limites des arbres de décision	45
Tableau 13 - Avantages et limites des machines à vecteurs de support (SVM)	47
Tableau 14 - Avantages et limites des réseaux de neurones.....	49
Tableau 15 - Avantages et limites des ensembles de classificateurs.....	55
Tableau 16 - Quelques méthodes de détection de dérives	58
Tableau 17 - Contrôlabilité des facteurs et piliers de l'économie durable.	62
Tableau 18 - Indicateurs de non-qualité	78
Tableau 19 - Indicateurs de flux.....	78
Tableau 20 - Type de variabilité en fonction du paramètre	83
Tableau 21 - Comparaison des concepts de cartographie visuelle et d'algorithme traditionnel.	91
Tableau 22 - Quelques algorithmes de clustering et leurs particularités.	92
Tableau 23 - Liste des défauts et de leurs causes possibles.....	101
Tableau 24 - Facteurs à suivre en fonction des causes possibles de l'apparition d'une coulure.....	102
Tableau 25 - Facteurs à analyser	102
Tableau 26 - Analyse fonctionnelle de la solution proposée	105
Tableau 27 - Résultats obtenus par le meilleur réseau de neurones, arbres, kpp et SVM.	114
Tableau 28 - Matrice de confusion des deux meilleurs classificateurs (NN et kpp)	114
Tableau 29 - Résultats obtenus en utilisant le meilleur ensemble de classificateurs.....	116
Tableau 30 - Matrice de confusion des deux meilleurs ensembles de classificateurs.....	117
Tableau 31 - Niveaux des facteurs du plan d'expérience.	126
Tableau 32 - Meilleure règle de pilotage en fonction des valeurs d'indicateurs imposés.....	127
Tableau 33 - Préconisation de travail en fonction des zones de saturation de l'atelier.	128

VI.3 REFERENCES

(s.d.). Récupéré sur <http://www.glasurit.com>.

Abouzahir, O., Gautier, R., & Gidel, T. (2003). Pilotage de l'amélioration des processus par les coûts de non-qualité. *10ème séminaire CONFERE*, 3-4.

AFGI. (1992). *Evaluer pour évaluer, les indicateurs de performance au service du pilotage industriel*.

Agard, B., & Kusiak, A. (2005). Exploration des bases de données industrielles à l'aide du datamining - Perspectives. *9ème colloque national AIP PRIMECA*.

- Aksela, M., & Laaksonen, J. (2006). Using diversity of errors for selecting members of a committee classifier. (Pergamon, Éd.) *Pattern Recognition*, 39(4), 608-623.
- Angelov, P., & Kasabov, N. (2005). Evolving computational intelligence systems. *Proceedings of the 1st international workshop on genetic fuzzy systems*, 76-82.
- Apley, D., & Shi, J. (2001). A Factor-Analysis Method for Diagnosing Variability in Multivariate Manufacturing Processes. *Technometrics*, 43, 84-95.
- Baena-Garcia, M., del Campo-Avila, J., Fidalgo, R., Bifet, A., Gavalda, R., & Morales-Bueno, R. (2006). Early drift detection method. *4ème International Workshop on Knowledge Discovery from Data Streams*.
- Barron, A. (1993). Universal approximation bounds for superpositions of a sigmoidal function. (IEEE, Éd.) *IEEE Transactions on Information Theory*, 39(3), 930-945.
- Basseville, M. (1986). On-line detection of jumps in mean. Dans S. B. Heidelberg (Éd.), *Detection of abrupt changes in signals and dynamical systems* (pp. 9-26).
- Bi, Y. (2012). The impact of diversity on the accuracy of evidential classifier ensembles. *Int. J. Approx. Reason.*, 53, 584-607.
- Bifet, A., & Gavalda, R. (2007). Learning from Time-Changing Data with Adaptive Windowing. *SDM*.
- Bifet, A., Holmes, G., Pfahringer, B., Kirkby, R., & Gavalda, R. (2009). New ensemble methods for evolving data streams. (ACM, Éd.) *Proceedings of the 15th ACM SIGKDD international conference on Knowledge discovery and data mining*, 139-148.
- Breiman, L. (1996). Bagging predictors. *Machine Learning*, 24, 123-140.
- Breiman, L. (1996). Bagging predictors. *Machine learning*, 24(2), 123-140.
- Breiman, L., Friedman, J., Stone, C., & Olshen, R. (1984). Classification and regression trees. (C. press, Éd.)
- Cheng, T. (1987). Optimal total-work-content-power due-date determination and sequencing. (Pergamon, Éd.) *Computers & Mathematics with Applications*, 14(8), 579-582.
- Clifford, M., & Duncan, S. (2002). Benefits of continuous on-line monitoring of mapping for CD control systems. *Pulp & Paper - Canada*, 103(8), 48-51.
- Colledani, M., & Tolio, T. (2006, décembre 31). Impact of quality control on production system performance. (Elsevier, Éd.) *CIRP Annals-Manufacturing Technology*, 55(1), pp. 453-456.

- Cortes, C., & Vapnik, V. (1995). Support-Vector Networks. *Machine Learning*, 20.
- Cover, T., & Hart, P. (1967). Nearest neighbor pattern classification. (IEEE, Éd.) *IEEE Transactions on Information Theory*, 13(1), 21-27.
- Cristianini, N., & Shawe-Taylor, J. (2000). An introduction to support vector machines and other kernel-based learning methods. (C. u. press, Éd.)
- Dai, Q. (2013). A competitive ensemble pruning approach based on cross-validation technique. (Elsevier, Éd.) *Knowledge-Based Systems*, 37, 394-414.
- Day, W., & Edelsbrunner, H. (1984). Efficient algorithms for agglomerative hierarchical clustering methods. (Springer-Verlag, Éd.) *Journal of classification*, 7-24.
- De Kok, A., Van Donselaar, K., & Van Woensel, T. (2008). A break-even analysis of RFID technology for inventory sensitive to shrinkage. (Elsevier, Éd.) *International Journal of Production Economics*, 112(2), 521-531.
- Dietterich, T. (2000). Ensemble methods in machine learning. (S. B. Heidelberg, Éd.) *Multiple classifier systems*, 1-15.
- Disney, S. M., & Towill, D. R. (2003, 6 30). On the bullwhip and inventory variance produced by an ordering policy. (Pergamon, Éd.) *Omega*, 157-167.
- Dreyfus, G., & al. (2002). Réseaux de neurones : Méthodologies et applications. Paris, France: Editions Eyrolles.
- Du, L., Ke, Y., & Su, S. (2012). The embedded Quality Control System of Product Manufacturing. *Advanced Materials Research*, 459, 510-513.
- Duffua, S., Khursheed, S., & Noman, S. (2004). Integrating Statistical Process Control, Engineering Process Control and Taguchi's Quality Engineering. *International Journal of Production Research*, 42, 4109-4118.
- El Haouzi, H. (2008). *Approche méthodologique pour l'intégration des systèmes contrôlés par le produit dans un environnement de juste-à-temps: Application à l'entreprise Trane*. Université Henri Poincaré-Nancy 1.
- Equitz, W. (1989). A new vector quantization clustering algorithm. (IEEE, Éd.) *Acoustics, Speech and Signal Processing*, 37(10), 1568-1575.
- Feigenbaum, A. (1951). *Quality Control*. Mc Graw-Hill.
- Freund, Y., & Schapire, R. (1997). A decision-theoretic generalization of on-line learning and an application to boosting. *Journal of computer and system sciences*, 55(1), 119-139.

- Gama, J., Medas, P., Castillo, G., & Rodrigues, P. (2004). Learning with drift detection. (S. B. Heidelberg, Éd.) *Advances in Artificial Intelligence - SBIA 2004*, pp. 286-295.
- Garvin, D. (1984). What does product quality really mean. *Sloan management review*, 26(1).
- Giacinto, G., & Roli, F. (2001). An approach to the automatic design of multiple classifier systems. (North-Holland, Éd.) *Pattern recognition letters*, 22(1), 25-33.
- Goldratt, E., & Cox, J. (1992). *The goal : A process of ongoing improvement*. Great Barrington, USA: North River Press, 2nd Revised edition.
- Grandvalet, J., & Canu, S. (2008). SVM and kernel methods. *matlab toolbox ASI-INSA de Rouan*.
- Grosjean, S., & Robichaud, D. (2010). Décider en temps réel : une activité située et distribuée mais aussi disloquée. (É. d. l'homme, Éd.) 134(4), 31-54.
- Guénoche, A., Hansen, P., & Jaumard, B. (1991). Efficient algorithms for divisive hierarchical clustering with the diameter criterion. (Springer-Verlag, Éd.) *Journal of classification*, 8(1), 5-30.
- Guo, L., & Boukir, S. (2013). Margin-based ordered aggregation for ensemble pruning. (North-Holland, Éd.) *Pattern Recognition Letters*, 34(6), 603-609.
- Haghighi, M., Vahedian, A., & Yazdi, H. (2011). Creating and measuring diversity in multiple classifier systems using support vector data description. (Elsevier, Éd.) *Applied Soft Computing*, 11(8), 4931-4942.
- Hajek, P., & Olej, V. (2010). Municipal revenue prediction by ensembles of neural networks and support vector machines. *WSEAS Transactions on Computers*, 9, 1255-1264.
- Herault, J., & Jutten, C. (1994). Réseaux neuronaux et traitement du signal. Hermès, Paris, France.
- Hernandez-Lobato, D., & Hernandez-Lobato, J. (2013). Learning feature selection dependencies in multi-task learning. *Advances in Neural Information Processing Systems*, 746-754.
- Herrera, C. (2011). *Proposition d'un cadrage générique de modélisation et de simulation de planifications logistiques dans un contexte de décisions partiellement distribuées*. Thèse de l'université de Lorraine.

- Hinkley, D. (1971). Inference about the change-point from cumulative sum tests. *Biometrika*, 58, 509-523.
- Ho, T. (1998). The random subspace method for constructing decision forests. (IEEE, Éd.) *Pattern Analysis and Machine Intelligence*, 20(8), 832-844.
- Hoc, J.-M., Mebarki, N., & Cegarra, J. (2004). L'assistance à l'opérateur humain pour l'ordonnancement dans les ateliers manufacturiers. (P. U. France, Éd.) 67(2), 181-208.
- Hornik, K. (1991). Approximation capabilities of multilayer feedforward networks. (Pergamon, Éd.) *Neural networks*, 4(2), 251-257.
- Hsu, C.-W., Chang, C.-C., & Lin, C.-J. (2003). A practical guide to support vector classification. 5.
- Ishikawa, K. (1982). *Guide to quality control* (Vol. 2). (A. P. Organization, Éd.)
- Jain, A., Murty, M., & Flynn, P. (1999). Data clustering: a review. (ACM, Éd.) *ACM computing surveys (CSUR)*, 31(3), 264-323.
- James, G., Witten, D., Hastie, T., & Tibshirani, R. (2013). *An Introduction to Statistical Learning*. New York: Springer.
- Joly, A., Frein, Y., Gauthier, D., & Bernier, V. (2004). Etude de l'impact des blocages sur le flux de production d'une usine terminale automobile. (Lavoisier, Éd.) *Journal européen des systèmes automatisés*, 38(3-4), pp. 291-313.
- Karp, R. (1992, 8). On-line algorithms versus off-line algorithms: How much is it worth to know the future? *IFIP Congress*, 12, 416-429.
- Keerthi, S., & Lin, C.-J. (2003). Asymptotic behaviors of support vector machines with Gaussian kernel. *Neural computation*, 15(7), 1667-1689.
- Khamassi, I., & Sayed-Mouchawed, M. (2014). Drift detection and monitoring in non-stationary environments. *Evolving and Adaptive Intelligent Systems (EAIS)*, 1-6.
- Kim, Y.-W., & Oh, I.-S. (2008). Classifier ensemble selection using hybrid genetic algorithms. (North-Holland, Éd.) *Pattern Recognition Letters*, 29(6), 796-802.
- Klein, T. (2008). *Le kanban actif pour assurer l'interopérabilité décisionnelle centralisé/distribué Application à un industriel de l'ameublement*. Université Henri Poincaré-Nancy 1.
- Ko, A., Sabourin, R., & Britto Jr, A. (2008). From dynamic classifier selection to dynamic ensemble selection. (Pergamon, Éd.) *Pattern Recognition*, 41(5), 1718-1731.

- Koksal, G., Batmaz, I., & Testik, M. (2011). A review of data mining applications for quality improvement in manufacturing industry. *Expert systems with applications*, 38(10), 13448-13467.
- Kotsiantis, S. (2007). Supervised machine learning: a review of classification techniques. *Informatica*, 31, 249-268.
- Kotsiantis, S. (2011). Combining bagging, boosting, rotation forest and random subspace methods. (S. Netherlands, Éd.) *Artificial Intelligence Review*, 35(3), 223-240.
- Kuhl, M., & Laubisch, G. (2004). A Simulation Study of Dispatching Rules and Rework Strategies in semiconductor manufacturing. (IEEE, Éd.) *Advanced Semiconductor Manufacturing*, 325-329.
- Kuncheva, L. (2002). A theoretical study on six classifier fusion strategies. (IEEE, Éd.) *IEEE Transactions on Pattern Analysis & Machine Intelligence*(2), 281-286.
- Kuncheva, L. (2004). *Combining pattern classifiers: methods and algorithms*. John Wiley & Sons.
- Kuncheva, L., & Whitaker, C. (2003). Measures of diversity in classifier ensembles and their relationship with the ensemble accuracy. (K. A. Publishers, Éd.) *Machine learning*, 51(2), 181-207.
- Kusiak, A. (2001). Rough set theory: a data mining tool for semiconductor manufacturing. *Electronics Packaging Manufacturing; IEEE Transactions on*, 24, 44-50.
- Li, L., Zhang, Y., Zou, L., Li, C., Yu, B., Zheng, X., & Zhou, Y. (2012). An ensemble classifier for eukaryotic protein subcellular location prediction using gene ontology categories and amino acid hydrophobicity. (P. L. Science, Éd.) *PloS one*, 7(1).
- Lopez, P. (1991). Approche énergétique pour l'ordonnancement de tâches sous contraintes de temps et de ressources. *Université Paul Sabatier-Toulouse III*.
- Love, P., Li, H., & Mandal, P. (1999). Rework: a symptom of a dysfunctional supply-chain. (Pergamon, Éd.) *European Journal of Purchasing & Supply Management*, 5(1), 1-11.
- Lughofer, E., & Angelov, P. (2011). Handling drifts and shifts in on-line data streams with evolving fuzzy systems. (Elsevier, Éd.) *Applied Soft Computing*, 11(2), 2057-2068.
- Lyonnet, B. (2010). Amélioration de la performance industrielle : Vers un système de production Lean adapté aux entreprises du pôle de compétitivité Arve Industries Haute-Savoie Mont-Blanc. *Thèse de l'université de Savoie*.

- Mac Farlane, D., Sarma, S., Chirn, J., Wonga, C., & Ashton, K. (2003). Auto ID systems and intelligent manufacturing control. *Engineering Applications of Artificial Intelligence*, 16, 365-376.
- McAuley, D. (1972). Machine grouping for efficient production. *The Production Engineer*, 51, 53-57.
- McClelland, J., & Rumelhart, D. (1986). Parallel distributed processing. *Explorations in the microstructure of cognition*, 2.
- McCulloch, W., & Pitts, W. (1943). A logical calculus of the ideas immanent in nervous activity. *Bulletin of Mathematics and Biophysics*, 5, 115-133.
- Mebarki, N. (1995). Une approche d'ordonnancement temps réel basée sur la sélection dynamique de règles de priorité. *Thèse de l'université de Lyon*.
- Mehta, M., Agrawal, R., & Rissanen, J. (1996). SLIQ: A fast scalable classifier for data mining. (S. B. Heidelberg, Éd.) *Advances in Database Technology—EDBT'96*, 18-32.
- Mercer, J. (1909). Functions of positive and negative type, and their connection with the theory of integral equations. (T. r. society, Éd.) *Philosophical transactions of the royal society of London. Series A, containing papers of mathematical or physical character*, 415-446.
- Meyer, D., Leisch, F., & Hornik, K. (2003). The support vector machine under test. *Neurocomputing*, 55, 169-186.
- Meyer, G., Främling, K., & Holmström, J. (2009). Intelligent products: A survey. (Elsevier, Éd.) *Computers in industry*, 60(3), 137-148.
- Mittler, M., & Schoening, A. (1999). Comparison of dispatching rules for semiconductor manufacturing using large facility models. (ACM, Éd.) *Proceedings of the 31st conference on Winter simulation: Simulation---a bridge to the future-Volume 1*, 709/713.
- Morel, G., Panetto, H., Zaremba, M., & Mayer, F. (2003). Manufacturing enterprise control and management system engineering paradigms and open issues. *Annual Reviews in Control*, 199-209.
- Muhl, E. (2002). *Contribution à la vision globale de l'ordonnancement d'un flux de véhicule : Vers au outil d'aide à la décision. Application pour un grand constructeur automobile*.

- Nguyen, D., & Widrow, B. (1990, juin 17). Improving the learning speed of 2-layer neural networks by choosing initial values of the adaptive weights. *IJCNN International Joint Conference on Neural Networks 1990*, 21-26.
- Noyel, M., Thomas, P., Charpentier, P., Thomas, A., & Brault, T. (2013, Octobre 28). Implantation of an on-line quality process monitoring. *International Conference on Industrial Engineering and Systems Management, IESM'13*.
- Noyel, M., Thomas, P., Thomas, A., & Beaupretre, B. (2012). Retour d'expérience industrielle sur le choix d'une technologie d'information portée par les produits. *9ème Conférence Internationale de Modélisation, Optimisation et Simulation, MOSIM'12*.
- Noyel, M., Thomas, P., Thomas, A., Charpentier, P., & Brault, T. (2014). Combinaison d'indicateurs pour évaluer la perturbation des flux dans les ateliers à fort taux de reprises. *10ème Conférence Internationale de Modélisation, Optimisation et Simulation, MOSIM'14*.
- Oakland, J. (2007). *Statistical Process Control*. Butterworth-Heinemann Ltd.
- Ohno, T., & Rosen, C. (1988). Toyota production systems: beyond large scale production. *Portland, USA, Productivity Press*.
- Oza, N. (2005). Online bagging and boosting. *IEEE international conference on Systems, man and cybernetics, 2005*, 3, 2340-2345.
- Page, E. (1954). Continuous Inspection Schemes. *Biometrika*, 41, pp. 100-115.
- Patel, M., & Panchal, M. (2012). A review on ensemble of diverse artificial neural networks. *Int. J. of Advanced Research in Computer Engineering and Technology*, 1(10), 63-70.
- Quinlan, J. (1993). Building Classification Models: ID3 i C4.5.
- Quinlan, J. (1996). Learning decision tree classifiers. (ACM, Éd.) *ACM Computing Surveys (CSUR)*, 28(1), 71-72.
- R. Segal, M. (2004). Machine learning benchmarks and random forest regression. *Center for Bioinformatics & Molecular Biostatistics*.
- Rabiee, M., Zandieh, M., & Jafarian, A. (2012). Scheduling of a no-wait two-machine flow shop with sequence-dependent setup times and probable rework using robust meta-heuristics. (I. J. Research, Éd.) *International Journal of Production Research*, 50(24), 7428-7446.

- Raquel, S., & João, G. (2009). A study on Change Detection Methods. <http://epia2009.web.ua.pt/onlineEdition/353.pdf>.
- Refke, A., Barbezat, G., & Loch, M. (2001). The benefit of an on-line diagnostic system for the optimization of plasma sprays devices and parameters. *Thermal Spray 2001: New Surfaces for a new millennium*, 765-770.
- Rodriguez, J., Kuncheva, L., & Alonso, C. (2006). Rotation forest: A new classifier ensemble method. (IEEE, Éd.) *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 28(10), 1619-1630.
- Rosenblatt, F. (1959). *Principles of neurodynamics*. Spartan Press, Washington, U.S.A.
- Roussew, P., & Leroy, A. (1987). *Robust regression and outlier detection*. New-York: Wiley.
- Ruta, D., & Gabrys, B. (2005). Classifier selection for majority voting. (Elsevier, Éd.) *Information fusion*, 6(1), 63-81.
- Sarac, A. (2010). *Modélisation et aide à la décision pour l'introduction des technologies RFID dans les chaînes logistiques*. Thèse.
- Sarle, W. (2002). Comp.ai.neural-nets FAQ, part 2, available on-line at: <ftp://ftp.sas.com/pub/neural/FAQ2.html>.
- Sayed-Mouchawed, M., & Lughofer, E. (2012). Learning in non-stationary environments: methods and application. *Springer Science & Business Media*.
- Sebastiao, R., & Gama, J. (2009). A study on change detection methods. *Progress in Artificial Intelligence, 14th Portuguese Conference on Artificial Intelligence, EPIA*, 12-15.
- Sebastião, R., Gama, J., Rodrigues, P., & Bernades, J. (2010). Monitoring incremental histogram distribution for change detection in data streams. (S. B. Heidelberg, Éd.) *Knowledge Discovery from Sensor Data*, pp. 25-42.
- Sematech, N. (2012). e-Handbook of Statistical Methods.
- Shafer, J., Agrawal, R., & Mehta, M. (1996). SPRINT: A scalable parallel classifier for data mining. *Proc. 1996 Int. Conf. Very Large Data Bases*, 544-555.
- Sick, B. (2002). On-line and indirect tool wear monitoring in turning with artificial neural networks: a review of more than a decade of research. (A. Press, Éd.) *Mechanical Systems and Signal Processing*, 16(4), 487-546.

- Singh, S., & Chauhan, N. (2011). K-means v/s K-medoids: A Comparative Study. *National Conference on Recent Trends in Engineering & Technology*, 13.
- Skidmann, A., Schweitzer, P., & Nof, S. (1985). Performance evaluation of a flexible manufacturing cell with random multiproduct feedback flow. (T. & Group, Éd.) *International Journal of Production Research*, 23(6), 1171-1184.
- Slama, I. (2008). Caractéristiques physico-mécaniques des composites bois-plastiques provenant de la valorisation des résidus des panneaux MDF: étude des possibilités de recyclage. (U. d. Abitibi-Témiscamingue, Éd.)
- Soto, J., Melin, P., & Castillo, O. (2013). Time series prediction using ensembles of neuro-fuzzy models with interval type-2 and type-1 fuzzy integrators. (IEEE, Éd.) *The 2013 International Joint Conference on Neural Networks (IJCNN)*, 1-6.
- Srivastava, B. (2004). The next revolution in SCM. *Business Horizons*, 47(6), 60-68.
- Stirl, T., & Skrzypek, R. (2003). Practical experiences and benefits with on-line monitoring systems for power transformers. *Proceedings of the 6th international conference on electrical machines and systems ICEMS 2003*, 1(2), 309-313.
- Taguchi, G. (1989). *Quality Engineering in Production Systems*. NY: MacGraw-Hill.
- Tague, N. (2004). *The Quality Toolbox, 2nd Edition*. ASQ Quality Press.
- Tang, E., Suganthan, P., & Yao, X. (2006). An analysis of diversity measures. (K. A. Publishers, Éd.) *Machine Learning*, 65(1), 247-271.
- Thierry, C., Thomas, A., & Bel, G. (2008). Simulation for product driven system. *Simulation for Supply Chain Management*, 221-255.
- Thomas, A. (2004). De la planification au pilotage pour les chaines logistiques. *HDR, Université Henri Poincaré - Nancy 1*.
- Thomas, A. (2009). RFID et nouvelles technologies de communication; enjeux économiques incontournables et problèmes d'éthique. *6ème Conférence Internationale Conception et Production Intégrées, CPI'2009*.
- Thomas, A., & Charpentier, P. (2005). Reducing simulation models for scheduling manufacturing facilities. (North-Holland, Éd.) *European Journal of Operational Research*(1), 111-125.
- Thomas, A., Trentesaux, D., & Valckenaers, P. (2012). Intelligents distributed production control. (S. V. (Germany), Éd.) *Journal of intelligent Manufacturing*, 23(6), 2507-2512.

- Thomas, P., & Thomas, A. (2009). How deals with discrete data for the reduction of simulation models using neural network. *13th IFAC Symp. On Information Control Problems in Manufacturing INCOM'09*, pp. 1177-1182.
- Thomas, P., Bloch, G., Sirou, F., & Eustache, V. (1999). Neurall modeling of an induction furnace using robust learning criteria. (I. Press, Éd.) *Integrated Computer-Aided Engineering*, 6(1), 15-26.
- Thomas, P., Suhner, M., Meuteler, B., & Brachotte, G. (2004). Quality monitoring of a high pressure die casting process based on bayesian and neural networks. *11th IFAC Symposium on automation in Mining Mineral and Metal processing MMM'04*, 38.
- Thorpe, S. (2012). Spike-based image processing: can we reproduce biological vision in hardware ? *Proceedings of the 12th international conference on Computer Vision*. doi:10.1007/978-3-642-33863-2_53
- Tibshirani, H., & Friedman, R. (2009). *The elements of statistical learning - data mining, inference and prediction*. (S.-V. N. Inc, Éd.) Springer series in statistics.
- Tsoumakas, G., Katakis, I., & Vlahavas, I. (2004). Effective voting of heterogeneous classifiers. Dans E. v. classifiers (Éd.), *Machine Learning: ECML 2004* (pp. 465-476).
- Tsoumakas, G., Partalas, I., & Vlahavas, I. (2009). An ensemble pruning primer. Dans S. B. Heidelberg (Éd.), *Applications of supervised and unsupervised ensemble methods* (pp. 1-13).
- Van Brussel, H., Wyns, J., Valckenaers, P., Bongaerts, L., & Peeters, P. (1998, Nov.). Reference architecture for holonic manufacturing systems: PROSA. *Comput. Ind.*, 37(3), pp. 255-274.
- Vapnik, V. (1999). An overview of statistical learning theory. (IEEE, Éd.) *Neural Networks, IEEE Transactions on*, 10(5), 988-999.
- Viennet, E. (2006). Réseaux à fonctions de base radiales. (Lavoisier, Éd.) 105.
- Villa-vialaneix, N., Olteanu, M., & Cierco-Ayrolles, C. (2013). Carte auto-organisatrice pour graphes étiquetés. *Atelier Fouilles de Grands Graphes (FGG)-EGC'2013*, 4.
- Vollmann, T., Berry, W., & Whybark, C. (s.d.). *Manufacturing Planning and Control Systems*. Dow Jones-Irwin.
- Voreux, C. (1987). *Un matériau riche d'utilisations : le panneau de fibres de moyenne densité (M.D.F.)*. (E. n. ENGREF, Éd.)

- Voreux, C. (1987). Un matériau riche d'utilisations: le panneau de fibres de moyenne densité (MDF). (E. n. ENGREF, Éd.)
- Weinberg, K. Q., & Saul, L. K. (2009). Distance metric learning for large margin nearest neighbor classification. (JMLR, Éd.) *The Journal of Machine Learning Research*, 10, 207-244.
- Windeatt, T. (2005). Diversity measures for multiple classifier system analysis and design. (Elsevier, Éd.) *Information Fusion*, 6(1), 21-36.
- Wozniak, M. (2007). Experiments with trained and untrained fusers. Dans S. B. Heidelberg (Éd.), *Innovations in Hybrid Intelligent Systems* (pp. 144-150).
- Wozniak, M. (2009). Evolutionary approach to produce classifier ensemble based on weighted voting. *World Congress on Nature & Biologically Inspired Computing, 2009. NaBIC 2009*, 648-653.
- Wozniak, M., Grana, M., & Corchado, E. (2014). A survey of multiple classifier systems as hybrid systems. (Elsevier, Éd.) *Information Fusion*, 16, 3-17.
- Yang, L. (2011). Classifiers selection for ensemble learning based on accuracy and diversity. (Elsevier, Éd.) *Procedia Engineering*, 15, 4266-4270.
- Zargar, A. (1995). Effect of rework strategies on cycle time. (Pergamon, Éd.) *Computers & Industrial Engineering*, 29(1), 239-243.
- Zhou, Z.-H., Wu, J., & Tang, W. (2002). Ensembling neural networks: many could be better than all. (Elsevier, Éd.) *Artificial intelligence*, 137(1), 239-263.

Partie

VII

VII ANNEXE

Cette annexe est dédiée à l'explication du plan d'expériences qui a été mené dans le but de réaliser une insertion dans un panneau qui puisse rester invisible dans le temps sous la laque.

Afin d'être exhaustif dans l'énumération des facteurs qui engendrent une insertion visible, nous utilisons le graphique Ishikawa présenté sur la Figure 50.

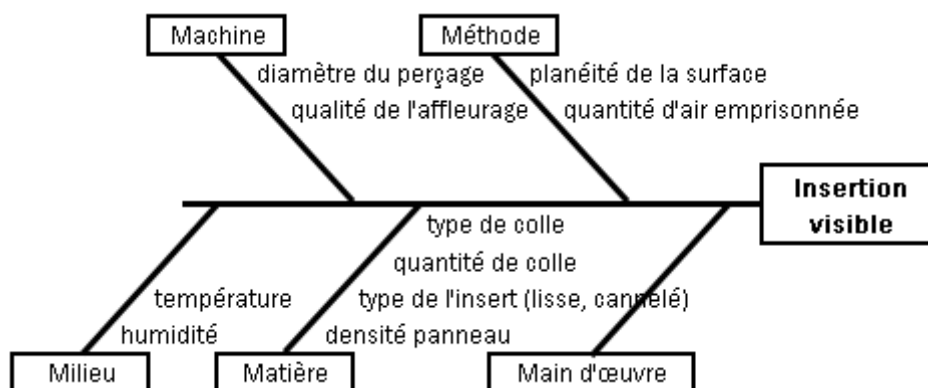


Figure 50 : Diagramme Ishikawa pour le problème de visibilité de l'insertion.

Sur la branche « machine », quand nous parlons du diamètre de perçage, nous tenons aussi compte des tolérances de perçage (dimensionnelles et tolérances de forme) d'après *Tolérances et écarts dimensionnels, géométriques et d'états de surface* de Jacques BOULANGER.

- Facteurs retenus, modalités et niveaux.

Modalité	niveau 1	niveau 2
Diamètre	8 mm	> à 8 mm
Planéité	oui	non
Densité panneau	e19	e50
Type de tourillon	lisse	cannelé
Colle	Non débordante	Débordante

- Résultats calculatoires

Le but n'étant pas de détailler entièrement le plan d'expériences, nous ne présentons ici que les résultats principaux ayant conduits à des réorientations du projet. Tout d'abord, l'importance des facteurs présentés ci-dessous qui confirme que le facteur le plus important n'est autre que l'état de surface en termes de planéité.

Planéité surface > Diamètre perçage > Type tourillon > Type panneau > Quantité de colle

Les valeurs des autres modalités présentées sur la Figure 51 ne font que confirmer que l'état de la jointure entre l'insert et le panneau doit être la plus nette et serrée possible.

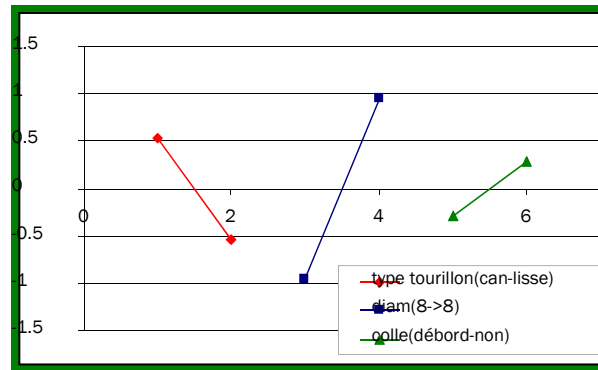


Figure 51 : Graphique représentant l'influence des 3 facteurs retenus.

Cependant, le meilleur réglage préconisé par la méthodologie de Taguchi ne permettait toujours pas d'obtenir une qualité d'insertion acceptable par le client. Deux nouvelles pistes ont alors été étudiées afin de vérifier que le problème ne venait pas de la différence de matériau entre l'insert et le panneau.

- Fabriquer des tourillons en tournant des lamelles de panneau.

Cette idée a pu être tirée des diverses utilisations possibles du panneau de MDF présentées dans (Voreux C. , 1987). Cependant, l'article aborde le tournage de grosses sections de panneau pour obtenir, par exemple, des montants pour les rampes d'escaliers. Tourner des sections de l'ordre de 8 mm de diamètre s'est révélé beaucoup plus compliqué et, malgré nos différents contacts, nous n'avons pu trouver personne en mesure de nous réaliser ces tourillons spéciaux. Nous tenons quand même à remercier la tournerie Patel pour leurs efforts en vue de nous aider.

- Injecter des tourillons en matériau composite dont les caractéristiques se rapprocheraient un maximum de nos panneaux.

D'après (Slama, 2008) il serait possible de réutiliser les chutes de panneaux afin de fabriquer notre propre matière à injecter. Cette perspective pouvait être très intéressante car elle nous permettrait d'agir sur nos déchets tout en envisageant de revendre le surplus de production. Cependant, elle nous oblige aussi à envisager l'investissement dans une ligne complète de fabrication de compounds et d'injection qui représente un corps de métier complet avec un savoir-faire totalement différent de celui de l'entreprise.

Ces nouvelles perspectives ont donné naissance à un nouveau plan d'expériences. Le premier des nouveaux facteurs concerne donc le matériau de l'insert avec pour modalités les points suivants :

- Bois de hêtre (tourillons lisses)
- Matériau injecté à base de bois uniquement (Arboform)
- Matériau injecté à base de bois/plastique (Promoplas)

Le deuxième facteur concerne l'affleurage et présente 3 modalités qui prennent en compte les constatations suite au premier plan d'expériences :

- Centre d'usinage
- Centre d'usinage + ponçage dur
- Affleurage de la plaqueuse de chants

Parce qu'il a été souligné que la colle pouvait, par sa dilatation, être à l'origine de la bosse à l'interface, et même si son importance n'avait pas été clairement identifiée comme primordiale dans le premier plan d'expériences, le temps de séchage peut quant à lui avoir une influence sur la dureté de cette bosse et une bosse plus dure sera plus difficile à poncer. La planéité serait donc plus difficile à garantir avec un ponçage standard. Les trois modalités suivantes seront donc étudiées afin d'être sûr des résultats, même si d'un point de vue processus de fabrication, un séchage de deux jours post-tagage serait inconcevable.

- Sans colle
- Séchage 1 heure
- Séchage 2 jours

Résumé :

Cette thèse CIFRE s'inscrit dans le cadre d'une collaboration entre Acta-Mobilier, fabricant de façades laquées haut de gamme, et le Centre de Recherche en Automatique Nancy. L'idée est de tirer parti du concept de Système Contrôlé par le Produit dans un environnement industriel perturbé par de nombreuses boucles de production et par un taux de reprises (non-qualités) non négligeable engendrant des pertes de pièces, le non-respect des délais, des charges de travail instables, etc... le lien impossible entre le produit et un identifiant infotronique rendant en plus la traçabilité difficile. Les travaux sur l'ordonnancement et son optimisation sont freinés par ces perturbations sur la chaîne de production qui rendent les plannings intenables. Le traitement prioritaire des pièces défectueuses permet d'assurer un taux de service qui reste remarquable au regard du pourcentage de pièces à réparer. Mais cela engendre aussi des pertes de pièces qui empêchent la livraison complète de la commande.

La problématique scientifique s'articule autour du pilotage des flux dans un contexte de production perturbé par les reprises et de la maîtrise de la qualité en évaluant son impact sur l'engorgement. L'enjeu de maîtrise de la qualité a été abordé à l'aide de réseaux de neurones capables de prévoir l'apparition du défaut auquel ils sont dédiés en fonction des paramètres de production et environnementaux. Cette anticipation permet de proposer une alternative de programme à utiliser ou à reporter la planification de la tâche. L'adaptation du modèle de prévision aux dérives du modèle physique au comportement considéré comme nerveux est réalisée « en-ligne » à l'aide de cartes de contrôle qui permettent de détecter la dérive et sa date de début.

Malgré cette simplification des flux, le pilotage reste complexe en raison des boucles normales de production et des non-qualités résiduelles. Il existe différents états de saturation du système pour lesquels la règle de pilotage la plus adaptée n'est pas toujours la même. Cette analyse est présentée sous forme de cartographie en deux dimensions dont chacun des axes présente un indicateur clé du taux de non-qualité et/ou de la perturbation des flux. Même si, contrairement aux algorithmes, la règle de pilotage la mieux adaptée ne sera pas toujours mise en évidence, cette cartographie présente d'autres avantages tels que la simplification du pilotage, la possibilité pour tous les utilisateurs d'avoir l'information importante sur l'état de l'atelier en un coup d'œil, ou encore la nécessité d'homogénéisation sur la globalité de l'unité de production.

Dans ce contexte, le container intelligent offre des perspectives intéressantes avec la volonté de tracer un groupe de produits ayant la même gamme de fabrication plutôt que des produits un à un, de partager des informations telles que sa date de livraison, son degré d'urgence, de connaître quels chemins ils doivent emprunter dans l'atelier et quelles sont les alternatives possibles ou encore de communiquer avec les machines et les autres systèmes dont celui de prévision de la qualité et retenir des informations au fil de la fabrication des produits. Le système proposé est donc interactif ou le conteneur est au cœur de la décision. Il signale sa présence au système d'ordonnancement seulement si les conditions qualité sont réunies, permettant ainsi de simplifier son travail autorisant alors un simple algorithme traditionnel de programmation linéaire à réaliser cette tâche particulièrement compliquée au premier abord. C'est en revanche à la charge de l'ordonnanceur de s'assurer de la règle de pilotage à utiliser et de demander les informations correspondantes aux lots disponibles.

La contribution de cette thèse est donc une méthodologie de simplification de problèmes complexes par une répartition des tâches entre différents sous-systèmes acteurs appliquée au cas d'une entreprise de fabrication de façades de cuisine laquées haut de gamme.

Abstract:

This CIFRE thesis is part of a collaboration between Acta-Mobilier, a manufacturer of high-finished lacquered panels, and the Centre de Recherche en Automatique de Nancy. The idea is to take advantage of Product Driven System in an industrial environment disturbed by many loops and a rework rate (non quality) causing significant loss of products, non-compliance deadlines, unstable workloads, etc ... impossible link between the product and identifying infotronic lead to more difficult traceability. Work on scheduling and optimization are hampered by these disturbances on the production line that make them untenable schedules. Priority processing on defective products ensures a service rate that remains outstanding compared to the percentage of products to repair. But it also leads to loss of products that prevent the full delivery of the order. The scientific problem revolves around the control of flow in a production context disturbed by the loops and the quality level by assessing its impact on congestion.

The quality-control issue has been addressed by using neural networks that can predict the occurrence of the defect to which they are dedicated from production and environmental parameters. This anticipation allows us to offer a program alternative to use or to plan to postpone the task. The adaptation of the forecasting model to the drift of the physical model with a behavior regarded as nervous is made "on line" using control charts that detect drift and its start date.

Despite this simplification of flows, the flow control remains complex due to normal production loops and residual non-qualities. There are different system saturation states for which the most suitable control rule is not always the same. This analysis is presented in a two-dimensional mapping which each axis has a key indicator on non-quality rate and / or disruption of flows. Although, unlike algorithms, the most suitable control rule will not always be highlighted, this mapping has other advantages such as the simplification of the control, the ability for all users to have important information about the workshop state, or the need for homogenization of the global state of the production unit.

In this context, the intelligent container offers interesting perspectives with the will to trace a group of products with the same routing sheet rather than products one by one, to share information such as its delivery date, the urgency degree, to know what paths they should take and what are the possible alternatives or to communicate with other machines and systems including the quality forecasting system and retain information over the manufacture of the products. The proposed system is so interactive where container is at the heart of the decision. It reported his presence to scheduling system only if the quality system requirements are met, and simplify this work while allowing a traditional linear algorithm to achieve this task seen as particularly complicated at first. It is however the responsibility of the scheduler to ensure the pilot rule to use and request the relevant information available to the lots. The contribution of this thesis is a methodology to simplify complex problems by a division of work between different subsystems actors applied to the case of a manufacturer of high-finished lacquered panels.