



Méthodologie pour le placement des capteurs à base de méthodes de classification en vue du diagnostic

Antonio Orantes Molina

► To cite this version:

Antonio Orantes Molina. Méthodologie pour le placement des capteurs à base de méthodes de classification en vue du diagnostic. Automatique / Robotique. INSA de Toulouse, 2005. Français. NNT : . tel-00010906

HAL Id: tel-00010906

<https://theses.hal.science/tel-00010906>

Submitted on 8 Nov 2005

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

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

THÈSE

Préparée au
Laboratoire d'Analyse et d'Architecture des Systèmes du CNRS

En vue de l'obtention du titre de
Docteur de l'Institut National des Sciences Appliquées de Toulouse

Spécialité
Systèmes Automatiques

Par
Antonio ORANTES MOLINA

METHODOLOGIE POUR LE PLACEMENT DES CAPTEURS A BASE DE METHODES DE CLASSIFICATION EN VUE DU DIAGNOSTIC

Version provisoire septembre 2005

Rapporteurs :	Joseba QUEVEDO Gilles TRYSTRAM
Directeur de thèse :	Marie-Véronique LE LANN
Examineurs :	Gilbert CASAMATTA Christophe GOURDON Cyrille LEMOINE
Invité(s) :	Joseph AGUILAR-MARTIN

Rapport LAAS N°

Table des matières

Introduction Générale.....	1
1. Diagnostic de procédés chimiques	5
1.1 Introduction	5
1.2 Détection de défauts et diagnostic.....	7
1.3 Classification des méthodes de diagnostic.....	11
1.3.1 Diagnostic Quantitatif	13
1.3.2 Diagnostic Qualitatif	17
1.4 Méthodes de classification	21
1.4.1 Algorithmes par hiérarchisation	21
1.4.2 Algorithmes par partitionnement.....	24
1.4.3 Algorithmes issus de la théorie des graphes.....	26
1.4.4 Algorithme de résolution mixte « Mixture-Resolving ».....	27
1.4.5 Classification par le plus proche voisin	27
1.4.6 Classification Floue.....	28
1.5 La méthode LAMDA	29
1.5.1 Les caractéristiques de LAMDA	30
1.5.2 Description de la méthodologie LAMDA	31
1.5.3 Les Concepts de LAMDA	33
1.5.4 L'algorithme de la classification	35

1.5.5 L'outil SALSA.....	39
1.6 Conclusions	40
2. Modélisation et simulation de procédés chimiques	43
2.1 Introduction	43
2.2 Concepts et caractéristiques du simulateur HYSYS	46
2.2.1 Concepts de base du simulateur HYSYS.....	46
2.2.2 Caractéristiques principales de HYSYS	49
2.3 Programmation de HYSYS	50
2.4 Modèle de l'unité de production du propylène glycol	53
2.5 Le nouveau réacteur OPR (<i>Open Plate Reactor</i>).....	59
2.6 Conclusions	61
3. Développement de la méthodologie d'aide au placement de capteurs en vue du diagnostic	63
3.1 Introduction	63
3.2 Méthodes de sélection de capteurs	64
3.3 Méthodes de sélection à partir du résultat de la classification de données	65
3.3.1 L'Entropie.....	66
3.3.2 Le gain de l'information.....	68
3.3.3 Critères de sélection des capteurs.....	69
3.4 Description de la méthodologie d'aide au placement de capteurs en vue du diagnostic.....	77
3.4.1 Paramètres du classificateur	78
3.4.2 Placement de capteurs.....	80
3.4.3 Modèle de comportement.....	81
3.4.4 Diagnostic en ligne	82
3.5 Conclusions	84
4. Placement de capteurs et diagnostic de procédés chimiques	87
4.1 Introduction	87
4.2 Application de la méthodologie à partir de l'utilisation du simulateur industriel HYSYS	89

4.2.1 Détection des défauts en utilisant tous les descripteurs : application au modèle de production de propylène glycol	91
4.2.2 Sélection des capteurs : application au modèle de la production de propylène glycol	94
4.2.3 Obtention du modèle du comportement	99
4.2.4 Reconnaissances de défauts inconnus.....	103
4.2.5 Utilisation des autres critères.....	115
4.3 Application de la méthodologie sur le nouveau procédé OPR	116
4.3.1 Réaction d'estérification	119
4.3.1.1 Sélection des descripteurs : Application à la réaction d'estérification	121
4.3.1.2 Conception du modèle de comportement pour la réaction d'estérification	124
4.3.1.3 Reconnaissances de défauts inconnus : Application à la réaction d'estérification	126
4.3.2 Réaction du thiosulfate.....	129
4.3.2.1 Sélection des descripteurs : Application à la réaction du thiosulfate mise en œuvre dans le réacteur « <i>Open Plate Reactor</i> ».....	130
4.3.2.2 Conception du modèle de comportement : Application à la réaction du thiosulfate mise en œuvre dans le réacteur « <i>Open Plate Reactor</i> ».....	132
4.3.2.3 Reconnaissances de défauts inconnus : Application à la réaction du thiosulfate	135
4.4 Conclusions	138
5. Conclusions et perspectives.....	139
Références	143
Annexes	
A. Exemple de calcul de la méthode LAMDA	151
B. programmation sous HYSYS	159
C. Propriétés de l'Entropie	165
D. Défauts Simultanés	169

Liste des figures

Figure 1-1 Différents types de défauts : a) défaut abrupt, b) défaut intermittent, c) défaut graduel	8
Figure 1-2 Schéma général d'un système de diagnostic	10
Figure 1-3 Classification des méthodes de diagnostic	12
Figure 1-4 Architecture générale de la détection de défauts à base de modèles	15
Figure 1-5 Différentes approches de classification non supervisée	21
Figure 1-6 Points dans trois classes	22
Figure 1-7 Dendrogramme obtenu en utilisant l'algorithme <i>single-link</i>	22
Figure 1-8 Clusters produits par les algorithmes de hiérarchisation sur un ensemble d'éléments contenant 2 classes.	23
Figure 1-9 Deux classes concentriques	23
Figure 1-10 Sensibilité de l'algorithme <i>k-Means</i> à la partition initiale	25
Figure 1-11 Classes formées en utilisant le MST (« <i>Minimal Spanning Tree</i> »)	27
Figure 1-12 Classes floues	29
Figure 1-13 Calcul de l'adéquation d'un élément	30
Figure 1-14 Algorithme général de LAMDA	38
Figure 1-15 Structure de SALSA hors ligne	39
Figure 1-16 Structure de SALSA en ligne	39
Figure 2-1 Exemple de <i>Main Flowsheet</i> dans HYSYS	48
Figure 2-2 Environnements de développement dans HYSYS	48
Figure 2-3 Organigramme des environnements dans la hiérarchie	50
Figure 2-4 Connexion en ligne du simulateur HYSYS et de l'outil de classification <i>SALSA</i>	52
Figure 2-5 Schéma de l'unité de production de propylène glycol	53
Figure 2-6 Architecture de l'ensemble des blocs de simulation utilisés sous HYSYS	54
Figure 2-7 Fenêtre des propriétés de l'oxyde de propylène du courant Prop Oxide	56
Figure 2-8 Sous-diagramme de la colonne de distillation	57
Figure 2-9 Schéma standard des régulateurs sur HYSYS	58
Figure 2-10 <i>Open Plate Reactor</i>	59
Figure 3-1 Concept du gain de l'information	68
Figure 3-2 Procédure probabiliste de sélection de capteurs par paires de classes	72

Figure 3-3 Valeurs de l'entropie probabilisé pour 2 classes	73
Figure 3-4 Procédure de sélection des capteurs par l'analyse de l'ensemble des classes	74
Figure 3-5 Valeurs de l'entropie pour 3 classes	75
Figure 3-6 Organigramme de la méthode basée sur l'entropie non probabiliste	76
Figure 3-7 Organigramme de la méthode de la variance	77
Figure 3-8 Schéma de la procédure d'aide au placement des capteurs	80
Figure 3-9 Conception du modèle comportemental	81
Figure 3-10 Schéma du diagnostic en ligne du procédé	82
Figure 3-11 Étapes de la conception du système de diagnostic	83
Figure 4-1 Diagramme des différentes étapes de la méthodologie de placement des capteurs et du diagnostic de procédés chimiques	88
Figure 4-2 Classification suite à la simulation d'une hausse d'oxyde de propylène	92
Figure 4-3. Profil des 2 classes « Normal » et « Défaut »	93
Figure 4-4 Résultat du gain relatif dans le cas de la défaillance « hausse d'oxyde »	96
Figure 4-5 Gains relatifs des descripteurs en prenant en compte tous les défauts	99
Figure 4-6 Identification de défauts avec les 6 descripteurs sélectionnés	101
Figure 4-7 Identification des 15 états fonctionnels avec 6 descripteurs	103
Figure 4-8 Scénario de deux défauts consécutifs critiques	104
Figure 4-9 Reconnaissance avec 6 descripteurs des défauts consécutifs : hausse d'oxyde et hausse de la température du fluide de refroidissement du réacteur	105
Figure 4-10 Scénario de deux défauts consécutifs non critiques	105
Figure 4-11 Reconnaissance avec 6 descripteurs des défauts consécutifs : baisse d'oxyde et baisse de la température du fluide de refroidissement du réacteur	106
Figure 4-12 Gain relatif pur l'ensemble des situations	108
Figure 4-13 Identification des défauts et états fonctionnels avec les 7 descripteurs sélectionnés	109
Figure 4-14 Reconnaissance avec 7 descripteurs des défauts consécutifs : baisse d'oxyde et baisse de la température du fluide de refroidissement du réacteur	111
Figure 4-15 Scénario : défauts simultanés	113
Figure 4-16 Identification d'états de défauts simultanés : scénario 1	113
Figure 4-17 Application de l'apprentissage supervisé	114
Figure 4-18 Reconnaissance des états de défauts simultanés après apprentissage supervisé	114
Figure 4-19 Représentation de l'OPR avec l'indice des cellules de modélisation	117
Figure 4-20 Perturbations réalisées sur la réaction d'estérification	120
Figure 4-21 Détection des défauts avec les 27 descripteurs : réaction d'estérification	121
Figure 4-22 Identification de défauts et d'alarmes avec les 8 descripteurs : réaction d'estérification	126
Figure 4-23 Classification de défauts inconnus : réaction d'estérification	127
Figure 4-24 Variations des conditions opératoires : réaction du thiosulfate	129
Figure 4-25 Détection de défauts avec 27 descripteurs : réaction du thiosulfate	130
Figure 4-26 Identification de défauts et alarmes avec 9 descripteurs : réaction du thiosulfate	133
Figure 4-27 Représentation des états fonctionnels du réacteur: réaction du thiosulfate	135
Figure 4-28 Classification des défauts inconnus : réaction du thiosulfate	136
Figure 4-29 Identifications des états fonctionnels : réaction du thiosulfate	137

REMERCIEMENTS

Je tiens tout d'abord à remercier mon directeur de thèse Mme **Marie-Véronique LE LANN** et mon co-directeur Mme **Tatiana KEMPOWSKY SANCHEZ** pour leur collaboration inestimable et leur disponibilité dans mon travail.

Je remercie Monsieur **Joseba QUEVEDO** et Monsieur **Gilles TRYSTRAM** d'avoir été rapporteurs de cette thèse. Ses grandes connaissances et ses critiques constructives ont permis d'affiner et de mieux situer ces travaux de thèse.

Je remercie aux Messieurs **Gilbert CASAMATTA** Président et Professeur à l'INPT, **Christophe GOURDON** Professeur ENSIACET/INPT, **Cyrille LEMOINE** Directeur Prog. Conduite Avancée de Procédés Anjou-Recherche et **Joseph AGUILAR-MARTIN** Directeur de Recherche LAAS-CNRS, pour avoir accepté de prendre part au jury. Je leur remercie tous pour l'intérêt qu'ils ont porté à ces travaux.

Je remercie **Jesús Alberto TEJEDA-RICARDEZ**, **Víctor Hugo GRISALES**, **David Garduño**, **Claudia ISAZA NARVAEZ** et **Héctor HERNANDEZ** qui ont toujours été prêts à répondre à mes interrogations de nature très diverse.

Je voudrais remercier à **Candita Gutiérrez** et à **Federico Miceli** du ITTG.

Je remercie l'amitié de Monsieur **André TITLI**, Professeur à l'INSA de Toulouse.

Je tiens à remercier Mme **Louise TRAVE-MASSUYES** et Monsieur **Joseph AGUILAR-MARTIN** de m'avoir accueilli au sein du groupe DISCO du LAAS-CNRS de TOULOUSE.

Je voudrais remercier à **CoSNET** et à **SFERE** par le financement de cette thèse et le suivi pédagogique, respectivement.

DEDICACE

*A **Dios** por darme una vida maravillosa*

*A mi padre **Antonio**, que es el mejor padre del mundo le dedico
especialmente esta tesis, gracias por ser mi
padre, gracias a ti he obtenido el mayor grado
de los estudios,*

*A mi madre **Lupita**, una mujer tierna y dulce siempre presente
para apoyarnos*

*A mi esposa **Martha Ilajaly** e hija **Marian**, que me acompañan
siempre y han dado gran sentido a mi vida,
ambas forman parte de mi existencia*

*A mis Abuelos **Panchita**, **Francisco** (†) y **Reynalda** (†)*

*A mis hermanos **Sergio**, **Gladys** y **Cecilia***

*A mis sobrinos: **Priscila**, **Sergito**, **Fatima** y **Aranza***

*A toda mi **familia política***

*A mis **colegas** de doctorado en Francia*

*A mis amigos **futboleros** en Toulouse*

*A mis compañeros y amigos del **ITTG***

INTRODUCTION GENERALE

Les systèmes industriels sont devenus de plus en plus complexes par l'automatisation de boucles de contrôle, par l'introduction de microprocesseurs à différents niveaux ou, plus récemment, par l'informatisation hiérarchisée et distribuée. De fait, cette évolution a complexifié la tâche de diagnostic des installations industrielles.

L'industrie chimique ou para-chimique renferme des procédés de plus en plus complexes qui peuvent être le plus souvent modélisés par des systèmes dynamiques hybrides (systèmes à variables continues et variables discrètes). De plus, ces systèmes sont fortement non-linéaires, leur fonctionnement est généralement étudié au travers de simulateurs dynamiques qui permettent de rendre compte des phénomènes mis en jeu dans ces procédés à l'aide des modèles dits de connaissance.

Ces simulateurs sont utilisés soit lors de la conception d'unités de production soit pour la formation d'opérateurs sur des installations existantes. Dans ce cadre, plusieurs sociétés ont développé des simulateurs dynamiques basés sur une description des phénomènes, parmi elles, la société *AEA technology* offre un simulateur dynamique HYSYS mondialement reconnu et utilisé dans le domaine industriel mais aussi dans celui de l'éducation. A ce jour, ce type de simulateur, n'est que très peu utilisé en ligne pour caractériser ou détecter un mauvais fonctionnement du processus. En effet, il n'existe pas encore de méthodologie générique qui permette de coupler ce simulateur à des techniques de supervision ou de diagnostic d'installations.

Il est généralement reconnu que les besoins pour la maintenance et le diagnostic devraient être pris en compte dès la phase de conception d'un système. Pour cette raison, les méthodes permettant d'analyser la diagnostabilité d'un système et de déterminer l'instrumentation nécessaire pour atteindre un certain degré de diagnostabilité sont hautement intéressantes.

L'objectif principal de ce travail est l'identification des capteurs et leur emplacement sur les procédés chimiques, capteurs qui seront nécessaires et suffisants pour pouvoir effectuer a posteriori la détection en-ligne de fautes et le diagnostic de dysfonctionnements. L'identification de capteurs est basée sur le résultat d'une technique de classification et sur une méthode de mesure de la quantité de l'information : l'Entropie. La méthodologie développée est générique et peut s'appuyer sur toute méthode de classification du moment que celle-ci fournisse une représentation des états de défaillance sous forme de classes caractérisées par une fonction de contribution des différents capteurs à ces classes. Pour illustrer cette méthodologie, le choix d'une méthode de classification s'est porté sur : «LAMDA» (*Learning Algorithm for Multivariate Data Analysis*), qui, comme son nom l'indique, est une technique d'Analyse Multidimensionnelle de Données avec Apprentissage développée au sein du groupe DISCO du LAAS et qui a été appliquée avec succès au diagnostic d'un certain nombre de processus. Récemment, elle a conduit à la réalisation du logiciel *SALSA* (*Situation Assessment using Lamda claSsification Algorithm*) développé dans le cadre du projet européen CHEM (Kempowsky, 2004). Cette méthode mêlant des concepts de représentation par modèle neuronal, de logique floue permet, sans changer de technique d'effectuer des classifications par apprentissage supervisé ou non supervisé, de manipuler des entrées (descripteurs) de nature différente : quantitatif ou qualitatif. La mesure de la quantité d'information d'un signal peut être effectuée selon différents critères : l'Entropie probabiliste ou non-probabiliste ou la variance. Une fois la sélection des capteurs effectuée sur la base d'un de ces critères, l'étape suivante de la méthodologie consiste à concevoir un modèle de comportement par classification à l'aide des seuls capteurs sélectionnés. Le but est d'obtenir un modèle du procédé à un niveau d'abstraction tel qu'il soit utilisable pour le diagnostic, c'est à dire qui rende explicite les états fonctionnels du procédé, leurs relations causales et les conditions de transition. Ce modèle est utilisé ensuite pendant la phase de reconnaissance en ligne des divers dysfonctionnements.

Ce mémoire de thèse est structuré en quatre chapitres.

Le premier chapitre propose un tour d'horizon des différentes méthodes de diagnostic de défaillances. Les diverses méthodes ont été classées suivant trois catégories : les méthodes basées sur les signaux, les méthodes basées sur les modèles mathématiques et les méthodes basées sur le raisonnement heuristique. Comme notre méthodologie de placement de capteurs et de diagnostic est basée sur l'analyse de données d'historique et de manière plus spécifique sur des résultats d'une classification (« *clustering* »), ce chapitre décrit, également,

différentes méthodes de classification. Nous présentons aussi de manière détaillée la méthodologie LAMDA, laquelle a été choisie comme technique de « *clustering* » spécifique. Finalement nous présentons l'outil *Salsa*, qui est basé sur la méthodologie LAMDA. Cet outil offre une interface modulaire, flexible et conviviale, qui permet l'interaction et le dialogue avec l'expert, lors de la construction du modèle de comportement.

Le deuxième est consacré à la simulation des procédés chimiques et aux deux procédés complexes étudiés lors de ce travail, sur lesquels a été appliquée la méthodologie. Plus spécifiquement, il présente les grandes caractéristiques du simulateur dynamique HYSYS utilisé pour simuler des scénarii de défaillances sur le premier procédé complexe étudié : l'unité de production de propylène glycol. Le deuxième procédé, l'OPR (« Open Plate reactor ») est un nouveau procédé en cours de conception dont une unité pilote est actuellement en fonctionnement au Laboratoire de Génie Chimique (UMR CNRS 5503). Cette dernière partie de notre travail s'est effectuée en collaboration avec l'équipe « Procédés de la Chimie Fine » dans le cadre d'un projet soutenu par l'Institut pour une Culture de Sécurité Industrielle (ICSI), projet intitulé : « *Nouvelles applications dans le domaine de la sécurité industrielle et du REX de systèmes d'aide à la conduite supervisée, d'abstraction d'informations, d'analyse et de classification de données* » et les scénarii de défaillance ont pu être simulés grâce au simulateur dynamique dédié qu'ils ont développé.

Dans le troisième chapitre, nous abordons les différentes approches du placement de capteurs développées par les différentes communautés de recherche. La plupart de ces approches s'appuient sur des modèles mathématiques de procédés, peu utilisent directement des données d'historique ou issues de simulations. Dans ces approches, nous incluons trois méthodes basées sur des données historiques pour sélectionner les capteurs qui utilisent le concept d'entropie et du gain d'information. Finalement, dans ce chapitre nous illustrons la totalité de la procédure permettant de mettre en œuvre la méthodologie développée au cours de notre travail.

Le chapitre 4 illustre largement les résultats obtenus lors de l'application de la méthodologie de placement de capteur et de diagnostic sur les deux procédés complexes : l'unité de production de propylène glycol simulé au travers du simulateur dynamique HYSYS et l'OPR.

Enfin, dans la conclusion générale nous reviendrons sur les atouts mais aussi les inconvénients de cette méthodologie et nous donnerons des perspectives d'amélioration.

1. DIAGNOSTIC DE PROCEDES CHIMIQUES

1.1 Introduction

Le domaine du contrôle des procédés a particulièrement explosé dans les 30 dernières années avec l'automatisation des processus complexes. Grâce aux progrès réalisés dans la commande distribuée, à la possibilité d'effectuer des prédictions basées sur des modèles (commande prédictive), les bénéfices acquis par les industries telles que, chimiques, pétrochimiques, de l'acier, de l'énergie ont été énormes. Toutefois, une tâche importante au sein des unités industrielles demeure une activité manuelle effectuée par l'opérateur humain. C'est la tâche de répondre aux événements anormaux qui surviennent sur le procédé. Ceci implique la détection à temps d'un événement anormal, en diagnostiquant ses causes et une prise de décision en retour, qui soit appropriée de sorte à ramener le procédé dans un état normal et sécurisé (supervision).

L'homme a été victime de problèmes survenus à cause d'erreurs de conception, de perturbations externes diverses ou encore, à cause d'un manque d'entretien des mécanismes.

Longtemps, l'observation humaine directe, c'est-à-dire sans utiliser d'instrumentations dédiées ou de capteurs spécifiques, a été le seul moyen pour détecter des déviations par rapport à l'état normal de fonctionnement. Les sensations perçues directement par l'opérateur humain (bruit, température, odeurs, vibrations...) étaient, et dans certains cas sont encore, les éléments de base de l'étape de détection des dysfonctionnements et de surveillance des installations. Ensuite sont apparus les capteurs et les moyens d'enregistrement ce qui lui a permis de percevoir la dimension dynamique des phénomènes qu'il observait.

Les systèmes industriels (nucléaire, chimique, pétrochimie, aérospatiale) sont également devenus de plus en plus complexes avec l'automatisation de boucles de régulation, l'introduction des microprocesseurs à différents niveaux et, plus récemment, l'information hiérarchisée et distribuée. De fait, cette évolution a complexifié la tâche de diagnostic des installations industrielles. Étant données ces difficultés, l'opérateur peut commettre des erreurs décisives et prendre de mauvaises actions. Des statistiques industrielles montrent que 70% des accidents industriels sont causés par des erreurs humaines. Ces événements anormaux ont eu des impacts d'ordre économique, mais surtout des impacts sur la sécurité et l'environnement. Malgré les avancées effectuées sur les commandes par ordinateur des unités chimiques, le fait de l'accident le plus récent sur l'usine AZF à Toulouse, avec 29 morts et plus de 50 personnes gravement blessées, pose le problème du développement de techniques de supervision. Les autres accidents récents les plus spectaculaires sont celui de l'Union Carbide à Bhopal en Inde et d'Occidental Petroleum à Piper Alpha [Lees, 1996] ou encore l'explosion d'une unité de Kuwait Petrochemical à la raffinerie de Mine Al-Ahmedi refinery en juin 2000, laquelle causa près de 100 millions de dollars de dommages.

De plus, des statistiques industrielles ont montré que de petits accidents sont très fréquents tous les jours, occasionnant des pertes économiques importantes [BLS, 1998 ; NSCI, 1999]. Ces pertes sont encore plus élevées dans le cas des industries pharmaceutiques et chimiques [Laser, 2000].

Les procédés chimiques sont des procédés complexes, faisant intervenir des réactions chimiques, dont les caractéristiques sont très difficiles à appréhender. Il arrive souvent que beaucoup de variables importantes du procédé ne soient pas accessibles à la mesure en ligne (compositions). De même, de nombreux paramètres sont mal connus et/ou susceptibles de varier au cours du temps. Les capteurs permettant d'accéder aux mesures de composition sont très coûteux et demandent souvent une maintenance lourde et onéreuse.

Néanmoins, lors de son fonctionnement, il faut gérer le procédé face à divers problèmes de fonctionnement, qu'il s'agisse de dysfonctionnements ou de pannes de capteurs, d'actionneurs (vannes, pompes, agitateurs...). Cette problématique fait appel à toutes les informations sur le procédé (qu'il s'agisse de celles qui proviennent de la modélisation, des capteurs physiques et logiciels ou de la commande).

L'objectif du diagnostic est de constater l'apparition d'un défaut, d'en trouver la cause puis d'en déduire la marche à suivre afin d'assurer la sûreté de fonctionnement du procédé.

Dans la section suivante, on donnera une revue non exhaustive de plusieurs méthodes de diagnostic à partir de différentes perspectives.

1.2 Détection de défauts et diagnostic

L'objectif de la maintenance préventive est de déterminer l'ensemble des actions à exercer sur le procédé afin de ne pas subir l'effet d'une défaillance. On peut à cet effet distinguer deux approches possibles : la maintenance préventive systématique et la maintenance préventive conditionnelle.

Dans le premier cas, les activités de maintenance sont planifiées et ont lieu selon un échancier basé sur le temps ou l'unité d'usage. Lors de ces interventions, les éléments sont remplacés même s'ils ne sont pas défaillants.

Dans le deuxième cas, les activités de maintenance sont déclenchées en fonction d'informations reflétant l'état de dégradation de l'équipement considéré. Dans ce deuxième cas, les éléments ne sont remplacés que si nécessaire.

On s'intéresse dans la suite à la maintenance conditionnelle, basée sur la surveillance en continu de l'évolution du système considéré. Ce type d'approche nécessite la conception d'un système de diagnostic permettant la détection précoce de déviations faibles par rapport à une caractérisation du système en fonctionnement nominal, afin de prévenir un dysfonctionnement avant qu'il n'arrive.

Lorsqu'il est appliqué à des systèmes industriels, une des difficultés du diagnostic réside dans la grande variété de défauts possibles, tant du point de vue de leur source que de leur amplitude et de leur fréquence.

Pour mieux étayer nos propos, nous introduisons quelques définitions utilisées dans le domaine du diagnostic [Toscano, 2005] :

- **Fonctionnement normal d'un système.** Un système est dit dans un état de fonctionnement normal lorsque les variables le caractérisant (variables d'état, variables de sortie, variables d'entrée, paramètres du système) demeurent au voisinage de leurs valeurs nominales. Le système est dit défaillant dans le cas contraire.
 - **Défaut.** C'est une déviation en dehors d'un intervalle acceptable, d'une variable observée ou d'un paramètre associé au procédé [Himmelblau, 1978]. C'est-à-dire, un
-

défaut est un processus anormal ou symptôme, tel que la hausse de température dans un réacteur ou la baisse de qualité du produit.

- **Défaillance.** C'est la cause d'une anomalie, telle qu'une panne d'une pompe de refroidissement ou d'un régulateur.
- **Détection de défauts.** La détection d'un défaut consiste à décider si le système se trouve ou non dans un état de fonctionnement normal.
- **Localisation d'un défaut.** A l'issue de la détection d'un défaut, il s'agit de déterminer le ou les éléments à l'origine du défaut.

Une classification des défauts à partir de leurs évolutions temporelles les définit comme :

- **abrupts** : la caractéristique principale de ce type de défauts est la discontinuité dans l'évolution temporelle de la variable. Cette évolution, si elle ne correspond pas aux évolutions dynamiques normales attendues pour la variable (changement de consigne), est caractéristique d'une panne brutale de l'élément en question : arrêt total ou partiel,
- **intermittents** : il s'agit d'un type de défauts caractéristiques de faux contacts ou de pannes intermittentes de capteurs. C'est un cas particulier de défaut brutal sur un capteur avec perte aléatoire de signal,
- **graduels** : ce type de défauts est caractéristique d'un encrassement ou d'une dérive dans les paramètres caractéristiques du procédé. Il s'agit de défauts très difficiles à détecter, car leurs évolutions temporelles sont les mêmes que celles d'une modification paramétrique lente représentant une non-stationnarité du procédé.

La Figure 1-1 résume cette classification des défauts à partir de leurs évolutions temporelles.

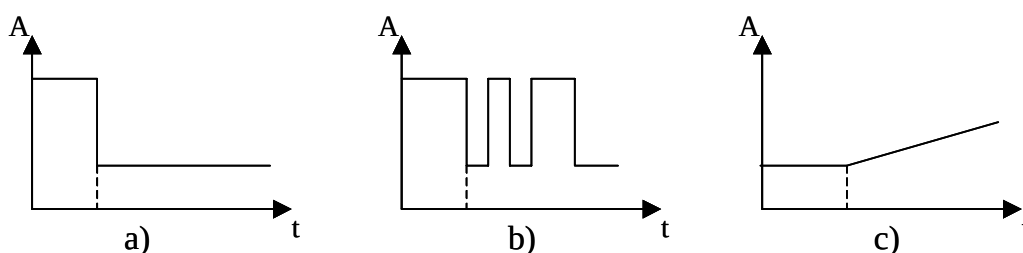


Figure 1-1 Différents types de défauts : a) défaut abrupt, b) défaut intermittent, c) défaut graduel

La présence de défauts multiples constitue une situation fréquente dans la pratique. Cette situation a lieu lorsqu'une ou plusieurs variables (ou paramètres) sortent de leurs évolutions normales à cause d'un dysfonctionnement simultané d'éléments d'une installation. En effet, la propagation d'une panne dans un procédé produit généralement des événements en cascade qui peuvent avoir ce type de conséquences.

Avec le développement de l'automatisation, les progrès techniques ont permis le développement de procédés de plus en plus complexes. Face à une défaillance, l'opérateur est par ailleurs soumis à une cascade d'informations qui met à rude épreuve ses capacités intellectuelles, sa réactivité et sa gestion du stress. Le diagnostic et la prise de décision deviennent des tâches qui ne sont pas aisées à réaliser car la détection de défauts doit être accomplie dans ces conditions réelles de fonctionnement des installations – donc en ligne – :

- Les caractéristiques temporelles des défauts sont inconnues.
- Le modèle nominal du système est incertain.
- Des bruits d'état et de mesure sont présents.
- Le signal de détection doit être obtenu dans un temps fini et est fonction des contraintes opératoires du procédé.

Pour remédier aux états de défaillance et trouver des solutions adaptées à chaque procédé, la détection précoce de défauts, le diagnostic automatisé et l'assistance aux opérateurs sont donc des domaines où la recherche scientifique est très active.

À cause de la difficulté d'effectuer en temps réel la détection de défauts et le diagnostic de procédés, plusieurs approches ont été développées. Citons comme techniques les plus récentes : les arbres de défaillances et digraphes, les approches analytiques, les systèmes à base de connaissance et encore plus récemment les réseaux de neurones.

Dans le cas où une modélisation du procédé est disponible ou est envisageable, il est possible d'utiliser des méthodes de diagnostic qui requièrent soit des modèles précis du procédé, soit des modèles semi-quantitatifs, voire même des modèles qualitatifs.

En revanche, il existe des méthodes qui ne nécessitent aucune modélisation proprement dite, mais qui se basent sur des informations/données historiques représentatives du fonctionnement du procédé lors de différents scénarii. La Figure 1-2 résume le contexte général du diagnostic de défauts.

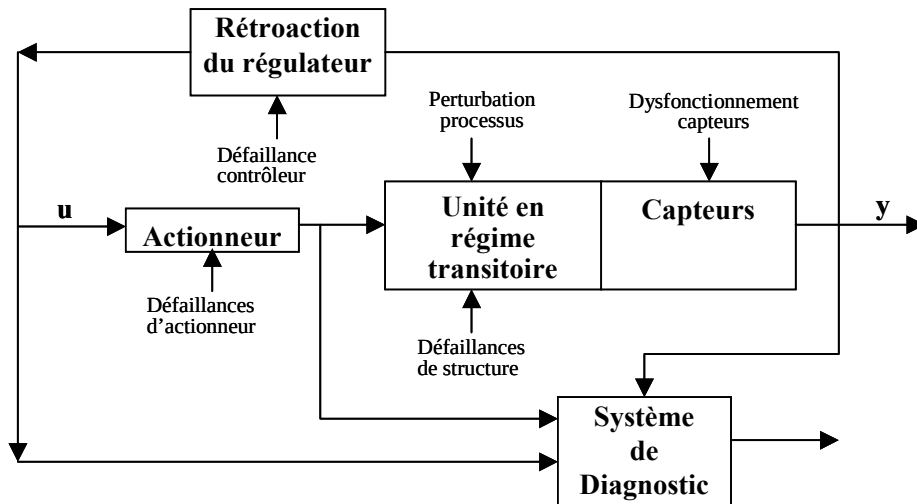


Figure 1-2 Schéma général d'un système de diagnostic

La figure présente un système de processus contrôlé et indique les différentes origines de ses défaillances. En général, il existe trois classes de défaillance ou de dysfonctionnement décrites comme suit [Venkatasubramanian et al., 2003] :

- **Changement paramétrique dans un modèle.** Un exemple de cette défaillance est le changement dans la concentration du réactif à partir de sa valeur normale ou état stable dans l'alimentation du réacteur.
- **Changement structurel.** On se réfère aux changements dans le processus même. Quelques exemples sont la défaillance d'un régulateur, la saturation d'une vanne, une fuite dans une canalisation.
- **Défaillances des capteurs et des actionneurs.** Ce peut être une déviation constante (positive ou négative) ou une défaillance de sortie de plage admissible.

Les autres défauts sont : les incertitudes de structure (défauts a priori non modélisés), bruits de processus (différence entre le processus actuel et les prédictions par le modèle) et le bruit de mesure (correspond généralement à l'addition d'un terme de haute fréquence à la sortie des capteurs).

L'ensemble des caractéristiques souhaitées qu'un système de diagnostic devrait posséder [Venkatasubramanian et al., 2003] est :

- a) Détection rapide et diagnostic.
- b) Isolation, c'est l'habilité pour différencier les défauts.
- c) Robustesse vis-à-vis de certains bruits et d'incertitudes.

- d) Identification de nouveauté, on se réfère à la capacité de décider si le processus est en état normal ou anormal. Dans le cas d'anomalie, il faut identifier s'il s'agit d'un défaut connu ou d'un nouveau défaut.
- e) Estimation de l'erreur de classification du défaut (diagnostic) en vue de sa fiabilité.
- f) Adaptabilité : le système de diagnostic devrait être adaptable aux changements de conditions du processus (perturbations, changements d'environnement).
- g) Facilité d'explication de l'origine des défauts et de la propagation de celui-ci. Ceci est important pour la prise de décision en ligne.
- h) Conditions de modélisation : pour le déploiement rapide et facile des classificateurs de diagnostic en temps réel, l'effort de modélisation devrait être aussi minimal que possible.
- i) Facilité de mise en œuvre informatique (faible complexité dans les algorithmes et leur implémentation) et capacité de stockage.
- j) Identification de multiples défauts : pour de grands processus, l'énumération combinatoire de multiples défauts est trop importante et ils ne peuvent être explorés de manière exhaustive.

1.3 Classification des méthodes de diagnostic

Les premières méthodes de diagnostic furent basées sur la redondance de matériels jugés critiques pour le fonctionnement du système. La redondance matérielle est très répandue dans les domaines où la sûreté de fonctionnement est cruciale pour la sécurité des personnes et de l'environnement, comme dans l'aéronautique ou le nucléaire. Les principaux inconvénients de la redondance matérielle sont liés aux coûts dus à la multiplication des éléments ainsi que l'encombrement et aux poids supplémentaires qu'elle génère.

Les spectaculaires progrès réalisés dans le domaine des calculateurs numériques combinés à une baisse des coûts permettent aujourd'hui la mise en œuvre, dans le milieu industriel, des méthodes modernes de l'automatique et de l'intelligence artificielle. Cette nouvelle approche permet d'éliminer en partie, voire même en totalité, la redondance matérielle pour le diagnostic des systèmes industriels. On peut globalement distinguer deux grandes familles dans les méthodes de diagnostic :

- Les méthodes basées sur une modélisation des systèmes ou des signaux, que nous dénommerons « diagnostic quantitatif ».
- Les méthodes basées sur l'intelligence artificielle que nous appellerons « diagnostic qualitatif ».

Le fait de distinguer ce qui est de l'ordre du quantitatif et du qualitatif, ne doit pas laisser penser que ces deux aspects sont disjoints. En réalité, ces deux types d'approche coexistent souvent au sein d'un même système de diagnostic. L'utilisation conjointe de méthodes quantitatives et qualitatives permet l'exploitation de l'ensemble des connaissances disponibles concernant le fonctionnement du système.

Nous présentons à la Figure 1-3 un panorama général des différentes méthodes de diagnostic dans l'une ou l'autre des catégories précédemment présentées.

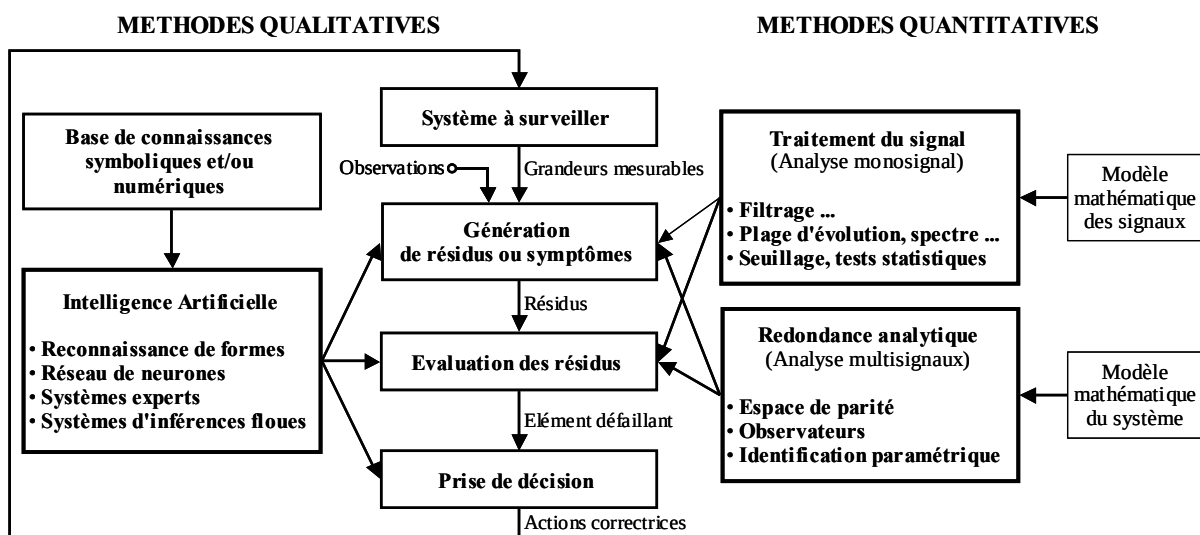


Figure 1-3 Classification des méthodes de diagnostic

En ce qui concerne le diagnostic quantitatif, on peut distinguer deux types d'approches suivant que l'on considère les mesures prises isolément les unes des autres ou qu'au contraire on présuppose des relations mathématiques les reliant. La première approche est connue sous le nom d'analyse mono signal, la deuxième est dénommée analyse multi-sigaux ou redondance analytique.

L'analyse mono-signal consiste dans l'étude des signaux, pris isolément les uns des autres, dans le but d'extraire des informations révélatrices de défauts. Cette approche repose sur les

techniques de traitement du signal. La méthode la plus simple consiste à fixer un ou plusieurs seuils limitant la plage d'évolution du signal considéré. D'autres approches, plus élaborées, sont basées sur les propriétés statistiques des signaux (moyenne, écart type, etc.), ou sur l'analyse spectrale.

L'analyse multi-sinaux met en jeu un ensemble de variables et les relations de causalité existant entre elles. Cette approche repose sur les techniques de l'automatique. En automatique, le concept clé est celui de modèle mathématique d'un processus physique. Ce modèle mathématique, censé représenter une réalité objective, est exploité afin d'élaborer des lois de commande permettant l'accroissement des performances dynamiques et statiques du système étudié. Cette connaissance analytique d'un système physique peut être mise à profit pour le diagnostic. En effet, le modèle étant censé décrire le comportement dynamique du système à surveiller, tout écart entre le comportement prévu par le modèle et le comportement mesuré traduit l'apparition d'un défaut. Dans la suite, nous entendrons essentiellement par diagnostic quantitatif, l'ensemble des méthodes fondées sur l'utilisation d'une représentation mathématique du système à surveiller.

Dans le cas du diagnostic qualitatif, les connaissances utilisables reposent sur le savoir et l'expérience d'opérateur humain ayant une parfaite maîtrise de l'installation à surveiller (base de connaissance symbolique) et/ou sur l'existence d'une base de connaissances numériques correspondant aux divers modes de fonctionnement de l'installation. Suivant cette approche, on trouve toutes les méthodes liées à l'intelligence artificielle. Elle inclut les systèmes experts, les systèmes d'inférence floue, la reconnaissance de forme, les réseaux neurones. Ces différentes approches sont alors utilisées pour construire un modèle de diagnostic qualitatif permettant de remonter des symptômes observés aux causes.

1.3.1 Diagnostic Quantitatif

A) Les méthodes basées sur l'approche mono-signal :

- **Redondance matérielle** : c'est une méthode employée dans des installations critiques (l'aérospatiale, le nucléaire). L'utilisation de plusieurs capteurs en vue d'obtenir la même information sur une variable permet de détecter les déviations par rapport à un état normal et de localiser un défaut de capteur. Les inconvénients de cette méthode sont l'accroissement du coût de l'installation et l'augmentation de la probabilité de pannes de capteurs – donc d'un besoin de maintenance supplémentaire.
-

- **Analyse spectrale** : Les signaux sont analysés en état normal de fonctionnement ; les hautes fréquences sont reliées au bruit et les basses fréquences aux évolutions propres de l'état du procédé. Ensuite, toute déviation des caractéristiques fréquentielles d'un signal est reliée à une situation de défaillance. Cette approche se révèle très utile pour analyser des signaux qui montrent des oscillations avec des périodes longues (les courants électriques, les débits, les pressions...). L'inconvénient est la sensibilité aux bruits de mesure quand ceux-ci coïncident avec la zone fréquentielle d'intérêt et la nécessité d'un échantillonnage fréquent pour permettre de reconstituer le signal de départ tout en minimisant la perte de fréquence. Les méthodes d'auto-corrélation, la densité spectrale des signaux, la transformée de Fourier, les ondelettes sont bien appropriées dans le cas où les fréquences représentatives de défauts sont connues. Dans le cas contraire, il est cependant préférable d'utiliser des modèles paramétriques des signaux qui permettent d'estimer en ligne les fréquences et les valeurs moyennes des paramètres.
- **Approches statistiques** : Ces approches se basent sur l'hypothèse de changements rapides (et non sur l'amplitude) des caractéristiques des signaux ou des paramètres des modèles par rapport à des dynamiques considérées comme étant lentes (procédés quasi stationnaires). Elles sont utilisées pour la détection de changements graduels avec des seuils de détection faibles. Les informations fournies par le nombre croissant de capteurs installés sur les procédés rendent très difficile l'analyse des résultats. L'analyse en composantes principales (*ACP*) et les moindres carrés partiels (*MCP*) permettent de réduire le nombre de variables à traiter par la détermination de relations linéaires entre elles (variables latentes), permettant ainsi d'expliquer la variance dans des séries de données.

B) L'approche multi-signaux

Ce sont de méthodes qui utilisent plus d'information que celles apportées par les seuls capteurs physiques. Ces informations peuvent provenir de la connaissance du comportement entrée/sortie d'un procédé ou des processus internes qui en gouvernent l'évolution.

Selon la méthode, différents types de modèle sont utilisés. Par exemple, pour les approches utilisant l'estimation d'état ou l'estimation paramétrique, on utilisera des modèles analytiques alors que pour l'approche systèmes experts, on recourra à des modèles de type base de connaissance. Le modèle servant directement de référence pour la détection de défauts, la

qualité du résultat dépend directement de la qualité des modèles. La mise en œuvre de ces méthodes nécessite donc une modélisation précise.

La détection de défaut basée sur l'utilisation de modèles peut être divisée en deux étapes principales : la génération de résidus et la prise de décision. Lors de la première étape, les signaux d'entrée et de sortie du système sont utilisés pour générer un résidu – c'est-à-dire un signal mettant en évidence la présence d'un défaut. En général, en régime de fonctionnement normal, ce signal est statistiquement nul et s'écarte notablement de zéro en présence de défaut. La génération de résidus est propre à la méthode utilisée. Durant la seconde étape, les résidus sont analysés pour décider s'il y a ou non présence de défaut, sur quelle composante du système il est intervenu (localisation) et pour déterminer la nature du défaut et sa cause (identification). La décision peut s'effectuer à l'aide d'un simple test de dépassement de seuil sur les valeurs instantanées ou les moyennes de résidus, en utilisant des fonctions floues, en faisant appel également à la reconnaissance de formes ou en utilisant des seuils adaptatifs qui évoluent en fonction du point de fonctionnement du processus surveillé.

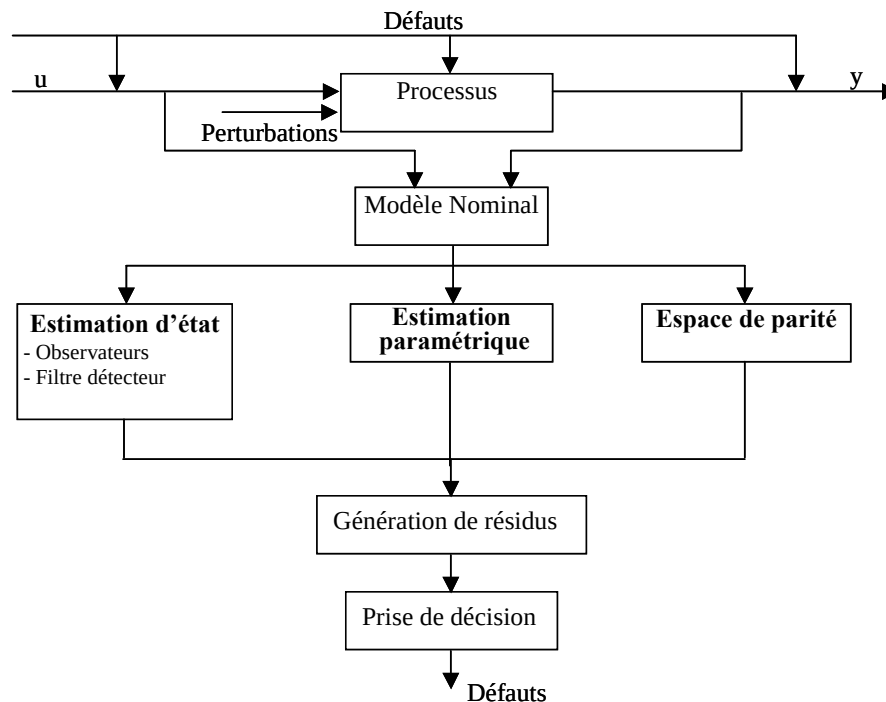


Figure 1-4 Architecture générale de la détection de défauts à base de modèles

La Figure 1-4 présente l'architecture générale de la détection de défauts basée sur l'utilisation de modèles [Isermann, 1984]. La génération de résidus s'effectue sur la base d'estimations

(états ou paramètres) ou à l'aide d'un raisonnement heuristique. L'étape de prise de décision permet ensuite de détecter la présence de changements, d'identifier la nature, l'amplitude du changement voire localiser les défauts éventuels.

Parmi les différentes méthodes de détection et de diagnostic utilisant des modèles mathématiques, nous trouvons principalement l'espace de parité, les observateurs et l'estimation paramétrique.

A) Espace de parité : Une relation de redondance analytique est une équation dans laquelle toutes les variables sont connues. La génération de telles relations permet d'engendrer des résidus. Pour la détection de défauts basée sur l'utilisation de modèles, un résidu est un signal temporel, fonction des entrées et des sorties du processus, indépendant (le plus possible) du point de fonctionnement de celui-ci. En l'absence de défauts, ce résidu est statistiquement nul. Lors de l'apparition d'un défaut, son amplitude évolue de manière significative.

L'approche la plus classique est celle dite de l'espace de parité. Les relations de parité utilisent la redondance directe au moyen de relations algébriques statiques liant les différents signaux ou la redondance temporelle issue de l'utilisation de relations dynamiques. Le terme « parité » a été emprunté au vocabulaire employé pour les systèmes logiques où la génération de bits de parité permet la détection d'erreur.

B) Observateurs : La génération de résidus à l'aide d'une estimation d'état consiste à reconstruire l'état ou, plus généralement, la sortie du processus à l'aide d'observateurs et à utiliser l'erreur d'estimation comme résidu. Cette méthode s'est beaucoup développée car elle donne lieu à la conception de générateurs de résidus flexibles.

C) Estimation paramétrique : L'approche d'estimation paramétrique considère que l'influence de défauts se reflète sur les paramètres et non pas uniquement, comme c'est le cas des observateurs, sur les variables du système physique. Le principe de cette méthode consiste à estimer en continu des paramètres du procédé en utilisant les mesures d'entrée/sortie et en évaluant la distance qui les sépare des valeurs de référence de l'état normal du procédé. L'estimation paramétrique possède l'avantage d'apporter de l'information sur l'importance des déviations. Toutefois, un des

inconvenients majeurs de la méthode réside dans la nécessité d'avoir un système physique excité en permanence. Ceci pose donc des problèmes d'ordre pratique dans le cas de procédés dangereux, coûteux ou fonctionnant en mode stationnaire. De plus, les relations entre paramètres mathématiques et physiques ne sont pas toujours inversibles de façon unitaire, ce qui complique la tâche du diagnostic basé sur les résidus.

Les principaux outils numériques d'analyse de résidus sont les méthodes de décision statistiques, visant à affirmer ou à infirmer qu'un résidu est nul, qui permettront d'engendrer les signatures expérimentales de défauts.

1.3.2 Diagnostic Qualitatif

Parmi les différentes méthodes de détection et de diagnostic utilisant des modèles de raisonnement heuristique, nous trouvons principalement les suivantes :

- A) **Modèles issus de l'intelligence artificielle (IA) :** De manière globale, trois types de modèle sont utilisés en IA dans le contexte de la surveillance :
- a. Dans les modèles associatifs, les situations significatives de dysfonctionnement sont associées à des schémas d'alarmes par l'expression directe des liens entre les symptômes et la situation caractéristique. Ils reposent sur l'expertise existante, qu'elle soit issue des opérateurs du procédé ou des ingénieurs de maintenance. De tels modèles sont en particulier utilisés dans les systèmes experts et dans la reconnaissance de scénarii de défaillance.
 - b. Les modèles prédictifs permettent de simuler les comportements possibles d'un système dans ces différents modes (normal, dégradé ou défaillant). Par leur capacité à calculer pas à pas les comportements prédits, ces modèles sont directement utilisables pour la détection de défauts (par comparaison des comportements prédits avec les observations). Ils peuvent s'obtenir par abstraction de modèles numériques, par réseaux bayésiens ou réseaux probabilistes temporels.
 - c. Les modèles explicatifs, s'appuient sur l'expertise déjà explicitée sur l'unité industrielle sous forme de catalogue de dysfonctionnements ou de graphes.
-

B) Réseaux de neurones artificiels (RNA): Un RNA est un système informatique constitué d'un nombre de processeurs élémentaires (ou nœuds) interconnectés entre eux qui traite – de façon dynamique – l'information qui lui arrive à partir des signaux extérieurs. De manière générale, l'utilisation des RNA se fait en deux phases. Tout d'abord, la synthèse du réseau est réalisée et comprend plusieurs étapes : le choix du type de réseau, du type de neurones, du nombre de couches, des méthodes d'apprentissage. L'apprentissage permet alors, sur la base de l'optimisation d'un critère, de reproduire le comportement du système à modéliser. Il consiste dans la recherche d'un jeu de paramètres (poids) et peut s'effectuer de deux manières : supervisé (le réseau utilise les données d'entrée et de sortie du système à modéliser) et non supervisé (seules les données d'entrée du système sont fournies et l'apprentissage s'effectue par comparaison entre exemples). Quand les résultats d'apprentissage obtenus par le RNA sont satisfaisants, il peut être utilisé pour la généralisation. Il s'agit ici de la deuxième phase où de nouveaux exemples – qui n'ont pas été utilisés pendant l'apprentissage – sont présentés au RNA pour juger de sa capacité à prédire les comportements du système ainsi modélisé.

Leur faible sensibilité aux bruits de mesure, leur capacité à résoudre des problèmes non linéaires et multi-variables, à stocker les connaissances de manière compacte, à « **apprendre** » en ligne et en temps réel, sont des propriétés qui rendent l'utilisation des RNA attrayante. Leur emploi peut alors se faire à trois niveaux :

- comme modèle du système à surveiller en état normal et générer un résidu d'erreur entre les observations et les prédictions,
- comme système d'évaluation de résidus pour le diagnostic,
- ou comme système de détection en une seule étape (en tant que classificateur), ou en deux étapes (pour la génération de résidus et le diagnostic).

C) Systèmes d'inférence flous (SIF) : La structure de base d'un SIF est constituée d' :

- a. un univers de discours qui contient les fonctions d'appartenance des variables d'entrée et de sortie à des classes. Ces fonctions peuvent avoir différentes formes, les plus usuelles étant les formes triangulaires, trapézoïdales et gaussiennes,

- b. une base de connaissance qui regroupe les règles liant les variables d'entrée et de sortie sous la forme « *SI...ALORS* »,
- c. un mécanisme de raisonnement qui base son fonctionnement sur la logique du *modus ponens* généralisé.

Les formalismes les plus utilisés pour les SIF sont ceux de Mandani [Mandani, 1977] et de Takagi-Sugeno [Takagi et al., 1985].

D) Reconnaissance de formes : La reconnaissance des formes regroupe l'ensemble des méthodes permettant la classification automatique d'objets, suivant sa ressemblance par rapport à un objet de référence. L'objectif est de décider, après avoir observé un objet, à quelle forme type celui-ci ressemble le plus. Une forme est définie à l'aide de p paramètres, appelés descripteurs, qui sont les composantes d'un vecteur forme que nous noterons $\mathbf{X}$.

La résolution d'un problème de reconnaissance de formes nécessite :

- La définition précise des l classes entre lesquelles va s'opérer la décision. Cette étape est liée à la nature des objets à classer et est donc spécifique au problème posé.
- Le choix d'un jeu de caractères pertinents pour la discrimination des vecteurs de formes. Malheureusement, il n'existe pas de méthodes systématiques permettant de choisir les paramètres les plus appropriés à la résolution d'un problème donné. Par conséquent, seule une bonne connaissance du problème à résoudre peut permettre un choix adéquat des paramètres de forme (ou descripteurs). Notons enfin que le nombre de caractères fixe la dimension de l'espace de représentation et, par conséquent, la quantité des calculs à mener. Ceci peut représenter une contrainte sévère dans un contexte de traitement en temps réel des objets à classer.
- L'élaboration d'un classificateur permettant l'affectation d'une forme observée à l'une des classes. Le classificateur est généralement établi à l'aide d'un ensemble d'apprentissage constitué de formes pour lesquelles on connaît l'appartenance aux différentes classes.

Les méthodes présentées précédemment sont complémentaires entre elles. Il n'est donc pas étonnant de trouver des applications où les différents niveaux de traitement (signaux, modèles et connaissance experte) soient intégrés dans une même structure.

D'autres articles traitent du domaine du diagnostic de défauts basé sur des modèles [Frank et al., 2000], sur le diagnostic de défauts par extraction de caractéristiques et classification de problèmes (défauts). La classification des défauts est examinée par Kramer et Mah [Kramer et al., 1993] selon trois catégories principales :

- i) Reconnaissance de formes : ici sont traitées les méthodes basées sur l'historique du processus.
- ii) Modèles basés sur le raisonnement : ils utilisent les techniques basées sur des modèles qualitatifs.
- iii) Les modèles causaux : dans ce cas, les techniques de recherche de symptômes en utilisant différents modèles sont traités.

Une autre méthodologie de diagnostic de défaillances à prendre en compte est celle basée sur des modèles à événements discrets (SED) [Lafortune, 1991 ; Garcia et al., 1998] et sur le concept de diagnostiqueur [Sampath et al., 1994], [Sampath et al., 1995 ; Garcia et al., 1999], lesquels permettent tant l'analyse **hors ligne**, qu'**en ligne** de la diagnosabilité des systèmes, relative à des défaillances qui peuvent se produire dans les processus. Le diagnostiqueur a comme objectif de détecter la cause de quelques situations (états fonctionnels) du procédé, en un temps fini, comme la conséquence d'une défaillance. Tant les SED que le diagnostiqueur sont des machines à états finis (MEF). La méthodologie du diagnostic à base de modèles à événements discrets [Ramadge, 1987] est applicable sur beaucoup de systèmes industriels. Toutefois, malgré son intérêt notable, elle présente l'inconvénient des systèmes partiellement observables, dans le sens de l'inévitable augmentation exponentielle des états produits au fur et à mesure qu'on incorpore de nouveaux modèles de dispositifs qui forment des parties du système ou que s'ajoutent de nouvelles situations de défaillance. Pour éviter l'augmentation des états, une conception modulaire du système a été adoptée, en divisant le système en sous-systèmes. La complexité des systèmes peut être considérée en tenant en compte du nombre de sous-systèmes liés entre eux. L'application de cette méthodologie est restreinte par la complexité du procédé.

1.4 Méthodes de classification

Différentes approches pour la classification non supervisée « *clustering* » de données peuvent être décrites avec l'aide de la structure hiérarchisée de la Figure 1-5 [Jain et al., 1988]. Sur ce graphe, il y a distinction entre les approches par hiérarchisation et par partitionnement. Les méthodes par hiérarchisation produisent une série de partitions liées, alors que les méthodes de partitionnement n'en produisent qu'une seule.

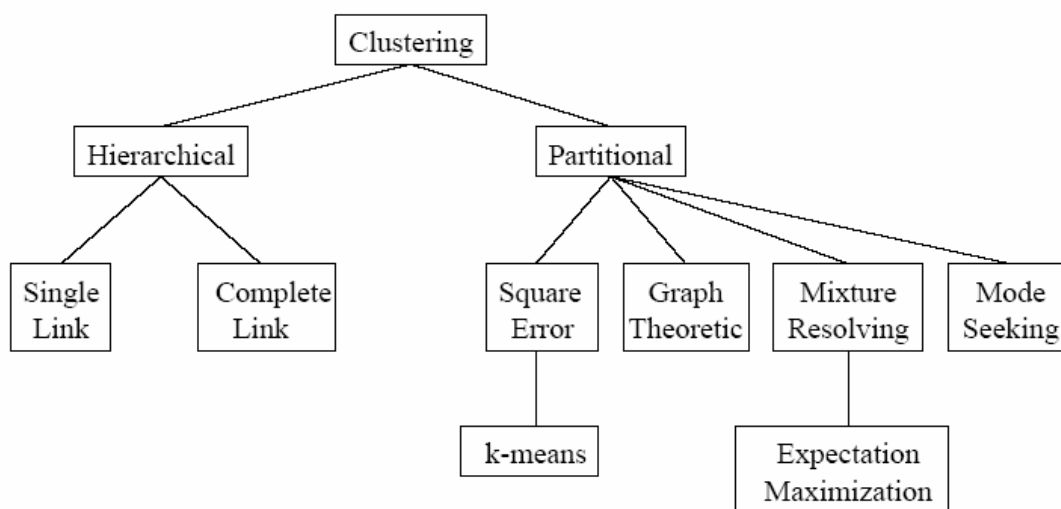


Figure 1-5 Différentes approches de classification non supervisée

Nous donnons, maintenant, une brève présentation des principales approches de classification non supervisée que nous considérons appropriées pour la mise en place de la stratégie proposée.

1.4.1 Algorithmes par hiérarchisation

L'algorithme par hiérarchisation crée une décomposition de données (ou d'éléments) en utilisant un critère. Le fonctionnement d'un algorithme par classification est illustré sur la Figure 1-6 en utilisant un ensemble de données en deux dimensions. Cette figure contient sept éléments étiquetés **A**, **B**, **C**, **D**, **E**, **F** et **G** dans trois classes. Un algorithme par hiérarchisation produit un dendrogramme qui représente le regroupement des éléments et les niveaux de similarité, permettant de faire figurer les changements de regroupement. Le dendrogramme qui correspond aux sept éléments de la Figure 1-6, obtenu à partir de l'algorithme *single-link*, est montré sur la Figure 1-7. Le dendrogramme est divisé en deux niveaux différents pour produire plusieurs classes de données.

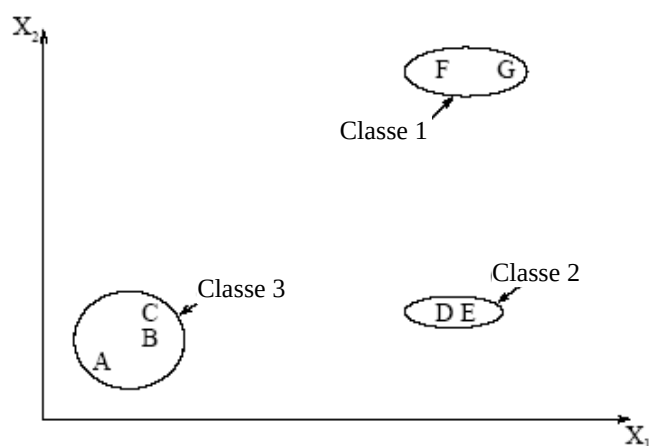
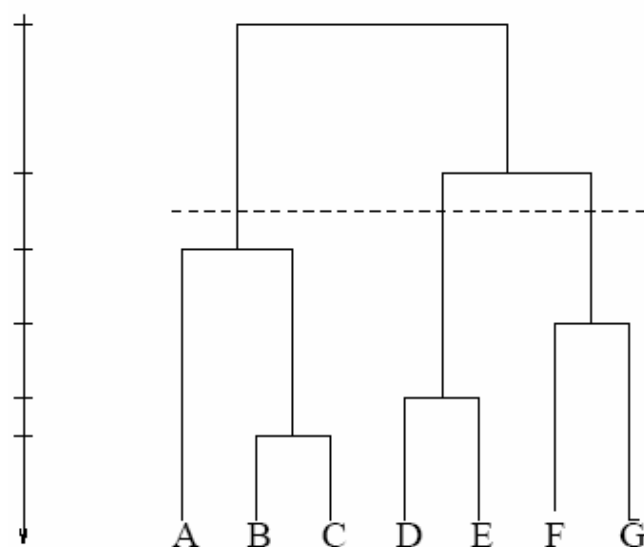


Figure 1-6 Points dans trois classes

Figure 1-7 Dendrogramme obtenu en utilisant l'algorithme *single-link*

Les autres algorithmes de classification similaires à celui développé par [Sneath et al., 1973] sont : l'algorithme *complete-link* par [King, 1967], et le *minimum de variance* [Ward, 1963 ; Murtagh, 1984]. Les deux algorithmes les plus populaires sont le *single-link* et le *complete-link*. Ces deux algorithmes se différencient dans la façon de caractériser la similarité entre une paire de classes. Dans la méthode *single-link*, la distance entre deux classes est le minimum des distances entre toutes les paires d'éléments appartenant à chaque classe (un élément dans la première classe, l'autre dans la seconde). Dans l'algorithme *complete-link*, la distance entre deux classes est le maximum des distances entre toutes les paires d'éléments sur les deux classes.

L'algorithme *complete-link* produit des classes compactes [Baeza-Yates, 1992]. Par contre, l'algorithme *single-link*, souffre d'un effet d'enchaînement [Nagy, 1968], c'est-à-dire, il tend à produire des classes allongées. L'algorithme *single-link* produit les classes telles qu'elles sont représentées par la Figure 1-8a, tandis que l'algorithme *complete-link* obtient les classes montrées sur la Figure 1-8b.

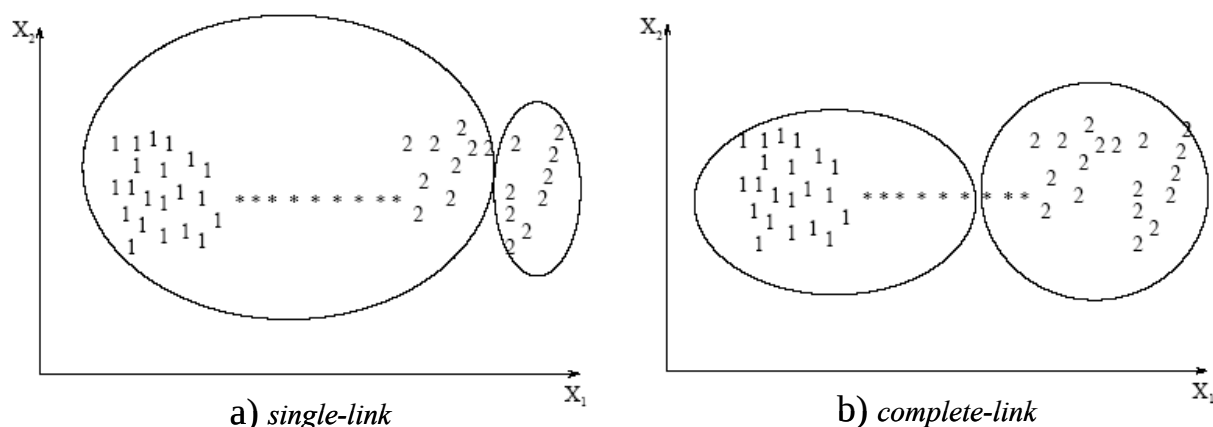


Figure 1-8 Clusters produits par les algorithmes de hiérarchisation sur un ensemble d'éléments contenant 2 classes.

Les classes obtenues par l'algorithme *complete-link* sont plus compactes que celles obtenues par l'algorithme *single-link*; la classe étiquetée comme un « 1 » obtenue en utilisant l'algorithme *single-link* (voir la Figure 1-8) est allongée à cause des éléments bruités étiquetés par étoiles « * ». Un avantage de l'algorithme *single-link* est qu'il peut extraire les classes concentriques telles que montrées dans la Figure 1-9, tandis que l'algorithme *complete-link* n'en est pas capable [Jain et al., 1988].

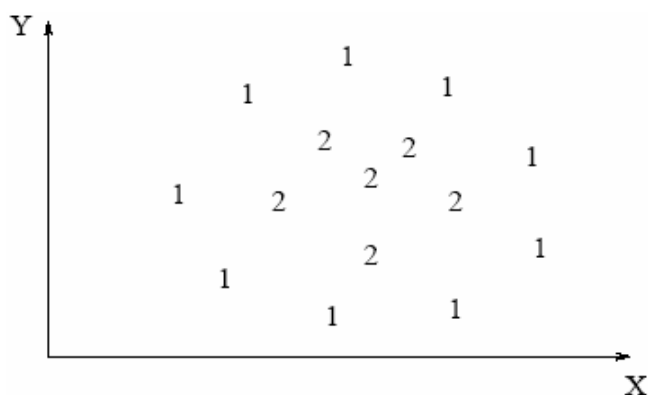


Figure 1-9 Deux classes concentriques

Procédure de l'algorithme *Single-Link* :

1. Placer chaque élément dans sa propre classe. Construire une liste de distances entre-éléments pour toutes les paires d'éléments non ordonnés et distincts, et ordonner cette liste par ordre ascendant.
2. A partir de la liste ordonnée des distances, former pour chaque valeur dissimilaire distincte d_k , un graphe sur les éléments où les paires d'éléments plus proches que d_k sont connectées par un graphe. Si tous les éléments sont membres d'un graphe connecté, alors arrêter. Sinon, répéter cette étape.
3. Le résultat de l'algorithme est un graphe hiérarchisé, lequel peut s'interrompre à un niveau de dissimilitude pour former une partition ou classe identifiée simplement à travers des composants connectés dans le graphe correspondant.

1.4.2 Algorithmes par partitionnement

Un algorithme de classification par partitionnement obtient une unique partition des données au lieu d'une structure hiérarchisée comme dans le cas précédent. Les méthodes de partitionnement ont l'avantage de manipuler une grande quantité d'ensembles de données, pour lesquelles, la construction d'un graphe hiérarchisé permettant de tout couvrir serait difficile. Un des problèmes relatifs aux algorithmes de partitionnement est le choix du nombre de classes désirées en sortie. Dans [Dubes, 1987] un guide pour choisir ce paramètre clé pour la conception est donné. Cette technique produit des classes en optimisant une fonction critère définie soit localement (sur un sous-ensemble d'éléments) soit globalement (définie sur l'ensemble des éléments). L'algorithme est réitéré de multiple fois avec différents nombres de classes de sortie, et la meilleure configuration obtenue est utilisée pour le choix de la classe de sortie.

- Algorithmes de minimisation de l'erreur quadratique

Cet algorithme est basé sur la minimisation d'un critère quadratique, critère qui est le plus utilisé dans les algorithmes de partitionnement, et qui est bien adapté dans le cas de classes séparées et compactes. L'erreur quadratique pour une classe L parmi K classes d'un ensemble d'éléments H est définie par :

$$e^2(H, L) = \sum_{j=1}^K \sum_{i=1}^{n_j} \|x_i^{(j)} - c_j\|^2 \quad (1.1)$$

où $x_i^{(j)}$ est le $i^{\text{ème}}$ élément qui appartient à la $j^{\text{ème}}$ classe et c_j est le centre de la $j^{\text{ème}}$ classe.

L'algorithme *k-means* est le plus simple et le plus connu des algorithmes de minimisation d'un critère quadratique [McQueen, 1967]. Il débute avec une partition initiale aléatoire et il déplace les éléments dans les classes sur la base de la similarité entre l'élément et le centre de la classe par minimisation du critère ((1.1) jusqu'à convergence (quand il n'y a pas de re-affectation d'élément d'une classe à une autre, ou lorsque l'erreur quadratique ne décroît plus significativement après un certain nombre d'itérations). L'algorithme *k-means* est populaire parce qu'il est facile à implanter. Un des problèmes majeurs de cet algorithme est qu'il est sensible au choix de la partition initiale et qu'il peut converger vers un minimum local de la fonction critère si la partition initiale n'est pas choisie de manière adéquate. La Figure 1-10 montre sept éléments en deux dimensions. Si nous commençons avec les éléments A, B et comme moyennes initiales les centres, trois classes sont construites, ceci conduit à la partition $\{\{A\}, \{B, C\}, \{D, E, F, G\}\}$ où les classes sont représentées par des ellipses. La valeur du critère de l'erreur quadratique est beaucoup plus grande pour cette partition que pour la meilleure partition : $\{\{A, B, C\}, \{D, E\}, \{F, G\}\}$ représentée par des rectangles, laquelle produit la valeur minimale globale de la fonction critère de l'erreur quadratique pour une classification sur trois classes. La solution correcte (trois classes) est obtenue en choisissant, par exemple, A, D et F comme les moyennes initiales des classes.

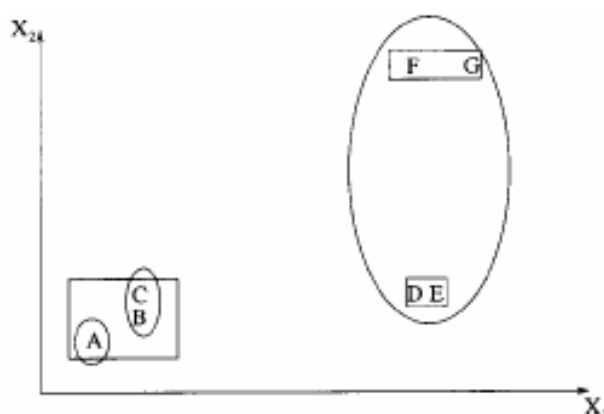


Figure 1-10 Sensibilité de l'algorithme *k-Means* à la partition initiale

Procédure de l'algorithme *k-Means*

1. Choisir k centres de classes pour coïncider avec les k points définis aléatoirement dans l'hyper volume contenant l'ensemble d'éléments.
2. Assigner chaque élément au centre de la classe la plus proche.
3. Recalculer les centres des classes en utilisant les nouveaux éléments de la classe.
4. Si un critère de convergence n'est pas satisfait, répéter l'étape 2. Ces critères de convergence typique sont : pas de nouvelle re-affectation d'éléments à un nouveau centre de classe, ou non-évolution du critère de l'erreur quadratique.

Plusieurs variantes de l'algorithme *k-Means* ont été rapportées dans la littérature [Anderberg, 1973]. Quelques algorithmes essaient de choisir une bonne partition initiale, pour que l'algorithme ait plus de possibilité de trouver la valeur du minimum global.

Une autre variante est de permettre de diviser et de fusionner les classes résultantes. Typiquement, une classe est divisée quand sa variance est au-dessus d'un seuil préétabli, et deux classes sont fusionnées quand la distance entre leurs centres est au-dessous d'un autre seuil pré-spécifié. En utilisant cette modification, il est possible d'obtenir une partition optimale en commençant par une partition initiale arbitraire, avec des valeurs de seuils correctement choisis.

L'algorithme ISODATA [Ball et al., 1965] utilise cette technique pour diviser et fusionner les classes. Si ISODATA prend comme partition initiale les ellipses de la Figure 1-10, il produira la partition optimale composée de trois classes. ISODATA d'abord fusionnera les classes {A} et {B, C} dans une classe parce que la distance entre leurs centres est la plus petite et ensuite divisera la classe {D, E, F, G} qui possède une grande variance, en deux classes : {D, E} et {F, G}.

1.4.3 Algorithmes issus de la théorie des graphes

Tout d'abord, définissons un concept important pour ce type d'algorithme. L'algorithme STA (« *Spanning Tree Algorithm* ») est un protocole de dialogue entre tous les ponts permettant de supprimer des liens redondants et de ne garder que les liens les plus intéressants (lignes de communications les plus rapides, les moins coûteuses, etc.).

Le plus connu de ces algorithmes est basé sur la construction du « MST » (*minimal spanning tree*) des données [Zahn, 1971]. Les segments du MST avec les plus grandes longueurs sont ensuite supprimés pour produire les classes. La Figure 1-11 représente le MST obtenu à partir de neuf points en deux dimensions. En supprimant le segment étiqueté CD avec une longueur de 6 unités (le segment avec la longueur Euclidienne maximale), deux classes ($\{A, B, C\}$ et $\{D, E, F, G, H, I\}$) sont obtenues. La deuxième classe peut être ultérieurement divisée en deux classes en supprimant la branche EF, laquelle a une longueur de 4.5 unités.

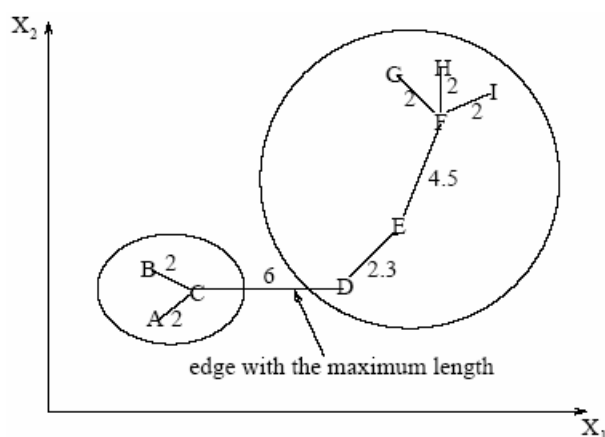


Figure 1-11 Classes formées en utilisant le MST (« *Minimal Spanning Tree* »)

1.4.4 Algorithme de résolution mixte « Mixture-Resolving »

Le fondement de l'approche « *Mixture-Resolving* » est que les éléments sont obtenus à partir de plusieurs distributions (par exemple, gaussiennes), et le but est d'identifier les paramètres ainsi que, leur nombre. Il consiste à obtenir (itérativement) une estimation de la probabilité maximale (Maximisation de l'espérance mathématique) du vecteur de paramètres des composants de densités [Jain et Dubes, 1988].

1.4.5 Classification par le plus proche voisin

Une procédure itérative a été proposée par Lu et Fu [1978]. Elle assigne chaque élément non étiqueté à la classe de son plus proche élément voisin étiqueté, si la distance au voisin étiqueté est au-dessous d'un seuil. Le processus continue jusqu'à ce que tous les éléments soient étiquetés ou qu'il ne se produise plus d'étiquetage.

1.4.6 Classification Floue

Les approches traditionnelles de classification génèrent des partitions; dans une partition, chaque élément appartient à une seule classe. La classification floue amplifie cette notion en associant chaque élément à chacune des classes en utilisant des fonctions d'appartenance [Zadeh, 1965].

- Procédure itérative de l'algorithme de classification floue

1. Sélectionner une partition floue initiale des N éléments dans K classes en sélectionnant la matrice d'appartenance U de dimension $N \times K$. Un élément u_{ij} de cette matrice représente le degré d'appartenance de l'élément x_i dans la classe c_j . Typiquement, $u_{ij} \in [0, 1]$.
2. Minimiser la valeur de la fonction critère floue, par exemple :

$$E^2(H, U) = \sum_{i=1}^N \sum_{k=1}^K u_{ik} \|x_i - c_k\|^2 \quad (1.2)$$

où $c_k = \sum_{i=1}^N u_{ik} x_i$ est le centre de la $k^{\text{ième}}$ classe, en modifiant de façon itérative la matrice U de manière à modifier l'assignation des éléments aux classes et ainsi faire décroître le critère $E^2(H, U)$.

3. Répéter l'étape 2 jusqu'à ce que la valeur de U ne change pas significativement.

Avec la classification floue, chaque classe est un ensemble flou de tous les éléments. La Figure 1-12 illustre cette approche. Les rectangles représentent deux classes typiques : $H_1 = \{1, 2, 3, 4, 5\}$ et $H_2 = \{6, 7, 8, 9\}$. Un algorithme de classification floue peut produire les deux classes floues F_1 et F_2 représentées par des ellipses. Les éléments auront des valeurs d'appartenance entre 0 et 1 pour chacune des classes. Par exemple, la classe floue F_1 pourrait être décrite comme :

$\{(1, 0.9), (2, 0.8), (3, 0.7), (4, 0.6), (5, 0.55), (6, 0.2), (7, 0.2), (8, 0.0), (9, 0.0)\}$

et F_2 comme :

$\{(1, 0.0), (2, 0.0), (3, 0.0), (4, 0.1), (5, 0.15), (6, 0.4), (7, 0.35), (8, 1.0), (9, 0.9)\}$

Les paires ordonnées (i, μ_i) pour chaque classe représentent le $i^{\text{ème}}$ élément et sa valeur d'appartenance à la classe μ_i .

Le plus connu des algorithmes de classification floue est le « *C-Means Flou* » (FCM). Cette méthode, développée par Dunn en 1973 [Dunn, 1973] et améliorée par Bezdek en 1981 [Bezdek, 1981], est souvent utilisée en reconnaissance de formes.

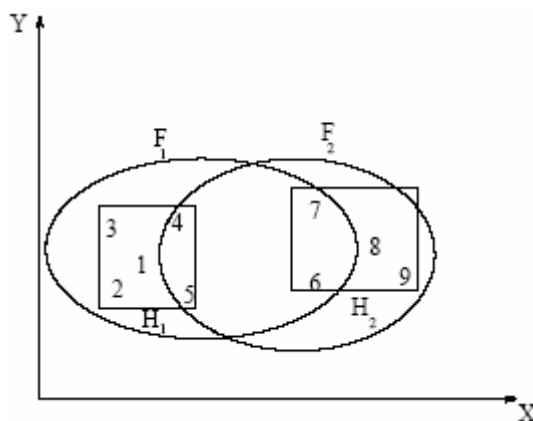


Figure 1-12 Classes floues

1.5 La méthode LAMDA

LAMDA (en anglais : « *Learning Algorithm for Multivariate Data Analysis* ») [J. Aguilar-Martin, 1980 ; Koodmen, 1989] est une technique de classification qui a été développée au Laboratoire d'Analyse et Architecture des Systèmes (LAAS). C'est un Algorithme d'Analyse de Données Multidimensionnelles par Apprentissage et Reconnaissance de Formes. La formation et la reconnaissance de classes dans cette méthode sont basées sur l'attribution d'un élément à une classe à partir de la règle heuristique appelée *adéquation maximale*. Cette règle d'adéquation a été proposée par Aguilar-Martin [J. Aguilar-Martin, 1980]. Chaque élément est représenté par des descripteurs. Le degré d'adéquation maximale d'un élément à une classe est déterminé à partir des degrés d'adéquation marginale de chaque descripteur dans les différentes classes.

LAMDA a été utilisé en reconnaissance de forme [Piera et Carreté et al., 1990], en analyse biomédicale [Chan et al., 1989], pour l'étude des processus de dépollution des eaux usées [Waissman-Vilanova et al. 2000] et pour les bio-procédés [Aguilar-Martin et al., 1999].

LAMDA est une technique de classification multi-source qui utilise des sous-ensembles flous et dont la démarche pour classer les éléments est la suivante :

1. Calcul du degré d'appartenance marginale pour chaque descripteur par l'utilisation d'une fonction d'appartenance. Ce degré est appelé Degré d'Adéquation Marginale (MAD).
2. Calcul d'un degré d'appartenance globale par l'utilisation d'un opérateur d'agrégation de l'information. Ce degré est appelé Degré d'Adéquation Globale (GAD).
3. Évaluation du GAD pour savoir s'il faut créer une nouvelle classe
 - a. Si c'est le cas, création d'une nouvelle classe et nouvelle évaluation des MAD et du GAD pour tous les points traités préalablement.
 - b. Sinon assignation de l'élément traité dans la classe ayant le plus grand GAD.

Le schéma général de la procédure d'attribution d'un élément à une classe est présenté sur la Figure 1-13 :

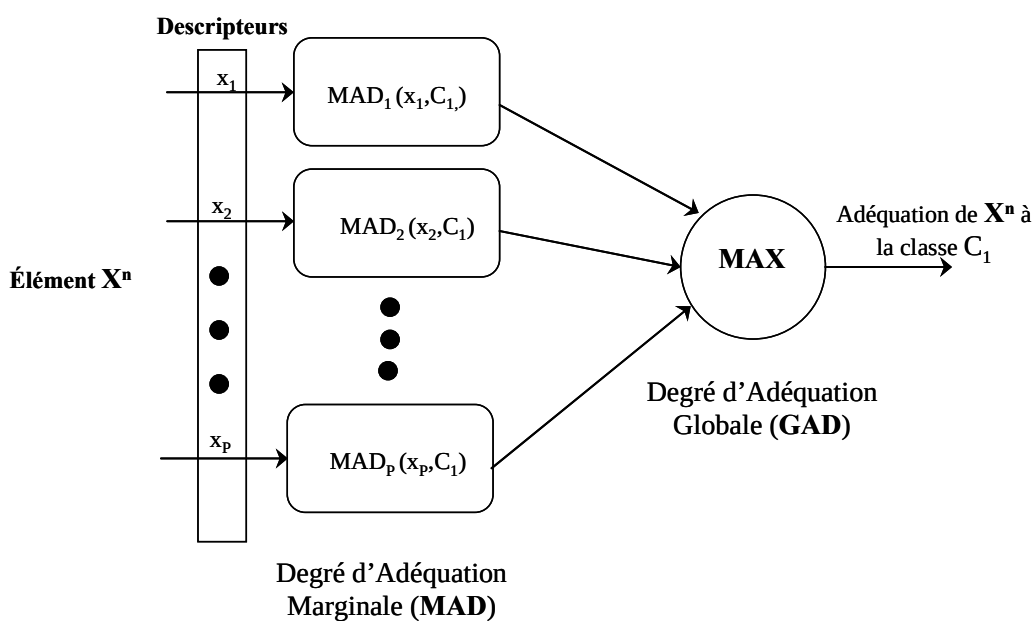


Figure 1-13 Calcul de l'adéquation d'un élément

1.5.1 Les caractéristiques de LAMDA

La méthode LAMDA est une méthode de classification issue de la logique floue, laquelle repose sur des notions de base comme la notion d'incertitude, d'imprécis, de vague ou de sous-ensemble flou [Zadeh, 1965]. Dans ce qui suit, les caractéristiques propres [Piera et al., 1989] sont présentées:

- **L'Adéquation Maximale (règle heuristique) :** LAMDA n'utilise pas de la similarité entre les éléments pour la classification, mais il calcule le *degré d'adéquation globale* entre l'élément et les classes déjà formées, et il affecte chaque élément à la classe dont le degré d'adéquation est maximal.
- **Le Concept de la classe non informative (NIC):** ce concept est à la base de la modélisation de l'homogénéité dans l'Univers, d'où proviennent tous les éléments. Initialement ces éléments sont issus d'une classe qui accepte tous les éléments avec le même degré d'adéquation. Cette classe est très importante pour le processus de formation d'une nouvelle classe. L'existence de cette classe agit comme une limitation ou un seuil : aucun élément ne sera assigné à une autre classe si son degré d'adéquation globale n'est pas supérieur à celui de la classe non informative.
- **Degré d'adéquation avec Connectifs :** Un élément sera affecté dans une classe non seulement si son degré d'appartenance à cette classe est maximal ou si au moins un de ses descripteurs a un degré d'adéquation marginale élevé, mais par l'intermédiaire d'un critère qui implique les connectifs « ET » et « OU »
- **La méthode LAMDA est issue du domaine de la Logique Floue**
- **Elle permet de gérer simultanément des variables qualitatives et quantitatives**
- **Apprentissage séquentiel :** Il peut s'adapter à une situation évoluant au cours du temps en raison d'un algorithme d'apprentissage séquentiel.

1.5.2 Description de la méthodologie LAMDA

LAMDA est une méthode de classification de données multidimensionnelles par apprentissage et reconnaissance de forme. Elle s'adresse à des séquences d'éléments décrits par un nombre fixe de composantes de type qualitatif ou quantitatif (nommés aussi descripteurs). Ces données peuvent être traitées dans un contexte probabiliste ou avec une approche floue. A chacune de ces deux méthodes est associé un connectif déterminé.

Il y a deux types de classification :

- Par Apprentissage (sans et avec initialisation) ou apprentissage non supervisé
 - Par Reconnaissance de Formes ou apprentissage supervisé.
-

A) Cas de l'apprentissage

L'« intelligence » de la machine se manifeste en apprenant les caractéristiques des classes associées à des partitions. Le paramétrage de ces classes est dynamique, leurs caractéristiques évoluent avec chaque affectation d'un nouvel élément.

Il y a deux modes de classification par apprentissage :

- Apprentissage sans initialisation : étant donné un ensemble d'éléments, on fait une classification qui donne la meilleure partition de ces éléments. En outre, il y a création des classes et conservation de leurs paramètres. *La classe « NIC » est la seule classe initialisée.*
- Apprentissage avec initialisation: Si on a une idée sur les caractéristiques de quelques classes, on peut en établir d'autres en se basant sur les classes déjà créées. Il y a possibilité de création de nouvelles classes et d'évolution des paramètres des anciennes classes. Dans le cas où il n'y aurait plus de possibilité d'ajout, les éléments non classés sont affectés à la classe NIC.

B) Cas de la Reconnaissance de Formes

On dispose a priori d'un certain nombre de classes dont les paramètres sont connus définitivement, et il s'agit de traiter un ensemble d'éléments par leur assignation à ces classes. Lors de l'affectation, les paramètres ne sont pas changés ce qui est normal puisqu'il s'agit d'une exploitation des connaissances et non pas de l'apprentissage. Les éléments non classés sont affectés à une classe appelée « classe résiduelle », qui a les mêmes caractéristiques que la classe NIC dans le cas de l'apprentissage.

Ces deux types de classification peuvent être complémentaires. On peut commencer par faire un apprentissage sans initialisation, puis effectuer une série de classifications avec initialisation en se basant sur les classes déjà formées pour les faire évoluer et arriver à une partition satisfaisante. Cet ensemble de classes est donc efficace pour faire, par la suite, la reconnaissance de formes.

Cet algorithme de classification n'est pas commutatif, c'est-à-dire que la partition finale dépendra de l'ordre dans lequel les données sont traitées. Dans les premiers pas d'un processus d'apprentissage, on peut affecter un élément à une classe non encore parfaitement

définie (vu qu'elle possède encore peu d'éléments) et à la fin, une fois que cette classe est mieux définie, si on refait le classement du même élément il peut avoir un degré d'appartenance très différent, supérieur ou inférieur par rapport à l'ancien.

Une caractéristique importante de LAMDA est qu'il modélise le concept de « Chaos », qui est la totale homogénéité. Il accepte tous les éléments de la même manière. Ce concept est représenté par une classe qui s'appelle « NIC » (non – informative – classe). Le degré d'appartenance à la classe NIC est le même pour tous les éléments. C'est le contexte qui détermine l'existence de cette classe, et ceci implique que la classe NIC est toujours présente dans l'espace des classes déjà établi. La classe NIC est invariante dans le temps.

1.5.3 Les Concepts de LAMDA

Pour mieux comprendre la méthode LAMDA, dans ce qui suit on présente les concepts et le vocabulaire qui lui sont spécifiques :

- **L'Univers de la Classification** est l'ensemble des possibilités ou éléments qui peut être représenté par une grille. Il représente le contexte dans lequel un élément peut être décrit.
- **Descripteurs (P)**: Ils peuvent être quantitatifs ou qualitatifs.

Descripteur quantitatif: Intervalle qui rassemble toutes les valeurs des éléments.

Descripteur qualitatif: Ensemble de modalités ou d'étiquettes qualitatives. A chacune est associée sa fréquence d'apparition.

- **Un ensemble de "N" éléments ($X_1, \dots, X_N$)**, où chacun est représenté par un vecteur $[x_1, \dots, x_p]$ tel que x_j est la caractéristique de l'élément pour le descripteur j .

$$X_1 = [x_1, \dots, x_p]$$

$$X_2 = [x_1, \dots, x_p]$$

$$\vdots$$

$$\vdots$$

$$\vdots$$

$$X_N = [x_1, \dots, x_p]$$

- **K Classes [$C_0, C_1, \dots, C_K$]** déjà existantes ou créées, où C_0 est la classe NIC. Chacune est définie par un vecteur de paramètres: $[\rho_1, \dots, \rho_p]$.

Si le descripteur est quantitatif, le paramètre ρ_j représente la moyenne des valeurs correspondantes aux éléments affectés à cette classe.

Si le descripteur est qualitatif, ρ_j sera la liste des modalités avec la probabilité reflétant son poids dans cette classe.

- **Connectifs**

Les connectifs sont reliés aux opérateurs logiques entre les ensembles, tels que l'union et l'intersection qui sont associées aux conjonctions linguistiques « *ET* » et « *OU* ».

Les connectifs sont associés à des fonctions de types T-norme et T-conorme. Parmi les familles de connectifs, les plus connues ont été choisies : le minimum (T-norme) et le maximum (T-conorme) de la théorie de la Logique Floue de Zadeh [Zadeh, 1978] et celles du produit et de la somme probabiliste qui présupposent un raisonnement de type statistique avec une indépendance entre les descripteurs.

Soit I l'intervalle $[0,1]$, si T est une fonction T-norme, la duale T-conorme correspondante est définie comme suit :

$$S(x, y) = I - T(I - x, I - y) \quad (1.3)$$

Les fonctions associées au contexte probabiliste et à l'approche floue sont les suivantes :

Tableau 1-A Fonctions T-Norme et T-Conorme

Connectif	T-Norme (intersection) T(X,Y)	T-conorme (union) S(X,Y)
Produit / Somme probabiliste	$X \cdot Y$	$1 - (1-X)(1-Y) = X+Y-X \cdot Y$
Minimum / Maximum	$\text{Min}(X, Y)$	$1 - \text{Min}(1-X, 1-Y) = \text{Max}(X, Y)$

Les connectifs mixtes constituent une extension à cette théorie se situant entre l'intersection et l'union. Ils dépendent d'un paramètre α appartenant à l'intervalle $[0,1]$:

$$CM(x, y) = \alpha T(x, y) + (1 - \alpha) S(x, y) \quad (1.4)$$

α sera maximum dans le cas de l'intersection et minimum dans le cas de l'union et est appelé « exigence ».

Le Tableau 1-B résume ce que l'on vient d'énoncer

Tableau 1-B Connectifs mixtes

Opérateurs Logiques		Intersection	Union
Conjonctions		ET	OU
Connectifs	Probabilistes	Produit	Somme
	Flous	Minimum	Maximum
	Mixtes	$\alpha = \quad 0 \quad \dots \quad 0.5 \quad \dots \quad 1$	

On peut remarquer dès maintenant que le résultat dépendra essentiellement du connectif.

- **Degré d'appartenance**

Le degré d'appartenance d'un élément à une classe donnée reflète la distance qui les sépare. Pour un élément donné, les caractéristiques par rapport à chaque descripteur interviennent dans le calcul du degré d'appartenance de cet élément à une classe par ce qu'on appelle « le degré d'adéquation marginale ». Pour chaque élément, on détermine un vecteur des degrés d'appartenance marginale. Ainsi, la distance de cet élément à la classe courante, qui est son degré d'appartenance, résume ou agrège toutes les informations données par les descripteurs. Donc il est fonction des appartenances marginales, et il s'appelle « degré d'adéquation globale ».

$GAD(X/C)$ est le degré d'appartenance globale d'un élément X à une classe C , $\mu_j = MAD(x_j/C)$ est le degré d'appartenance marginale (ou partielle) par rapport au descripteur j , et $[\mu_1, \dots, \mu_j \dots, \mu_p]$ est le vecteur des appartenances marginales.

1.5.4 L'algorithme de la classification

Dans l'annexe A de cette thèse, on peut trouver un exemple simple du développement de l'algorithme en utilisant des données quantitatives, permettant de mieux comprendre cette méthode. Dans ce qui suit, nous donnons de façon détaillée les différentes étapes de l'algorithme de classification.

Soit un élément X et les classes $C_0, C_1, \dots, C_K$, une classification se déroule de la façon suivante:

Calculer les degrés d'appartenance globale (GAD) de l'élément X par chacune des classes $C_1, \dots, C_K$ et C_0 la classe vide ou résiduelle, notés $[\rho_1, \dots, \rho_p]$. Pour ce faire, on calcule des

degrés d'appartenance marginale (μ_j) par rapport à chaque descripteur. Le calcul du degré d'appartenance marginale ou partielle dépend du type de descripteur correspondant :

A) Cas des descripteurs qualitatifs :

Un descripteur qualitatif est caractérisé par un ensemble non ordonné de modalités. Lors de la classification, on procède par le calcul des fréquences de chaque modalité à l'intérieur d'une classe. Les numéros associés à chacune des différentes modalités sont les valeurs qui représentent le mieux la classe. On les utilise pour le calcul de la fonction d'appartenance d'un élément selon la modalité observée pour chaque composante.

B) Cas des descripteurs quantitatifs :

Les descripteurs quantitatifs sont tels que les valeurs associées peuvent se mettre dans un ensemble ordinal discret ou continu. Cet ensemble se présente donc comme un intervalle $[x_{\min}, x_{\max}]$ et peut être réduit par la formule de normalisation suivante :

$$x_j = \frac{x_j - x_{\min}}{x_{\max} - x_{\min}} \quad (1.5)$$

Il existe plusieurs algorithmes pour calculer l'appartenance d'un descripteur. Dans ce qui suit, nous donnons les 5 algorithmes que nous avons utilisés : Binomiale, Binomiale-Centré, Binomiale-Distance, Binomiale-Distance au Carré et Gauss.

1. **Binomiale:** Dans ce cas, le MAD se calcule par la formule suivante :

$$\mu(x_j / C_i) = \rho_{i,j}^{x_j} (1 - \rho_{i,j})^{(l-x_j)} \quad (1.6)$$

2. **Binomiale-Centré:** ce cas est une variante du cas précédent, en introduisant un centre :

$$par = \rho_{i,j}^{x_j} \left(1 - \rho_{i,j} \right)^{(l-x_j)} \quad (1.7)$$

$$des = x_j^{x_j} (1 - x_j)^{(l-x_j)} \quad (1.8)$$

$$\mu(x_j / C_i) = \frac{par}{des} \quad (1.9)$$

3. **Binomiale-Distance:** Comme son nom l'indique, on fait intervenir une pseudo-distance :

$$a = \max[\rho_{i,j}^{x_j}, (1 - \rho_{i,j})] \quad (1.10)$$

$$x_{dist} = 1 - \text{abs}(x_j - \rho_{i,j}) \quad (1.11)$$

$$\mu(x_j / C_i) = (a)^{x_{dist}} (1 - a)^{(1 - x_{dist})} \quad (1.12)$$

4. **Binomiale-Distance au Carré:** Par rapport à l'algorithme 3, on introduit une normalisation :

$$a = \max[\rho_{i,j}^{x_j}, (1 - \rho_{i,j})] \quad (1.13)$$

$$x_{dist} = 1 - \text{abs}(x_j - \rho_{i,j}) \quad (1.14)$$

$$b = x_{dist}^{x_{dist}} (1 - x_{dist})^{(1 - x_{dist})} \quad (1.15)$$

$$\mu(x_j / C_i) = \frac{a^{x_{dist}} (1 - a)^{(1 - x_{dist})}}{b} \quad (1.16)$$

5. **Gauss:** Dans ce cas, les relations utilisées sont à rapprocher de celles donnant la moyenne et l'écart type d'une distribution gaussienne :

$$\mu(x_j / C_{ij}) = e^{-\frac{1}{2\sigma_{ij}^2}(x_j - \mu_{ij})^2} \quad (1.17)$$

Où μ_{ij} et σ_{ij}^2 correspondent, respectivement, à la valeur moyenne et à la variance du descripteur j pour la classe i .

L'étape suivante consiste, à l'aide du connectif, à déterminer le degré d'appartenance globale (GAD) de l'élément X à la classe C_j .

Les appartenances marginales pour chaque descripteur, permettent, à l'aide d'un connectif au choix (produit, somme, minimum, maximum), de calculer l'appartenance d'un élément à chacune des classes.

Cet élément est assigné à la classe la plus proche dont le degré d'appartenance globale correspondant est maximal. S'il y en a plusieurs, on choisit la première trouvée.

Enfin, l'actualisation des paramètres associés aux descripteurs quantitatifs se fait de la façon suivante:

$$\rho_{i,j} = \rho_{i,j} + \frac{x_j - \rho_{i,j}}{N + 1} \quad (1.18)$$

Pour procéder séquentiellement, il est nécessaire de connaître le nombre d'éléments (N) ayant servi au calcul des fréquences dans la classe correspondante.

On peut créer aussi, à l'aide de la classe vide (NIC) une nouvelle classe qui va être caractérisée par l'affectation d'un élément à cette classe. X est le premier élément d'une nouvelle classe C_{K+1} et la représentation de cette nouvelle classe dépendra de cet élément. On prendra un paramètre fictif correspondant au «nombre d'éléments de la classe « NIC ».

$$\rho_{i,j} = \rho_{i,j} + \frac{x_j - \rho_{i,j}}{N_0 + 1} \quad (1.19)$$

Le paramètre N_0 détermine l'initialisation de l'apprentissage. Sa valeur peut être choisie arbitrairement puisqu'elle n'influe pas sur le résultat de la classification (mais elle doit être positive).

Dans le cas de la reconnaissance de formes, l'élément est attribué à une classe significative ou rejetée dans la classe résiduelle. Dans le cas de l'apprentissage, s'il est affecté à une classe significative il y a modification des paramètres de cette classe. En revanche, si la classe vide est la plus proche, une nouvelle classe doit être créée pour contenir cet élément. Il y a rejet si la classe vide est la plus proche et qu'il n'y a pas possibilité de création de nouvelle classe parce que le nombre maximum des classes créées est atteint par exemple.

L'algorithme général de la classification est donné sur la Figure 1-14 :

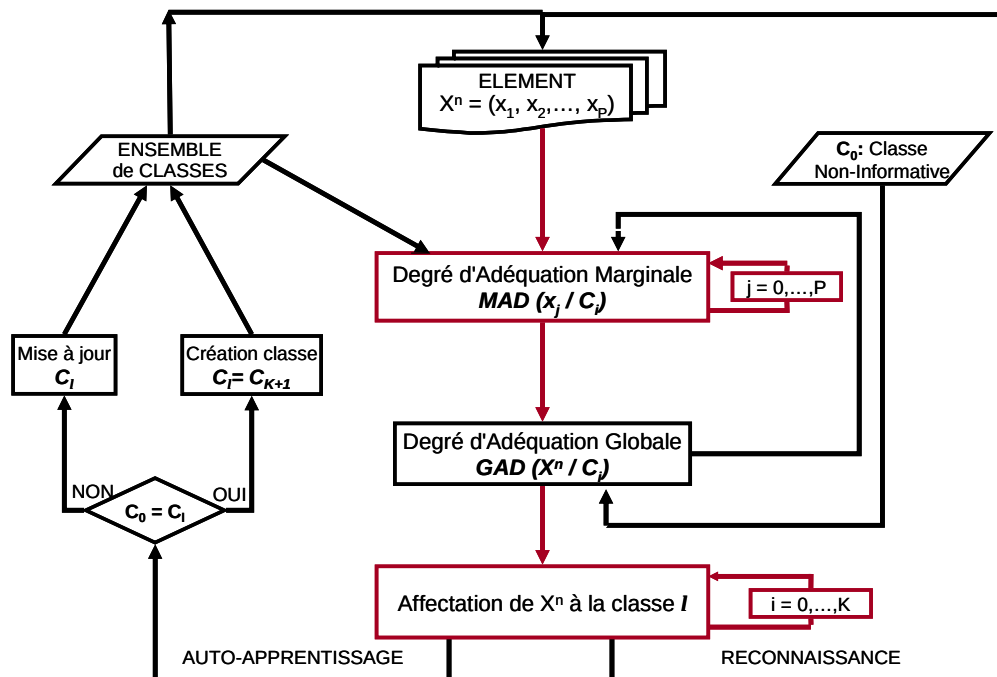


Figure 1-14 Algorithme général de LAMDA

Le paragraphe suivant est consacré à une brève description de l'outil *SALSA* développé sur la base de la méthode LAMDA au cours du projet CHEM [CHEM].

1.5.5 L'outil SALSA

SALSA (*Situation Assessment using Lamda claSsification Algorithm*) est un outil développé sous LabWindows par [Kempowsky, 2004b] pour déterminer l'état fonctionnel (situation) d'un processus en traitant des données en ligne et hors ligne. L'objectif de la phase hors ligne est la conception d'un modèle de comportement du processus à partir du résultat de la classification effectuée sur des données historiques ou expérimentales. La phase en ligne est effectuée pour reconnaître et déterminer l'état fonctionnel actuel d'un processus en temps réel, en utilisant le modèle de comportement obtenu durant la phase hors ligne. La Figure 1-15 et la Figure 1-16 schématisent ces deux phases d'utilisation :

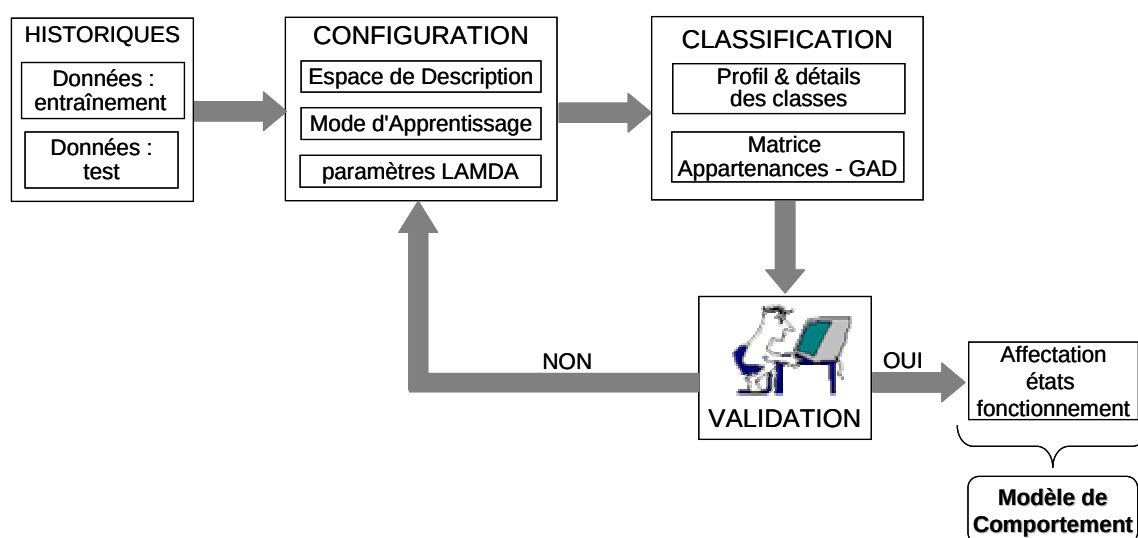


Figure 1-15 Structure de SALSA hors ligne

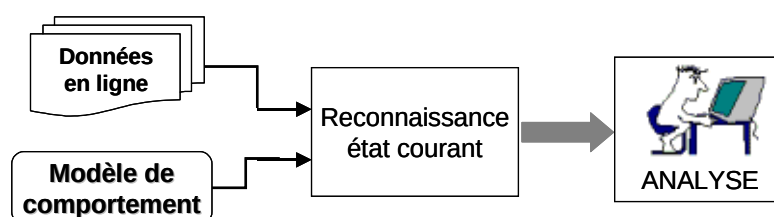


Figure 1-16 Structure de SALSA en ligne

L'avantage de l'outil Salsa pour faire le diagnostic est qu'il n'a pas besoin d'un modèle initial ni analytique ni issu de l'intelligence artificielle (logique floue, réseau neuronal). En revanche, il nécessite l'avis d'un expert pour valider l'affectation des états de fonctionnement du processus à des classes et obtenir ainsi, le modèle de comportement.

Comme caractéristiques principales de SALSA, nous trouvons qu'il :

- Accepte aussi bien de l'information qualitative que quantitative.
- Permet indifféremment l'apprentissage non supervisé (« *clustering* ») et l'apprentissage supervisé (classes prédéfinies).
- Permet d'obtenir plusieurs classifications depuis un ensemble de données.
- Nécessite un nombre minimum de paramètres à régler par l'opérateur.
- Dispose d'un algorithme simple et rapide pour la reconnaissance en ligne.
- Est facile pour l'installation et la configuration.

L'outil Salsa a été développé dans le cadre du projet européen CHEM (*Advanced Decision Support Systems for Chemical and Petrochemical Manufacturing Processes*) dont l'objectif principal a été le développement d'une plateforme générique d'outils intégrés basés sur des méthodologies avancées pour la surveillance, la supervision, la détection de défauts et le diagnostic des procédés chimiques, pétrochimiques et de raffinage, afin d'en améliorer la sécurité, la qualité des produits et la qualité de leur fonctionnement. L'outil SALSA a été intégré à cette plateforme [CHEM].

1.6 Conclusions

L'objectif du diagnostic est de détecter rapidement les divers défauts existant sur un procédé (défaut d'instrumentation, de paramètres, structurel) pour éviter la dégradation de ses performances et augmenter la sécurité des opérateurs et de l'environnement.

Ce chapitre a fourni une présentation succincte des différentes techniques de diagnostic classiquement utilisées. Il est clair qu'un très grand nombre de développements a concerné les approches basées sur des modèles quantitatifs du processus utilisant la technique du calcul des résidus. Cependant différents facteurs tels que la complexité du système [Bailey, 1984], la dimension élevée, la non linéarité des processus ont souvent rendu très difficile le développement d'un modèle mathématique précis.

Cette difficulté limite l'application de cette approche aux processus industriels réels. En particulier, un modèle trop simple ne permet pas de générer des résidus représentatifs uniquement de défaillances, mais ils intègrent le plus souvent les erreurs de modélisation ou de dérive normale des paramètres.

Parmi les méthodes de diagnostic ne nécessitant pas de modèle mathématique explicite, se situent les techniques de classification. Un système de diagnostic peut être considéré en effet comme un classificateur, puisque la tâche de diagnostic peut être vue comme une classification de problèmes.

Parmi les classifications floues existantes, nous avons choisi d'utiliser la méthode LAMDA qui cherche à modéliser un certain type de raisonnement humain et dont les caractéristiques de flexibilité, à savoir la possibilité d'utiliser indifféremment et simultanément un apprentissage supervisé ou non, des données qualitatives ou quantitatives, lui confère de réels avantages.

Après avoir présenté les différentes phases suivies lors d'une classification floue effectuées par LAMDA, les caractéristiques de l'outil *SALSA* ont été décrites. Cet outil développé dans le cadre du projet européen CHEM, donne à l'opérateur un support pour une prise de décision en temps réel. Il peut aussi être utilisé comme outil d'entraînement hors ligne.

Il est à la base (comme technique de classification particulière) de la méthodologie pour le placement des capteurs. Néanmoins, toute technique de classification floue pourrait être utilisée.

L'application de la méthodologie pour le placement de capteurs a été effectuée dans le cas de procédés chimiques complexes. Pour le faire, elle s'appuie sur la simulation dynamique de scénarii de défaillances effectuées grâce à des simulateurs de procédés dont on présente les caractéristiques dans le chapitre suivant.

2. MODELISATION ET SIMULATION DE PROCÉDES CHIMIQUES

2.1 Introduction

La simulation est un outil utilisé dans différents domaines de l'ingénierie et de la recherche en général, permettant d'analyser le comportement d'un système avant de l'implémenter et d'optimiser son fonctionnement en testant différentes solutions et différentes conditions opératoires. Elle s'appuie sur l'élaboration d'un modèle du système, et permet de réaliser des scénarii et d'en déduire le comportement du système physique analysé. Un modèle n'est pas une représentation exacte de la réalité physique, mais il est seulement apte à restituer les caractéristiques les plus importantes du système analysé. Il existe plusieurs types de modèle d'un système physique : allant du modèle de représentation qui ne s'appuie que sur des relations mathématiques traduisant les grandes caractéristiques de son fonctionnement, jusqu'au modèle de connaissance complexe issu de l'écriture des lois physiques régissant les phénomènes mis en jeu. Le choix du type de modèle dépend principalement des objectifs poursuivis.

Les simulateurs de procédés chimiques utilisés classiquement dans l'industrie chimique ou para-chimique, peuvent être considérés comme des modèles de connaissance. Ils sont basés sur la résolution de bilans de masse et d'énergie, des équations d'équilibres thermodynamiques, ... et sont à même de fournir l'information de base pour la conception. Ils sont principalement utilisés pour la conception de nouveaux procédés (dimensionnement d'appareil, analyse du fonctionnement pour différentes conditions opératoires, optimisation), pour l'optimisation de procédés existants et l'évaluation de changements effectués sur les conditions opératoires.

Avant même de parler de modèles d'opération de transformation de la matière, il faut des modèles pour prédire les propriétés physiques de la matière. C'est pourquoi ces simulateurs disposent tous d'une base de données thermodynamiques contenant les propriétés des corps purs (masse molaire, température d'ébullition sous conditions normales, paramètres des lois de tension de vapeur, ...). Cette base de données est enrichie d'un ensemble de modèles thermodynamiques permettant d'estimer les propriétés des mélanges.

Tout simulateur industriel de procédés chimiques est organisé autour des modules suivants :

- Une base de données des corps purs et un ensemble de méthodes pour estimer les propriétés des mélanges appelés aussi modèles thermodynamiques.
- Un schéma de procédé permettant de décrire les liaisons entre les différentes opérations unitaires constituant l'unité (PFD pour *Process Flow Diagram*).
- Des modules de calcul des différentes opérations unitaires contenant les équations relatives à leur fonctionnement : réacteur chimique, colonne de distillation, colonne de séparation, échangeurs de chaleur, pertes de charges, etc.
- Un ensemble de méthodes numériques de résolution des équations des modèles.

Avec ce type de logiciel, les ingénieurs peuvent à partir de la donnée des corps purs présents dans le procédé et du schéma de procédé, développer un modèle du processus reposant sur la mise en commun des équations décrivant les différentes opérations unitaires, les réactions chimiques, les propriétés des substances et des mélanges, qui puisse aussi communiquer avec d'autres applications comme Excel, Visual Basic, Matlab,

Il y a deux modes de fonctionnement dans un simulateur : statique (ou stationnaire) et dynamique. Les simulateurs statiques résolvent des équations statiques qui traduisent le fonctionnement en régime permanent (à l'équilibre), tandis que les simulateurs dynamiques permettent d'évaluer l'évolution des variables dans le temps à partir de la résolution de systèmes d'équations différentielles. Les simulateurs industriels les plus connus mondialement sont :

- Statiques : Aspen Plus (Aspen Technologies), Design II (de WinSim), HYSYS (Hyprotech), , PRO/II (Simulation Sciences), Prosim
 - Dynamiques : HYSYS, Aspen Dynamics (Aspen Technologies), Design II (de WinSim), Dymosym (Simulation Sciences Inc.)
-

Selon *Winter* (Winter, 1992) les simulateurs dynamiques sont en passe de se substituer aux simulateurs en régime permanent. Par exemple, HYSYS (*Hyprotech*) peut passer de la simulation d'un régime permanent à celle d'un régime transitoire (dynamique) par un seul « click » sur un bouton.

Néanmoins, tout procédé ne peut être simulé à l'aide de ces simulateurs industriels. En effet, dans le cas de la mise au point de nouveau procédé, il est généralement nécessaire de disposer de son propre simulateur. Le concept est le même : sur la base des propriétés thermodynamiques des corps purs impliqués dans l'opération et des modèles thermodynamiques, il y a résolution des équations de bilan de matière et d'énergie et des relations d'équilibre constituant le modèle. La différence vient du fait que généralement seules les propriétés des corps présents dans le procédé chimique considéré ne sont détaillées et que l'environnement de développement est moins convivial. On parlera de simulateur dédié (spécifique à un procédé donné). Il a l'avantage de pouvoir avoir une totale maîtrise sur la façon d'écrire les équations du modèle et de les résoudre.

Dans ce travail, nous avons utilisé les deux types de simulateurs :

- Généraliste (industriel) : nous avons constitué à l'aide du simulateur HYSYS (développé par *AEA Technology Engineering Software Products*), un modèle du fonctionnement dynamique d'une unité de production de propylène glycol.
- Dédié : dans le cadre du projet « *Nouvelles applications dans le domaine de la sécurité industrielle et du REX de systèmes d'aide à la conduite supervisée, d'abstraction d'informations, d'analyse et de classification de données* » soutenu par l'Institut pour une Culture de Sécurité Industrielle (ICSI), nous avons collaboré avec l'équipe « Procédés de la Chimie Fine » du Laboratoire de Génie Chimique (UMR CNRS 5503) de Toulouse qui développe, en partenariat avec ALFA-LAVAL, un nouveau réacteur appelé « *Open Plate Reactor* » et utilisé les résultats fournis par leur simulateur dédié (*Data Processing Tool*).

L'objectif de ce chapitre est de décrire ces deux modèles de procédé sur lesquels la méthodologie développée lors de ce travail a été appliquée. La première partie est consacrée à la description des concepts du simulateur HYSYS et de ses caractéristiques et fonctionnalités qui lui permettent d'être un puissant simulateur industriel et à l'explication détaillée du

modèle du procédé chimique étudié (production de propylène glycol). Dans la deuxième partie, nous présenterons le nouveau procédé développé par le Laboratoire de Génie Chimique, ainsi que les caractéristiques principales de son simulateur qui a été validé à partir d'expérimentations effectuées sur une première version pilote de ce procédé.

2.2 Concepts et caractéristiques du simulateur HYSYS

2.2.1 Concepts de base du simulateur HYSYS

HYSYS est un simulateur de conception orientée-objets. Tout changement spécifié sur un élément est répercuté dans tout le modèle.

Dans ce qui suit, on définit les principaux concepts de base et vocabulaires associés, qui sont utilisés pendant les étapes de construction d'un modèle dans le simulateur HYSYS (voir la Figure 2-1) [HYSYS, 2002a] :

- « *Flowsheet* » : c'est un ensemble d'objets « *Flowsheet Elements* » (courants de matière, d'énergie, d'opérations unitaires, de variables opératoires) qui constituent tout ou une partie du procédé simulé et qui utilisent la même base de données thermodynamique « *Fluid Package* ». Ce simulateur possède une Architecture Multi-Flowsheet : il n'y a pas de limite par rapport au nombre de *Flowsheets*. On peut préalablement construire des *Flowsheets* pour les utiliser dans une autre simulation, ou organiser la description de procédés complexes en le scindant en sous-*Flowsheets* qui sont des modèles plus concis (ceci permet de hiérarchiser un processus très complexe). Il possède un certain nombre d'entités particulières : un « *Process Flow Diagram* » (PFD), un « *Workbook* ».
- « *Fluid Package* » : il permet de définir les composants chimiques présents dans le procédé simulé et leurs affecte les propriétés chimiques et physiques contenues dans la base de données des corps purs. Il permet aussi de définir les modèles thermodynamiques qui seront utilisés pour le calcul des propriétés des mélanges et de définir les cinétiques des réactions chimiques mises en jeu dans le procédé.
- « *Process Flow Diagram* » : ce diagramme permet de visualiser les courants et les opérations unitaires, représentées par des symboles dans le « *Flowsheet* », ainsi que la connectivité entre les courants, les opérations unitaires et les tableaux des propriétés des courants.

- « *Workbook* » : il permet d'avoir accès à l'information sur les courants et les opérations unitaires sous forme de tableau de données.
- « *Desktop* »: c'est l'espace principal de HYSYS pour visualiser les fenêtres lors de la conception.
- « *Property view* » : il contient l'information décrivant un objet (opération ou courant)
- « *Simulation Case* » (fichier de simulation) : c'est l'ensemble des « *Fluid Packages* », « *Flowsheets* » et « *Flowsheet Elements* » qui constituent le modèle.

HYSYS est un logiciel de simulation interactif intégrant la gestion d'événements (« *Event driven* ») : c'est-à-dire qu'à tout moment, un accès instantané à l'information est possible, de même que toute nouvelle information est traitée sur demande et que les calculs qui en découlent s'effectuent de manière automatique. Deuxièmement, il allie le concept d'opérations modulaires à celui de résolution non-séquentielle. Non seulement toute nouvelle information est traitée dès son arrivée mais elle est propagée tout au long du *Flowsheet*.

Il existe 5 environnements de développement pour manipuler et mettre en forme l'information dans le simulateur (voir la Figure 2-2) :

1. **Environnement « *Basis Manager* »**: cet environnement permet de créer et modifier le « *Fluid Package* ».
2. **Environnement « *Oil Characterization* »**: il est utilisé pour caractériser les fluides de type pétrolier.
3. **Environnement « *Main Flowsheet* »**: il permet de définir la topologie du *Flowsheet* principal de la simulation. Il est utilisé pour placer et définir les différents courants, opérations unitaires et « *Sub-Flowsheets* » qui constituent le procédé simulé.
4. **Environnement « *Sub-Flowsheet* »**: il permet de définir la topologie d'un sous-ensemble particulier du schéma principal (un courant ou une opération particulière et des autres *Sub-Flowsheets*).
5. **Environnement « *Column* »**: c'est un objet particulier permettant de définir la topologie de l'opération unitaire colonne à distiller. Il possède ses propres « *Flowsheet* », « *Fluid Package* », « PFD » et « *Workbook* ».

Dans la Figure 2-2, les flèches montrent que seuls l'environnement « *Column* » et le « *sub-Flowsheet* » sont accessibles depuis l'environnement principal « *Main Flowsheet* ». Toutefois, en utilisant l'*Object Navigator* (voir la Figure 2-1) on peut se déplacer directement d'un *Flowsheet* à autre.

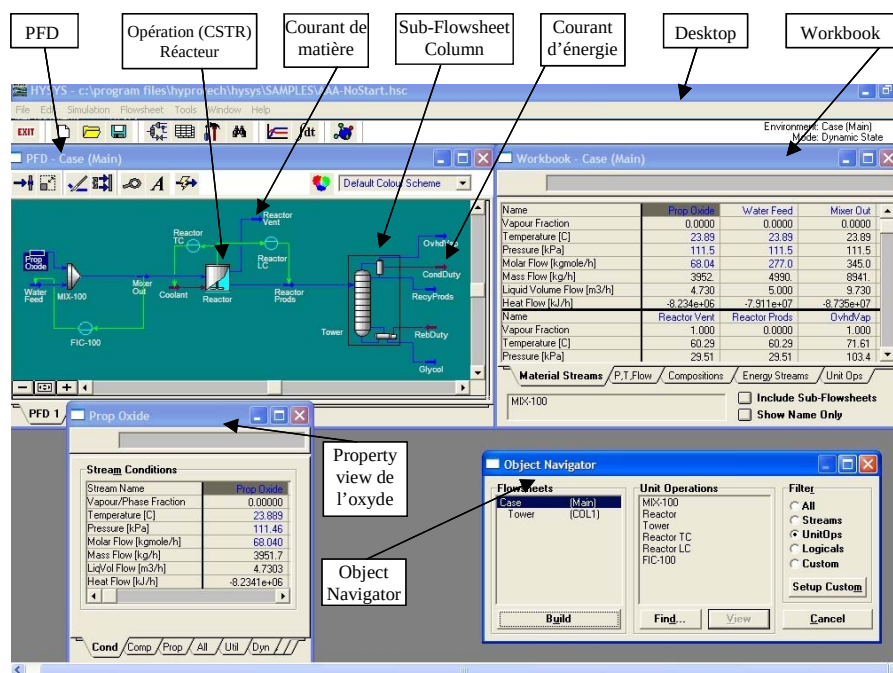


Figure 2-1 Exemple de *Main Flowsheet* dans HYSYS

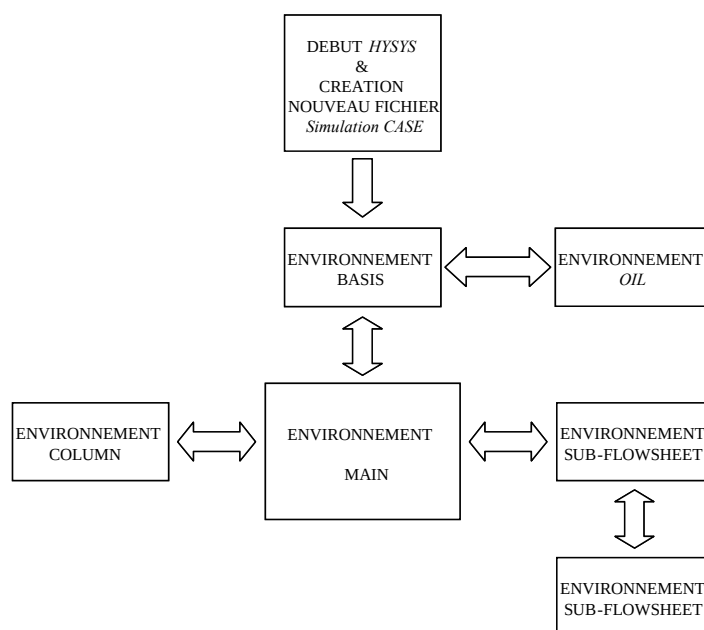


Figure 2-2 Environnements de développement dans HYSYS

2.2.2 Caractéristiques principales de HYSYS

Cette partie décrit brièvement les caractéristiques importantes qui font de HYSYS une plateforme de simulation et de développement très puissant.

- (*The Integrated Engineering Environment*) : Toutes les applications nécessaires sont utilisées dans un environnement de simulation commun.
- Il intègre la possibilité d'une modélisation dans un état stable ou stationnaire et en régime dynamique : la modélisation dans un état stable et l'optimisation étant utilisées lors de la conception des procédés ; la simulation en régime dynamique étant réservée aux études de contrôlabilité de procédés et au développement de stratégies de contrôle.
- Programmation de HYSYS : HYSYS contient un *Internal Macro Engine* qui supporte la même syntaxe que *Microsoft Visual Basic*. On peut automatiser différentes tâches dans HYSYS sans avoir besoin d'un autre programme.

Voici quelques caractéristiques de HYSYS sur la manière dont sont réalisés les calculs :

- Gestion des événements (*Event Driven*): HYSYS combine le calcul interactif (les calculs sont exécutés automatiquement chaque fois que l'on fournit une nouvelle information) avec un accès instantané à l'information (à tout moment on peut avoir accès à l'information depuis n'importe quel environnement de simulation).
- Gestion intelligente de l'information (*Built-in Intelligence*): Les calculs des propriétés thermodynamiques s'effectuent instantanément et automatiquement dès qu'une nouvelle information est disponible.
- Opérations Modulaires: Chaque courant ou unité d'opération peut réaliser tous les calculs nécessaires, en utilisant l'information soit indiquée dans l'opération ou communiquée depuis un courant. L'information est transmise dans les deux directions à travers les *Flowsheets*.
- Algorithme de résolution non séquentielle : on peut construire des *Flowsheets* dans n'importe quel ordre.

Voici les caractéristiques de HYSYS sur comment opèrent les environnements :

Lorsque l'on effectue des développements dans un *Flowsheet* particulier, seul ce *Flowsheet* et les autres situés au-dessous dans la description hiérarchique, seront modifiés. Par exemple, si l'on considère la Figure 2-3 et que l'on suppose que l'on désire faire des changements dans le

SubFlowsheet D, on se place dans son environnement pour y effectuer ces changements. Puisque D est au-dessus de E dans la hiérarchie, tous les *Flowsheets* autres que D et E resteront inchangés. Dès que les calculs dans D seront effectués, il est possible alors de se déplacer dans l'environnement *Main Flowsheet* pour recalculer toutes les autres parties du modèle contenues dans les autres *SubFlowsheets*.

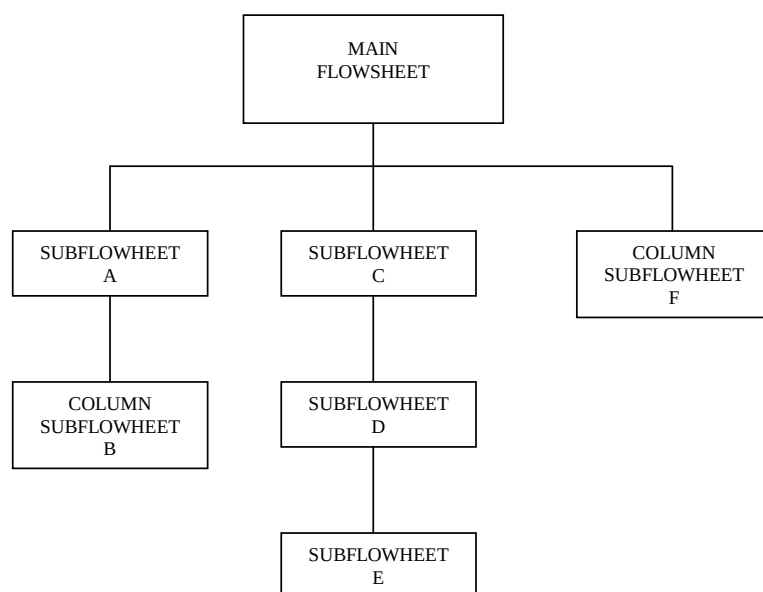


Figure 2-3 Organigramme des environnements dans la hiérarchie

2.3 Programmation de HYSYS

Il y a plusieurs outils CAPE (*Computer-Aided Process Engineering*) utilisés pour différentes activités d'ingénierie, par exemple, en simulation, en estimation des propriétés physiques, en conception de procédés, en optimisation, etc. Il y a eu un effort important des concepteurs de logiciels généraux de simulation de procédés pour intégrer, développer et adapter leurs outils aux besoins spécifiques de leurs clients. En particulier, la capacité d'utilisation de simulateurs de procédés avec des applications externes est de plus en plus privilégiée. Le simulateur de procédés HYSYS contient une structure avancée orientée-objets et utilise la technologie OLE (*Objet Linking and Embedding*) pour envoyer et recevoir de l'information vers ou depuis d'autres applications telles que Matlab, Excel, Word, C++ ou Visual Basic. Ceci permet à l'utilisateur de créer ses propres applications et communiquer avec le simulateur HYSYS.

Un objet comprend un ensemble de fonctions et variables. Généralement, les fonctions d'un objet sont nommées méthodes et les variables sont appelées propriétés. Chaque propriété est

une variable qui a une valeur associée à l'objet et les méthodes sont des fonctions et procédures de calcul associées à l'objet.

Considérons un simple ballon comme un objet. Le ballon a un ensemble de propriétés telles que : la couleur, la taille, l'état de gonflage, etc. Le ballon peut avoir des méthodes telles que : gonfler, dégonfler, etc. En utilisant les propriétés et les méthodes associées au ballon, il est possible de définir, manipuler et interagir avec l'objet.

Un objet peut contenir un autre objet lequel est un sous-ensemble de l'objet principal. Par exemple, l'objet voiture peut contenir d'autres objets tels que le moteur ou les pneus. Ces objets peuvent avoir leurs propres propriétés et méthodes. Un moteur peut avoir comme propriétés le nombre de soupapes et la taille des pistons. Les pneus peuvent avoir la propriété du type de modèle de pneu.

Le développement d'applications spécifiques à un utilisateur peut se faire selon deux méthodes qui permettent l'adaptation de HYSYS aux besoins propres de l'utilisateur: *l'Automation* et *l'Extensibilité*. Avec *l'Automation*, il est possible de manipuler une application à partir d'une autre, au travers d'une relation Client-Serveur. *L'Extensibilité* permet la création de bases de données, de modèles réactionnels et d'opérations unitaires, propres à l'utilisateur, lesquels font partie des éléments de la simulation et fonctionnent comme des objets construits dans HYSYS.

Au cours de notre travail, nous avons utilisé la méthode *Automation* pour manipuler les objets de HYSYS en utilisant le *HYSYS Macro Language Editor* (c'est un outil permettant de développer, tester et exécuter des macro-instructions interactives dans l'environnement de HYSYS) et le langage de programmation Visual Basic. Cette méthode repose sur la technologie OLE, laquelle permet de communiquer de manière interactive avec une application à travers des objets conçus par les développeurs d'une application. Pour manipuler les objets, l'utilisateur n'a pas besoin de comprendre le code généré dans HYSYS, seulement de connaître les noms des objets qui sont disponibles [HYSYS, 2003].

A titre d'exemple, le code pour récupérer la température sur le réacteur est le suivant :

Dim Simcase As Object

Set Simcase = ActiveCase

Var_1 = Simcase.Flowsheet.Operations.Item(1).VesselTemperature

La première instruction définit la variable « *simcase* » comme un objet. Dans la ligne, l'instruction « *Set* » assigne le fichier de simulation actuel (« *ActiveCase* ») à la variable de

type objet « *simcase* ». La dernière instruction donne le chemin complet pour accéder à la valeur de la température du réacteur CSTR (*Continuous Stirred Tank Reactor*), laquelle est assignée à la variable « Var_1 ». « Item » est un ensemble d'objets qui contient toutes les opérations du *Flowsheet* du fichier de simulation actuel (*Simcase*). Dans ce cas, « item(1) » représente le réacteur CSTR. Finalement, la propriété « VesselTemperature » contient la valeur recherchée.

La Figure 2-4 présente la structure de la connexion en ligne du simulateur *HYSYS* et de l'outil de classification *SALSA* [Kempowsky 2004b]. Le bloc **Automation** est en charge de récupérer les mesures disponibles sur le simulateur *HYSYS* et de créer le **Fichier des descripteurs** en contenant l'élément à classer. Le bloc **Acquisition** est chargé de la lecture du **Fichier de descripteurs** pour le traitement avec l'outil de classification *SALSA*. Le bloc de **communication** est chargé de la synchronisation du système ; l'outil *SALSA* attend la création du **Fichier des descripteurs** pour le traitement de la classification, et à la fin de cette classification, l'outil génère un **Fichier de sortie**, lequel indique la disponibilité de l'outil pour recevoir une nouvelle donnée à partir de *HYSYS*. Dès que *HYSYS* a lu ce **Fichier de sortie**, il est prêt pour envoyer un nouveau **Fichier des descripteurs**.

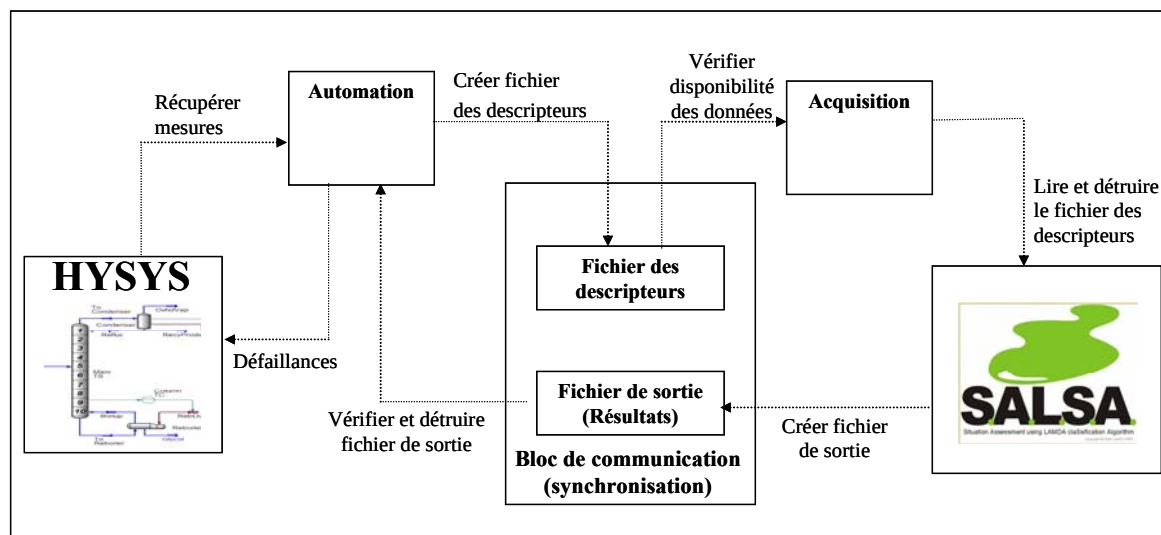


Figure 2-4 Connexion en ligne du simulateur *HYSYS* et de l'outil de classification *SALSA*

L'annexe B montre le programme de récupération de mesures et de génération de défaillances sur le simulateur *HYSYS*.

2.4 Modèle de l'unité de production du propylène glycol

Le propylène glycol est un liquide clair, sans couleur, légèrement sirupeux à la température ambiante. Il peut exister dans l'air sous forme gazeuse, bien que le propylène glycol doive être chauffé ou vivement secoué pour produire un gaz. Le propylène glycol est pratiquement inodore et insipide. Ce composé est employé comme antigel dans les solutions de dégivrage pour les voitures, les avions, les bateaux; utilisé pour la fabrication de composés polyester et comme dissolvant dans les industries de peinture et de plastiques. *The Food and Drug Administration* (FDA) a classé le propylène glycol comme additif sûr pour un usage en alimentation. Il est employé pour absorber l'excès d'eau et pour maintenir l'humidité dans certains médicaments, produits de beauté ou produits alimentaires. C'est un dissolvant pour certains colorants alimentaires et arômes. Le propylène glycol est également employé pour créer la fumée ou le brouillard artificiel utilisé dans la formation de lutte contre l'incendie et dans des productions théâtrales. Le propylène glycol affecte la chimie du corps humain en augmentant la quantité d'acide dans le corps, ayant pour résultat des problèmes métaboliques. Cependant, de grandes quantités de propylène glycol sont nécessaires pour causer cet effet (*Agency for Toxic Substances and Disease Registry*) [ATSDR].

Le procédé simulé lors de ce travail est celui permettant de produire du propylène glycol par réaction de l'oxyde de propylène et de l'eau.

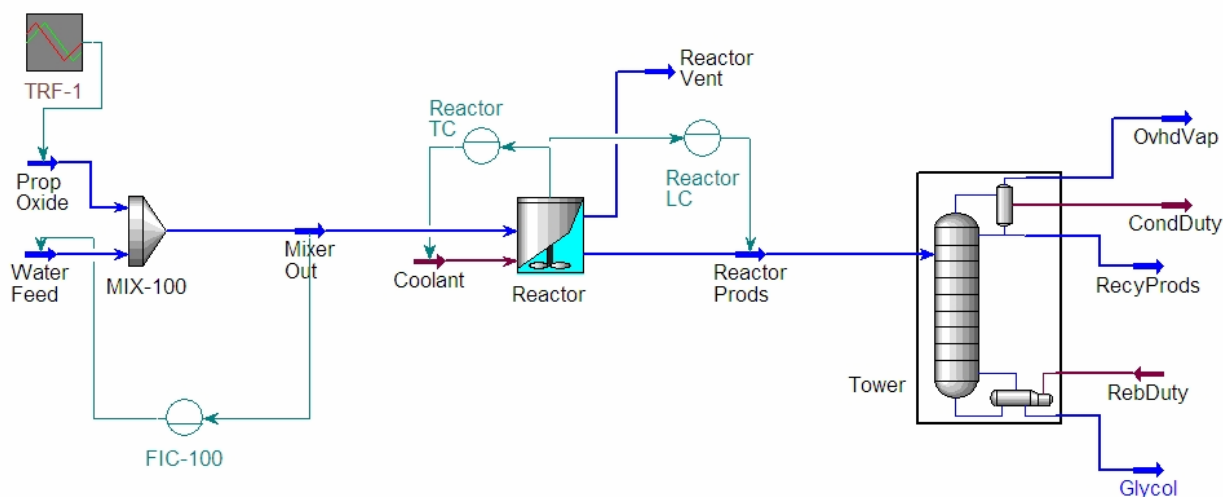


Figure 2-5 Schéma de l'unité de production de propylène glycol

Le procédé dont le schéma de procédé construit sous HYSYS est donné sur la Figure 2-5, est composé d'un mélangeur MIX-100 qui permet de mélanger l'oxyde de propylène et l'eau. Ce

mélange alimente ensuite le réacteur chimique où il y a production de propylène glycol suivant la réaction :



Les produits et réactifs liquides non consommés présents dans le courant (ReactorProds) sont envoyés dans la colonne de distillation pour y être séparés. Cette colonne a pour but de récupérer au maximum (en pied de colonne) le propylène glycol. Elle comporte 10 étages théoriques. L'alimentation de cette colonne s'effectue sur le plateau 5 et est donc composée de 3 constituants : l'eau, le propylène glycol, l'oxyde de propylène. Le courant de sortie en tête de colonne (RecyProds) est composé en grande partie d'eau et d'oxyde. Celui de fond (Glycol), même s'il est composé en grande partie de propylène glycol, contient encore de l'eau et de l'oxyde. La qualité du produit obtenu est donnée par la fraction de propylène glycol dans ce courant, l'objectif étant de récupérer le maximum de propylène en sortie (fraction molaire voisine de 100%)

La Figure 2-6 montre l'architecture de l'ensemble des blocs de simulation utilisés lors de la conception du modèle sous HYSYS. On peut constater que le modèle consiste en un sous-diagramme « Colonne » connecté à l'Environnement Principal (*Main Flowsheet*) renfermant le schéma du procédé tel que présenté sur la Figure 2-1.

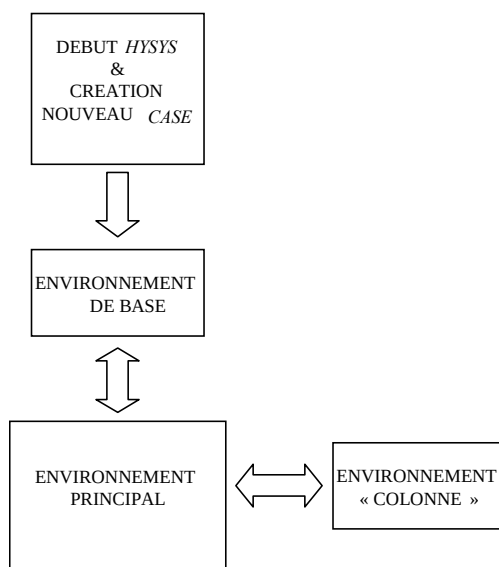


Figure 2-6 Architecture de l'ensemble des blocs de simulation utilisés sous HYSYS

Comme il a été précisé précédemment, il existe deux modes de simulation : le mode régime permanent et le mode régime transitoire ou dynamique. La simulation en mode dynamique

s'effectue à partir d'un point de régime permanent. La construction du modèle pour une utilisation suivant ces deux modes est donnée ci-après.

A) Mode régime permanent.

En suivant le diagramme de procédé de la Figure 2-6, il faut tout d'abord définir l'Environnement de Base, puis l'Environnement Principal et finalement le sous-diagramme « Colonne »

1. Environnement de Base :

Dans l'Environnement de Base, la première étape consiste à créer le « *Fluid Package* », lequel contient les composants chimiques présents dans l'exemple traité. La deuxième étape concerne la donnée du « *Property Package* », c'est-à-dire du modèle thermodynamique qui sera utilisé pour calculer les propriétés des composés et des mélanges dans les conditions (température, pression) calculées au cours de la simulation. Enfin, la réaction chimique, dont la vitesse est supposée suivre la loi d'Arrhenius. Elle est décrite par la donnée de sa stœchiométrie, du coefficient pré-exponentiel et de l'énergie d'activation :

$$r = k * f(Basis) \quad (2.2)$$

(en mole de réactif clé qui a réagi par unité de temps et par unité de volume).

$f(basis)$ permet d'intégrer les différentes caractéristiques de la cinétique : composants impliqués, nature de la réaction (phase liquide, phase gazeuse), température min et max. La fonction $f(Basis)$ fait apparaître généralement le produit des concentrations ou pressions partielles des réactifs élevées à des puissances respectives. Le plus souvent ces puissances sont les valeurs des coefficients stœchiométriques [HYSYS, 1995].

Dans notre cas, la réaction s'effectue en phase liquide et la vitesse a pour équation :

$$r = k[H_2O]^\alpha [C_3H_6O]^\beta \quad (2.3)$$

Avec $\alpha = \beta = 1$ (coefficients stœchiométriques), et

$$k = A e^{\frac{-E_A}{RT}} \quad (2.4)$$

où E_A est l'énergie d'activation exprimée en kJmol^{-1} , A est le facteur de fréquence, T la température dans le réacteur en K et R la constante des gaz parfaits (8.3144 J/K/mole). Les

valeurs nominales de l'énergie d'activation et du facteur de fréquence sont : $E_A = 75362$ KJ/Kmole et $A = 1.70E+13$, respectivement.

Les composés chimiques présents dans le procédé simulé sont l'Oxyde de Propylène (C_3H_6O), l'Eau (H_2O) et le Propylène Glycol ($C_3H_8O_2$). Chaque courant et chaque opération (objets) contenus dans HYSYS ont leur propre fenêtre de propriétés, laquelle contient de multiples pages d'information qui décrivent l'objet. La Figure 2-7 montre la fenêtre des propriétés de l'Oxyde de Propylène dans le courant « *Prop Oxide* » de l'alimentation.

Stream Conditions		
	Prop Oxide	Liquid Phase
Stream Name	Prop Oxide	
Vapour/Phase Fraction	0.00000	1.0000
Temperature [C]	23.889	23.889
Pressure [kPa]	111.46	111.46
Molar Flow [kgmole/h]	68.040	68.040
Mass Flow [kg/h]	3951.8	3951.8
LiqVol Flow [m3/h]	4.7304	4.7304
Heat Flow [kJ/h]	-8.2341e+06	-8.2341e+06

Figure 2-7 Fenêtre des propriétés de l'oxyde de propylène du courant **Prop Oxide**

Dans le cas du mélange Eau - Propylène Glycol- Oxyde de Propylène, le modèle thermodynamique préconisé [HYSYS, 2002a] est le modèle **UNIQUAC**.

2. Environnement Principal (Environnement de Simulation)

L'Environnement Principal permet de construire le diagramme de procédés en définissant les différents courants de matière et d'énergie et les opérations unitaires. Les opérations unitaires utilisées dans cet exemple sont un mélangeur (« *MIXER* »), un réacteur agité continu (« *CSTR* ») et une colonne de distillation (rectification).

3. Environnement sous-diagramme « Colonne »

L'opération unitaire « colonne de distillation » doit être définie dans un sous-diagramme spécifique (voir la Figure 2-8). La colonne est constituée de dix plateaux (numérotés de haut en bas) et son alimentation s'effectue sur le plateau numéro 5. Elle possède un condenseur et un bouilleur (« *Reboiler* ») et deux régulateurs PID (pression et température).

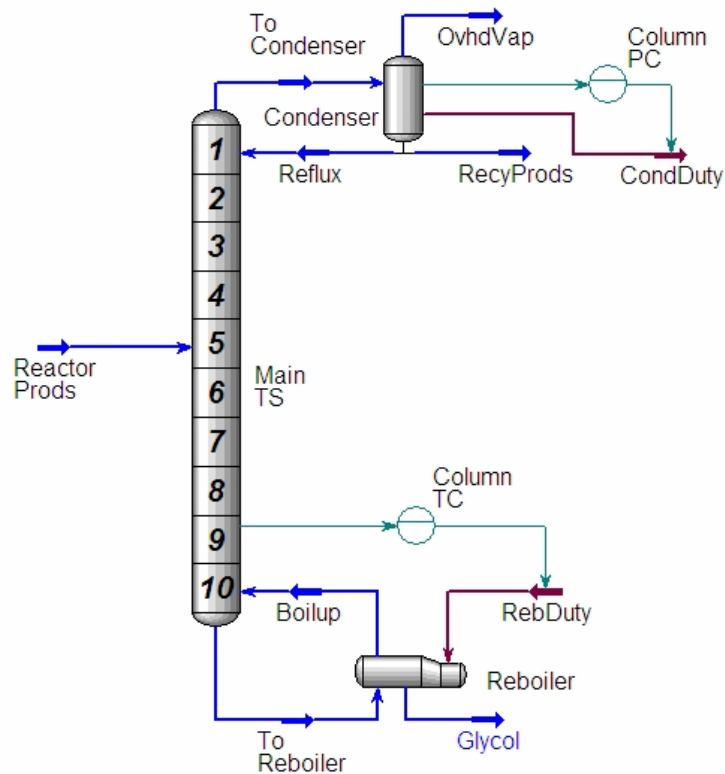


Figure 2-8 Sous-diagramme de la colonne de distillation

B) Mode régime transitoire

Une fois que la simulation en régime permanent a été correctement exécutée, il est possible de basculer en mode dynamique. Pour pouvoir exécuter une simulation en régime transitoire, il faut définir les éléments propres au fonctionnement dynamique du procédé, c'est-à-dire en pratique les régulateurs PID ainsi que les éléments de type fonction de transfert (TRF-1) (voir la Figure 2-5) qui, dans l'exemple présenté, ont été utilisés pour introduire des perturbations non instantanées sur des courants du procédé chimique. Dans le cas présent, cette perturbation se situe sur l'alimentation principale en Oxyde de Propylène. Sur la Figure 2-5 et la Figure 2-8, les régulateurs sont schématisés par :

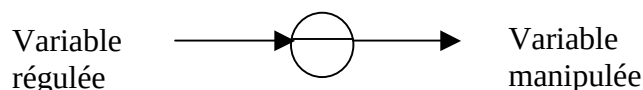


Figure 2-9 Schéma standard des régulateurs sur HYSYS

Les sigles « TC, FC, LC, PC » sont les sigles normalisés représentant respectivement un régulateur de température, de débit, de niveau et de pression.

Le Tableau 2-A donne de manière synthétique la variable régulée et la variable manipulée pour chaque régulateur.

Le contrôle de la température sur le plateau 9 reproduit ce qui est fait en pratique : en absence de capteur de concentration, il permet un pseudo-contrôle de la qualité de la séparation. L'emplacement de ce capteur de température sur le plateau 9 (avant-dernier plateau) a été choisi après une analyse du fonctionnement de la colonne en régime permanent qui a montré que la sensibilité de la température à une variation de la puissance fournie au bouilleur était la plus importante à cet endroit.

Tableau 2-A Les différents régulateurs PID utilisés sur le procédé

Nom du régulateur	Action
Reactor LC	Contrôle le niveau de liquide dans le réacteur en manipulant le débit du courant de sortie
Reactor TC	Contrôle la température dans le réacteur en manipulant le débit du fluide de refroidissement
FIC-100	Permet de fixer le débit molaire total à l'entrée du réacteur
Column PC	Contrôle la pression dans la colonne de distillation en manipulant la puissance échangée au condenseur avec le fluide de refroidissement
Column TC	Contrôle la température de la colonne de distillation en manipulant la puissance délivrée au bouilleur RebDuty

La simulation a été effectuée autour d'un point de régime permanent qui est aussi un point d'équilibre dans le cas d'une installation telle que celle-ci qui fonctionne en continu (il n'y a pas de changement de consignes).

2.5 Le nouveau réacteur OPR (*Open Plate Reactor*)

La mise en oeuvre de réactions chimiques dans des réacteurs discontinus (batch) ou semi-continus est fortement limitée par des contraintes liées à l'évacuation de la chaleur engendrée par les réactions lorsqu'elles sont exothermiques. Un nouveau concept de réacteur intégrant les caractéristiques fondamentales des échangeurs de chaleur devrait permettre d'améliorer grandement la maîtrise des échanges thermiques dans le cas de réacteur fonctionnant en continu. C'est dans cet objectif que l'*Open Plate Reactor* a été développée par le Laboratoire de Génie Chimique en partenariat avec ALFA-LAVAL.

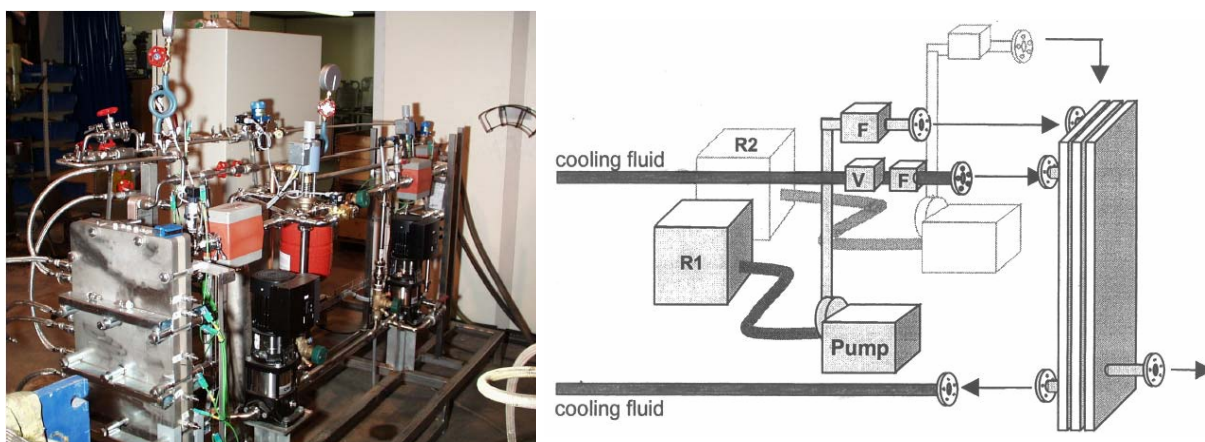


Figure 2-10 *Open Plate Reactor*

La Figure 2-10 montre l'unité pilote conçue et construite dans ce laboratoire autour de ce principe. Ce réacteur reprend les caractéristiques d'un échangeur de chaleur à plaques. Le réacteur consiste en trois blocs avec différentes configurations thermiques : un échange thermique de type co-courant pour le premier et troisième bloc et un échange thermique à contre-courant pour le deuxième bloc. Deux systèmes d'alimentation assurent l'introduction des réactifs (R1 et R2) à la température ambiante. Chaque système se compose d'une pompe, d'un système de mesure de débit (F) et d'une boucle de régulation de débit. Le débit du liquide de refroidissement est commandé par l'intermédiaire d'une vanne (V). Des équipements de mesure de température sont installés sur les admissions et les sorties du fluide. La température du procédé à l'extrémité du premier bloc est de même disponible. Le dispositif de mesure de pression est installé sur l'alimentation générale du réacteur. Toutes les mesures sont enregistrées sur un PC et une représentation graphique en ligne des variations de toutes les variables enregistrées est effectuée. La procédure de mise en régime de l'unité est la

suivante : le fluide de refroidissement doit avoir circulé pendant une heure avant le début de l'expérience. Le réacteur est rempli d'eau distillée et une circulation est assurée pour éliminer les bulles d'air résiduelles. Les réactions sont effectuées avec un excès d'un des réactifs qui est présent au préalable dans le réacteur. Quand l'état d'équilibre est atteint, le deuxième réactif est injecté.

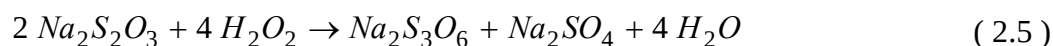
La méthodologie développée lors de ces travaux a été appliquée initialement sur des données expérimentales, résultats d'expériences classiques effectuées sur ce type d'unité pilote. Quelques dysfonctionnements ou changements de conditions opératoires étaient présents dans ce scénarii mais pas totalement représentatifs de défaillances majeures pouvant affecter le fonctionnement de cette unité. C'est pourquoi, pour pouvoir diagnostiquer un certain nombre de défaillances dont certaines difficilement réalisables en pratique pour des raisons de sécurité, nous avons généré ces scénarii à l'aide d'un simulateur du fonctionnement de l'unité pilote qui a été au préalable validé par le LGC [Devatine et al., 2005].

Plusieurs réactions chimiques ont été étudiées :

- La réaction d'oxydation du thiosulfate de sodium ($\text{Na}_2\text{S}_2\text{O}_3$) par l'eau oxygénée (H_2O_2), en milieu homogène.
- L'estérification de l'anhydride propionique par le 2Butanol.

La réaction d'oxydation du thiosulfate présente les différentes caractéristiques : une stœchiométrie connue, elle est irréversible, rapide et fortement exothermique. L'eau oxygénée est utilisée en excès pour éviter toute réaction secondaire. L'énergie générée lors de la réaction est connue et vaut : $\Delta H = -586.2 \text{ KJ.mol}^{-1}$

Le schéma réactionnel est le suivant :



L'estérification de l'anhydride propionique par le 2-Butanol, est largement citée dans les publications relatives à la sécurité des procédés physico-chimiques. Cette réaction est utilisée dans le cadre de la validation expérimentale d'études théoriques concernant par exemple la prédiction d'emballement thermique au sein de réacteur batch ou la détermination rapide de données cinétiques [Galvan et al., 1996].

Cette synthèse a également été mise en œuvre dans le cadre de travaux purement expérimentaux. Ainsi, à l'aide d'outils de régulations et d'appareils d'analyses

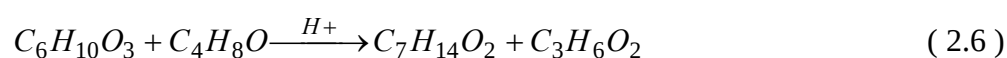
calorimétriques, le profil d'addition de réactif dans un procédé semi-batch a pu être optimisé en ligne pour le cas particulier d'une réaction d'estérification exothermique du second ordre [Ubrich et al., 1999 ; Ubrich et al., 2001].

Enfin, plus récemment, cette réaction a également été utilisée pour tester une nouvelle génération du simulateur dynamique permettant la mise en place d'une stratégie de contrôle thermique performante.

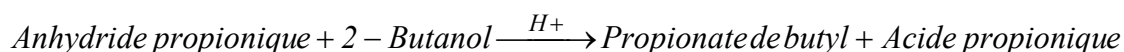
Cette réaction présente certains avantages :

- Elle s'effectue en milieu homogène [Benuzzi et al, 1991].
- Elle est relativement simple à mettre en œuvre.
- Elle est modérément exothermique.
- Le comportement cinétique change selon que l'on opère avec ou sans catalyseur : la vitesse de réaction est du second ordre en l'absence d'acide fort et est de type auto-catalysée en présence d'acide sulfurique.

Le schéma réactionnel est la suivante :



Ou



2.6 Conclusions

Dans ce chapitre, nous avons décrit les deux procédés chimiques complexes sur lesquels nous avons appliqué la méthodologie d'aide au placement de capteurs en vue du diagnostic.

Le premier est celui de la production de propylène glycol. Il a été simulé à l'aide d'un simulateur dynamique mondialement connu et très utilisé dans l'industrie. Ce choix, même s'il a nécessité le besoin de s'investir sur la prise en main, la compréhension, la maîtrise, l'adaptation d'un outil aussi puissant qu'HYSYS, a permis de démontrer la faisabilité de la procédure développée alliée à l'utilisation de simulateurs dynamiques de procédés commerciaux et donc, son caractère générique. Dans ce chapitre, nous avons décrit les

principales caractéristiques et concepts d'HYSYS en insistant sur ceux liés à la construction du modèle de simulation que nous avons développé au cours de ce travail.

Le deuxième exemple est celui d'un procédé en cours de développement : *l'OPR*. L'unité pilote existante a permis de mettre en évidence l'apport de la classification pour son diagnostic. La mise en œuvre de la procédure de placement de capteurs proprement dite a été effectuée grâce à la simulation. En effet, il a été plus facile et surtout plus sûr de générer des scénarii de défaillance à partir du simulateur spécifique à cette installation développé par le Laboratoire de Génie Chimique, que de les réaliser sur l'installation réelle. Deux types de réactions chimiques ont été étudiés, permettant de couvrir un large domaine de possibilités d'utilisation de ce type de réacteur.

3. DEVELOPPEMENT DE LA METHODOLOGIE D'AIDE AU PLACEMENT DE CAPTEURS EN VUE DU DIAGNOSTIC

3.1 Introduction

Pour pouvoir effectuer correctement le suivi en ligne, détecter les anomalies de fonctionnement, il est nécessaire de recueillir des informations qui doivent être les plus pertinentes possibles. Une solution consisterait à placer le maximum de capteurs. Cependant, en plus du coût prohibitif, un afflux d'informations non pertinentes peut être préjudiciable pour la facilité d'analyse de la situation. La technique de placement de capteurs proposée vise à répondre aux deux questions suivantes : Quelle sont les meilleures localisations pour placer des capteurs en nombre limité sur le procédé ? Quel est le type de capteur qui donne la meilleure information pour surveiller le système ? Le placement de ces capteurs est crucial pour réaliser ultérieurement des mesures de bonne qualité utilisées comme base pour la surveillance et le diagnostic de défaillances.

Un premier paragraphe fait tout d'abord le point sur les différentes méthodes développées pour résoudre le problème de placement de capteurs sur des processus dynamiques : les méthodes utilisant des modèles mathématiques, les méthodes par analyse en composantes principales et la méthode développée au cours de ces travaux qui utilise les résultats d'une classification effectuée à partir de données historiques.

Pour cette dernière technique, nous proposons plusieurs critères pour évaluer la quantité d'information donnée par un capteur et deux procédures pour déterminer les capteurs les plus pertinents.

La fin de ce chapitre est consacrée à la description de la méthodologie développée qui se divise en 3 parties : le placement des capteurs, l'établissement d'un modèle du comportement et le diagnostic en ligne.

3.2 Méthodes de sélection de capteurs

Le problème du placement optimal des capteurs est crucial lors de la phase de conception des systèmes. Ayant comme objectif de réduire les coûts d'instrumentation et d'accroître la pertinence des informations reçues, les approches développées généralement se basent sur des estimateurs d'états ou l'analyse des corrélations entre capteurs. Les différentes approches pour le placement de capteurs peuvent être classées suivant [Basseville et al., 1987] :

- **Placement optimal de capteurs par reconstruction d'états**

Le problème est de trouver une matrice de mesure H laquelle optimise un critère qui reflète la performance de l'estimateur d'état pour le système dynamique linéaire considéré. Plusieurs auteurs ont utilisé une mesure indirecte de la performance, telle que la matrice d'information de Fisher [Qureshi et al., 1980], la matrice d'observabilité [Dochain et al., 1996], ou la matrice de covariance d'un filtre de Kalman [Kumar et al., 1978]

- **Placement optimal de capteurs pour les systèmes à paramètres distribués.**

Beaucoup de systèmes en science et technologie sont distribués en espace et en temps, et par conséquent, ils peuvent être décrits par des équations différentielles partielles non linéaires. Ces modèles mathématiques contiennent généralement divers paramètres inconnus dont les valeurs numériques doivent être obtenues à partir des données expérimentales.

Dans ses travaux [Wouwer et al., 2000] par rapport au placement optimal de capteurs dans des systèmes à paramètres distribués, Wouwer distingue les objectifs d'estimation d'états, des objectifs d'estimation de paramètres [Papadimitriou, 2003]. Dans le premier cas, le critère d'optimisation est basé sur une mesure d'indépendance entre les réponses des capteurs, tandis que dans le deuxième cas, l'objectif est de déterminer le placement des capteurs pour mieux identifier les paramètres d'un modèle, en recourant à une mesure d'indépendance entre les fonctions de sensibilité paramétrique.

- Placement optimal de capteurs pour la détection de défaillances

Ce problème a été abordé par Watanabe et al. (1985) dans le cas des systèmes non-linéaires. La détection et le diagnostic des défaillances sont réalisés en considérant les résultats de l'estimation d'états et/ou de plusieurs observateurs de résidus. Le problème du placement optimal des capteurs est ensuite résolu par recherche exhaustive pour minimiser le coût d'observation associé à chaque ensemble de mesures qui sont possibles pour cette stratégie de détection de défaillances.

- Placement optimal de capteurs par Analyse en Composantes Principales (ACP)

Une méthodologie a été proposée pour identifier le nombre et la localisation des mesures de température à utiliser comme entrées pour estimer le profil de compositions dans une colonne de distillation discontinue [Zamproga et al., 2004]. Cette méthodologie basée sur l'Analyse en Composantes Principales (ACP) a été proposée pour sélectionner les variables du processus secondaires les plus pertinentes (on appelle variable secondaire toute variable autre qu'une composition). Dans cette approche, une matrice mesurant la sensibilité instantanée de chaque variable secondaire par rapport aux variables primaires (les compositions que l'on cherche à estimer) est définie. Les variables les plus sensibles sont extraites de cette matrice en exploitant les propriétés de l'ACP, ensuite elles sont utilisées comme variables d'entrée pour le développement d'un modèle convenable pour l'implantation en ligne.

3.3 Méthodes de sélection à partir du résultat de la classification de données

Il y a beaucoup de publications sur le diagnostic, l'isolation et l'identification de défaillances. Ces méthodes ont souvent besoin d'un grand nombre de mesures. C'est la raison pour laquelle on propose une technique de placement de capteurs pertinents en vue de son diagnostic, sans avoir besoin d'un modèle mathématique du processus.

Cette technique repose sur le concept de l'entropie de SHANNON et du gain d'information dont nous donnons les définitions et propriétés dans ce qui suit.

3.3.1 L'Entropie

Le concept d'entropie se réfère à deux domaines : la physique et la théorie de l'information. En physique, l'entropie est une mesure du désordre de l'énergie et elle augmente naturellement. Le désordre d'un système est le nombre d'états dans lesquels le système peut se maintenir. Supposons le cas de l'automobile et regardons le Tableau 3-A :

Tableau 3-A Utilisation de l'énergie dans une automobile

Entrées d'énergie	Sorties d'énergie
Combustible 100%	Pertes à cause de l'eau de refroidissement : 36%
	Sortie d'énergie du moteur : 26%
	Pertes dans l'échappement : 38%

Seulement 26% de l'énergie sont « utiles », le reste est de l'énergie dissipée, qui s'ajoute au désordre mondial. De plus, les 26% ne sont pas totalement utilisés dans l'accélération du véhicule :

Sortie d'énergie du moteur 26% = énergie utilisée dans l'accélération : 5% + énergie dissipée due à la friction des pneus : 7% + énergie pour les accessoires : 4% + pertes dues à la résistance de l'air : 7% + pertes dans la transmission de puissance 3%.

Dans la théorie de l'information, l'Entropie est la grandeur qui mesure l'information contenue dans un flux de données, c'est-à-dire ce que nous apporte une donnée ou un fait concret. Par exemple : si nous disons « les rues sont mouillées » en sachant qu'il vient de pleuvoir, ceci nous apporte peu d'information, puisque c'est très logique. En revanche, si nous disons « les rues sont mouillées » en sachant qu'il ne vient pas de pleuvoir, ceci nous apporte beaucoup plus d'information. La quantité d'information est différente avec pourtant le même message.

Nous utilisons les concepts des deux approches pour faire les calculs de l'entropie, c'est-à-dire, l'entropie dépend du nombre d'états dans un système, et de la probabilité avec laquelle un élément du système appartient à chaque état.

La théorie de l'information a été proposée par Claude SHANNON [Shannon, 1948], c'est un domaine des probabilités, dans lequel on propose un nouveau modèle mathématique des systèmes de communication. Un postulé de base de cette théorie est que « l'information » peut se traiter comme une quantité physique mesurable, telle que la densité ou la masse.

Dans cette théorie, SHANNON a montré qu'il était possible de quantifier la capacité d'information en introduisant par une valeur numérique, le concept d'Entropie.

Dans la théorie probabiliste, un système complet d'événements signifie un ensemble d'événements tels qu'un et seulement un, peut se produire à chaque essai. Soit $E_1, E_2, E_3, \dots, E_n$

les événements d'un système complet et $p_1, p_2, \dots, p_n$ $\left(p_i \geq 0, \sum_{i=1}^n p_i = 1 \right)$ leurs probabilités,

alors on possède un **espace probabilisé** T [Kninchin, 1957]:

$$T = \begin{pmatrix} \mathbf{E1} & \mathbf{E2} & \dots & \mathbf{En} \\ \mathbf{p1} & \mathbf{p2} & \dots & \mathbf{pn} \end{pmatrix}$$

Chaque espace probabilisé décrit un état d'incertitude. Lors d'une expérience, le résultat ne peut être qu'un des événements $E_1, E_2, E_3, \dots, E_n$ et on connaît seulement la probabilité avec laquelle cet événement peut se produire. Suivant cette incertitude, on aura un espace probabilisé différent. On peut démontrer ceci avec l'exemple suivant. Soient deux espaces probabilisés :

$$\begin{pmatrix} \mathbf{E_1} & \mathbf{E_2} \\ \mathbf{0.5} & \mathbf{0.5} \end{pmatrix} \quad \begin{pmatrix} \mathbf{E_1} & \mathbf{E_2} \\ \mathbf{0.99} & \mathbf{0.01} \end{pmatrix}$$

La première alternative représente beaucoup plus d'incertitude que la deuxième (l'Entropie est maximale à cause de l'équiprobabilité). Dans le deuxième cas, la probabilité d'avoir le résultat de l'événement E_1 est plus grande, tandis que dans le premier cas on doit éviter de réaliser une prédiction.

SHANNON a proposé que la quantité d'information d'un signal (ou événement) est inversement proportionnelle à sa probabilité :

$$h(s_k) = h(p[s_k]) \propto \frac{1}{p[s_k]} \quad (3.1)$$

En sachant que l'information est additive, l'information de deux signaux sera la somme de ses mesures d'information h :

$$h(s_m, s_n) = h(p[s_m]) + h(p[s_n]) \quad (3.2)$$

Comme il s'agit d'un espace probabilisé, alors :

$$h(p[s_m]) + h(p[s_n]) = h(p[s_m] p[s_n]) \quad (3.3)$$

SHANNON a conclu que la fonction h devait être logarithmique à cause de la propriété des logarithmes : $\log_b x + \log_b y = \log_b(xy)$

On a donc $h(s_k) = -\log(p[s_k])$ et $h(s_k)$ peut être interprété :

- a priori, par l'incertitude qui règne sur la réalisation de s_k
- a posteriori, par l'information apportée par la réalisation de s_k .

On peut introduire une quantité pour mesurer l'incertitude associée à un espace probabilisé donné. En appelant p_i la probabilité du $i^{\text{ème}}$ événement et en prenant les logarithmes de base « **a** » (on notera désormais $\log$ le logarithme en base **a**), la mesure de l'information moyenne H est :

$$H = \sum_{i=1}^n p_i h(s_i) = \sum_{i=1}^n p_i \log\left(\frac{1}{p_i}\right) = - \sum_{i=1}^n p_i \log(p_i) \quad (3.4)$$

Cette équation représente l'Entropie de SHANNON. La base logarithmique peut être arbitraire, si on prend le logarithme base 2, alors l'unité d'entropie s'appelle *BIT*. Lorsqu'on utilise les logarithmes en base 10 l'unité s'appelle *HARTLEY* et *NITS* si on utilise les logarithmes népériens [Silviu, 1977].

Dans l'annexe C, nous donnons les propriétés principales de l'Entropie.

3.3.2 Le gain de l'information

En appliquant ce concept dans le contexte du diagnostic, l'Entropie de SHANNON peut être vue comme une mesure de la qualité du descripteur pour la discrimination de fautes.

Selon le travail de FERNANDEZ R [Fernandez, 2002], le gain de l'information est l'outil qui permet de quantifier l'information fournie par un descripteur (d) et permet de résoudre le problème de sélection des descripteurs les plus représentatifs. Il est défini à partir de la différence entre l'Entropie maximale de SHANNON ($H(C)$) et l'entropie connaissant la valeur du descripteur ($H(C|d_k)$) (voir la Figure 3-1).

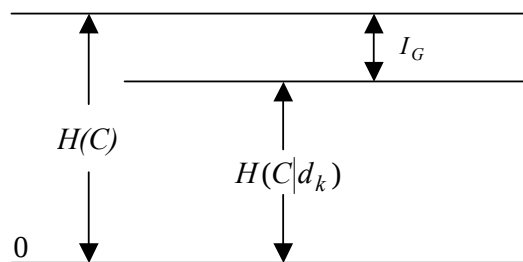


Figure 3-1 Concept du gain de l'information

Soient n descripteurs, le gain d'information est défini par :

$$I_G = H(C) - H(C|d_k) \quad (3.5)$$

où $I_G \geq 0$, $H(C) = \log n$, $H(C|d_k) = -\sum_{i=1}^n (p_{iprob}) \log(p_{iprob})$ et p_{iprob} est la probabilité qu'un élément appartienne à la classe C_i .

On peut définir aussi un gain relatif de l'information :

$$I_{GR} = \frac{I_G}{H(C)} \times 100\% \quad (3.6)$$

L'objectif de chaque descripteur est de maximiser le gain pour avoir une meilleure contribution dans la détection du défaut.

3.3.3 Critères de sélection des capteurs.

Dans ce paragraphe, nous proposons trois critères différents pour la sélection de capteurs les plus pertinents pour la détection, l'isolation et le diagnostic des défaillances: l'entropie probabiliste, l'entropie non-probabiliste et la variance.

Les données d'entrée pour le calcul de ces critères proviennent de la matrice contenant les profils des classes obtenues par classification. Dans notre cas la méthode de classification est la méthode de classification floue LAMDA (logiciel *SALSA*) mais la démarche développée est générique et peut donc s'appliquer à toute technique de classification donnant comme résultat une matrice liant les classes aux différents descripteurs. Dans cette matrice, chaque ligne représente une classe de défaillance du processus et les colonnes représentent l'ensemble des descripteurs disponibles. Le but de ces méthodes est de choisir les descripteurs les plus pertinents par rapport à chaque classe de défaillance du processus ; c'est-à-dire, le descripteur qui donne le plus d'information pour la détection, la localisation et le diagnostic de la défaillance étudiée.

A. Critère de l'Entropie probabiliste

En partant de la matrice du profil des classes (résultat de la classification), nous proposons deux types de procédure pour la sélection des capteurs : procédure au moyen de paires de classes de défaillance et la procédure qui examine l'ensemble des classes de défaillance.

- Procédure par paires

Ce type de procédure consiste à former des paires de classes, où chaque paire est composée de la classe de fonctionnement normal du processus et une classe de défaillance ou opération anormale (voir la Figure 3-2). Le nombre total de paires formées est égal au nombre de défaillances appliquées au processus (F). Le résultat de la classification pour une « paire » est le profil de classes tel qu'il est donné sur le Tableau 3-B.

Tableau 3-B Résultats de la classification pour une paire état normal/défaillance

	Descripteur 1	Descripteur 2	...	Descripteur P
Classe 1 (état normal)	A_{11}	A_{12}	...	A_{1p}
Classe 2 (défaillance)	A_{21}	A_{22}	...	A_{2p}

Il comprend la valeur moyenne normalisée de chaque descripteur j à la classe i à laquelle il appartient (A_{ij})

La méthodologie développée repose sur cinq étapes distinctes :

Étape 1. Pour appliquer le concept d'entropie, nous devons manipuler un espace probabilisé car la somme des colonnes dans le Tableau 3-B : $\sum_{i=1}^2 A_{ij}$ est généralement différente de 1. Pour se ramener au cas probabiliste, chaque élément est normalisé par rapport à la somme totale sur toutes les classes des valeurs moyennes du descripteur considéré. La formule de normalisation est la suivante :

$$\tilde{A}_{ij} = \frac{A_{ij}}{\sum_{i=1}^2 A_{ij}} \quad (3.7)$$

où « j » identifie le descripteur et « i » la classe.

Étape 2 : L'Entropie maximale H_{max} de la classification de défauts résultante est alors calculée comme suit [Walter et Pronzato, 1997] :

$$H_{\max} = \text{Log } 2 \quad (3.8)$$

L'Entropie maximale correspond en effet au cas où l'élément a la même probabilité d'appartenir à l'une ou l'autre des classes ($p_i=1/2$). L'Entropie de SHANNON s'interprète comme l'erreur de classification d'un élément à la classe. C'est-à-dire, c'est l'information dont on a besoin pour résoudre le problème de classification. Toutefois, elle ne dit pas comment obtenir cette information à partir de l'analyse des descripteurs [Costes-Albrespic, 1984].

Étape 3. L'entropie probabiliste de SHANNON par rapport à chaque descripteur $H(A_j)$ est ensuite calculée de la façon suivante :

$$H(\tilde{A}_j) = -\sum_{i=1}^2 \tilde{A}_{ij} \text{Log } \tilde{A}_{ij} \quad (3.9)$$

Étape 4. On calcule le gain d'information de chaque descripteur j . Le gain d'information est l'outil qui permet de quantifier l'information amenée par un descripteur spécifique. Il est défini comme la différence entre l'Entropie maximale et l'Entropie par rapport au descripteur j , lequel est toujours ≥ 0 :

$$I_G = H_{\max} - H(\tilde{A}_j) \quad (3.10)$$

Étape 5. Finalement, on obtient le gain relatif d'un descripteur en utilisant la relation suivante:

$$I_{G \text{ rel}} = \frac{I_G}{H_{\max}} \times 100\% \quad (3.11)$$

Le calcul des gains relatifs d'information pour tous les descripteurs fournit un tableau. Le descripteur (capteur) le plus pertinent ou représentatif pour un défaut donné est celui qui fournit le gain relatif le plus élevé. Cette procédure est répétée pour toutes les paires : état normal/défaillance. Un organigramme est donné sur la Figure 3-2.

.

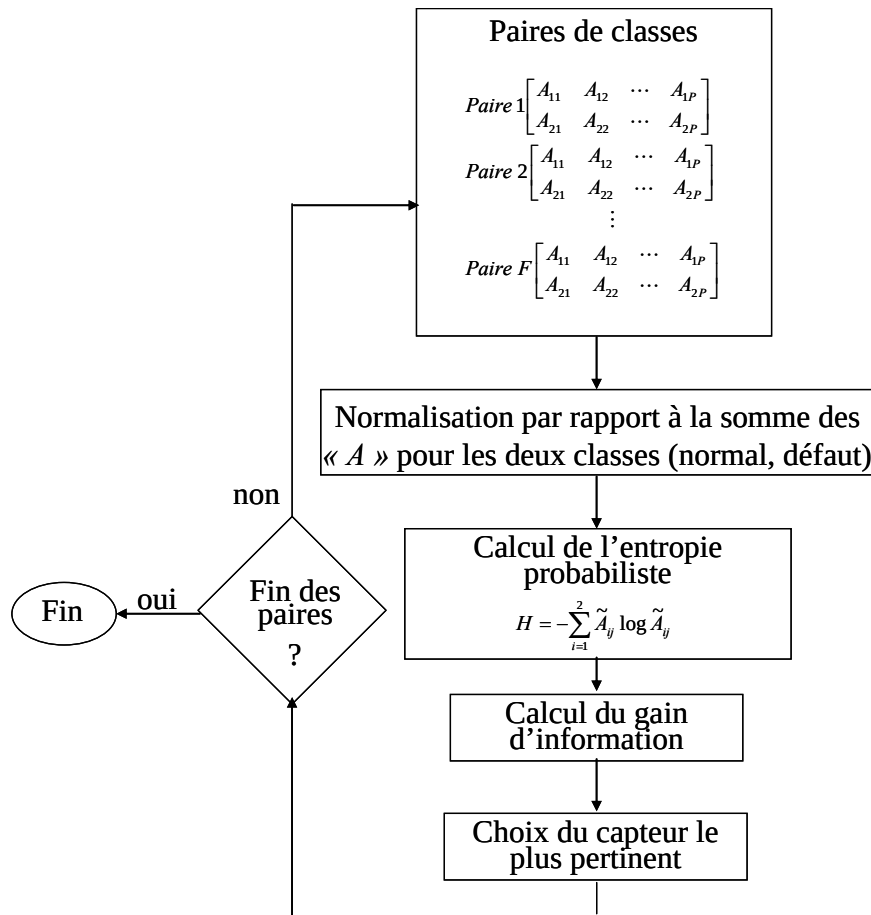


Figure 3-2 Procédure probabiliste de sélection de capteurs par paires de classes

L'avantage de cette procédure est qu'on peut obtenir directement les descripteurs les plus pertinents à chaque défaillance, puisque, l'analyse pour chaque paire est faite de façon indépendante.

Pour mieux comprendre la méthodologie développée, expliquons plus en détails le calcul de l'entropie par chaque paire de classes :

Si un descripteur quelconque de la classe normale a une valeur égale à celle prise dans la classe de défaillance (c'est-à-dire qu'il n'y a pas de changement de valeur à la suite de l'application d'une défaillance pour ce descripteur), alors il aura une valeur de 0.5, puisqu'il s'agit d'un espace probabilisé, et ceci implique qu'il aura une valeur d'entropie maximale (voir la Figure 3-3) et pourtant un gain d'information nulle. En revanche, si un descripteur quelconque de la classe normale a une valeur très différente de celle prise dans la classe de défaillance (par exemple 0 ou 1), il aura une valeur d'entropie minimale (voir la Figure 3-3) et pourtant un gain d'information maximal. Ceci signifie que ce descripteur permet de bien

détecter cette défaillance ; en effet il permet de mieux faire la différence entre le fonctionnement normal et la défaillance.

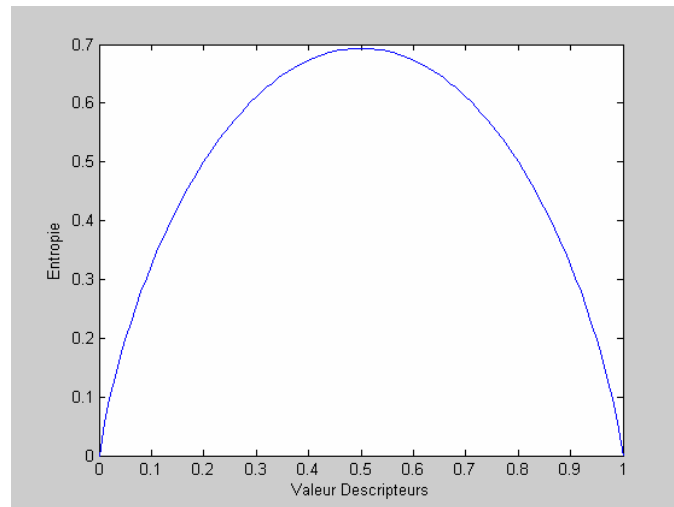


Figure 3-3 Valeurs de l'entropie probabilisée pour 2 classes

- Procédure d'analyse de l'ensemble des classes de défaillance

Dans ce cas on considère l'ensemble des classes de défaillance (F) et la classe normale. C'est-à-dire, cette procédure s'applique une seule fois sur la matrice composée de $(F+1)$ situations ou classes, et de « P » descripteurs (voir le Tableau 3-C).

Tableau 3-C. Profil des classes

	Descripteur 1	Descripteur 2	...	Descripteur P
Classe 1	A_{11}	A_{12}	...	A_{1P}
Classe 2	A_{21}	A_{22}	...	A_{2P}
$\vdots$	$\vdots$	$\vdots$	...	$\vdots$
Classe F	A_{F1}	A_{F2}	...	A_{FP}
Normal	$A_{(F+1)1}$	$A_{(F+1)2}$	...	$A_{(F+1)P}$

Le calcul de la normalisation et de l'entropie pour chaque descripteur est réalisé de la même manière que précédemment, mais en prenant en compte les $F+1$ classes dans la matrice (au lieu de deux comme dans la procédure précédente). Pour le calcul du gain d'information, il faut considérer que la valeur de l'entropie maximale est : **$\log (F+1)$** . L'organigramme de la procédure est donné sur la Figure 3-4. L'inconvénient de ce type d'analyse est qu'il faut

définir un seuil pour choisir le nombre final de descripteurs pertinents, de sorte qu'ils permettent de détecter toutes les défaillances étudiées préalablement.

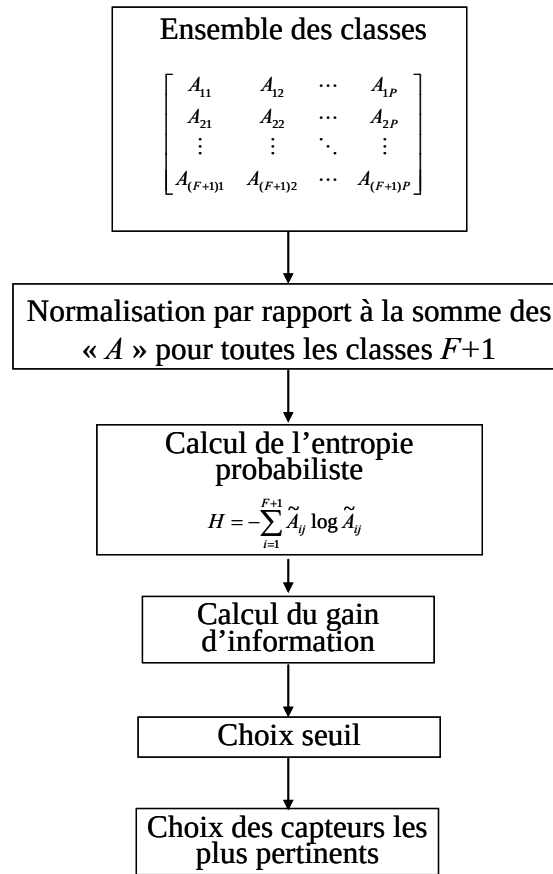


Figure 3-4 Procédure de sélection des capteurs par l'analyse de l'ensemble des classes

Pour mieux interpréter le résultat final de la sélection des capteurs à l'aide de cette procédure, prenons un exemple de calcul de l'entropie d'un descripteur en fonction de trois classes, afin d'observer graphiquement les diverses valeurs de l'entropie (voir la Figure 3-5) :

- Soient p_1 , p_2 et p_3 les probabilités d'un descripteur pour les classes 1, 2 et 3, respectivement, telles que $\forall p_1, p_2, p_3 \in [0, 1]$, les p_i satisfont : $p_1 + p_2 + p_3 = 1$. Le calcul de l'entropie pour 3 classes est donné par:

$$H = -\sum_{i=1}^3 p_i \log p_i = -[(p_1 \log p_1 + p_2 \log p_2 + p_3 \log p_3)] \quad (3.12)$$

- Si la valeur probabilisée d'un descripteur à une classe est égale à zéro (en prenant $0 \cdot \log 0 \sim 0$), le graphique est identique à celui obtenu dans le cas de deux classes (voir la Figure 3-3).

- La valeur maximale de l'entropie s'obtient quand $p_1 = p_2 = p_3 = 1/3$, c'est-à-dire dans une espace équiprobable, et elle est égale à : $H = \log(3) = 1.098$
- La valeur minimale de l'entropie (c'est-à-dire zéro) s'obtient quand deux des probabilités sont égales à zéro, ce qui implique que la troisième a sa valeur égale à 1. Alors on peut conclure que le descripteur avec le gain d'information le plus grand est très pertinent par rapport à une classe de défaillance et qu'il reste invariant pour les autres classes.

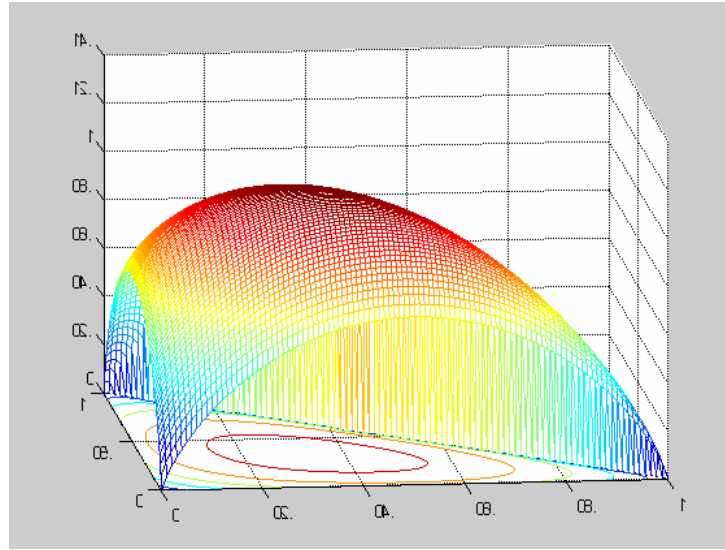


Figure 3-5 Valeurs de l'entropie pour 3 classes

B. Critère de l'entropie non probabiliste

Pour cette méthode, on calcule l'entropie non probabiliste défini par [De Luca et al., 1972]. Dans ce cas, il n'y a pas de normalisation des données de la matrice des profils des classes (voir la Figure 3-6), c'est-à-dire que les valeurs de la matrice (A_{ij}) sont prises directement dans la matrice résultante de la classification.

La relation donnant l'entropie non probabiliste est :

$$H = - \sum_{i=1}^F A_{ij} \log A_{ij} + (1 - A_{ij}) \log(1 - A_{ij}) \quad (3.13)$$

où K est le nombre de classes de défaillances. La valeur maximale pour H correspond à $A_{ij} = 0.5$, pour $\forall A_{ij} \in [0,1]$; c'est-à-dire :

$$H_{max} = F \log 2 \quad (3.14)$$

$F = 2$ dans la procédure par paire.

Cette valeur sera utilisée pour le calcul du gain de l'information par rapport à chaque descripteur.

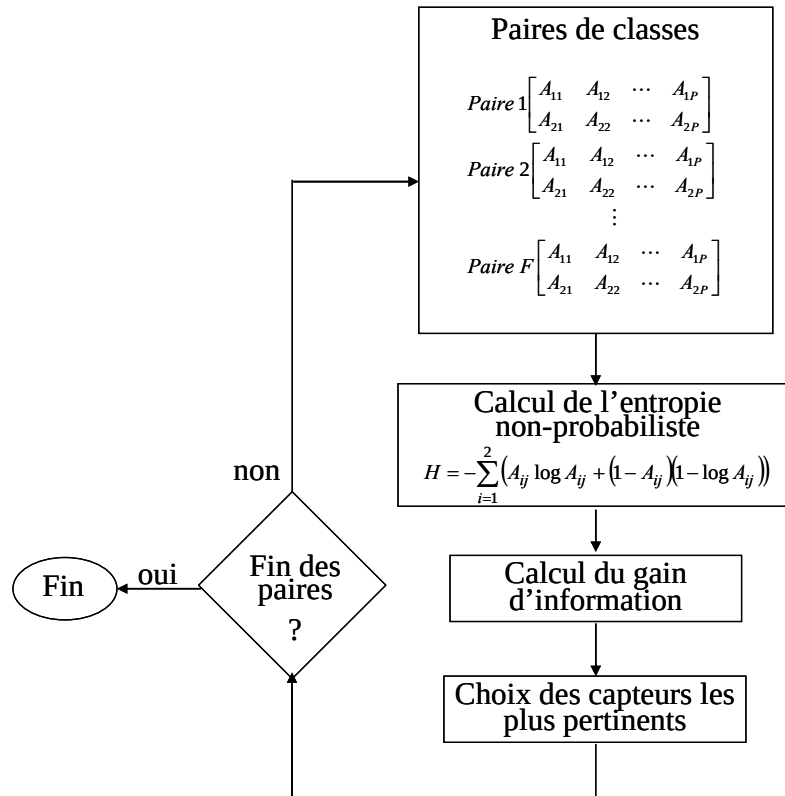


Figure 3-6 Organigramme de la méthode basée sur l'entropie non probabiliste

Cette méthode est intéressante quand les valeurs normalisées des descripteurs dans la classe normale sont très proches de 0.5, puisque toute autre valeur proche de 1 ou 0 donnera une valeur significative du gain d'information.

C. Critère de la Variance

Cette méthode est la plus simple puisqu'il s'agit de calculer la variance de chaque valeur probabilisée du descripteur pour les différentes classes. Comme précédemment, on peut définir des paires : classe normale/classe de défaillance. La Figure 3-7 illustre la procédure que nous venons d'expliquer, où μ est la moyenne des valeurs de la colonne pour chaque paire de classes et σ^2 est la valeur de la variance.

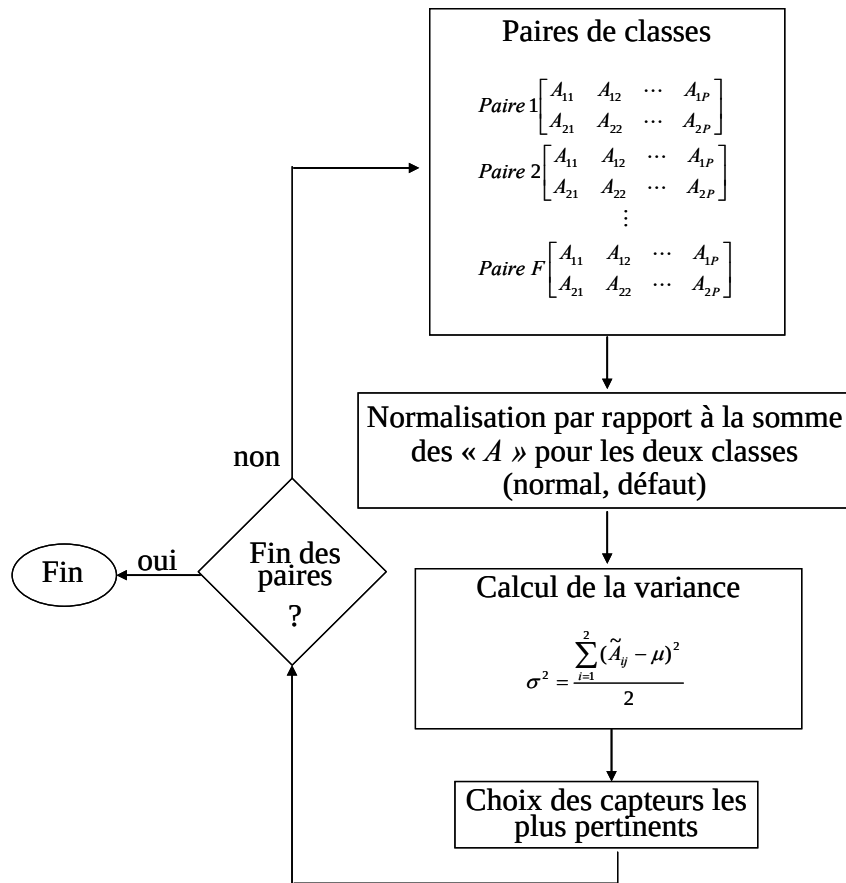


Figure 3-7 Organigramme de la méthode de la variance

3.4 Description de la méthodologie d'aide au placement de capteurs en vue du diagnostic

Dans ce paragraphe, nous allons montrer les différentes étapes de conception du système de diagnostic proposé. La méthode que nous proposons peut être considérée comme une approche intermédiaire entre les méthodes basées sur la reconnaissance et les méthodes basées sur les signaux.

Nous avons divisé le processus de la procédure de diagnostic en 3 parties : le placement des capteurs, l'établissement d'un modèle du comportement et le diagnostic en ligne. Les sections suivantes s'attachent à mieux expliquer chacune de ces parties. Tout d'abord, on donne les critères de paramétrage de la méthodologie LAMDA utilisée pour obtenir la classification à l'aide de tous les capteurs possibles puis pour la construction du modèle de comportement du procédé en vue de son diagnostic.

3.4.1 Paramètres du classificateur

Le premier aspect à considérer dans l'élaboration du système de surveillance est le paramétrage du classificateur en mode auto-apprentissage. Ces paramètres sont relatifs aux réglages nécessaires pour l'obtention du modèle de comportement à partir d'informations issues directement des mesures du processus ou bien par le biais d'outils de filtrage ou d'extraction d'information (FFT, pente, analyses statistiques) et de la méthode de classification sélectionnée.

Ce paramétrage est fait hors ligne et nécessite une intervention active de l'expert, de façon à intégrer ses connaissances dans toutes les étapes. Pour cela, nous devons fournir à l'expert des outils lui permettant de faire ce réglage en lui donnant des critères de choix objectifs.

A) Contexte

Dans le cas de la méthode LAMDA, l'expert doit définir l'intervalle de fonctionnement, c'est-à-dire les valeurs minimale et maximale pour la normalisation des descripteurs quantitatifs. Ainsi, l'expert pourra établir plusieurs combinaisons de configuration pour chaque descripteur dans l'espace de représentation. Le fait de pouvoir ajuster l'intervalle des descripteurs quantitatifs permet de définir jusqu'à quel point leurs variations peuvent être considérées. De même, le réglage de l'intervalle peut déterminer l'influence des descripteurs sur la caractérisation des classes. Kempowsky a donné quelques exemples de normalisation en utilisant différentes valeurs maximal et minimal lors d'une phase d'apprentissage par la méthodologie LAMDA au travers de l'outil *SALSA* [Kempowsky, 2004a].

B) Mode d'apprentissage

La classification de défauts par la méthodologie LAMDA peut être réalisée par auto-apprentissage (apprentissage non supervisé) ou imposée par l'expert (apprentissage supervisé).

- a) Auto-apprentissage : Il consiste dans la création, à partir des informations contenues dans un ensemble de données, de classes caractérisant les différents modes de fonctionnement. De cette façon, la méthode de classification propose une première partition à l'expert qui, ensuite, doit valider ces résultats en prenant en compte sa connaissance ainsi que les lois physiques régissant le comportement du processus.
-

- b) Apprentissage supervisé : Lorsque l'expert connaît a priori les modes réels de fonctionnement contenus dans l'ensemble d'apprentissage, nous parlons d'un apprentissage dirigé ou imposé par l'expert. En effet, l'expert étiquette les observations qui d'après lui, représentent au mieux les différentes situations. De cette façon, les paramètres de chaque classe seront une combinaison des caractéristiques des observations imposées par l'expert.

C) Réglage des paramètres du classificateur

Il faut choisir et ajuster certains paramètres de façon à obtenir une partition satisfaisante. Concernant la méthode LAMDA, il y a trois paramètres à régler : la fonction d'adéquation marginale (ou d'appartenance), les connectifs mixtes d'association pour l'agrégation des contributions de chaque descripteur et l'indice d'exigence.

Concernant la sélection de la fonction d'adéquation marginale [Kempowsky, 2004a] propose trois cas. Si la classification résultante n'est pas satisfaisante, il est possible de choisir une autre fonction quelconque qui permettra de modifier cette classification. Lors de nos travaux, nous avons utilisé la fonction Binomiale (Lamda1 sur l'outil SALSA) pour l'exemple de la production de glycol et la fonction Gauss (Gauss1 sur l'outil SALSA) pour celui sur le réacteur OPR. Ces choix ont été faits de façon heuristique, puisqu'il est difficile de sélectionner de manière analytique ou statistique l'une ou l'autre des méthodes lors de la présence de beaucoup de descripteurs (21 descripteurs sur la production de glycol et 27 descripteurs sur l'OPR).

Il est possible de régler l'exigence du classificateur à l'aide de l'indice d'exigence α . Ainsi, le classificateur le plus exigeant accepte un élément dans une classe uniquement si tous ces descripteurs présentent un Degré d'Adéquation Marginale (MAD) élevé à cette classe; ce qui correspond à avoir $\alpha=1$ (T-Norme (intersection)). De la même manière, le classificateur le moins exigeant accepte un élément dans une classe si au moins l'un de ses descripteurs possède une appartenance marginale élevée à cette classe ; dans ce cas il s'agit de la T-Conorme ($\alpha=0$). Donc, à partir du connectif mixte, l'expert peut avoir une série de partitions plus ou moins strictes de l'espace de description.

Deux autres paramètres importants à prendre en compte, qui ne sont disponibles dans l'outil Salsa qu'en mode auto-apprentissage sont : le pourcentage maximal de variation désiré (la variabilité) et le nombre maximal d'itérations permises. Ces itérations sont parfois nécessaires

parce que dans l'apprentissage non supervisé, les paramètres des classes varient ; ainsi pour obtenir la stabilité, la même population est classée plusieurs fois jusqu'à un pourcentage maximal de changements d'éléments d'une itération à une autre. Toutefois, le temps pour atteindre la stabilité peut être très élevée, c'est la raison pour laquelle un nombre maximal d'itérations a été introduit.

3.4.2 Placement de capteurs

Cette phase de la procédure consiste à rechercher les capteurs les plus pertinents sur le processus en vue de son diagnostic, à partir de l'analyse des signaux de ces capteurs qui, au départ, constituent l'ensemble des capteurs possibles (facilement disponibles). Cette procédure de sélection de capteurs, basée sur l'utilisation de la simulation ou d'expériences effectuées sur des installations pilotes permet d'établir dès la conception du procédé les points clés de mesure permettant un diagnostic aisé de son fonctionnement ; trop souvent, l'ajout d'un capteur supplémentaire une fois la construction de l'installation achevée ou en cours d'exploitation est pratiquement impossible. Cette procédure permet de réduire les coûts d'équipement mais aussi de permettre de déterminer où focaliser l'attention des opérateurs par exemple en réduisant le bruit (création exhaustive de classes). La Figure 3-8 explicite sous forme schématique l'enchaînement des différentes étapes constituant la procédure d'aide au placement de capteurs.

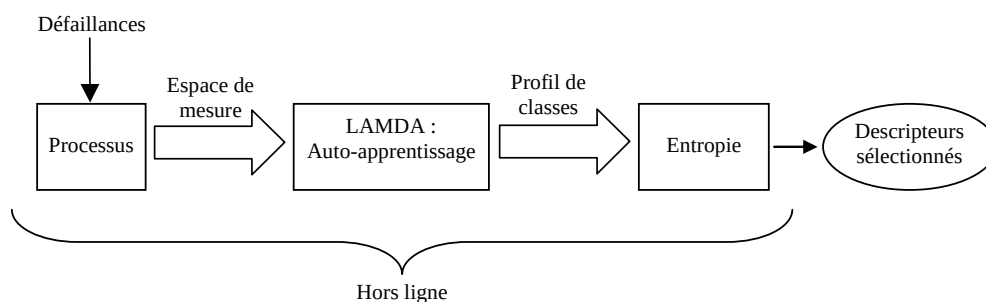


Figure 3-8 Schéma de la procédure d'aide au placement des capteurs

Pour effectuer cette phase de la méthodologie, il est nécessaire de posséder un nombre suffisant de données (hors ligne) pour décrire le comportement du processus en état normal et dans tous les états de défaillance possibles. En général, suffisamment de données sont

disponibles pour modéliser le fonctionnement normal du processus ; toutefois, d'une part il est difficile de se baser sur des données historiques pour modéliser correctement les fonctionnements anormaux du processus et d'autre part ce travail de placement de capteur doit être fait au plus tôt dans la phase de conception du procédé. Dans notre cas, nous avons eu accès à de puissants simulateurs dynamiques (commercial comme HYSYS, dans le cas de production de glycol, ou dédié comme le simulateur du réacteur OPR développé par le LGC [Elgue,2004]), grâce auxquels la génération de données a été facile à obtenir.

On a simulé divers défauts sur le processus pour classer les situations de défauts stabilisés. L'espace de mesure sur la figure précédente est l'ensemble des données composées des descripteurs $d_1, d_2, \dots, d_p$ qui jouent le rôle d'entrées du système d'aide au placement de capteurs. Un descripteur est défini comme un signal direct en provenance d'un capteur, un signal traité (e.g. un signal filtré) ou une donnée qualitative (suggérée par l'expert).

L'étape suivante de la méthodologie, est l'application de la méthode de classification LAMDA. L'objectif de cette phase est de créer une classification des diverses situations du processus (état normal et anormal).

3.4.3 Modèle de comportement

Cette phase correspond à la conception du modèle de comportement. Il s'agit de la construction à partir des données historiques issues du processus, d'un classificateur caractérisé par un ensemble de classes qui permettent d'identifier un ensemble de situations. Ceci correspond à l'obtention d'un modèle de comportement. La Figure 3-9 illustre schématiquement la phase de conception du modèle comportemental.

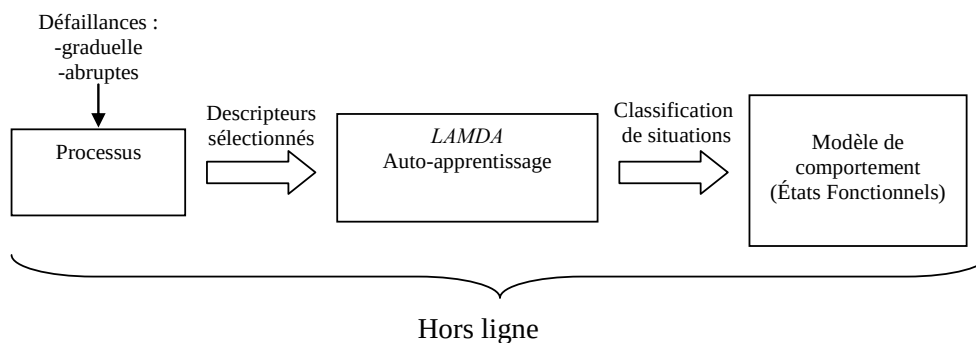


Figure 3-9 Conception du modèle comportemental

De même que dans l'étape précédente, il faut posséder des données historiques ou d'observations de différentes situations (normal et anormales). Grâce à des simulateurs, nous pouvons disposer d'un nombre suffisant de données pour chaque état de défaillance et chaque état de fonctionnement normal du processus. En revanche, il est fort probable que tous les défauts n'aient pas été simulés a priori, à cause de la difficulté d'appréhender l'exhaustivité des défauts dans un processus complexe ou dans le cas où deux ou plusieurs défaillances combinées provoquent des effets qu'aucune d'elles, prise isolément, ne pourrait produire (synergie).

Une fois le modèle construit, le suivi des situations attendues, identifié comme étant la phase de reconnaissance, vise à associer toute nouvelle observation à l'une des classes déterminées.

3.4.4 Diagnostic en ligne

Comme nous l'avons mentionné dans la section précédente, rien ne permet d'assurer l'exhaustivité de la connaissance du comportement du processus car les données historiques ne peuvent pas couvrir la vie entière du processus. De plus, au fur et à mesure de la présentation de nouvelles observations, de nouvelles situations peuvent apparaître dans la structure initiale. Pour cette raison, il est nécessaire que le système de surveillance présente un caractère adaptatif au moment de l'identification des nouvelles situations. Pour cela, deux principes d'apprentissage sont prévus : un apprentissage hors ligne et une reconnaissance en ligne. L'apprentissage hors ligne a été expliqué dans la section précédente (3.4.3), lequel est utilisé dans l'étape de création du modèle de comportement, mais aussi pour caractériser de nouvelles situations. La reconnaissance en ligne permet de détecter en continu un défaut en se basant sur le modèle de comportement, afin de détecter, localiser et identifier une défaillance. La Figure 3-10 montre le diagramme de blocs associé au diagnostic en ligne.

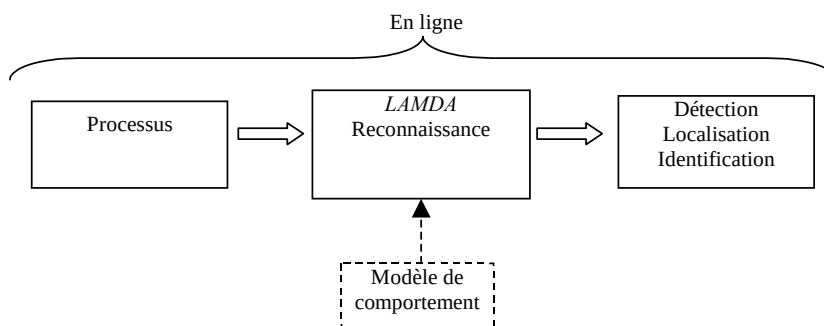


Figure 3-10 Schéma du diagnostic en ligne du procédé

Afin d'expliquer de manière générique l'algorithme de diagnostic en ligne proposé, nous avons divisé la phase de diagnostic en trois parties (voir la Figure 3-11) : l'espace de mesure, l'espace de reconnaissance et l'espace de diagnostic. Dans ce qui suit, nous détaillons chacune de ces étapes.

- a) **Espace de mesure** : C'est l'ensemble des descripteurs $d_1, d_2, \dots, d_p$ sélectionnés dans la phase d'aide au placement de capteurs (section 3.4.2). Ce sont les entrées du système de diagnostic en ligne.

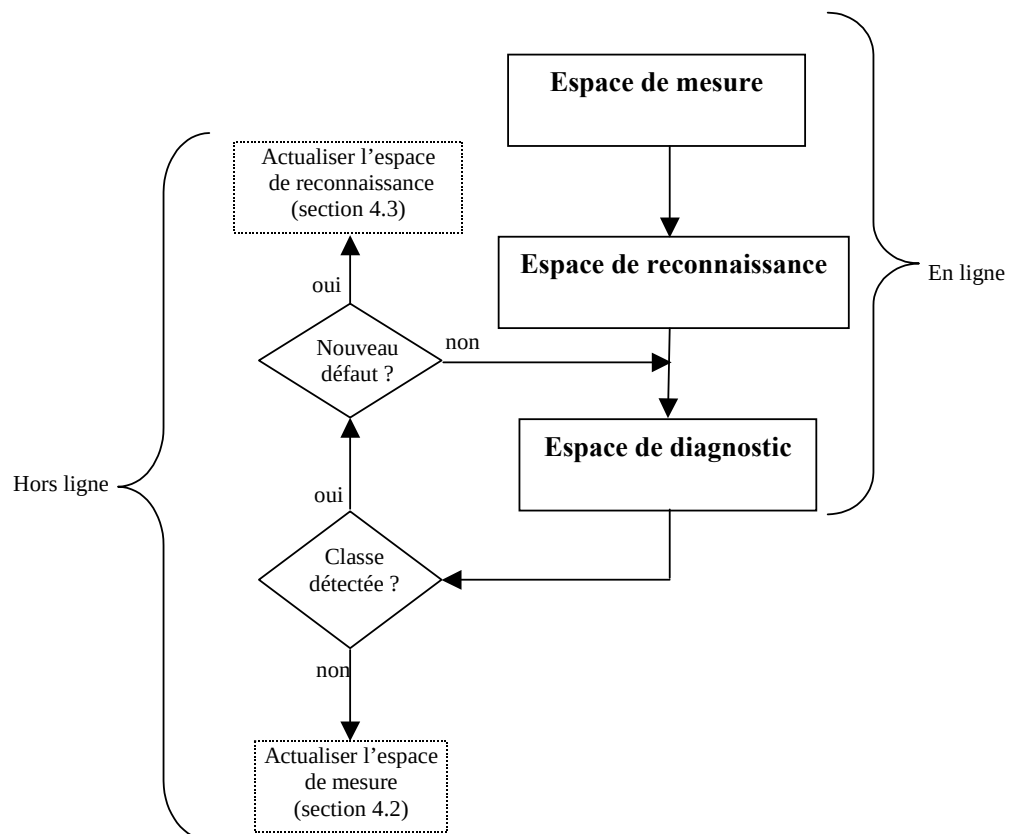


Figure 3-11 Étapes de la conception du système de diagnostic

- b) **Espace de reconnaissance**. Cette étape consiste à déterminer, à chaque instant, dans quel état fonctionnel se trouve le procédé lorsqu'une nouvelle observation est présentée. Il est basé sur le modèle de comportement construit hors ligne.
- c) **L'espace de diagnostic**. Cette étape explique le type de défaillance qui a causé le défaut sur le procédé. Il faut prendre en compte que tous les modes de fonctionnement ne sont pas forcément identifiés lors de l'étape de création du

modèle (en mode apprentissage). Si un nouveau défaut détecté arrive, alors il faut actualiser le modèle de comportement, en ajoutant une nouvelle classe qui représente cette défaillance.

Il est possible que, des opérations anormales dans un procédé n'aient pas été détectées, à cause d'un manque de descripteurs pour en faire le diagnostic. Ceci veut dire que, le système de diagnostic n'est capable de détecter le défaut (il confond ce défaut avec l'opération normale du procédé) mais que la détection ait été réalisée par l'opérateur. Dans ce cas, il faut revenir à l'étape de placement de capteurs et créer une nouvelle classe qui représente le défaut non détecté.

Le plus important dans l'identification d'un nouveau défaut est de ne pas le confondre avec un autre dysfonctionnement connu : il vaut mieux dire à l'opérateur : « je ne sais pas ce que c'est , mais ce n'est pas le fonctionnement normal » que « c'est ce type de défaillance » et se tromper. Dans le cas d'une confusion, il faut refaire la conception du modèle en modifiant les paramètres du système de diagnostic.

3.5 Conclusions

Dans ce chapitre, nous avons cité diverses références avec différentes approches pour optimiser le placement de capteurs. La plupart fait appel à un modèle explicite du procédé, comme le placement optimal par reconstruction d'états pour les systèmes linéaires ou systèmes à paramètres distribués. Une seule technique autre que celle issue de la classification, consiste à utiliser une analyse par composantes principales.

La méthode que nous avons développée s'appuie sur le résultat d'une classification (quelle qu'elle soit) et donc l'utilisation de données historiques de fonctionnement ou de la simulation du procédé. Elle est basée sur le concept de l'Entropie de SHANNON, permettant de quantifier l'apport d'un descripteur à la discrimination d'une classe de défaillance par rapport à l'état normal.

Après avoir rappelé les principaux concepts et propriétés de l'entropie de SHANNON et donné plusieurs critères de sélection (entropie probabiliste, entropie non probabiliste ou variance), deux procédures de sélection des capteurs ont été définies : la procédure « par paire de classes » qui consiste à former des paires classe de défaillance/classe de fonctionnement

normal et de sélectionner le capteur donnant le plus fort gain d'information pour chacune des paires et la procédure d'analyse des descripteurs sur l'ensemble des classes ; la sélection des capteurs se faisant dans ce cas par rapport à un seuil préalablement fixé. Cette dernière procédure permet de mieux discriminer toutes les défaillances comme il sera vu dans le chapitre suivant qui montre l'application de cette méthodologie à deux exemples : un procédé chimique complexe incluant plusieurs opérations unitaires et un nouveau procédé en cours de développement : *l'Open Plate Reactor*.

Nous avons présenté les différents paramètres de la méthode LAMDA à régler pour avoir une classification acceptable. Les résultats de cette classification sont primordiaux tant pour le placement de capteurs que pour la conception du modèle de comportement du système sur lequel est basé le diagnostic.

4. PLACEMENT DE CAPTEURS ET DIAGNOSTIC DE PROCÉDES CHIMIQUES

4.1 Introduction

L'objectif de ce chapitre est de présenter et discuter les résultats obtenus grâce à la méthodologie développée au cours de nos travaux concernant la détermination des capteurs les plus pertinents pour effectuer un diagnostic en ligne. A partir d'un ensemble de capteurs possibles (en excluant les capteurs de concentration), un sous-ensemble de capteurs est déterminé de manière à pouvoir détecter en ligne les états fonctionnels associés à des défaillances lors d'une phase de reconnaissance. Deux exemples viennent illustrer l'application de cette méthodologie. Le premier (production de propylène glycol) permet de montrer comment grâce au couplage de cette méthodologie à l'utilisation d'un simulateur dynamique de procédés du commerce, il est possible de déterminer les capteurs les plus adaptés pour la supervision. Le second exemple porte sur une unité pilote en cours de développement (l'OPR). Dans chacun des exemples, la démarche adoptée a été identique et s'est déroulée selon 4 étapes. Les résultats des 4 étapes successives sont présentés selon 4 sections (voir la Figure 4-1) :

1. Détection des défauts du procédé en utilisant tous les descripteurs disponibles

Le but de cette étape est de classer de manière stable et concise tous les défauts du procédé à partir d'une classification floue (*LAMDA*) et de déterminer grâce aux profils des classes, la valeur moyenne normalisée de chaque descripteur dans la classe correspondante.

2. Sélection des capteurs en utilisant l'entropie

La finalité de cette partie est d'utiliser les concepts de l'entropie et du gain de l'information comme bases pour la sélection des capteurs les plus pertinents sur le procédé afin d'effectuer ultérieurement le diagnostic de fautes.

3. Obtention du modèle de comportement

L'objectif de cette section est d'obtenir le modèle de comportement du procédé à partir des descripteurs les plus pertinents, qui permet d'identifier un ensemble de situations.

4. Reconnaissance de défauts inconnus

Cette dernière partie est consacrée à la reconnaissance en-ligne de défauts dont la date et l'amplitude ne sont pas connues au préalable, ceci pour vérifier et valider le modèle de comportement ainsi que la pertinence du choix des capteurs retenus.

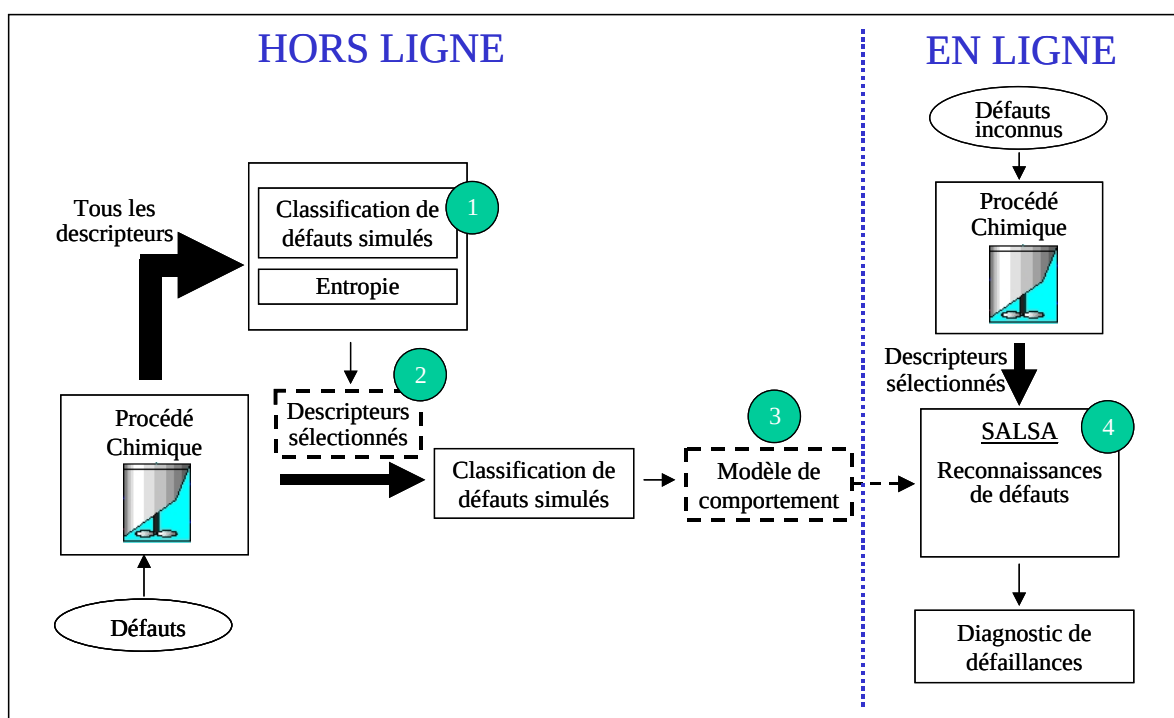


Figure 4-1 Diagramme des différentes étapes de la méthodologie de placement des capteurs et du diagnostic de procédés chimiques

4.2 Application de la méthodologie à partir de l'utilisation du simulateur industriel HYSYS

La première étape de l'application de la méthodologie consiste à définir l'ensemble des capteurs qui pourront être installés sur l'installation. Dans l'exemple présenté, nous avons choisi de ne considérer que des capteurs usuels : température, pression, niveau. Les capteurs de concentration ne font pas partie de cette liste puisqu'en pratique, il est très difficile et onéreux de trouver un moyen de mesurer **en-ligne** la concentration d'un constituant dans un mélange multiconstituant. Les premiers capteurs déjà présents sur l'unité sont ceux utilisés par les 4 régulateurs utilisés dans le réacteur et dans la colonne de distillation (voir la Figure 2-5 et la Figure 2-8). A ces capteurs, nous avons ajouté les valeurs de sortie des PID qui sont les valeurs des variables manipulées (ou commandes délivrées par les PID). A ces 8 variables ont été rajoutées des mesures de pression et de température. En particulier, il nous est apparu intéressant de pouvoir utiliser comme information des mesures de température à l'intérieur de la colonne de distillation, c'est pourquoi il a été considéré que tout plateau pouvait être muni d'un capteur de température. Ainsi, au total il y a 21 points de mesure possibles qui constituent les descripteurs en entrée de la technique de classification. Ces 21 descripteurs ainsi que leur nature sont regroupés dans le Tableau 4-A.

Tableau 4-A Nom et nature des descripteurs utilisés pour la classification

Numéro Descrip.	Nom	Numéro Descrip.	Nom	Numéro Descrip.	Nom
1	Pression Réacteur (Kpa)	8	Température plateau 2 (°C)	15	Température plateau 9 (°C) (Variable régulée)
2	Température Réacteur (°C) (Variable régulée)	9	Température plateau 3 (°C)	16	Température plateau 10 (°C)
3	Niveau Liquide Réacteur (%) (Variable régulée)	10	Température plateau 4 (°C)	17	Puissance délivrée au bouilleur (Kj/h) (Variable manipulée)
4	Vanne pour régler le débit de Refroidissement Réacteur (%) (Variable manipulée)	11	Température plateau 5 (°C)	18	Vanne pour régler le débit de Refroidissement Condenseur (%) (Variable manipulée)
5	Vanne pour régler le Niveau du Liquide Réacteur (%) (Variable manipulée)	12	Température plateau 6 (°C)	19	Pression Condenseur (Variable régulée)
6	Température de sortie du fluide Refroidissement Réacteur (°C)	13	Température plateau 7 (°C)	20	Température de sortie du fluide de refroidissement du condenseur (°C)
7	Température plateau 1 (°C)	14	Température plateau 8 (°C)	21	Température sortie du Glycol (°C)

La première phase de la méthodologie pour déterminer les capteurs pertinents, est la génération de simulations suivant des scénarii de défaillance. Chaque perturbation ou

défaillance simulée comporte deux phases : la hausse et la baisse de la variable perturbée. Dans le Tableau 4-B, nous avons regroupé les perturbations simulées, la valeur nominale de la variable perturbée et ses valeurs maximale et minimale à partir desquelles le système est capable de retourner à l'opération normale. Il est clair qu'une perturbation trop importante peut rendre le procédé instable (le procédé est dit instable si son évolution ne tend plus vers un point de régime permanent). Cette situation n'est pas recherchée puisque le diagnostic de l'unité doit permettre d'identifier au plus tôt les dérives pour pouvoir agir en conséquence et revenir à un état stationnaire. La gamme des valeurs des perturbations pour lesquelles le système reste stable a été obtenue de façon empirique en utilisant le simulateur dynamique HYSYS.

Tableau 4-B Gamme des valeurs des perturbations pour lesquelles le système reste stable

Défaillance	Gamme de valeurs			Unités
	Minimum	Nominal	Maximum	
Variation de débit d'oxyde de propylène	50	68.04	80	Kgmole/h
Variation de la température du débit d'alimentation du fluide de refroidissement du réacteur	11	15	25	°C
Modification de la vitesse de réaction: Variation du facteur de fréquence	1.20E+13	1.70E+13	2.20E+13	Moles de réactif consommé/s/unité de volume
Variation de la température du fluide de refroidissement du condenseur	11	15	19	°C

La détermination de ces valeurs minimales et maximales a permis de définir les bornes de variation des différents descripteurs. Ainsi, le contexte dans lequel la classification a été effectuée ne repose pas sur les valeurs minimales et maximales des descripteurs lors d'une simulation particulière d'une défaillance, mais bien sur un intervalle de variation physiquement cohérent. La normalisation des descripteurs a donc été effectuée à l'aide des valeurs limites pour chaque descripteur après avoir réalisé toutes les simulations des scénarii étudiés. Le réglage de l'intervalle peut déterminer l'influence des descripteurs sur la caractérisation des classes (la définition de cet intervalle permet de définir le contexte de la classification et peut donc influencer les résultats de la classification et le poids des différents descripteurs sur le profil de la caractérisation des classes). Si nous choisissons un intervalle

étroit, les oscillations des signaux sont agrandies, par contre si l'intervalle est très large, les petites modifications des signaux peuvent être non observables.

Les signaux en provenance du procédé tendent à osciller et /ou à diverger quand les valeurs des défaillances sont en dehors des intervalles donnés dans le Tableau 4-B. Avant d'appliquer les méthodologies de placement de capteurs et de diagnostic, on peut traiter les signaux du procédé en utilisant les méthodes d'abstraction de signaux. De même, des descripteurs quantitatifs peuvent être obtenus à partir du pré-traitement des signaux bruts : par exemple par filtrage ou en appliquant une FFT (Transformée de Fourier Rapide) [Orantes, 2003 ; Orantes, 2004].

4.2.1 Détection des défauts en utilisant tous les descripteurs : application au modèle de production de propylène glycol

Le but de cette étape dans le processus de détermination des capteurs pour le diagnostic, est d'obtenir une classification des états fonctionnels du procédé correspondant aux différentes défaillances simulées. Cette classification effectuée grâce à la méthode LAMDA (sur laquelle repose le logiciel *SALSA*) doit être concise (classe bien définie) et stable (sans oscillation). Elle permet d'obtenir la valeur moyenne normalisée de chaque descripteur à la classe correspondante (information obtenue à partir du profil des classes).

A partir des gammes des perturbations données dans le tableau 4-B, 8 défauts ont été simulés correspondant aux 8 valeurs extrêmes des perturbations. La classification a été réalisée par auto-apprentissage de ces 8 défauts en utilisant tous les descripteurs disponibles (21 au total) sur le procédé. Dans ce cas, nous avons utilisé la fonction *Binomiale* et le connectif *Max-Min* puisqu'ils ont donné les meilleurs résultats de façon empirique, ainsi qu'un niveau d'exigence (α) de 0.8 et une variabilité maximale de 0.1.

La fonction *Binomiale* fait une partition de l'espace de représentation par rapport aux valeurs extrêmes, c'est-à-dire qu'il y a une tendance vers la bipolarisation « réussite-échec », identifiées à 1 et 0. Les travaux de [Kempowsky, 2004a] montrent quelques partitions de l'espace de description pour diverses fonctions MAD pour le cas quantitatif, en utilisant deux descripteurs.

Pour permettre une meilleure visualisation de la procédure développée, nous nous intéresserons dans ce qui suit au cas de la détection d'une défaillance particulière à savoir la variation de l'alimentation en oxyde. Les évolutions normalisées des 21 descripteurs sont reportées sur le graphe de la Figure 4-2 (graphe du bas). Ces résultats ont été ensuite traités

par l'outil de classification *SALSA* [Kempowsky, 2004b] pour obtenir 2 classes (la classe de fonctionnement normal et la classe correspondant à la défaillance établie). On rappelle que l'apprentissage se fait en mode automatique, c'est-à-dire sans connaissance des classes prédéfinies. La Figure 4-2 montre le défaut hausse de propylène d'oxyde depuis la valeur nominale (68.04 kgmoles/h) jusqu'à 80 kgmoles/h (valeur maximale). Dans cette simulation de défaillance, la défaillance appliquée est de type graduel : incrément de 0.04 par minute sur le débit d'oxyde pendant 300 minutes.

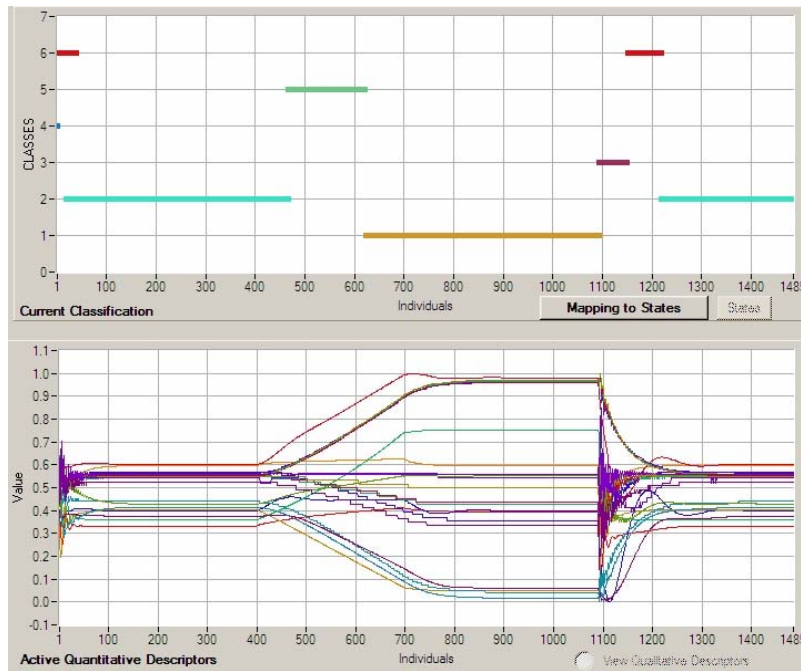


Figure 4-2 Classification suite à la simulation d'une hausse d'oxyde de propylène

Sur le graphe supérieur de la Figure 4-2, on peut constater la création de 6 classes. La classe 1 correspond à la défaillance lorsque celle-ci s'est établie. La classe 2 est rattachée à l'état de fonctionnement normal. La classe 5 représente le régime transitoire entre un état de fonctionnement normal et la défaillance établie. Les autres classes 3, 4 et 6 représentent divers états de transition liés soit à l'établissement du régime de fonctionnement normal soit à la phase transitoire représentant le retour à un fonctionnement normal depuis un état de défaillance. A partir des résultats de cette classification, on peut vérifier que le défaut est bien détecté (classe 1) et peut-être affecté à une seule classe et que l'état de fonctionnement normal est lui aussi bien reconnu (classe 2). Si on s'intéresse aux deux classes particulières : 1 et 2, nous pouvons grâce aux résultats de la classification construire la matrice du profil de ces 2 classes :

Tableau 4-C Profil de deux classes

	d1	d2	d3	d4	d5	d6	d7	d8	d9	d10	d11	d12	d13	d14	d15	d16	d17	d18	d19	d20	d21
NORMAL	0.33	0.59	0.40	0.36	0.60	0.40	0.56	0.55	0.55	0.55	0.56	0.56	0.56	0.56	0.54	0.44	0.37	0.43	0.56	0.41	0.52
DEFAUT	0.39	0.60	0.41	0.73	0.96	0.07	0.38	0.44	0.45	0.50	0.92	0.92	0.92	0.93	0.55	0.09	0.10	0.55	0.56	0.06	0.35

Chaque élément de la matrice représente la valeur moyenne normalisée du descripteur à la classe correspondante (voir le Tableau 4-A pour la description des descripteurs). Ces valeurs ont été obtenues à partir du calcul par la moyenne générale (voir le paragraphe descriptif de la méthode LAMDA au chapitre 1). La Figure 4-3 donne l'histogramme des profils des deux classes formées lors de l'étape de l'apprentissage (Normal et Défaut).

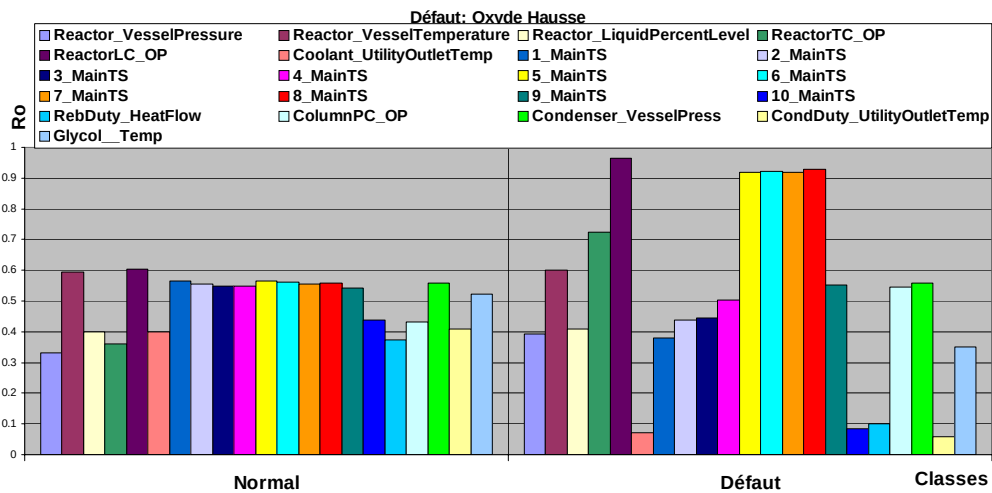


Figure 4-3. Profil des 2 classes « Normal » et « Défaut »

Cette procédure peut être répétée pour chaque défaut (voir le Tableau 4-B). On obtient ainsi 8 profils de classes au total.

Grâce à l'application d'un défaut graduel (sous forme de pente), on vérifie que l'on peut bien détecter un état de transition depuis l'opération Normale du procédé jusqu'à l'établissement de la défaillance. La classe qui représente ce défaut ne sera pas prise en compte pour le placement des capteurs, par contre cet état de transition est très important pour la reconnaissance en ligne de la défaillance pour la détecter préalablement. Cette classe est appelée « alarme » de la défaillance.

Il faut remarquer que cette phase de la procédure, préalable à la phase de sélection des capteurs pertinents proprement dite, est très importante car elle doit permettre d'obtenir des classes bien définies et concises de telle sorte que les valeurs obtenues à partir du profil de classes soient les plus précises possibles, et ainsi, permettre d'avoir une meilleure sélection

des capteurs. A partir du résultat de la classification de la défaillance, l'expert doit valider que la classe caractérise correctement la défaillance. Si ce n'est pas le cas, on doit modifier les paramètres de l'exigence (α) et de la variabilité pour obtenir la bonne caractérisation.

Une fois qu'une classification a été validée par l'expert et à partir du profil de classes obtenues, la deuxième phase de la procédure qui consiste à appliquer le concept de **l'entropie et du gain d'information** pour sélectionner les descripteurs les plus pertinents, peut être initiée.

4.2.2 Sélection des capteurs : application au modèle de la production de propylène glycol

L'objectif principal de nos travaux a été de développer une procédure de placement optimal des capteurs grâce à l'utilisation du concept de l'entropie de *SHANNON* couplé avec la classification. Pour illustrer cette procédure, repartons du profil de classes obtenu à la suite de la classification effectuée à l'aide de tous les descripteurs possibles, telle celle obtenue dans la section précédente. La méthodologie développée repose sur cinq étapes distinctes :

Étape 1 : Espace probabilisé. Cette étape initiale est nécessaire pour pouvoir appliquer le concept de l'entropie d'un signal comme il a été défini au chapitre 3. En effet, ce concept repose sur la possession d'un espace probabilisé. Comme nous travaillons avec des valeurs moyennes normalisées non probabilisées, la somme sera différente de 1. Pour se ramener au cas probabiliste, chaque élément du profil de classes est normalisé par rapport à la somme totale des valeurs moyennes normalisées des classes formées. A partir de la matrice (Tableau 4-C) définissant le profil des classes, nous obtenons alors la matrice suivante :

Tableau 4-D Profil normalisé de deux classes

$$\begin{pmatrix} 0.46 & 0.50 & 0.49 & 0.33 & 0.39 & 0.85 & 0.60 & 0.56 & 0.55 & 0.52 & 0.38 & 0.38 & 0.38 & 0.38 & 0.50 & 0.84 & 0.79 & 0.44 & 0.50 & 0.88 & 0.60 \\ 0.54 & 0.50 & 0.51 & 0.67 & 0.61 & 0.15 & 0.40 & 0.44 & 0.45 & 0.48 & 0.62 & 0.62 & 0.62 & 0.62 & 0.50 & 0.16 & 0.21 & 0.56 & 0.50 & 0.12 & 0.40 \end{pmatrix}$$

Étape 2 : Calcul de l'Entropie Maximale. L'Entropie maximale $H(C)$ de la classification résultante est alors calculée comme suit [Walter et Pronzato, 1997] :

$$H(C) = \log n \quad (4.1)$$

Dans ce cas, le logarithme utilisé a été le logarithme népérien. L'Entropie maximale correspond en effet au cas où l'élément a la même probabilité d'appartenir à l'une ou l'autre des classes ($P_i=1/2$ dans le cas de 2 classes). L'Entropie maximale du système de classification $H(C)$ est alors égale à 0.6931. On rappelle que l'Entropie de SHANNON s'interprète comme l'incertitude de classification d'un élément à la classe.

Étape 3. L'Entropie du descripteur. L'entropie probabiliste par rapport à chaque descripteur est ensuite calculée (voir l'équation 3.4). Elle représente la disparité entre les valeurs des classes étudiées (« normal » et « défaut »). La valeur de l'Entropie de classification pour les descripteurs d_k , ($H(C|d_k)$) est :

Tableau 4-E Valeurs de l'Entropie par rapport à chaque descripteur

(0.69 0.69 0.69 0.64 0.67 0.43 0.67 0.69 0.69 0.69 0.66 0.66 0.66 0.66 0.69 0.44 0.52 0.69 0.69 0.38 0.67)

Étape 4. Calcul du « Gain d'information » des descripteurs. On calcule le « gain de l'information » comme étant la différence entre l'Entropie maximale et l'Entropie par rapport à chaque descripteur (voir l'équation 3.5):

Tableau 4-F Valeurs du Gain d'information par rapport à chaque descripteur

(0.00 0.00 0.00 0.06 0.03 0.27 0.02 0.01 0.01 0.00 0.03 0.03 0.03 0.03 0.00 0.25 0.17 0.01 0.00 0.32 0.02)

Etape 5. Calcul du Gain Relatif des descripteurs. Finalement, on obtient le gain relatif par rapport au gain maximal (voir l'équation 3.6) :

Tableau 4-G Valeurs du Gain d'information Relatif par rapport à chaque descripteur

(0.53 0.00 0.01 8.33 3.83 38.68 2.82 0.97 0.77 0.12 4.18 4.33 4.39 4.46 0.00 36.06 24.97 0.97 0.00 45.86 2.80)

La Figure 4-4 montre les gains relatifs de chacun des descripteurs pour le défaut considéré. Dans cet exemple, on voit que le descripteur *CondDuty_UtilityOutletTemp* (température du

fluide de refroidissement en sortie du condenseur) donne la meilleure information discriminante pour ce défaut avec une valeur de 46% de gain relatif.

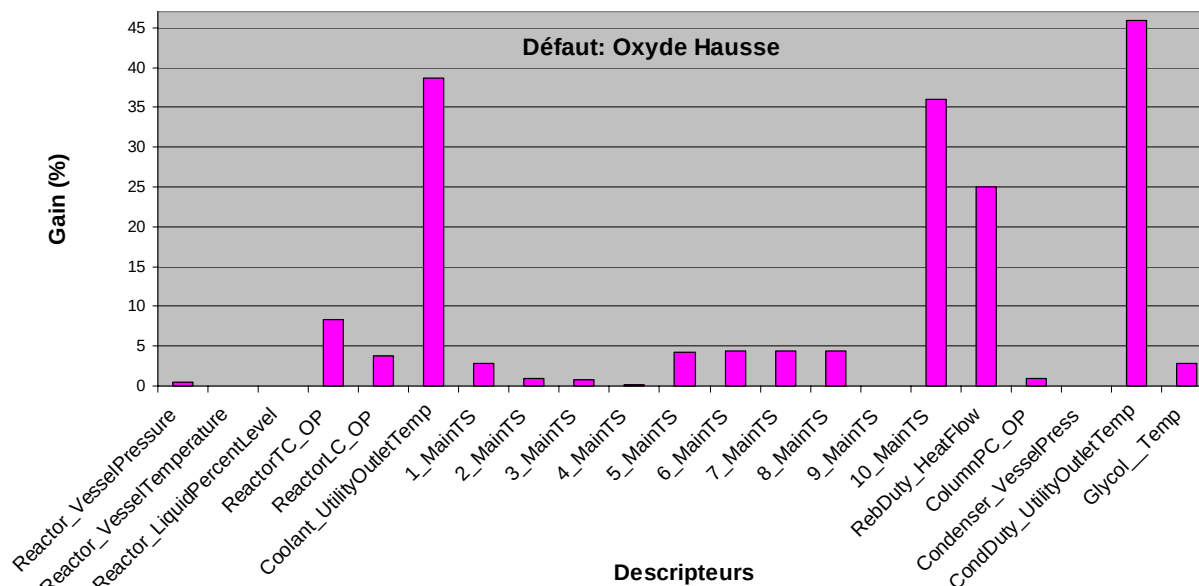


Figure 4-4 Résultat du gain relatif dans le cas de la défaillance « hausse d'oxyde »

Cette procédure qui a été décrite pour la détection du défaut « hausse d'oxyde » à partir d'un fonctionnement normal, peut être effectuée pour chacun des défauts répertoriés sur le procédé (8 défauts au total). A chaque fois, à partir des résultats de la simulation d'un scénario comprenant une phase de fonctionnement normal suivi d'une défaillance, la classification est réalisée et les calculs des étapes 1, 2, 3, 4, 5 sont effectués. On obtient la matrice des gains relatifs par rapport à toutes les classes représentant un défaut établi, qui est donnée ci-dessous (Tableau 4-H) : chaque colonne représente un descripteur et chaque ligne une défaillance (voir le Tableau 4-B pour la description des défaillances appliquées). Dans cette matrice, on peut constater que la dernière ligne correspond aux résultats de la défaillance « hausse d'oxyde » dont nous venons d'expliquer les calculs du gain relatif.

Tableau 4-H Valeurs de GAINS Relatifs appliqués à toutes les défaillances

13.6	0.00	0.00	0.16	0.00	0.19	3.66	1.18	1.26	1.28	0.62	0.17	0.16	0.14	0.00	0.15	0.21	10.9	0.00	71.4	0.02
43.2	0.00	0.00	0.05	0.00	0.06	0.74	0.21	0.32	0.33	0.16	0.04	0.05	0.04	0.00	0.11	0.14	23.1	0.00	4.65	0.02
0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	4.38	0.00	0.00	0.00
0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	3.86	0.00	0.00	0.00
0.00	0.00	0.00	1.43	0.00	1.07	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00
0.00	0.01	0.00	14.3	0.00	23.8	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.00	0.02	0.00	0.02	0.00
3.50	0.00	0.04	31.6	52.2	11.9	1.76	0.65	0.68	0.12	27.4	30.1	30.1	32.7	0.03	7.48	8.17	2.71	0.00	10.7	1.77
0.53	0.00	0.01	8.33	3.83	38.7	2.82	0.97	0.77	0.12	4.18	4.33	4.39	4.46	0.00	36.1	25.0	0.97	0.00	45.9	2.80

Les valeurs de cette matrice peuvent être portées sur un histogramme (Figure 4-5). Chaque défaut est repérable par une couleur. Par exemple, sur cette figure, le défaut « hausse d'oxyde » est représenté en couleur rose. Nous retrouvons sur ce graphe que, pour ce défaut, le descripteur fournissant le gain relatif le plus élevé est le **descripteur 20** (voir le Tableau 4-A), c'est-à-dire la température du fluide de refroidissement en sortie du condenseur. Pour chaque défaut, on peut sélectionner de cette manière le descripteur qui fournit un gain relatif maximal. On obtient ainsi les descripteurs qui fournissent le plus d'information pour la classification exprimée comme la disparité entre les valeurs des classes étudiées (« normal » et « défaut ») et donc les plus pertinents pour le diagnostic du procédé. Comme le montre le Tableau 4-I, ils sont au nombre de 6 descripteurs, chaque défaut correspond à une couleur différente. Il est intéressant à ce stade de revenir sur l'interprétation physique de ces choix et en particulier d'examiner s'ils sont cohérents avec la connaissance que l'on a des phénomènes mis en jeu.

La sortie du régulateur assurant le contrôle de la pression de la colonne (*ColumnPC OP*) a été sélectionnée comme « meilleur » capteur pour identifier des défauts sur la température d'alimentation du fluide de refroidissement au condenseur. La pression dans le condenseur étant régulée par action sur le débit de liquide de refroidissement au condenseur, une baisse ou une hausse de sa température d'alimentation aura comme conséquence une augmentation ou une diminution de son débit par le régulateur pour compenser ces variations et maintenir une pression constante. Le choix de cette variable (*ColumnPC OP*) est donc cohérent avec la connaissance des phénomènes physiques mis en jeu.

La baisse d'oxyde à l'alimentation est identifiable à partir de la variable de commande du régulateur de niveau : on alimente moins d'où une modification de cette variable pour maintenir le niveau.

Le défaut hausse du facteur de fréquence a pour conséquence une augmentation de la quantité de glycol formée et donc du niveau ou pression dans le réacteur.

La température de sortie du fluide de refroidissement du condenseur (*CondDuty_UtilityOutletTemp*) permet d'évaluer si la composition en tête de colonne a changé. En effet, un changement de composition affecte les propriétés thermodynamiques du mélange et donc la quantité d'énergie à échanger avec le fluide de refroidissement pour condenser le mélange. Cette modification se traduit physiquement par une modification de la température du fluide de refroidissement en sortie du condenseur. Cette variation de la composition en tête de colonne peut être causée par une modification de la composition de l'alimentation du procédé (défaut : hausse oxyde) ou par un changement de la cinétique (en pratique changement de la réactivité des réactifs, changement de l'activité du catalyseur, ...) simulé dans notre cas par une variation du facteur pré-exponentiel (défaut baisse A). Ce qui est intéressant de noter c'est que ces deux défauts ont la même conséquence : une augmentation de la composition d'oxyde dans le courant d'alimentation de la colonne.

Un défaut sur la température du liquide de refroidissement du réacteur (*TempInletCool*) est identifié dans le cas d'une hausse par la mesure de la température du fluide de refroidissement en sortie (*Coolant_UtilityOutletTemp*) : celle-ci augmente. Une baisse de cette température d'alimentation est en revanche plus détectable à partir de la variable de commande générée par le régulateur de température (débit du fluide de refroidissement). La distinction entre ces deux informations suivant le sens de la perturbation est cohérente. En effet, lors d'une augmentation de la température du fluide de refroidissement, le régulateur va augmenter le débit jusqu'à atteindre un débit maximum et c'est ensuite la température du fluide de refroidissement en sortie qui augmentant, contiendra le plus d'information. Dans le cas d'une baisse de la température d'alimentation du fluide de refroidissement, le régulateur diminuera le débit en ayant plus de marge de manoeuvre que dans le cas précédent donc la variable (*ReactorTC OP*) sera porteuse d'information.

Il y a donc bien cohérence entre les capteurs choisis et les connaissances a priori des phénomènes mis en jeu.

La Figure 4-5 confirme ce qui pouvait être appréhendé à savoir que les variables régulées des divers régulateurs dans le procédé n'apportent que très peu d'information. En effet, le fait qu'elles doivent être maintenues à des valeurs constantes, leur variation n'est pas synonyme d'un changement de point de fonctionnement. En revanche, ce changement de point de fonctionnement se voit sur les valeurs des variables manipulées. Il est bien entendu que les

variables sur lesquelles ont porté les perturbations (par exemple la température du fluide de refroidissement) n'ont pas été retenues comme descripteurs car dans ce cas, la détection d'un dysfonctionnement est immédiate.

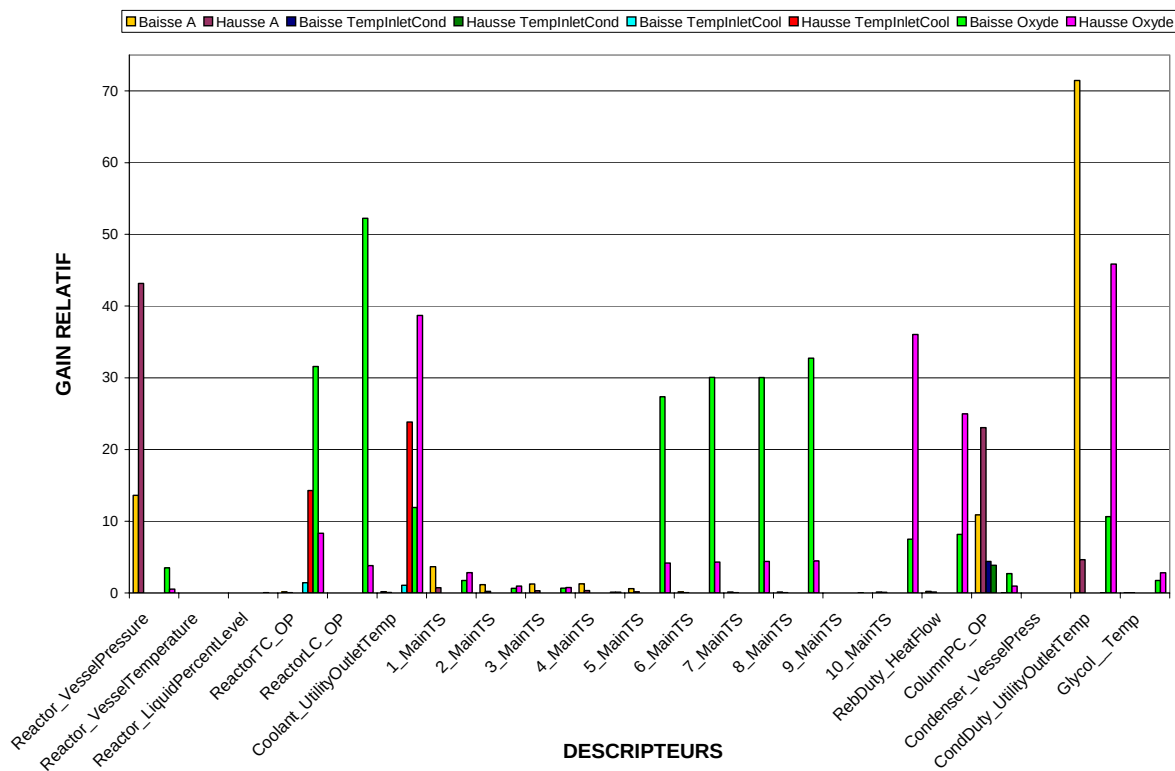


Figure 4-5 Gains relatifs des descripteurs en prenant en compte tous les défauts

Tableau 4-I Résultat des descripteurs sélectionnés

NOM DÉFAUT	DESCRIPTEURS SÉLECTIONNÉS
Baisse A (Fréq. Naturelle)	CondDuty_UtilityOutletTemp
Hausse A (Fréq. Naturelle)	Reactor_VesselPressure
Baisse TempInletCond	ColumnPC_OP
Hausse TempInletCond	ColumnPC_OP
Baisse TempInletCool	ReactorTC_OP
Hausse TempInletCool	Coolant_UtilityOutletTemp
Baisse Oxyde	ReactorLC_OP
Hausse Oxyde	CondDuty_UtilityOutletTemp

4.2.3 Obtention du modèle du comportement

L'étape suivante est de trouver le modèle qui soit capable de caractériser le comportement du procédé. Elle consiste à construire, à partir des descripteurs sélectionnés dans la section précédente, un classificateur caractérisé par un ensemble de classes qui permettent d'identifier un ensemble de situations.

Tout d'abord, un apprentissage non supervisé (auto-apprentissage) est effectué avec ces descripteurs pour déterminer les classes correspondant aux défauts, ainsi qu'aux diverses alarmes. La Figure 4-6 donne les résultats obtenus après l'auto-apprentissage à partir de la simulation d'un scénario incluant 8 défauts (les évolutions des 6 descripteurs sélectionnés dans l'étape précédente lors de cette simulation sont données sur le graphe du bas). L'ordre d'application des défauts est le suivant : Variation de l'oxyde, variation de la température du liquide de refroidissement du réacteur, variation de la température du liquide de refroidissement du condenseur et variation de la vitesse de réaction (variation du facteur de fréquence dans la loi d'Arrhenius). Les paramètres utilisés pour effectuer cet auto-apprentissage sont pratiquement identiques à ceux utilisés lors de la première classification effectuée avec tous les descripteurs : niveau d'exigence de 0.87, variabilité maximale de 0.1, la fonction d'appartenance *Binomiale* et le connectif *Max-Min*. Seul le niveau d'exigence a été augmenté par rapport à la première classification, ce qui a pour effet de permettre d'être un peu plus sévère sur le nombre de classes créées.

On peut constater que la nouvelle classification permet de détecter les 8 défauts de manière satisfaisante. Le Tableau 4-J montre la correspondance des classes aux défauts.

Tableau 4-J Correspondance des défauts et numéro des classes issues de l'auto-apprentissage avec 6 descripteurs

NOM DÉFAUT	NUMÉRO DE CLASSE
Hausse d'oxyde	2
Baisse d'oxyde	4
Hausse de la température du fluide de refroidissement du réacteur	9
Baisse de la température du fluide de refroidissement du réacteur	8
Hausse de température du fluide de refroidissement du condenseur	12
Baisse de température du fluide de refroidissement du condenseur	1
Hausse de la vitesse de réaction	14
Baisse de la vitesse de réaction	16

Les transitions de l'opération normale à une situation de défaillance sont bien représentées par les classes 17, 18, 19 et 22, grâce à l'application graduelle des défauts.

Il convient dans l'étape suivante de synthétiser les résultats de cette classification de telle sorte à définir des états fonctionnels à partir des différentes classes créées. Un état fonctionnel n'est qu'un regroupement de classes qui représentent un état spécifique de fonctionnement du

procédé (soit un état normal ou un défaut particulier). C'est l'expert qui effectue l'assignation des états aux classes. Ce travail s'accompagne généralement d'une analyse implicite du fonctionnement du procédé, elle permet aussi de mieux comprendre l'évolution du procédé. On peut par exemple regrouper dans un même état des classes qui sont liées à un phénomène identique.

L'ensemble des correspondances état/classe est donné dans le Tableau 4-K. On peut observer que l'état correspondant au fonctionnement normal du procédé est composé des classes 13, 21, 24, 28 et 31. Ces classes ont été regroupées sous le même état numéro 10 et cette information sera celle présentée à l'opérateur lors de la phase de reconnaissance (diagnostic en ligne).

Les états d'alarme sont très importants pour la détection anticipée des défaillances. On peut voir dans le tableau du modèle de comportement que les classes 17 et 22 constituent l'état d'alarme de hausse de la température du fluide de refroidissement dans le réacteur.

Quelques classes correspondent à l'oscillation des signaux comme conséquence de la récupération brutale des défaillances, par les régulateurs par exemple. Ce type des classes n'a pas été affecté à un état de fonctionnement, à cause de son instabilité comme c'est le cas des classes 23, 25, 29 et 30.

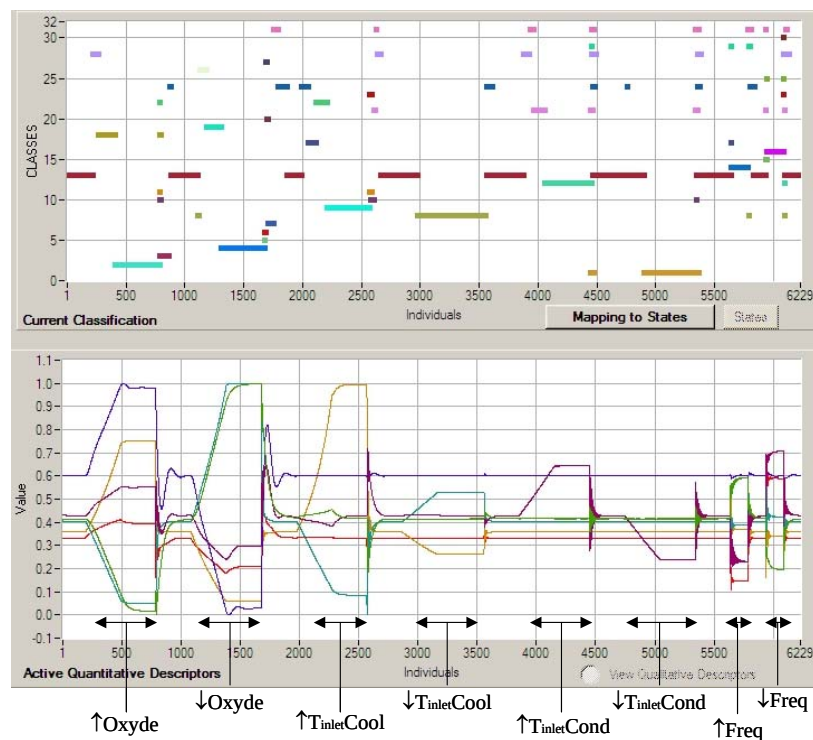


Figure 4-6 Identification de défauts avec les 6 descripteurs sélectionnés

Tableau 4-K Assignment des classes aux états à partir des 6 descripteurs sélectionnés

Numéro de classe	Numéro d'état	Etat Fonctionnel	Description d'état	Numéro de classe	Numéro d'état	Etat Fonctionnel	Description d'état
1	1	B-TCOND	Baisse de la température du fluide de refroid. du Condenseur	16	12	B-FREQV	Baisse facteur fréquence (vitesse de réaction)
2	2	H-OXYDE	Hausse du débit d'oxyde	17	13	AL_H-TCOOL	Alarme de Hausse de la Température du fluide de refroid. du réacteur
3	3	REC_H-OXY_N	Récupération de la hausse du débit d'oxyde	18	14	AL_H-OXYDE	Alarme de hausse du débit d'oxyde
4	4	B-OXYDE	Baisse du débit d'oxyde	19	15	AL_B-OXYDE	Alarme de baisse du débit d'oxyde
5	5	REC_B-OXY_N	Récupération de la baisse du débit d'oxyde	20	5	REC_B-OXY_N	Récupération de baisse d'oxyde vers état Normal
6	5	REC_B-OXY_N	Récupération de la baisse du débit d'oxyde	21	10	NORMAL	Opération Normale du procédé
7	5	REC_B-OXY_N	Récupération de la baisse du débit d'oxyde	22	13	AL_H-TCOOL	Alarme de Hausse de la Température du fluide de refroid. du réacteur
8	6	B-TCOOL	Baisse de la Température du fluide de refroid. du réacteur	23		NoState	pas d'assignation d'état
9	7	H-TCOOL	Hausse de la Température du fluide de refroid. du réacteur	24	10	NORMAL	Opération Normale du procédé
10	8	REC-HAUSSE	Récupération dans l'état Normal	25		NoState	pas d'assignation d'état
11	8	REC-HAUSSE	Récupération dans l'état Normal	26	15	AL_B-OXYDE	Alarme de baisse du débit d'oxyde
12	9	H-TCOND	Hausse de la température du fluide de refroid. du Condenseur	27	5	REC_B-OXY_N	Récupération de baisse du débit d'oxyde
13	10	NORMAL	Opération Normale du procédé	28	10	NORMAL	Opération Normale du procédé
14	11	H-FREQV	Hausse facteur fréquence (vitesse de réaction)	29		NoState	pas d'assignation d'état
15	11	H-FREQV	Hausse facteur fréquence (vitesse de réaction)	30		NoState	pas d'assignation d'état
				31	10	NORMAL	Opération Normale du procédé

La Figure 4-7 montre les différents états fonctionnels détectés en utilisant les 6 descripteurs sélectionnés dans le cas de la simulation des différents défauts présentés précédemment. On peut constater la détection correcte de tous les défauts, ainsi que, la détection des différents états d'alarmes.

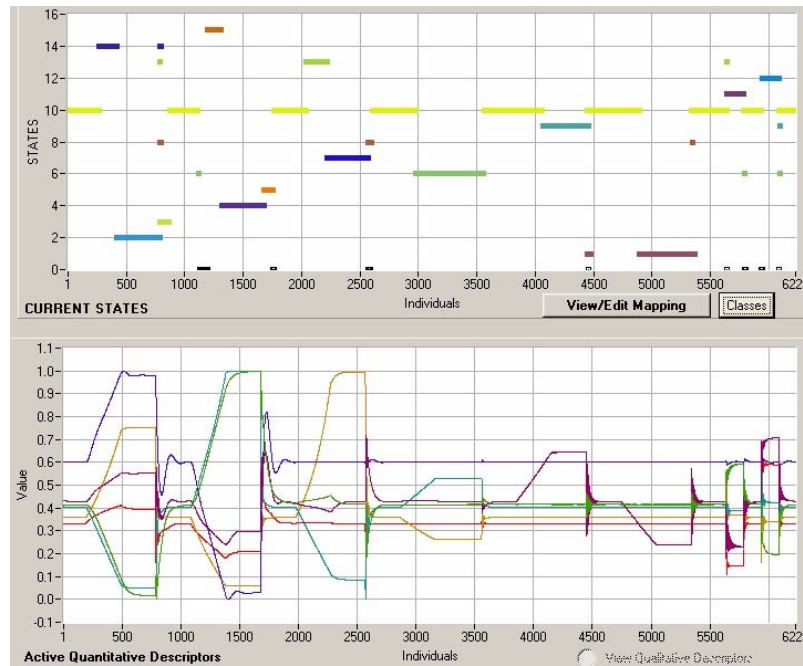


Figure 4-7 Identification des 15 états fonctionnels avec 6 descripteurs

4.2.4 Reconnaissances de défauts inconnus

Il est évident que l'étape du suivi repose sur le résultat de la phase de modélisation du comportement (section 4.2.3) du procédé chimique, ce qui pose le problème de la validité du modèle de comportement. Il est probable que toutes les classes ne soient pas connues a priori, l'expert peut avoir une certaine connaissance sur quelques situations présentes dans l'ensemble d'apprentissage, mais il peut s'avérer difficile d'en identifier d'autres. Il est possible que les classes ne soient pas correctement représentées par les observations disponibles initialement; et il se peut aussi que, pour certaines situations, l'acquisition de données soit difficile ; il s'agit, souvent, de situations critiques qui ne se produisent pas fréquemment. Dans ces conditions, rien ne permet d'assurer l'exhaustivité de la connaissance du comportement du processus car les données historiques ne peuvent pas couvrir la vie entière du processus [Kempowsky 2004a].

A) Défauts inconnus continus

Dans cette dernière partie, nous avons fait la reconnaissance de diverses situations de défauts inconnus (magnitude et dérive différentes, simultanéité de défauts) afin de vérifier et valider notre modèle comportemental.

D'abord nous avons fait la simulation de deux défauts de façon continue avec amplitude et dérivées différentes de celles utilisées pour construire le modèle de comportement. Ces défauts sont : la hausse de l'oxyde de propylène à 74 Kgmole/h et la hausse de la température du fluide de refroidissement du réacteur à 18 °C. La Figure 4-8 illustre le scénario simulé, avec les deux défauts consécutifs.

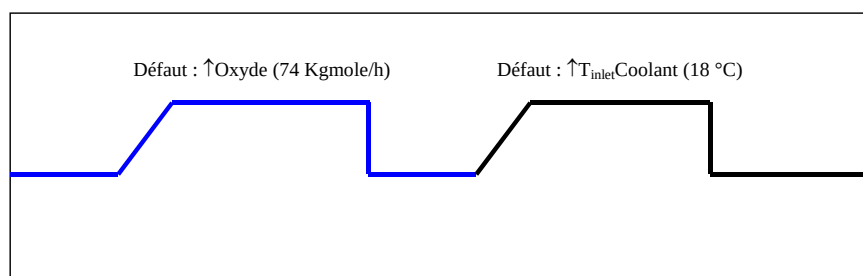


Figure 4-8 Scénario de deux défauts consécutifs critiques

La Figure 4-9 montre le résultat de la reconnaissance en utilisant les 6 descripteurs sélectionnés et le modèle des états fonctionnels construit (voir tableau précédent). On peut constater qu'il y a création de 6 états et l'état zéro (état de non-reconnaissance). En regardant le tableau du modèle de comportement, nous pouvons vérifier que l'état 10 correspond à l'opération normale du procédé, l'état 14 correspond à l'alarme de hausse d'oxyde, l'état 13 correspond à la hausse de la température du fluide de refroidissement du réacteur, l'état 3 correspond à la récupération de la hausse d'oxyde et les états 8 et 9 sont des transitoires liés à la phase de mise en régime permanent du procédé. On peut observer sur cette figure, que les états d'alarme des deux défaillances sont bien identifiés, ce qui confère plus de sécurité au processus.

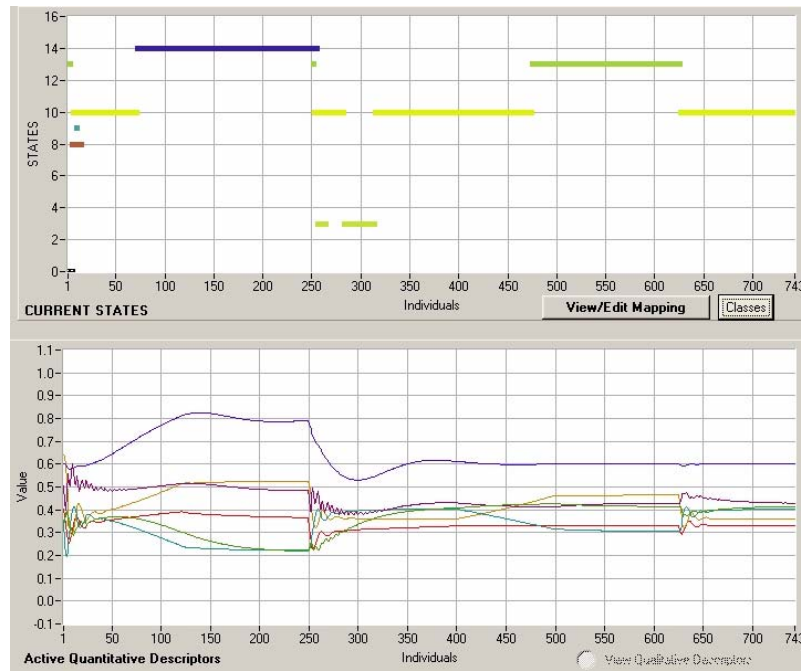


Figure 4-9 Reconnaissance avec 6 descripteurs des défauts consécutifs : hausse d'oxyde et hausse de la température du fluide de refroidissement du réacteur

Dans ce qui suit, nous avons fait la simulation de deux défauts consécutifs moins critiques par rapport aux précédents avec une amplitude et des dérives différentes de celles utilisées lors de la conception du modèle de comportement. Ces défauts sont : une baisse de l'oxyde de propylène à 62 Kgmole/h et une baisse de la température du fluide de refroidissement du réacteur à 13 °C. Le scénario simulé est illustré sur la Figure 4-10, avec les deux défauts consécutifs.

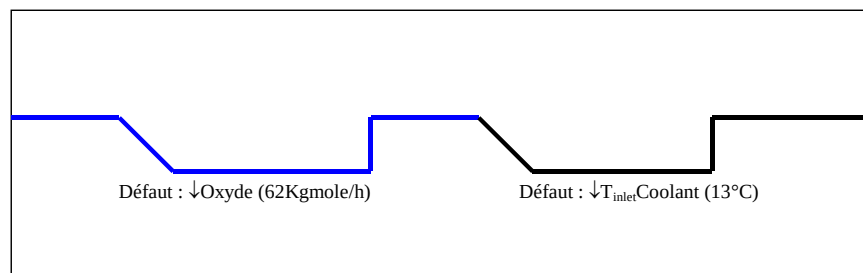


Figure 4-10 Scénario de deux défauts consécutifs non critiques

Sur la Figure 4-11, on montre le résultat de la reconnaissance en utilisant les 6 descripteurs sélectionnés et le modèle des états fonctionnels. On peut constater qu'il y a création de 6 états et l'état zéro (état de non-reconnaissance). L'état 10 correspond à l'opération normale du

procédé, l'état 15 correspond à l'alarme de la baisse d'oxyde, l'état 6 correspond à la baisse de la température du fluide de refroidissement du réacteur, et les états 13, 9 et 8 sont des transitoires. On peut observer sur cette figure, que l'état 6 est répété lors de la phase de détection de l'alarme de baisse d'oxyde, mais c'est un problème sans importance puisque le nombre d'échantillons que contient cet état est très petit (3 éléments). En revanche, un problème plus grave est que l'état d'alarme de la baisse d'oxyde (état 15) est reconnu très tardivement et, en plus, une grande partie de ce défaut est restée inconnue (état zéro).

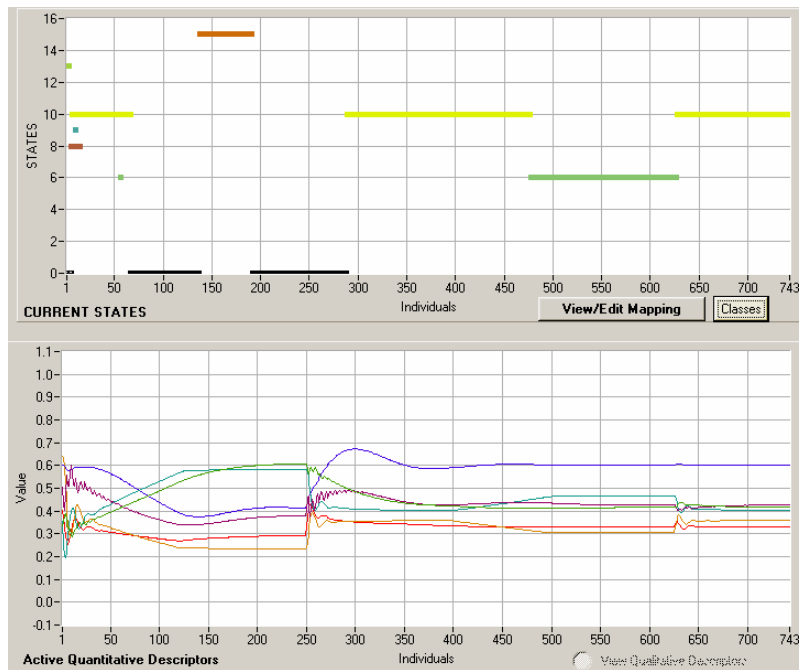


Figure 4-11 Reconnaissance avec 6 descripteurs des défauts consécutifs : baisse d'oxyde et baisse de la température du fluide de refroidissement du réacteur

La solution à ce problème est l'ajout d'un descripteur pour améliorer la détection des états fonctionnels étudiés. Ce descripteur peut être ajouté en utilisant une nouvelle analyse de l'entropie.

Cette analyse est similaire à celle effectuée lors des étapes présentées dans la section 4.2.2 seule différence est que, dans cette méthode, on a pris l'ensemble des défauts en même temps pour les calculs de l'entropie (voir la section 3.3.3). D'abord, on obtient le profil des classes à partir de la classification. Chaque ligne représente une classe (soit l'opération normale ou soit un défaut) et chaque colonne un descripteur.

Tableau 4-L Profil des classes de la classification à partir de 6 descripteurs

	d1	d2	d3	d4	d5	d6	d7	d8	d9	d10	d11	d12	d13	d14	d15	d16	d17	d18	d19	d20	d21
Normal	0.33	0.59	0.40	0.36	0.60	0.40	0.56	0.55	0.55	0.55	0.56	0.56	0.55	0.56	0.54	0.44	0.38	0.43	0.56	0.41	0.52
Défaut	0.83	0.60	0.40	0.33	0.60	0.45	0.36	0.43	0.42	0.42	0.47	0.51	0.51	0.51	0.54	0.48	0.43	0.97	0.56	0.02	0.54
Défaut	0.05	0.60	0.40	0.38	0.60	0.38	0.69	0.62	0.63	0.63	0.62	0.59	0.58	0.58	0.54	0.41	0.34	0.12	0.56	0.69	0.51
Défaut	0.33	0.59	0.40	0.36	0.60	0.40	0.57	0.56	0.55	0.55	0.56	0.56	0.55	0.56	0.54	0.44	0.37	0.26	0.56	0.41	0.52
Défaut	0.33	0.60	0.40	0.36	0.60	0.40	0.57	0.56	0.55	0.55	0.56	0.56	0.56	0.56	0.54	0.44	0.37	0.69	0.56	0.41	0.52
Défaut	0.33	0.59	0.40	0.27	0.60	0.52	0.57	0.56	0.55	0.55	0.56	0.56	0.55	0.56	0.54	0.44	0.37	0.43	0.56	0.41	0.52
Défaut	0.33	0.61	0.40	0.92	0.60	0.11	0.57	0.56	0.55	0.55	0.57	0.56	0.56	0.56	0.54	0.44	0.37	0.42	0.56	0.42	0.52
Défaut	0.21	0.59	0.38	0.08	0.07	0.94	0.78	0.67	0.67	0.59	0.14	0.13	0.13	0.12	0.52	0.86	0.75	0.29	0.56	0.93	0.72
Défaut	0.39	0.60	0.41	0.73	0.96	0.07	0.38	0.44	0.45	0.50	0.92	0.92	0.92	0.93	0.55	0.09	0.10	0.55	0.56	0.06	0.35

A partir de ce profil, on normalise (pour obtenir un espace probabilisé), on calcule l'entropie, le gain d'information du système en utilisant la valeur maximale de l'entropie dans le système ($\log 9 = 2.197$), et finalement, le gain relatif en pourcentage. Le résultat final des calculs est le suivant :

Tableau 4-M Valeurs du gain relatif

(6.82 0.00 0.01 7.06 4.11 7.86 1.15 0.40 0.41 0.24 3.04 3.14 3.15 3.31 0.00 4.22 3.78 5.87 0.00 10.64 0.63)

La Figure 4-12 montre le résultat sous forme graphique. Le résultat principal de ce critère de sélection est qu'il convient d'ajouter le descripteur de la température sur le plateau 10 de la colonne de distillation (10_MainTS), puisque la valeur du gain de ce descripteur dans ce critère est plus grande que la valeur du descripteur pourcentage d'ouverture de la vanne pour contrôler le niveau de liquide dans le réacteur ($ReactorLC_OP$), qui a été choisi grâce à l'autre critère développé dans la section 4.2.2.

Ce rajout d'une variable supplémentaire est cohérent puisqu'il permet d'avoir une information supplémentaire sur la composition du mélange au bouilleur de la colonne de rectification, reflet d'une variation de la composition à l'alimentation de cette colonne, qui elle-même peut être consécutive à une baisse de l'oxyde à l'alimentation du procédé.

A partir de cette constatation, nous avons 7 descripteurs pour la reconnaissance des diverses situations inconnues.

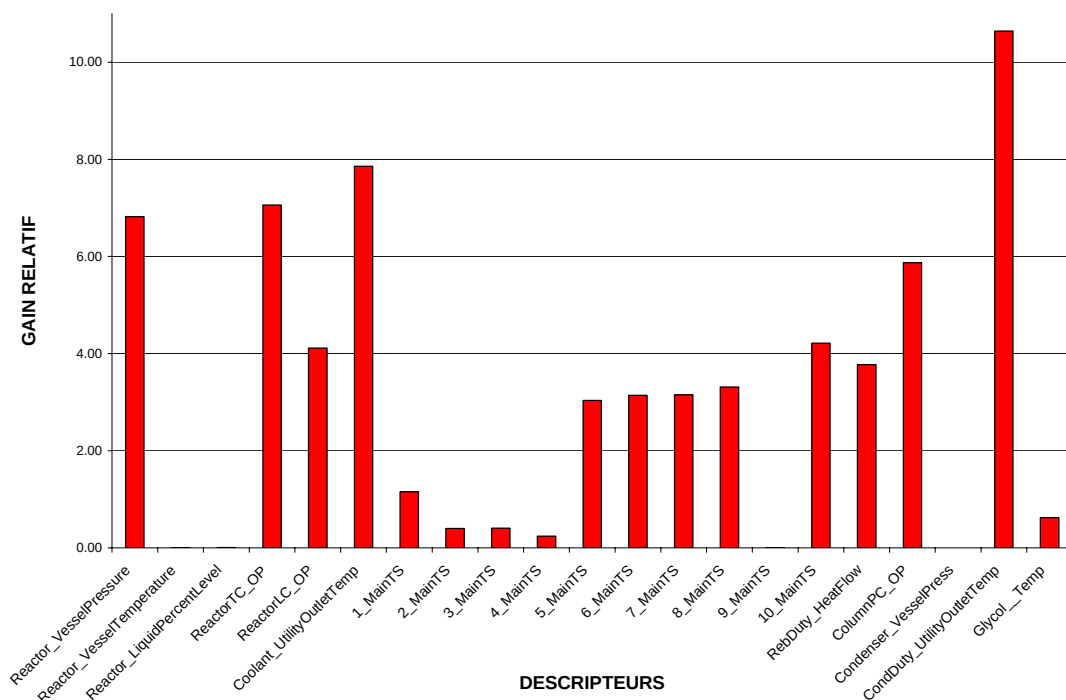
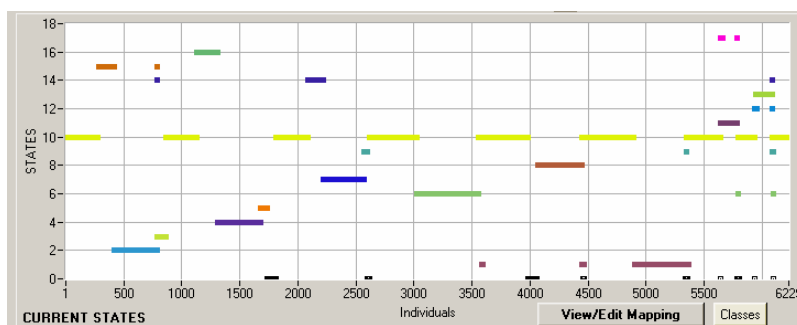
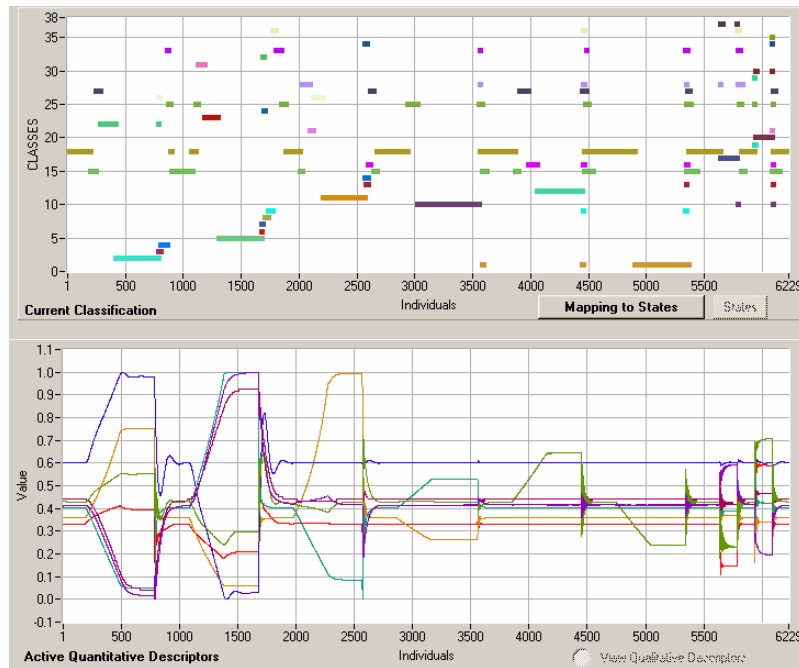


Figure 4-12 Gain relatif pur l'ensemble des situations

La Figure 4-13a montre les résultats obtenus après l'auto-apprentissage des 8 défauts, en utilisant les 7 descripteurs sélectionnés. On peut constater que la nouvelle classification peut détecter les 8 défauts de manière satisfaisante. Les paramètres utilisés pour cette classification sont : un niveau d'exigence de 0.9, une variabilité maximale de 0.1, la fonction d'appartenance *Binomiale* et le connectif *Min-Max*.



b)



a)

Figure 4-13 Identification des défauts et états fonctionnels avec les 7 descripteurs sélectionnés

L'étape suivante est la définition des états fonctionnels représentatifs en agrégeant les classes créées telles que représentées sur la Figure 4-13b. L'assignation des classes aux états a été effectuée de la façon suivante (Tableau 4-N).

On peut observer que l'état d'opération normale du procédé est composé des classes 15 et 18.

En comparant ce modèle avec le modèle de comportement avec 6 descripteurs (Tableau 4-K), nous pouvons vérifier l'existence d'un plus petit nombre de classes (37 classes) et d'états (17 états) dans ce modèle. En plus, on peut voir dans ce nouveau modèle, la disparition de l'état REC-HAUSSE, et la création de trois nouveaux états : REC_H-TCOOL, AL_B-FREQV et AL-H-FREQV. En général, nous pouvons dire que ce dernier modèle est meilleur que le précédent, puisqu'il permet de mieux discriminer les défaillances et de donner plus d'états d'alarmes.

Cette étape de la procédure a donc permis de créer un modèle de comportement du procédé basé sur la classification des 7 descripteurs sélectionnés.

Tableau 4-N Modèle de comportement avec 7 descripteurs

Numéro de classe	Numéro d'état	Etat fonctionnel	Description d'état	Numéro de classe	Numéro d'état	Etat fonctionnel	Description d'état
1	1	B-TCOND	Baisse de la température du fluide de refroid. du Condenseur	15, 18, 25, 27, 28, 33	10	NORMAL	Opération Normale du procédé
2	2	H-OXYDE	Hausse du débit d'oxyde	17	11	H-FREQV	Hausse facteur fréquence (vitesse de réaction)
3, 4	3	REC_H-OXY_N	Récupération de la hausse du débit d'oxyde	19, 31, 35	12	AL_B-FREQV	Alarme de Baisse facteur fréquence (vitesse de réaction)
5	4	B-OXYDE	Baisse du débit d'oxyde	20	13	B-FREQV	Baisse facteur fréquence (vitesse de réaction)
6, 7, 8, 24, 32	5	REC_B-OXY_N	Récupération de la baisse du débit d'oxyde	21, 26	14	AL_H-TCOOL	Alarme de Hausse de Température du fluide de refroidissement réacteur
10	6	B-TCOOL	Baisse de la Température du fluide de refroid. du réacteur	22	15	AL_H-OXYDE	Alarme de hausse de débit d'oxyde
11	7	H-TCOOL	Hausse de la Température du fluide de refroid. Du réacteur	23, 30	16	AL_B-OXYDE	Alarme de baisse du débit d'oxyde
12	8	H-TCOND	Hausse de la température du fluide de refroid. du Condenseur	37	17	AL_H-FREQV	Alarme de Hausse facteur fréquence (vitesse de réaction)
13, 14, 34	9	REC_H-TCOOL	Récupération de la hausse de la Température du fluide de refroid. Du réacteur	9, 16, 29, 36	0	NoState	pas d'assignation d'état

La Figure 4-13b montre les résultats de l'identification des états fonctionnels du procédé en utilisant les 7 descripteurs sélectionnés. On peut constater que tous les défauts, ainsi que, les états d'alarmes sont correctement détectés.

Rappelons que l'ajout d'un descripteur supplémentaire avait pour but d'améliorer l'identification d'états inconnus. Reprenons le scénario simulé précédemment incluant les deux défauts méconnus (Figure 4-10). En effectuant la classification en mode « reconnaissance » avec cette fois les 7 descripteurs, on obtient les résultats représentés sur la Figure 4-14. On constate que, même si l'on retrouve une toute petite période au tout début de l'apparition de la baisse d'oxyde où il n'y a pas reconnaissance immédiate du défaut (on rappelle que lors de l'apprentissage le défaut simulé avait une amplitude et une dérivée beaucoup plus importante), grâce à l'utilisation du descripteur supplémentaire, on a pu détecter ce défaut beaucoup plus tôt et seul le tout début n'est pas reconnu contrairement à ce qui avait été obtenu avec les 6 descripteurs.

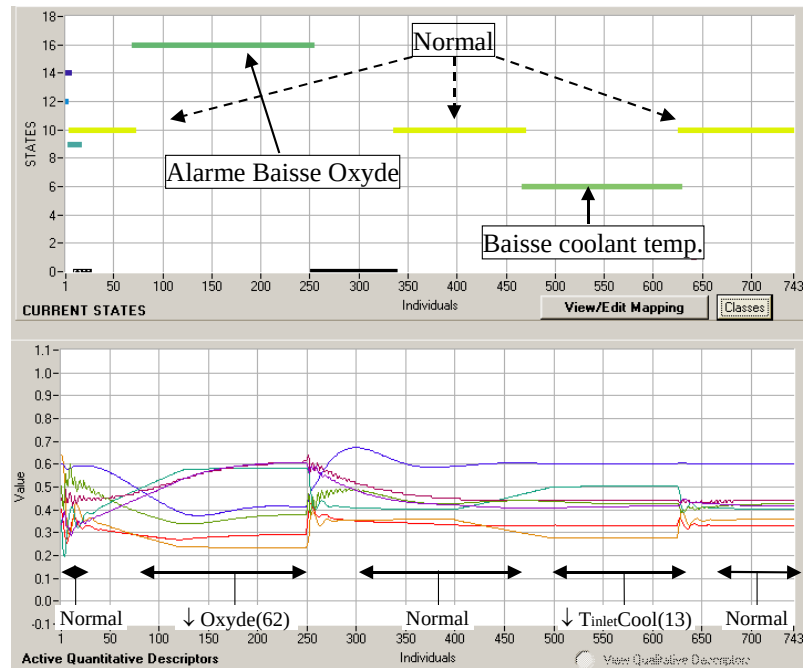


Figure 4-14 Reconnaissance avec 7 descripteurs des défauts consécutifs : baisse d'oxyde et baisse de la température du fluide de refroidissement du réacteur

Ce paragraphe a permis de montrer que dans le cas où un des défauts est mal reconnu, il est possible d'améliorer sa détection en rajoutant un descripteur supplémentaire en regardant ceux qui donnent les gains d'information relatifs les plus élevés non plus uniquement dans le cas où l'on considère la détection d'un défaut particulier par rapport à l'état de fonctionnement normal mais en considérant l'ensemble des défauts.

B) Défauts inconnus simultanés

L'ultime validation de la procédure développée concerne la détection de défauts simultanés. Il est clair qu'en pratique cette parfaite simultanéité est peu probable, néanmoins, il est intéressant d'observer le comportement de l'outil de diagnostic face à cette situation inconnue. Le scénario utilisé pour ce test est la superposition de deux défauts :

- i) Hausse de la température du liquide de refroidissement du réacteur : la valeur de la pente d'incrément de température est de $0.035\text{ }^{\circ}\text{C}/\text{min}$ faisant passer la température de sa valeur nominale ($15\text{ }^{\circ}\text{C}$) jusqu'à la valeur finale de la défaillance ($22\text{ }^{\circ}\text{C}$).

- ii) Hausse de la vitesse de réaction : ce défaut est engendré en modifiant le facteur de fréquence de la loi d'Arrhenius. Le facteur de fréquence est modifié depuis sa valeur nominale ($1.70E+13$) et passe à $2.00E+13$.

Il est important de mentionner que ces défauts ont des valeurs et s'établissent selon des dérives différentes de celles des scénarii étudiés lors de l'apprentissage qui a permis de développer le modèle d'états fonctionnels. L'objectif principal est de retrouver à travers le modèle d'états construit précédemment les différents défauts isolés, ainsi que les états d'alarme de ces défauts et de vérifier la non-reconnaissance de défauts simultanés (états de fonctionnement non modélisés) puisqu'ils n'ont jamais été « appris ».

Au fur et à mesure de la présentation de nouvelles observations, des nouvelles situations peuvent apparaître dans la structure initiale. De même, les classes existantes peuvent aussi évoluer au cours du temps. Pour ces raisons, il est nécessaire que le système de surveillance présente un caractère **adaptable**, non seulement au moment de la construction du modèle de comportement (phase d'apprentissage) mais, aussi, au moment de l'identification des nouvelles situations. Pour ceci, deux principes d'apprentissage sont prévus : un apprentissage hors ligne et un apprentissage en ligne. L'apprentissage hors ligne est utilisé dans l'étape de caractérisations de situations, mais aussi pour essayer de caractériser des nouvelles situations qui peuvent se présenter lorsqu'un nombre suffisant d'observations n'est plus reconnu, compte tenu des classes existantes. L'apprentissage en ligne permet de modifier en continu les classes existantes et, même, d'en créer des nouvelles, tout au long de l'étape du suivi.

Dans ce qui suit, nous donnons le résultat obtenu lors de la reconnaissance d'un scénario de simultanéité de défauts, à partir du modèle de comportement avec 7 descripteurs:

Scénario : Dans ce cas le défaut : hausse de la vitesse de réaction, est créé exactement en même temps que le défaut : hausse de la température du liquide de refroidissement du réacteur, mais il est supprimé plus tôt (voir la Figure 4-15)

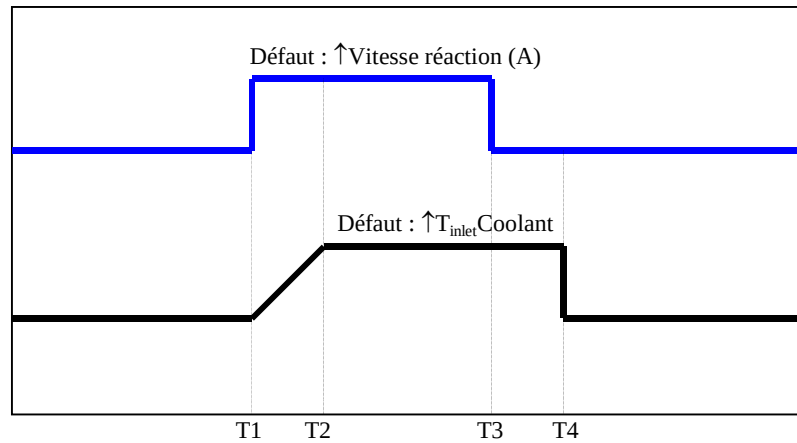


Figure 4-15 Scénario : défauts simultanés

On peut vérifier sur la Figure 4-16 qu'il y a création de 4 états et de l'état zéro (état de non-reconnaissance). L'état 10 représente l'opération normale du procédé, l'état 14 correspond au défaut d'alarme de hausse de la température du fluide de refroidissement du réacteur, les états 11 et 17 correspondent au défaut de hausse de la vitesse de réaction et à l'alarme de ce même défaut, respectivement (voir le modèle de comportement du Tableau 4-N).

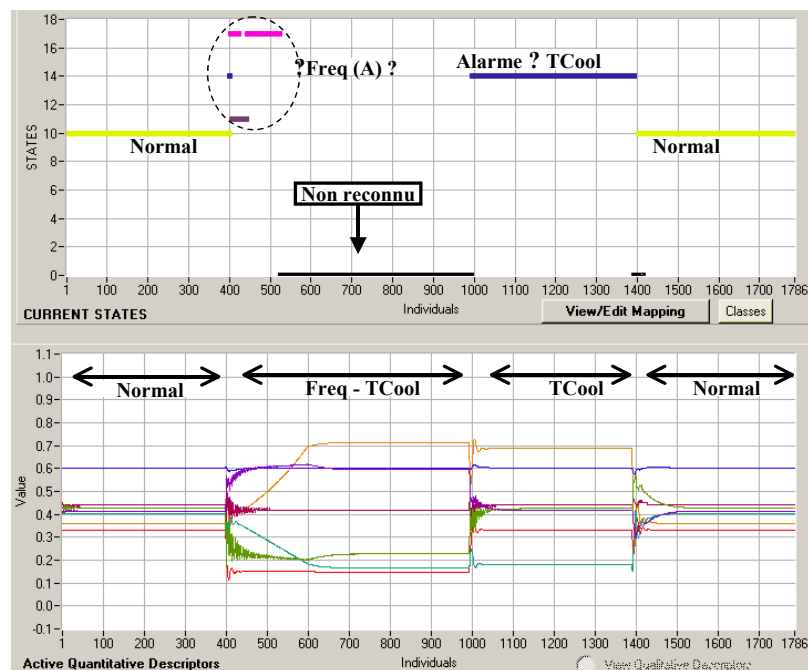


Figure 4-16 Identification d'états de défauts simultanés : scénario 1

L'état zéro représente les défauts simultanés et, il n'est pas reconnu, puisqu'il n'a pas été considéré pendant la construction du modèle de comportement. Pour le reconnaître, il faut

faire un apprentissage supervisé (*supervised learning*) pour affecter la classe NIC (*Non Information Classe*) à une classe représentant ces défauts simultanés (voir la Figure 4-17).

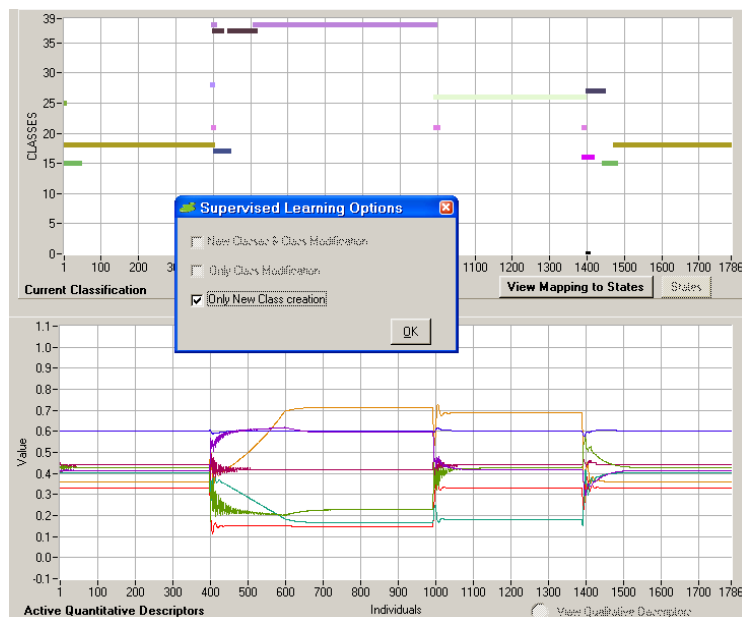


Figure 4-17 Application de l'apprentissage supervisé

Finalement, après cet apprentissage on peut faire la reconnaissance des états, et on peut constater la bonne reconnaissance de l'état numéro 18 (c'est un nouvel état créé) pour représenter la nouvelle défaillance (Figure 4-18). Nous pouvons constater dans le Tableau 4-N que la classe 18 n'existait pas dans le modèle construit préalablement.

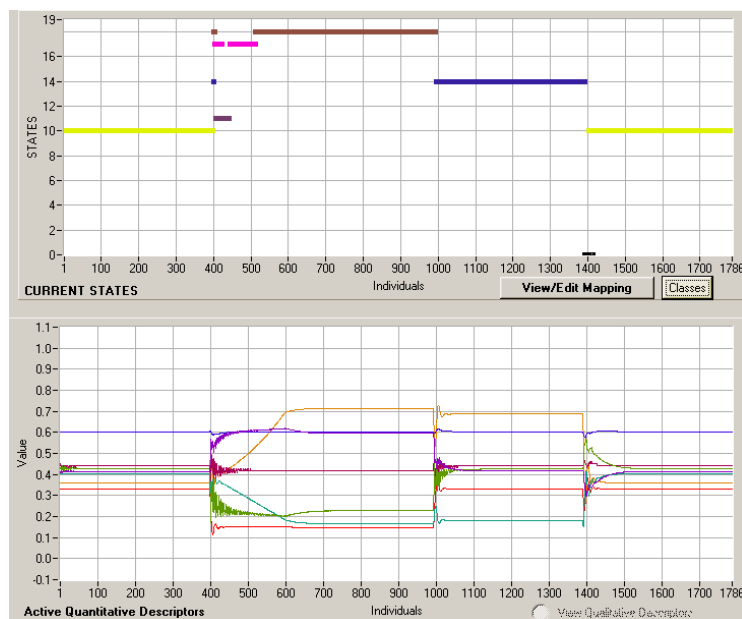


Figure 4-18 Reconnaissance des états de défauts simultanés après apprentissage supervisé

L'annexe D montre les résultats des autres scénarii de défauts simultanés étudiés.

4.2.5 Utilisation des autres critères

Comme il a été vu au chapitre 3, il existe deux autres critères pour mesurer l'information : l'entropie non probabiliste et la variance. Sans reprendre toutes les étapes du processus permettant de déterminer les capteurs à partir de la classification réalisée sur l'ensemble des capteurs possibles, nous donnons dans le Tableau 4-O les résultats en appliquant ces deux critères dans le cas de la procédure dite par paire. Comme nous le voyons sur ce tableau, les capteurs sélectionnés par le critère de la variance sont parfaitement identiques à ceux sélectionnés par le critère de l'entropie probabiliste. Ce résultat conforte ceux obtenus précédemment. En revanche, le critère de l'entropie non probabiliste ne donne pas sensiblement les mêmes résultats. Si l'on regarde l'ensemble final des capteurs sélectionnés, c'est le même que pour les deux autres critères : au total on a les mêmes 6 descripteurs. Ce qui différencie les résultats obtenus avec ce critère des autres est la correspondance capteur sélectionné/défaut. Toutes les permutations (sauf une) sont physiquement explicables : par exemple utiliser la valeur de l'ouverture de la vanne de régulation de la température dans le réacteur comme capteur pour un défaut de hausse de la température du fluide de refroidissement paraît correct, d'ailleurs c'est ce qui a été choisi pour la baisse de la température par tous les critères. Le seul cas sujet à discussion est celui de défaut de hausse de la température du fluide de refroidissement au condenseur qui serait mieux identifié avec la pression dans le réacteur comme capteur ; ce qui paraît difficile à expliquer physiquement puisque la colonne de distillation est en aval du réacteur.

De ces résultats et de ceux obtenus sur l'autre procédé (l'OPR), il ressort que les critères Entropie probabiliste et Variance donnent les mêmes résultats. En revanche, certaines différences peuvent apparaître avec le critère d'*entropie non-probabiliste* en particulier sur des couplages capteur/défaut difficilement explicables. C'est pourquoi dans tout ce qui suit, nous donnerons uniquement les résultats obtenus avec le critère de l'Entropie probabiliste.

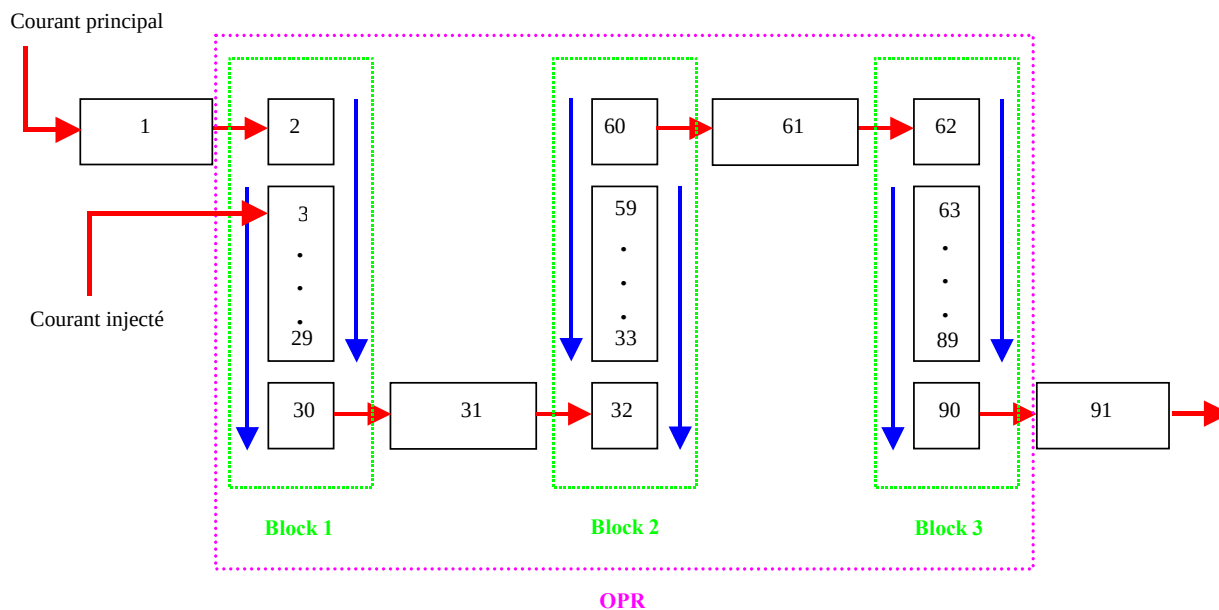
Tableau 4-O Capteurs sélectionnés suivant les trois critères

NOM DÉFAUT	DESCRIPTEURS SÉLECTIONNÉS		
	PROBABILISTE	VARIANCE	NON-PROBABILISTE
Baisse A (Freq Naturelle)	CondDuty_UtilityOutletTemp	CondDuty_UtilityOutletTemp	CondDuty_UtilityOutletTemp
Hausse A (Freq Naturelle)	Reactor_VesselPressure	Reactor_VesselPressure	Reactor_VesselPressure
Baisse TempInletCond	ColumnPC_OP	ColumnPC_OP	ColumnPC_OP
Hausse TempInletCond	ColumnPC_OP	ColumnPC_OP	Reactor_VesselPressure
Baisse TempInletCool	ReactorTC_OP	ReactorTC_OP	ReactorTC_OP
Hausse TempInletCool	Coolant_UtilityOutletTemp	Coolant_UtilityOutletTemp	ReactorTC_OP
Baisse Oxyde	ReactorLC_OP	ReactorLC_OP	Coolant_UtilityOutletTemp
Hausse Oxyde	CondDuty_UtilityOutletTemp	CondDuty_UtilityOutletTemp	ReactorLC_OP

4.3 Application de la méthodologie sur le nouveau procédé OPR

Les synthèses des produits chimiques ou pharmaceutiques, largement effectuées dans des réacteurs discontinus ou semi-continus, sont souvent fortement limitées par des contraintes liées à la dissipation de la chaleur produite par les réactions. Un nouveau concept de réacteurs d'échange thermique, « *OPEN PLATE REACTOR (OPR)* », développé par *ALFA LAVAL AB* permet la mise en oeuvre de réactions chimiques complexes avec une commande thermique très précise [Haugwitz et al., 2004 ; Haugwitz et al., 2005].

L'« *OPR* » consiste en 3 blocs avec différentes configurations thermiques. Le réacteur *OPR* est formé d'un nombre de plaques dans lesquelles les réactifs sont mélangés et la réaction se produit. De chaque côté de la plaque où s'effectue la réaction, il y a une plaque de refroidissement, à travers laquelle circule de l'eau froide. La Figure 4-19 montre un schéma de l'*OPR* avec les numéros des cellules de modélisation.



X_MAIN représente la fraction molaire de chaque composé dans le fluide principal. X_INJ celle dans le fluide secondaire appelé aussi fluide injecté. Attention, ces valeurs sont données uniquement pour information, elles ne constituent pas des descripteurs car elles ne peuvent pas être mesurées en-ligne sur le réacteur.

Les données acquises pour l'analyse de l'entropie et du diagnostic ont été obtenues à partir d'un simulateur développé en Fortran (*Data Processing Tool*) par Sébastien ELGUE (*Alfa Laval Vicarb*). La formulation dynamique du modèle mène à un système d'équations différentielles et algébriques (*DAE*). La résolution de ce système est obtenue au moyen d'un solveur d'équations algébriques et différentielles : DISCo [Elgue et al., 2004 ; Sargousse et al., 1999].

L'OPR peut-être amplement équipé de capteurs. Expérimentalement, il y a plusieurs capteurs de température internes pour donner une information sur l'état actuel de la réaction. Grâce au simulateur dédié, une étude exhaustive sur ces capteurs a pu être possible. Le Tableau 4-P montre les différents variables susceptibles d'être utilisées comme mesure. Parmi ces variables, seules celles qui sont facilement mesurables en ligne ont été retenues : toutes les concentrations (12) n'ont été considérées. De même toutes les variables sur lesquelles ont été effectuées les perturbations n'ont pas été gardées puisque l'identification d'une perturbation sur ces variables devenait triviale. Ces variables sont les températures d'alimentation (courant principal et injecté) les débits des deux courants d'alimentation, le débit et la température

d'alimentation de l'utilité : soit 6 variables. Il y a donc 27 capteurs possibles qui constituent l'ensemble initial sur lequel portera la recherche.

La conception et la validation du modèle de comportement pour le réacteur OPR ont été étudiées sur deux réactions : la réaction du thiosulfate et la réaction d'estérification (voir la section 2.5). Ces deux réactions présentent deux cas extrêmes, la première est une réaction hautement exothermique tandis que, la réaction d'estérification est très lente et faiblement exothermique.

Tableau 4-P Description des variables : modèle OPR

NOM VARIABLE	DESCRIPTION	No. DE CELLULE DE CALCUL
TP_IN_1	temp. alim courant principal	1
TP_IN_2	temp. alim courant injecté	3
TP_B1_1	temp. fluide procédé bloc 1 cellule 1	3
TP_B1_2	temp. fluide procédé bloc 1 cellule 2	4
TP_B1_3	temp. fluide procédé bloc 1 cellule 3	5
TP_B1_4	temp. fluide procédé bloc 1 cellule 4	6
TP_B1_5	temp. fluide procédé bloc 1 cellule 5	7
TP_B1_10	temp. fluide procédé bloc 1 cellule 10	8
TP_B1_15	temp. fluide procédé bloc 1 cellule 15	17
TP_B1_20	temp. fluide procédé bloc 1 cellule 20	22
TP_B1_27	temp. fluide procédé bloc 1 cellule 27	29
TP_B2_1	temp. fluide procédé bloc 2 cellule 1	33
TP_B2_5	temp. fluide procédé bloc 2 cellule 5	37
TP_B2_10	temp. fluide procédé bloc 2 cellule 10	42
TP_B2_15	temp. fluide procédé bloc 2 cellule 15	47
TP_B2_20	temp. fluide procédé bloc 2 cellule 20	52
TP_B2_27	temp. fluide procédé bloc 2 cellule 27	59
TP_B3_1	temp. fluide procédé bloc 3 cellule 1	63
TP_B3_5	temp. fluide procédé bloc 3 cellule 5	67
TP_B3_10	temp. fluide procédé bloc 3 cellule 10	72
TP_B3_15	temp. fluide procédé bloc 3 cellule 15	77
TP_B3_20	temp. fluide procédé bloc 3 cellule 20	82
TP_B3_27	temp. fluide procédé bloc 3 cellule 27	89
FP_IN_1	débit alimentation courant principal	1
FP_IN_2	débit alimentation courant injecté	3
FP_OUT	débit fluide procédé en sortie	91
PP_IN_1	pression alimentation courant principal	1
PP_IN_2	pression alimentation courant injecté	3
PP_OUT	pression fluide procédé en sortie	91
FU_IN	débit alimentation fluide utilité	2
TU_B1_IN	temp. alimentation fluide utilité	2
TU_B1_OUT	temp. fluide utilité sortie cellule 1	30
TU_OUT	temp. sortie fluide utilité	90

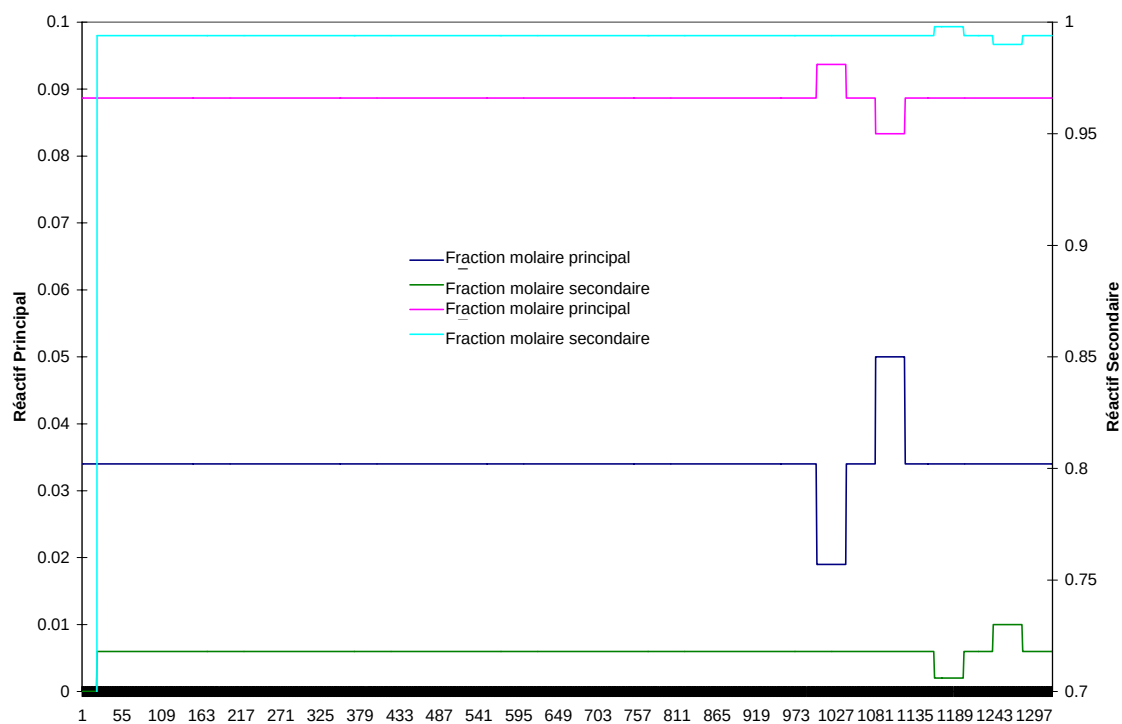
X_MAIN_1	fraction molaire alim. courant principal, composé 1	1
X_MAIN_2	fraction molaire alim. courant principal, composé 2	1
X_MAIN_3	fraction molaire alim. courant principal, composé 3	1
X_MAIN_4	fraction molaire alim. courant principal, composé 4	1
X_MAIN_5	fraction molaire alim. courant principal, composé 5	1
X_MAIN_6	fraction molaire alim. courant principal, composé 6	1
X_INJ_1	fraction molaire alim. courant injecté, composé 1	3
X_INJ_2	fraction molaire alim. courant injecté, composé 2	3
X_INJ_3	fraction molaire alim. courant injecté, composé 3	3
X_INJ_4	fraction molaire alim. courant injecté, composé 4	3
X_INJ_5	fraction molaire alim. courant injecté, composé 5	3
X_INJ_6	fraction molaire alim. courant injecté, composé 6	3

4.3.1 Réaction d'estérification

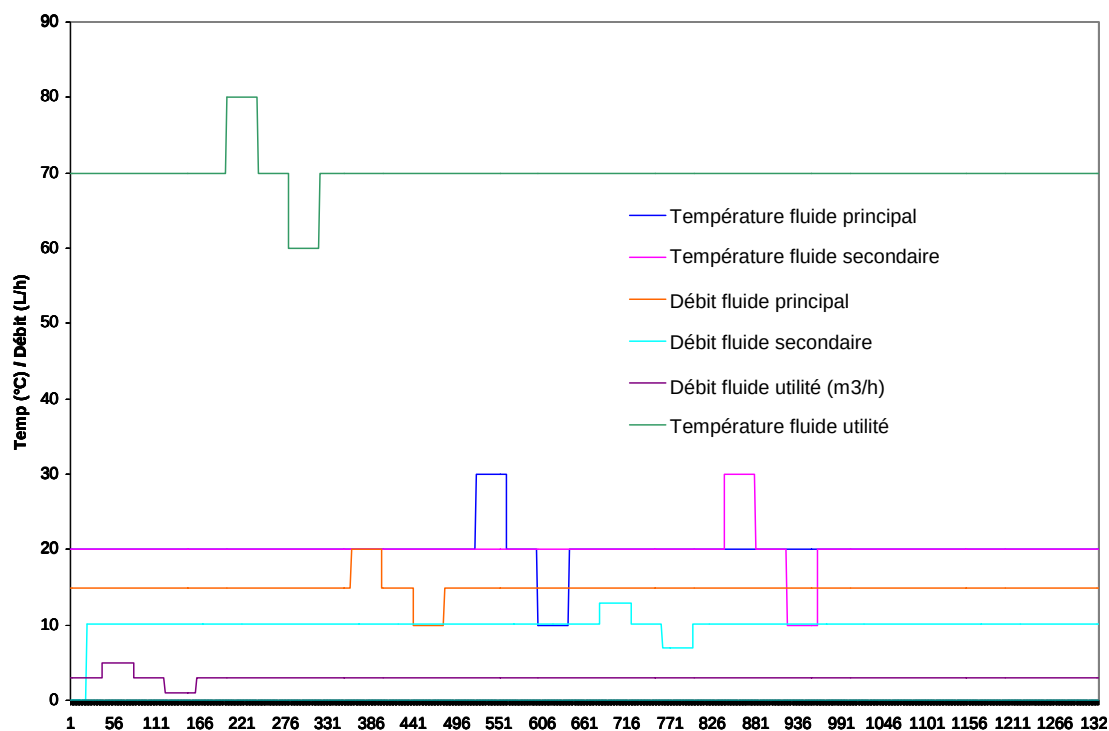
La réaction d'estérification est une réaction lente et faiblement exothermique. Pour l'accélérer, il faut chauffer le milieu réactionnel. Dans ce cas, le fluide utilité sert de fluide de chauffage. Il y a moins de contraintes de sécurité que dans le cas de la réaction du thiosulfate.

Les perturbations dans l'OPR pour la réaction d'estérification ont été réalisées à partir de variations en simulation sur les températures, les débits et les compositions des diverses variables (réactif principal, réactif injecté, débit d'utilité). Au total, ce sont 16 défauts qui ont été appliqués au réacteur (voir la Figure 4-20).

Les défaillances dans l'OPR ont été simulées sous forme de perturbations sur les températures et les débits du réactif principal (C_4H_8O), du réactif secondaire ou injecté ($C_6H_{10}O_3$) et sur le système de refroidissement (système utilité), ainsi que sur la composition dans les réactifs principal et secondaire.



b) Perturbations de composition



a) Perturbations sur le fluide utilité et sur les températures

Figure 4-20 Perturbations réalisées sur la réaction d'estérification

4.3.1.1 Sélection des descripteurs : Application à la réaction d'estérification

Pour obtenir le modèle de comportement, tout d'abord nous avons réalisé un auto-apprentissage en utilisant comme base d'apprentissage, les évolutions de tous les descripteurs suite aux variations de conditions opératoires. Les résultats de cette classification sont regroupés sur la Figure 4-21. Les paramètres utilisés pour cette classification sont : un niveau d'exigence de 0.8, une variabilité maximale de 0.087, la fonction d'appartenance *Gauss* 1 et le connectif *Min-Max*.

Nous pouvons constater la création de 46 classes ; le résultat le plus important est que nous avons détecté de façon stable et concise chaque défaut survenu sur le réacteur, étape préalable à l'application du gain d'information pour la sélection des descripteurs les plus pertinents.

Le Tableau 4-Q donne les classes qui représentent les divers défauts et qui ont été utilisées pour la sélection des descripteurs.

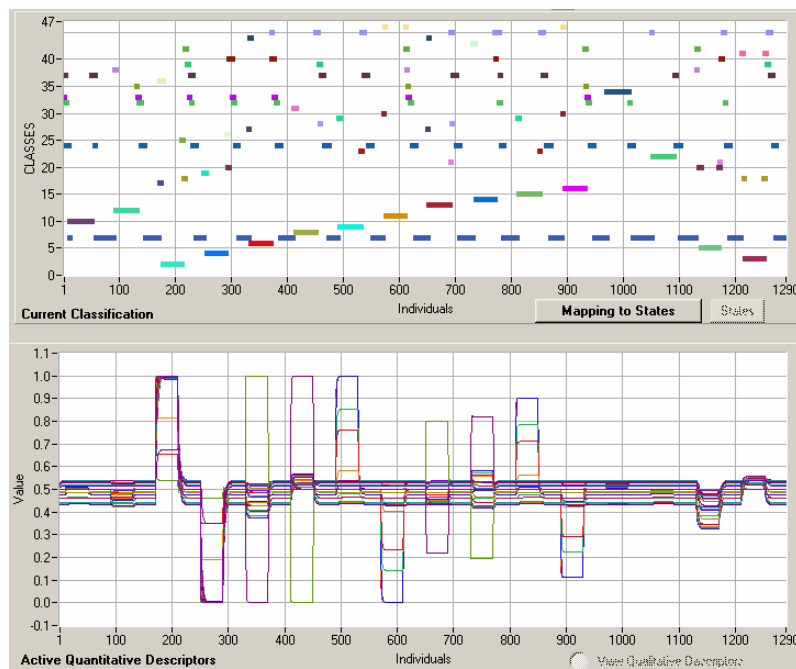


Figure 4-21 Détection des défauts avec les 27 descripteurs : réaction d'estérification

Tableau 4-Q Classes qui représentent les divers défauts : réaction d'estérification

NUMÉRO DE CLASSE	NOM DÉFAUT
2	H-TU_B1_IN
3	H-X_INJ
4	B-TU_B1_IN
5	B-X_INJ
6	H-FP_IN1
7	NORMAL
8	B-FP_IN1
9	H-TP_IN1
10	H-FU_IN
11	B-TP_IN1
12	B-FU_IN
13	H-FP_IN2
14	B-FP_IN2
15	H-TP_IN2
16	B-TP_IN2
22	H-X_MAIN
34	B-X_MAIN

Le résultat de l'auto-apprentissage est le profil des classes, lequel contient la valeur moyenne normalisée de chaque descripteur pour chaque classe qui représente un défaut. A partir de ce résultat, le gain d'information peut être calculé pour chaque défaut, en prenant comme référence l'état normal. Comme précédemment, nous avons formé des paires de situations de la façon suivante : **Normal-Défaut**, avec la finalité de développer les calculs du gain d'information comme ceci a été expliqué dans le chapitre précédent.

Le Tableau 4-R donne les descripteurs sélectionnés après avoir appliqué la procédure de sélection des capteurs grâce au calcul du gain d'information. Comme précédemment, le critère de choix a été de prendre le descripteur avec un gain d'information le plus élevé par rapport à chaque défaut sur le réacteur (voir le Tableau 4-P pour la description des descripteurs).

Tableau 4-R Descripteurs sélectionnés : réaction d'estérification

PAIRE DE SITUATION	DESCRIPTEUR SÉLECTIONNÉ	DESCRIPTION
NORMAL-H-TU_B1_IN	TP_B2_5	temp. fluide procédé bloc 2 cellule 5
NORMAL-H-X_INJ	TP_B2_5	temp. fluide procédé bloc 2 cellule 5
NORMAL-B-TU_B1_IN	TU_B1_OU	temp. fluide utilité sortie cellule 1
NORMAL-B-X_INJ	TP_B2_10	temp. fluide procédé bloc 2 cellule 10
NORMAL-H-FP_IN1	PP_OUT	pression fluide procédé en sortie
NORMAL-B-FP_IN1	FP_OUT	débit fluide procédé en sortie
NORMAL-H-TP_IN1	TP_B1_2	temp. fluide procédé bloc 1 cellule 2
NORMAL-H-FU_IN	TP_B1_3	temp. fluide procédé bloc 1 cellule 3
NORMAL-B-TP_IN1	TP_B1_1	temp. fluide procédé bloc 1 cellule 1
NORMAL-B-FU_IN	TP_B1_3	temp. fluide procédé bloc 1 cellule 3
NORMAL-H-FP_IN2	PP_OUT	pression fluide procédé en sortie
NORMAL-B-FP_IN2	FP_OUT	débit fluide procédé en sortie
NORMAL-H-TP_IN2	TP_B1_3	temp. fluide procédé bloc 1 cellule 3
NORMAL-B-TP_IN2	TP_B1_1	temp. fluide procédé bloc 1 cellule 1
NORMAL-H-X_MAIN	FP_OUT	débit fluide procédé en sortie
NORMAL-B-X_MAIN	PP_OUT	pression fluide procédé en sortie

Dans ce tableau nous pouvons constater la présence de 8 descripteurs différents.

Sans revenir sur chacun des choix, essayons d'en dégager les grandes caractéristiques. Cinq des 8 descripteurs sont des mesures de température du fluide procédé le long du réacteur. Puisqu'il n'y a pas de mesure en ligne des concentrations, la source d'information sur un changement du fonctionnement du réacteur reste le profil de température le long du réacteur. Ces points de mesure sélectionnés sont positionnés à la sortie des premières cellules du premier bloc (les 3 premières cellules) et deux cellules du bloc 2. La localisation de ces points dans le réacteur correspond bien aux points où la réaction chimique est la plus intense, c'est-à-dire où la vitesse de la réaction est maximale. La température du fluide utilité en sortie du premier bloc sera une variable affectée par une modification de la température de ce même fluide à l'alimentation. Un défaut sur la quantité produite sera perceptible sur les variables comme la pression fluide procédé ou le débit du fluide procédé en sortie. Cette quantité peut être altérée suite à une modification des débits à l'alimentation (fluide injecté ou fluide principal, ou une variation de la concentration du réactif principal).

La section suivante est consacrée au développement du modèle de comportement en utilisant seulement ces descripteurs sélectionnés.

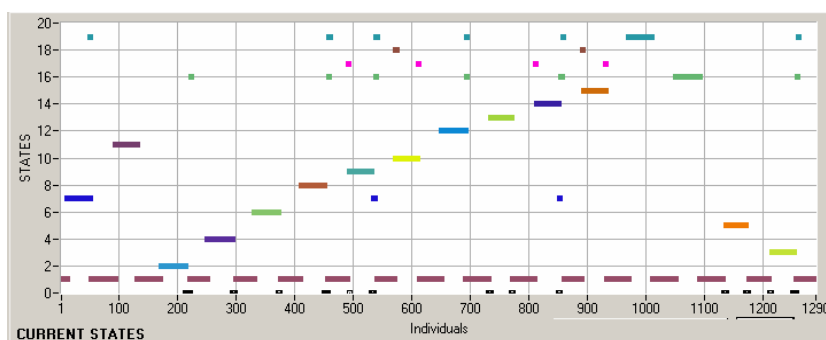
4.3.1.2 Conception du modèle de comportement pour la réaction d'estérification

Le modèle de comportement doit renfermer les classes qui définissent les défauts, ainsi que, les situations d'alarmes de telle sorte qu'il permette d'anticiper la détection des défauts. Sur la Figure 4-22a, on montre le résultat d'identification des défauts et des alarmes avec les 8 descripteurs sélectionnés. Il y a création de 47 classes à cause de la valeur très petite de la variabilité maximale (0.087). Les autres paramètres utilisés pour cette classification sont les suivants : un niveau d'exigence de 0.8, la fonction d'appartenance *Gauss* 1 et le connectif *Min-Max*. Le modèle de comportement résultant pour cette réaction est donné dans le Tableau 4-S avec le nom du défaut, l'ensemble des classes qui composent l'état du procédé et le numéro de l'état. Dans la colonne des classes, outre les différentes classes regroupées, figurent les éventuels pré-défaut et post-défaut. Ils représentent les états d'alarme avant et après que le défaut se soit établi. Par exemple, dans le défaut « baisse du débit d'alimentation du fluide d'utilité » ($\downarrow FU_IN$) la classe 12 correspond à ce défaut stabilisé, et il n'y a pas de pré-défaut ni de post-défaut. Le défaut « hausse de température de l'alimentation de fluide utilité » est représenté par la classe 2, tandis que, les classes 17, 34, 45 sont des alarmes de pré-défaillance, et les classes 18 et 45 correspondent à la récupération du procédé vers l'état normal, à partir de ce défaut. La création des états 17 ($\uparrow TP_IN$) et 18 ($\downarrow TP_IN$) a été effectuée pour rendre compte d'un même phénomène « hausse ou baisse de la température » affectant le fluide principal ou le fluide injecté. Sur la Figure 4-22a, on observe que les classes 28 et 25 correspondent au début du défaut (pré-défaut) « hausse de la température du fluide principal » et au pré-défaut « hausse de la température du fluide injecté », mais aussi à la remontée de la température lors du retour à l'opération normale depuis une baisse de la température du fluide principal ou du fluide injecté (post-défaut). Dans ce cas, il nous est apparu cohérent de regrouper ces classes comme l'état d'« alarme de hausse de la température du fluide d'alimentation » (soit principal ou injecté) ($\uparrow TP_IN$). La situation est similaire pour les classes 42 et 46. Dans ce dernier cas, nous pouvons définir ces classes comme l'état d'« alarme de baisse de la température du fluide d'alimentation » (soit principal ou injecté) ($\downarrow TP_IN$).

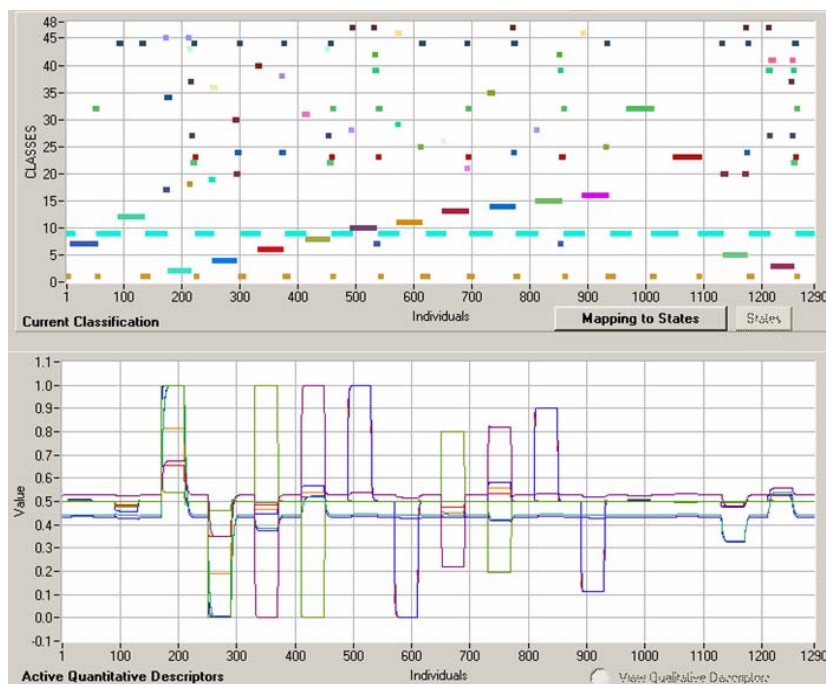
En prenant comme base le modèle du Tableau 4-S, on peut construire le modèle d'états fonctionnels comme le montre la Figure 4-22b. Nous pouvons constater la création de 19 états fonctionnels. Par exemple, l'état numéro 1 correspond à l'opération Normale et le défaut « baisse du fluide utilité » ($\downarrow FU_IN$) correspond à l'état numéro 11.

Tableau 4-S Modèle de comportement de la réaction d'estérification

	DÉFAUT	CLASSES	ÉTAT
	NORMAL	1, 9, 44	1
1	↑FU_IN	7	7
2	↓FU_IN	12	11
3	↑TU_B1_IN	2 Alarme : 17, 34, 45 - Retour : 45, 18	2
4	↓TU_B1_IN	4 Alarme : 19, 36 - Retour : 30	4
5	↑FP_IN_1	6 Alarme : 40 - Retour : 38	6
6	↓FP_IN_1	8 Alarme : 31	8
7	↑TP_IN_1	10 Alarme : 18	9
8	↓TP_IN_1	11	10
9	↑FP_IN_2	13 Alarme : 26 - Retour : 21	12
10	↓FP_IN_2	14	13
11	↑TP_IN_2	15	14
12	↓TP_IN_2	16	15
13	↑X_MAIN	32	19
14	↓X_MAIN	23	16
15	↑X_INJ	5	5
16	↓X_INJ	3 Transitoire : 41	3
	Tendance ↑TP_IN	Transitoire : 25, 28	17
	Tendance ↓TP_IN	Transitoire : 42, 46	18



b)



a)

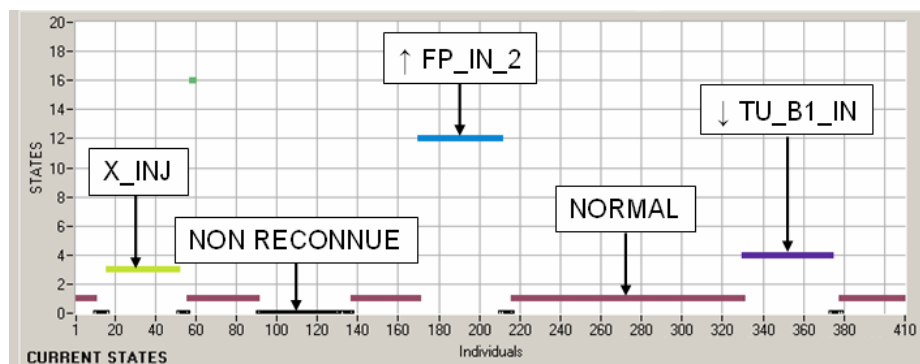
Figure 4-22 Identification de défauts et d'alarmes avec les 8 descripteurs : réaction d'estérification

4.3.1.3 Reconnaissances de défauts inconnus : Application à la réaction d'estérification

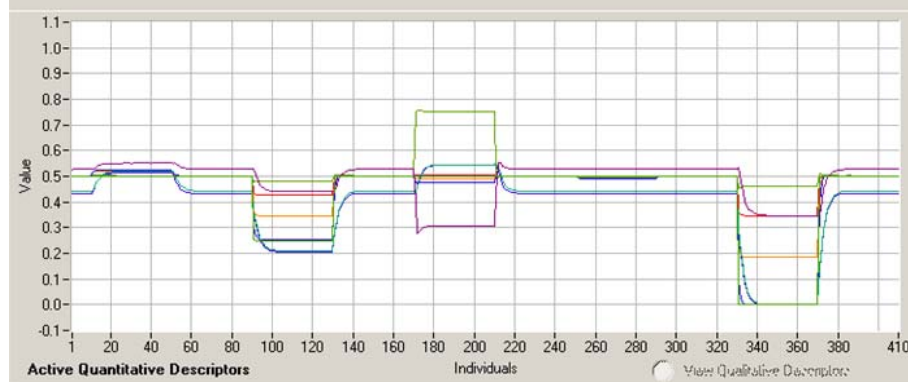
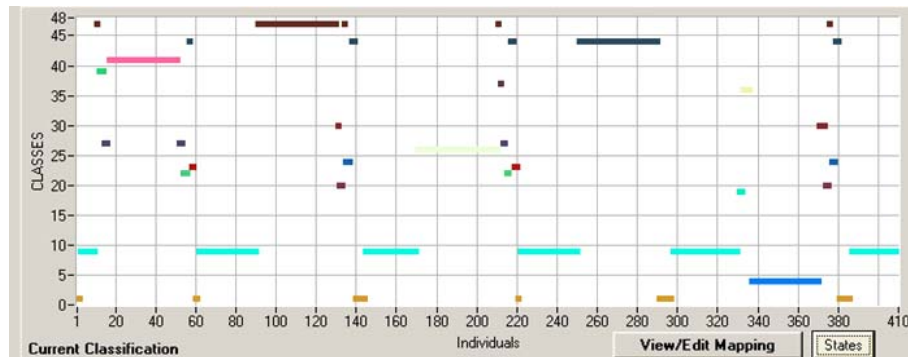
La validation de cette méthodologie a été effectuée dans le cadre du projet ICSI, en collaboration avec l'équipe « Procédés de la Chimie Fine » du LGC et elle s'est voulue réaliste, c'est-à-dire correspondre à la reconnaissance a priori par l'opérateur des défaillances survenant sur le procédé. Pour faire la reconnaissance en utilisant le modèle construit dans la section précédente, un fichier de données acquises à partir de la simulation de scénarii effectuée sur le simulateur de réacteur nous a été fourni sans préciser le type de perturbation effectué ni s'il y avait existence de défauts simultanés. Le seul renseignement était le nombre de perturbations (5) et leur date :

Perturbation	Début	Fin
Perturbation 1	10	50
Perturbation 2	90	130
Perturbation 3	170	210
Perturbation 4	250	290
Perturbation 5	330	370

Sur la Figure 4-23a nous pouvons visualiser la reconnaissance des défauts inconnus en utilisant le modèle de comportement. Les paramètres utilisés pour cette phase de reconnaissance sont : un degré d'exigence plus faible (0.57) que celui utilisé pour la construction du modèle, afin d'être moins strict et ainsi assigner le défaut à une classe. Avec l'aide du tableau donnant les états du modèle de comportement, nous pouvons décrypter ce qui s'est passé dans le réacteur. Les classes 19 et 36 représentent des alarmes de pré-défaillance correspondant à un début de « baisse de la température du fluide utilité » ($\downarrow$ TU_B1_IN) (voir le Tableau 4-P).



b)



a)

Figure 4-23 Classification de défauts inconnus : réaction d'estérification

La Figure 4-23b présente les résultats de l'identification des états fonctionnels. On peut y noter la création de la classe zéro (correspondant à des éléments non reconnus), ainsi que les autres états fonctionnels.

Nous pouvons remarquer que les perturbations 2 et 3 ne sont pas identifiées. En ce qui concerne la quatrième perturbation, l'amplitude de cette perturbation étant très faible les éléments ont été assignés à la classe normale.

Connaissant les perturbations qui ont été appliquées, essayons de dégager quelques conclusions. La première perturbation a été effectuée sur la fraction massique de 2-butanol (réactif injecté) qui est passée de 0.994 à 0.991. Cette perturbation est bien reconnue. La deuxième perturbation concerne une faible variation sur la température d'alimentation de l'utilité dans le bloc 1 ($\downarrow TU_B1_IN$) qui passe de 70°C à 65°C. Cette baisse n'a pas pu être identifiée lors de la phase de reconnaissance. La perturbation 4 concerne la baisse du fluide utilité ($\downarrow FU_IN$) qui passe de 3 à 2 m³/h. Cette variation est là aussi trop faible pour être détectée : la réaction étant très faiblement exothermique, l'impact d'une telle variation n'a que peu d'incidence sur la réaction, c'est pour cela que l'état identifié est l'état normal. Les perturbations 3 et 5 sont des défauts simultanés et donc jamais présents dans la base d'apprentissage lors de la phase de conception du modèle de comportement. La perturbation 3 correspond en réalité à une hausse du débit du réactif injecté ($\uparrow FP_IN_2$) et en même temps d'une baisse de la fraction massique de 2-butanol avec la même amplitude que pour la perturbation 1. Ces deux changements de conditions opératoires vont dans un sens opposé et donc leurs effets s'annulent : on alimente plus mais plus dilué. La dernière perturbation correspond aux défauts simultanés « baisse de la température d'utilité » et « baisse du débit de fluide utilité » ($\downarrow TU_B1_IN$ et $\downarrow FU_IN$) mais l'amplitude de la variation sur la température de l'utilité est deux fois plus importante (de 70°C à 60°C) ; celle sur le débit restant inchangée par rapport à la perturbation 4. Dans ce cas, l'état reconnu est celui correspondant à la défaillance sur la température de l'utilité ; les deux effets des perturbations étant dans ce cas dans le même sens (un défaut de chauffage). On voit donc que, lorsque la variation sur la température est assez importante et du même ordre de grandeur que celle du scénario utilisé pour construire le modèle de comportement, le système peut le reconnaître.

On peut noter que malgré un échec, grâce au modèle de comportement on a pu identifier la plupart des défauts et que dans le cas de défauts simultanés, il a été possible d'identifier au moins l'un des deux. Dans l'exemple suivant, nous montrons comment le modèle peut être modifié lorsque de nouveaux états ne sont pas reconnus (en particulier les défauts simultanés).

4.3.2 Réaction du thiosulfate

La réaction du thiosulfate est une réaction très rapide et fortement exothermique. Elle nécessite un refroidissement permanent par un fluide utilité (de l'eau). La forte exothermicité entraîne de nombreuses contraintes de sécurité.

Les défaillances dans l'OPR pour la réaction du thiosulfate ont été simulées sous forme de perturbations sur les mêmes variables que dans l'exemple précédent : températures et les débits du réactif principal ($\text{Na}_2\text{S}_2\text{O}_3$), du réactif secondaire ou injecté (H_2O_2) et sur le système de refroidissement (système utilité), ainsi que sur la composition dans les réactifs principal et secondaire, mais les amplitudes des variations sont plus importantes. La différence majeure avec la réaction précédente c'est qu'il y a arrêt total du fluide utilité. Au total ce sont 17 défauts qui ont été appliqués au réacteur (la dernière perturbation entre 20800 et 21400 seconds correspond à un arrêt total du fluide utilité) (voir la Figure 4-24).

Les réactions sont mises en œuvre avec un excès de H_2SO_4 (réactif principal), lequel est d'abord introduit dans le réacteur avec un débit fixé à 40 l/h. Quand l'état stable est atteint, le réactif secondaire NaOH est injecté à raison de 10 l/h.

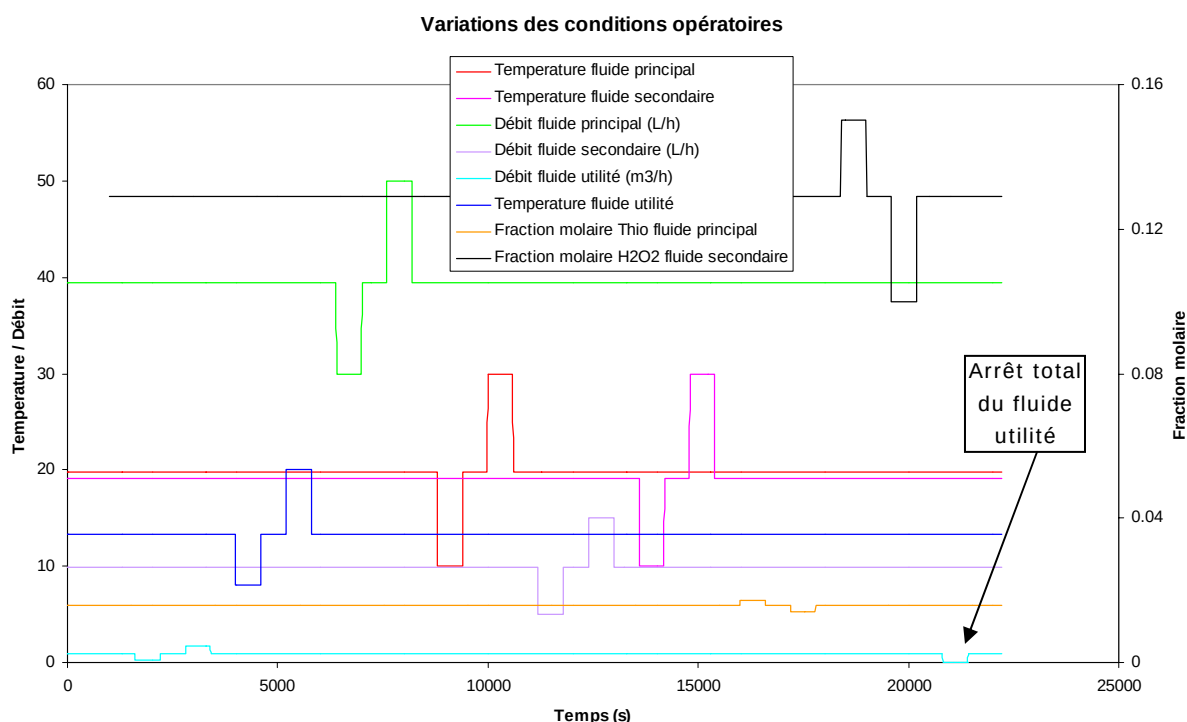


Figure 4-24 Variations des conditions opératoires : réaction du thiosulfate

4.3.2.1 Sélection des descripteurs : Application à la réaction du thiosulfate mise en œuvre dans le réacteur « *Open Plate Reactor* »

Comme première étape de la conception du modèle de comportement, nous avons effectué un auto-apprentissage à partir de la base d'apprentissage constituée de toutes les évolutions des descripteurs lors des variations des conditions opératoires. Le but de cet auto-apprentissage est d'obtenir une classification stable de chaque défaut simulé, pour avoir un profil de classes bien définies. La Figure 4-25, donne le résultat de cette classification. Les paramètres utilisés pour cette classification sont : un niveau d'exigence de 0.6, une variabilité maximale de 0.1, la fonction d'appartenance *Gauss 1* et le connectif *Max-Min*.

Nous pouvons constater la création de 55 classes ; le résultat le plus important c'est que nous avons détecté de façon stable et concise chaque défaut sur le réacteur.

Le Tableau 4-T donne les classes obtenues à partir de la base d'apprentissage constituée des variations de tous les descripteurs (27) suite aux variations des conditions opératoires.

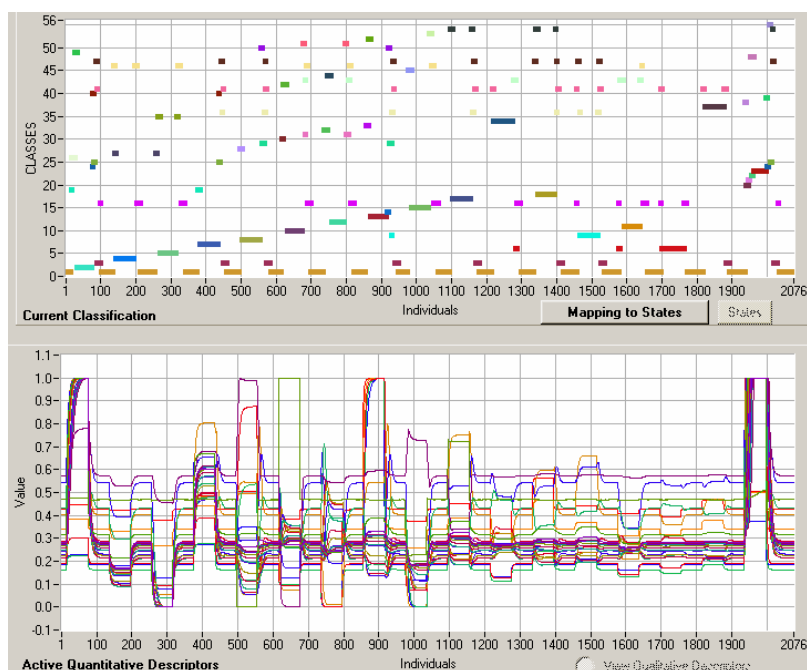


Figure 4-25 Détection de défauts avec 27 descripteurs : réaction du thiosulfate

Tableau 4-T Défauts représentés par les classes

NUMÉRO DE CLASSE	NOM DÉFAUT
1	NORMAL
2	Baisse Débit_Utilité
4	Hausse Débit_Utilité
5	Baisse Temp_Utilité
6	Hausse Fracc_Mol_Second
7	Hausse Temp_Utilité
8	Baisse Débit_Principal
9	Hausse Fracc_Mol_Principal
10	Hausse Débit_Principal
11	Baisse Fracc_Mol_Principal
12	Baisse Temp_Principal
13	Hausse Temp_Principal
15	Baisse Débit_Second
17	Hausse Débit_Second
18	Hausse Temp_Second
23	Arret utilité
34	Baisse Temp_second
37	Baisse Fracc_Mol_Second

Le résultat de l'auto-apprentissage est le profil de classes, lequel contient la valeur moyenne normalisée de chaque descripteur pour la classe qui représente le défaut. A partir de ce résultat nous pouvons calculer l'entropie pour chaque défaut, en prenant comme référence l'état normal. C'est-à-dire, nous avons formé les paires de situations de la façon suivante : **Normal– Défaut**, avec la finalité de développer les calculs de l'entropie comme nous l'avons fait dans la section 4.2.2 de ce chapitre.

Le Tableau 4-U présente les descripteurs sélectionnés après avoir réaliser les calculs du gain d'information. Le critère de choix a été de prendre le descripteur avec un gain d'information le plus grand par rapport à chaque situation.

Dans ce tableau, nous pouvons constater qu'il y a 9 descripteurs différents (voir le Tableau 4-P pour la description de chaque descripteur).

On retrouve comme précédemment principalement des capteurs de température dans les deux premiers blocs. La réaction est cette fois-ci très exothermique et comme il n'y a pas de régulation de température au sein du réacteur, l'évolution des températures le long des blocs est porteuse d'information : plus la réaction est rapide plus il y a production de chaleur, plus il y a échauffement du fluide procédé (puisque la température augmente et que le débit de fluide utilité reste inchangé, plus la réaction s'accélère). Donc le fait de trouver des mesures de température du fluide procédé en sortie des blocs 1 et 2, en plus des températures dans les premières cellules est cohérent avec ce phénomène d'exothermicité. On remarque que le profil de température dans le premier bloc contient à lui tout seul beaucoup : 4 sur 9 au total.

La section suivante est consacrée à la construction du modèle de comportement en utilisant seulement ces descripteurs sélectionnés.

Tableau 4-U Descripteurs sélectionnés : réaction du thiosulfate

PAIRE DE SITUATION	DESCRIPTEUR SÉLECTIONNÉ	DESCRIPTION
Normal - Baisse Débit Utilité	TP_B2_10	temp. fluide procédé bloc 2 cellule 10
Normal - Hausse Débit Utilité	TP_B2_10	temp. fluide procédé bloc 2 cellule 10
Normal - Baisse Temp Utilité	TU_B1_OU	temp. fluide utilité sortie cellule 1
Normal - Hausse Fracc Mol Second	TP_B1_10	temp. fluide procédé bloc 1 cellule 10
Normal - Hausse Temp Utilité	TP_B2_20	temp. fluide procédé bloc 2 cellule 20
Normal - Baisse Débit Principal	FP_OUT	débit fluide procédé en sortie
Normal - Hausse Fracc Mol Principal	TP_B1_10	temp. fluide procédé bloc 1 cellule 10
Normal - Hausse Débit Principal	PP_OUT	pression fluide procédé en sortie
Normal - Baisse Fracc Mol Principal	TP_B1_10	temp. fluide procédé bloc 1 cellule 10
Normal - Baisse Temp Principal	TP_B1_1	temp. fluide procédé bloc 1 cellule 1
Normal - Hausse Temp Principal	TP_B1_4	temp. fluide procédé bloc 1 cellule 4
Normal - Baisse Débit Second	TP_B1_10	temp. fluide procédé bloc 1 cellule 10
Normal - Hausse Débit Second	TP_B1_10	temp. fluide procédé bloc 1 cellule 10
Normal - Hausse Temp Second	TP_B1_5	temp. fluide procédé bloc 1 cellule 5
Normal - Arrêt utilité	TP_B1_4	temp. fluide procédé bloc 1 cellule 4
Normal - Baisse Temp second	TP_B1_10	temp. fluide procédé bloc 1 cellule 10
Normal - Baisse Fracc Mol Second	TP_B1_5	temp. fluide procédé bloc 1 cellule 5

4.3.2.2 Conception du modèle de comportement : Application à la réaction du thiosulfate mise en oeuvre dans le réacteur « *Open Plate Reactor* »

Le modèle de comportement doit comprendre les classes qui définissent les défauts, ainsi que, les situations d'alarmes de telle sorte qu'il soit capable d'anticiper la détection des défauts. La Figure 4-26 montre le résultat d'identification des défauts et des alarmes avec 9 descripteurs pour l'exemple de la réaction du thiosulfate. Il y a création d'un grand nombre de classes (61 au total) à cause de la valeur très petite de la variabilité maximale (0.07). Une valeur supérieure de ce paramètre n'a pas permis de trouver les classes que correspondent aux alarmes des défauts. Les autres paramètres utilisés pour cette classification sont : un niveau d'exigence de 0.51, la fonction d'appartenance *Gauss 1* et le connectif *Max-Min*.

Le Tableau 4-V donne le modèle de comportement du réacteur lors de la mise en oeuvre de la réaction du thiosulfate. Dans ce tableau, nous avons reporté le nom du défaut, l'ensemble des classes qui composent l'état du procédé et le numéro de l'état. Pour chaque classe liée à un défaut, nous avons fait figurer un pré-défaut et un post-défaut (alarme et retour, respectivement); ils représentent les états d'alarme avant et après respectivement, que le défaut se soit manifesté. Par exemple pour le défaut « baisse du débit d'alimentation du fluide utilité » ($\downarrow FU_IN$) la classe 2 correspond à la défaillance lorsqu'elle est parfaitement établie,

tandis que, la classe 39 correspond à l'apparition de ce défaut, et les classes 19 et 55 correspondent à la récupération du procédé vers l'état normal, à partir de ce défaut établi. Une autre particularité dans ce modèle est l'existence de l'état numéro 16 qui représente une récupération du défaut « Baisse de fluide utilité » (B-POST-FU) et est formée par les classes 19 et 55.

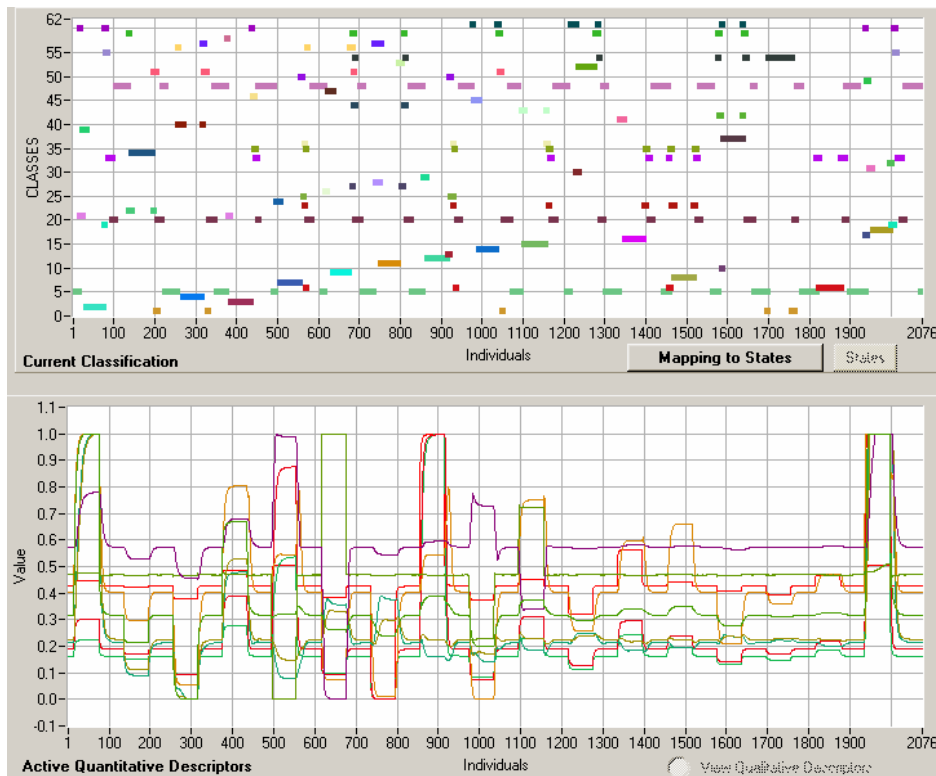


Figure 4-26 Identification de défauts et alarmes avec 9 descripteurs : réaction du thiosulfate

Tableau 4-V Modèle de comportement de la réaction du thiosulfate

DÉFAUT		CLASSES	ÉTAT
	NORMAL	5, 20, 48	4
1	↓FU_IN	2 Pré-défaut : 39, Post-défaut : 19, 55	1
2	↑FU_IN	34 Pré-défaut 22 : Post-défaut : 22	17
3	↓TU_B1_IN	4 Pré-défaut : 40 Post-défaut : 40	3
4	↑TU_B1_IN	3 Pré-défaut : 58 Post-défaut : 46	2

DÉFAUT		CLASSES	ÉTAT
5	↓FP_IN_1	7 Pré-défaut : 24	6
6	↑FP_IN_1	9 Pré-défaut : 26, 47	8
7	↓TP_IN_1	11 Pré-défaut : 28 Post-défaut : 53	10
8	↑TP_IN_1	12 Pré-défaut : 29 Post-défaut : 13	11
9	↓FP_IN_2	14 Pré-défaut : 45	12
10	↑FP_IN_2	15 Pré-défaut : 43 Post-défaut : 43	13
11	↓TP_IN_2	52 Pré-défaut : 30	18
12	↑TP_IN_2	16 Pré-défaut : 41	14
13	↑X_MAIN	8	7
14	↓X_MAIN	37 Pré-défaut : 10, 42 Post-défaut : 42	9
15	↑X_INJ	54	19
16	↓X_INJ	6	5
17	Arrêt FU_IN	18 Pré-défaut : 17, 31, 49 Post-défaut : 32	15
	B-POST-FU	19, 55	16

En prenant comme base le modèle donné dans le Tableau 4-V, on peut représenter les états fonctionnels du réacteur (réaction du thiosulfate) et obtenir la Figure 4-27. Nous pouvons constater la création de 19 états fonctionnels. A gauche de la figure se trouve la liste des états existants. Par exemple, l'état numéro 1 correspond au défaut 1 (DF1), en regardant le Tableau 4-V ce défaut correspond à la baisse du fluide utilité (↓FU_IN) et l'état numéro 4 correspond à l'opération normale du réacteur.

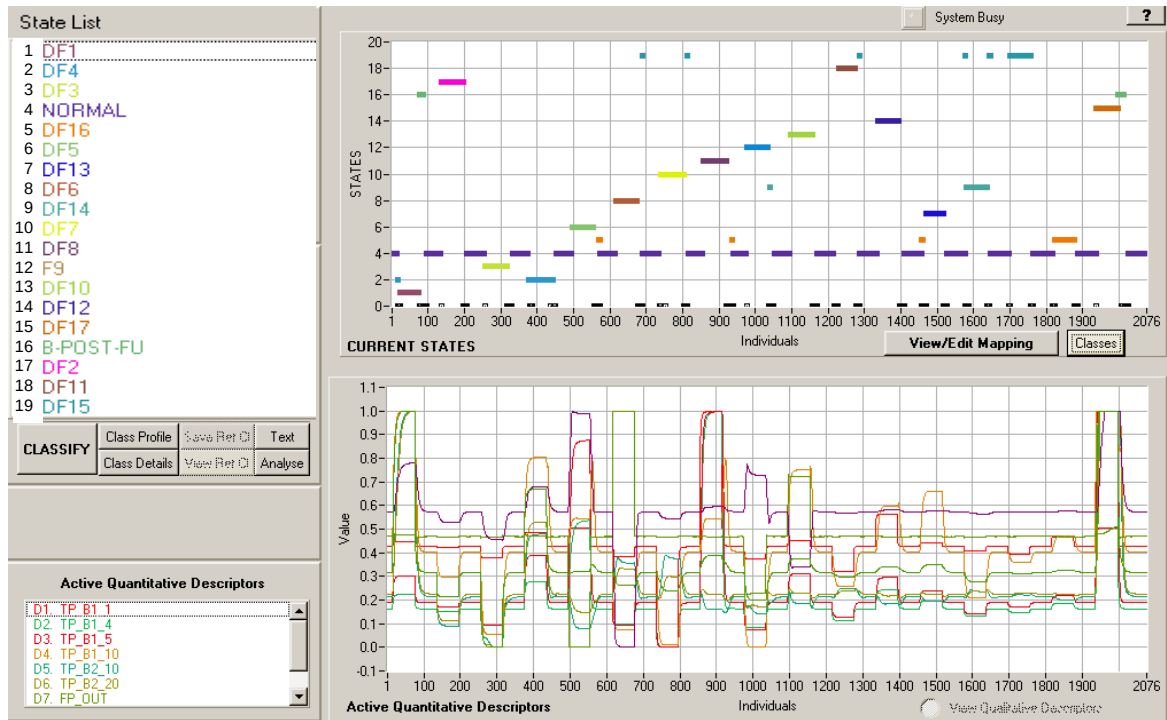


Figure 4-27 Représentation des états fonctionnels du réacteur: réaction du thiosulfate

4.3.2.3 Reconnaissances de défauts inconnus : Application à la réaction du thiosulfate

Comme précédemment, la validation de cette méthodologie a été effectuée en réalisant une reconnaissance en utilisant le modèle construit dans la section précédente sur des données de simulation de scénarii, qui nous ont été fournies sans préciser le type de perturbation effectuée ni s'il y avait existence de défauts simultanés. Le seul renseignement était le nombre de perturbations et leur date :

Perturbation	Début	Fin
Perturbation 1	15	75
Perturbation 2	135	195
Perturbation 3	255	315
Perturbation 4	375	435
Perturbation 5	495	555
Perturbation 6	615	675

La Figure 4-28 permet de visualiser les résultats obtenus lors de la phase de reconnaissance de ces défauts non connus au préalable en utilisant le modèle de comportement de la réaction du thiosulfate. Avec l'aide du tableau du modèle, nous pouvons décrypter ce qui s'est passé dans

le réacteur. Par exemple, la classe 8 correspond à une hausse de la composition du réactif principal. La classe 39 représente une alarme de pré-défaillance de baisse du débit d'utilité. C'est un cas similaire pour la classe 58 qui correspond à l'alarme de pré-défaillance de hausse de la température du fluide utilité, mais dans ce cas, il existe un problème de reconnaissance car la même classe est présentée au début de la classification. En revanche, la perturbation numéro 5 est assignée à la classe 61, laquelle est une classe inexistante dans le modèle. Ce peut être un nouveau défaut non pris en compte pour la conception du modèle ou une simultanéité de défauts.

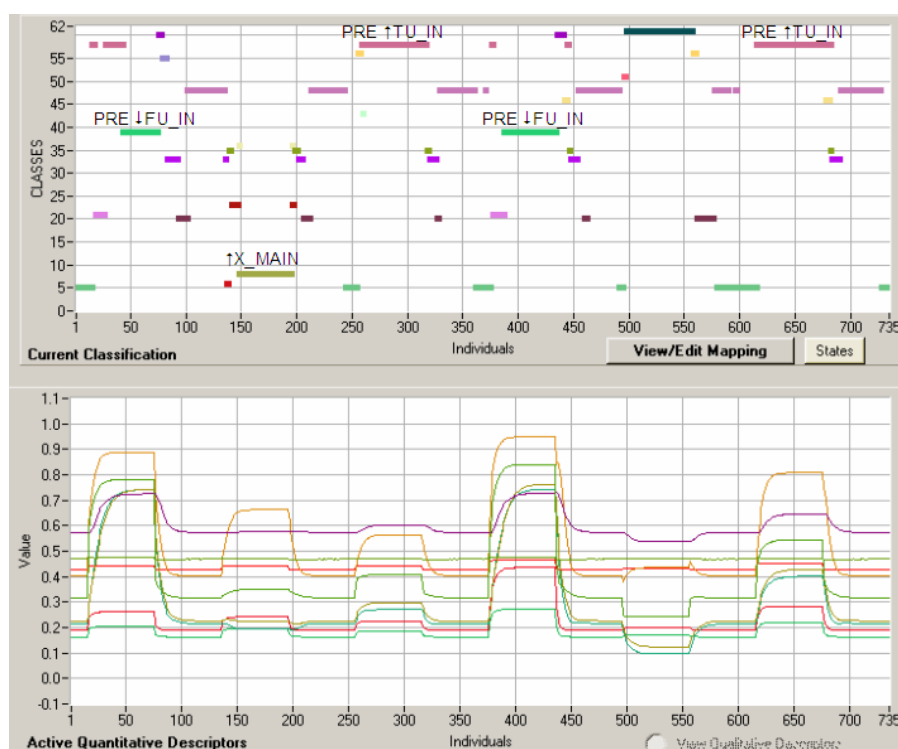


Figure 4-28 Classification des défauts inconnus : réaction du thiosulfate

Reprenons maintenant les défaillances qui ont été réellement simulées.

Perturbation 1 : Débit d'utilité de 0.916 m³/h à 0.3 (↓FU_IN)

Perturbation 2 : Fraction massique du thiosulfate de 0.0137 à 0.017 (↑X_MAIN)

Perturbation 3 : Température de l'utilité de 13.37 à 15°C (↑TU_B1_IN)

Perturbation 4 : Défauts simultanés : hausse de la fraction massique du thiosulfate de 0.0137 à 0.0175 et baisse du débit d'utilité de 0.916 m³/h à 0.3 (↑X_MAIN + ↓FU_IN).

Perturbation 5 : Défauts simultanés : hausse de la fraction massique du thiosulfate de 0.0137 à 0.0165 et baisse du débit d'utilité de 0.916 m³/h à 1.5 (↑X_MAIN + ↓FU_IN).

Perturbation 6 : 3 défauts simultanés : hausse de la fraction massique du thiosulfate de 0.0137 à 0.0165 et baisse du débit d'utilité de 0.916 m³/h à 0.6 et hausse de la température de l'utilité de 13.37 à 14.8°C (↑X_MAIN + ↓FU_IN + ↑TU_B1_IN).

Il y a donc 3 défaillances qui sont bien reconnues, celles correspondant aux perturbations 1, 2 et 3. Dans les cas 1 et 3, on a trouvé 2 alarmes puisque dans les deux perturbations, les variations du débit d'utilité et de la température de l'utilité sont plus petites que celles présentées lors de la conception du modèle et ont donc été assimilées à des pré-défauts.

Dans le cas de défaillances simultanées, seule une des défaillances est reconnue : la baisse du débit d'utilité dans le cas n°4 et la hausse de la température de l'utilité dans le cas n°6. La perturbation n°5 n'est pas identifiée comme étant un état de défaillance : les effets conjugués des 2 perturbations s'annulent : on alimente avec un réactif plus concentré mais on refroidit plus ce qui correspond bien à une opération dite normale (par exemple s'il y avait un régulateur pour maintenir la température en sortie du bloc 1, face à une augmentation de la fraction massique d'un réactif, il augmenterait le débit d'utilité).

Sur la Figure 4-29, on présente le résultat de l'identification des états fonctionnels. Ici on peut remarquer l'état zéro correspondant à la classe zéro ou classe NIC (non reconnue), ainsi que les autres états fonctionnels.

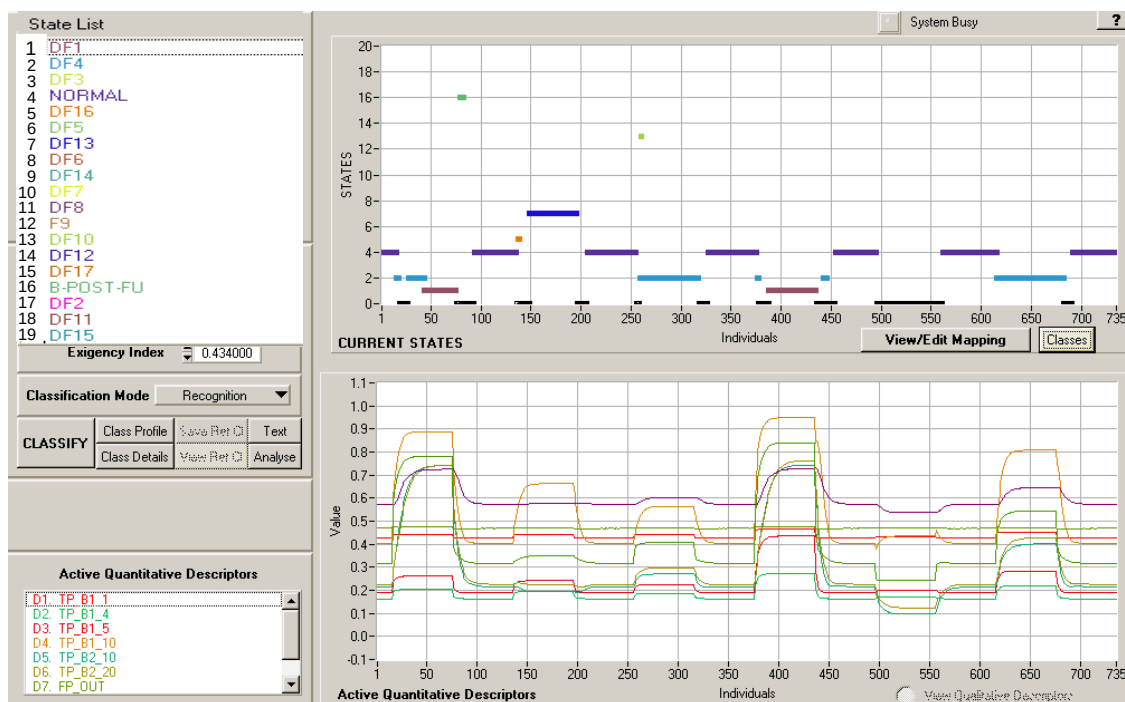


Figure 4-29 Identifications des états fonctionnels : réaction du thiosulfate

4.4 Conclusions

Nous avons développé une méthodologie générique, qui permet au final de déterminer le nombre de capteurs pertinents et leur emplacement à partir du résultat d'une classification effectuée en utilisant tous les descripteurs possibles.

La technique de classification utilisée est la technique de classification floue *LAMDA*, mais tant autre technique aurait pu être choisie du moment qu'elle fournit un profil de classes sur lequel peut s'appuyer la procédure de calcul de la quantité d'information (basé sur l'entropie) que fournit tout descripteur.

Pour la création du modèle de comportement du procédé de production de glycol simulé sur HYSYS, nous avons fait la simulation de diverses perturbations avec des dérives (changements lents) jusqu'à la valeur maximale permise de telle sorte que, la situation du procédé soit récupérable vers un état de fonctionnement. Nous avons modélisé les divers états fonctionnels du procédé par rapport à chaque défaut, et nous avons détecté de façon satisfaisante les différents états de fonctionnement (défauts avec dérives et amplitudes différentes que celles utilisées lors de la phase de conception du modèle de comportement). En revanche, il n'a pas été possible de reconnaître les défauts simultanés. Néanmoins, nous avons montré qu'il était possible de refaire un apprentissage et d'affecter un nouvel état à cette situation non présente dans la base d'apprentissage initiale.

En ce qui concerne le réacteur OPR, nous avons montré les résultats obtenus avec 2 réactions chimiques différentes : une réaction d'estérification qui est une réaction faiblement exothermique et la réaction du thiosulfate, laquelle est une réaction très exothermique. Pour concevoir le modèle de comportement, nous avons initialement simulé des perturbations de type changement brusque. À cause de cela, nous n'avons pas pu reconnaître quelques défauts de très faibles amplitudes lors de la phase de reconnaissance. De même, il a été parfois difficile de reconnaître des défauts simultanés non présents dans la base d'apprentissage initiale.

5. CONCLUSIONS ET PERSPECTIVES

Le résultat le plus important dans ce mémoire de thèse, est que l'on dispose maintenant des outils génériques qui permettent de localiser l'emplacement des capteurs qui vont permettre de caractériser les différents défauts ou états de fonctionnement d'un procédé et d'identifier correctement ces défauts en ligne, ces capteurs étant des capteurs de mesure facilement disponible sur un procédé (température, pression, débits, pourcentage d'ouverture de vanne).

Nous avons développé une méthode de placement de capteurs qui permet la redondance matérielle tout en permettant de réduire les coûts dans la conception d'une unité industrielle. Cette méthodologie s'appuie sur le concept de gain d'information qui est basé sur l'entropie de SHANNON qui permet d'effectuer la discrimination entre plusieurs situations et qui a été utilisée dans diverses approches comme la planification expérimentale. Comme le montrent les résultats dans le cas du modèle de la production de glycol construit sur le simulateur HYSYS, grâce aux mesures de signaux facilement disponibles (températures, pression, débit, pourcentage d'ouverture de vanne), il a été possible de détecter des défauts soit dans l'alimentation du procédé chimique soit dans une des unités (réacteur), avec un nombre réduit de descripteurs sélectionnés. Ces défauts ont une incidence sur la qualité des composés produits en sortie de l'installation. La détection de ces défauts de qualité a été effectuée sans avoir recours à un capteur de composition (indisponible en pratique la plupart du temps) mais en tenant compte d'une mesure d'une température dans une section correctement identifiée de la colonne de rectification. Cette information, si elle est correctement traitée permet donc de remonter à la source du dysfonctionnement : défaut sur l'alimentation principale de l'installation ou défaut sur le réacteur. Dans ce cas précis, la position des capteurs à l'intérieur de la colonne (choix des plateaux sur lesquels positionner les capteurs de température) est importante pour maximiser l'information et donc améliorer ou tout

simplement permettre ultérieurement le diagnostic. Ce que l'on peut aussi mettre en exergue c'est l'apport, en terme de richesse d'information, des variables de commande par rapport aux variables régulées qui, si les régulateurs fonctionnent correctement restent toujours dans une marge relativement proche des consignes. Il est intéressant de noter que cette mesure de discrimination a bien montré que les variables régulées des divers régulateurs ne sont pas porteuses d'information pertinente sur le procédé et qu'il vaut mieux suivre l'évolution des variables manipulées ou des autres variables dites « libres » qui portent elles beaucoup plus d'information indirecte sur la défaillance.

La validation de cette méthodologie s'est effectuée grâce à la classification en phase de reconnaissance de défauts inconnus ; phase de reconnaissance en « ligne » dans le cas du procédé de production de propylène glycol simulé sur HYSYS grâce au développement d'une interface de communication entre l'outil de diagnostic SALSA et HYSYS. Il en est ressorti que grâce à la possibilité sous HYSYS de simuler des défaillances de manière graduelle, le système de diagnostic, au travers du modèle de comportement, a pu ensuite détecter les alarmes de défauts (pré-défauts), ce qui est très important pour la sécurité et pour la récupération rapide du procédé vers un état d'opération normale, et ainsi éviter les dommages sur le produit et l'unité elle-même.

Concernant les résultats obtenus pour le réacteur « OPR », à cause du type de défaillances appliquées (très abruptes), on n'a pas pu complètement détecter les diverses alarmes de défaillances. Le résultat le plus important est que le système de diagnostic a été capable de détecter les défauts de composition en utilisant seulement 8 des 27 descripteurs disponibles sur le réacteur sans avoir recours à une mesure directe de la composition et ce malgré la faiblesse de cette défaillance. La validation sur l'unité pilote installée au Laboratoire de Génie Chimique est une des phases suivantes du projet ICSI.

Aussi bien la méthode de placement de capteurs que la méthode de diagnostic ont besoin de défaillances étudiées préalablement sur le procédé, celles-ci peuvent être générées soit

- à l'aide de simulation et alors la stratégie de placement de capteurs peut être incérée comme une étape particulière de la procédure de conception d'une unité de production. L'agrégation de défaillances est facile, par contre, rien ne permet d'assurer l'exhaustivité de la connaissance du comportement du processus.

- à partir d'historiques de défaillances et alors la méthodologie de placement de capteur peut être utilisée comme un moyen de synthétiser l'information à donner à l'opérateur (ne donner que l'information essentielle)

Quel que soit le domaine d'application, il est clair qu'il est impossible d'envisager que toutes les défaillances aient été étudiées de manière exhaustive. Il aurait fallu par exemple faire l'analyse d'autres défaillances comme celles affectant les actionneurs vannes, pompes, agitateurs ou de défauts intermittents pouvant simuler par exemple la défaillance d'une vanne de régulation. Nous avons montré qu'il était possible d'adapter le modèle de comportement au fur et à mesure de la présentation de nouvelles observations, de nouvelles situations. On peut imaginer aussi qu'un retour sur la simulation puisse permettre une nouvelle application de la méthodologie de placement de capteurs et la conception d'un nouveau modèle de comportement **en ligne** pour le diagnostic et donc l'installation de nouveaux capteurs (lorsque cela est possible sur une installation existante) : cette phase de maintenance à des fins de diagnostic pourrait être intégrée dans la vie du processus.

La détection de défauts inconnus (assignation à la classe **NIC** pendant la reconnaissance **en ligne**) a été démontrée dans certains cas. En revanche, il est arrivé qu'il ne soit pas possible de détecter un nouveau défaut : la non-détection de défauts (confondre la défaillance avec l'opération normale du processus) est un problème grave puisque ceci peut affecter la sécurité de l'environnement (opérateur, unité, produit). Un autre problème qui peut se présenter pendant la reconnaissance de défauts **en ligne**, c'est de confondre un défaut avec un autre, et de ne pas identifier l'importance de la nouvelle défaillance. Ces problèmes nous conduisent à la question suivante :

Est-il possible de concevoir une technique pour indiquer les capteurs ou descripteurs manquants dans un procédé pour en faire le diagnostic et la supervision ?

Pour permettre l'application de cette technique de manière plus générale, nous proposons comme futurs travaux de développer une méthodologie qui permette au final :

- d'extraire des informations plus pertinentes que la valeur brute du signal : pente, pics, tendances, fréquences, ... suivant les caractéristiques recherchées (défaillance de capteur, dérives, ...)
- de sélectionner des informations de type qualitative qui seraient données de manière spécifique par l'opérateur.

Pour ce faire, il est important de pouvoir juger de la validation de la classification obtenue grâce à l'interprétation des classes obtenues. L'établissement de critères géométriques ou autres sur la distribution des classes devrait permettre de mieux savoir quelle technique de classification ou algorithme utiliser (dans ce travail nous avons utilisé la fonction « binomiale » pour faire la classification de défauts sur l'exemple simulé avec HYSYS et « gauss » dans le cas du réacteur « OPR ») et ainsi aider l'expert pendant la phase de conception du modèle de comportement.

Concernant ce modèle de comportement, des travaux sont actuellement menés dans le cadre du projet ICSI en collaboration avec Mme Irène Gaillard (maître de Conférences à l'IPST, effectuant sa recherche au CERTOP (Centre d'Etudes et de Recherches « Travail, Organisation et Pouvoir », spécialiste d'ergonomie du travail). L'objectif est tout d'abord d'évaluer l'apport d'un tel outil pour les opérateurs lors de la surveillance d'unités de production complexes mais aussi d'extraire des expériences (qui devraient être menées sur un panel d'opérateurs) des caractéristiques nouvelles qu'un outil de diagnostic de ce type devraient fournir.

De manière générale, on peut dire qu'il existe encore de travaux à mener dans le domaine du diagnostic des procédés complexes. En effet, un système de diagnostic dans l'absolu devrait pouvoir accepter et mêler des informations de type mesures comme c'était le cas dans ce travail mais aussi issues de la connaissance de la physique des phénomènes mis en jeu (résidus générés grâce à des simulations dynamiques en ligne des processus) et/ou de « macro-connaissance » de type qualitative construite à partir d'une modélisation simple du procédé.

RÉFÉRENCES

- [ATSDR] *Agency for Toxic Substances and Disease Registry* disponible sur:
<http://www.atsdr.cdc.gov/>
- [Aguilar-Martin et al., 1999] Aguilar-Martin J., Waissman-Vilanova, J., Sarrate-Estruch, R., et Dahou, B. Knowledge based measurement fusion in bio-reactors, *IEEE EMTECH*, 1999.
- [Aguilar-Martin, 1980] Aguilar-Martin J., Balssa M., R.D.M. Estimation récursive d'une partition. Exemple d'apprentissage et auto-apprentissage dans **R**. Technical Report No. 880139, LAAS-CNRS, 1980.
- [Anderberg, 1973] Anderberg, M. R., Cluster Analysis for Applications. Academic Press, Inc., New York, NY., 1973.
- [Bailey, 1984] Bailey, S. J., From desktop to plant floor, a CRT is the control operators window on the process. *Control Engineering*, vol. 31, no. 6, pp. 86-90, 1984.
- [Ball et al., 1965] Ball, G. H., D. J., ISODATA, a novel method of data analysis and classification. Technical report Stanford University, Stanford, CA., 1965.
- [Basseville, 1987] Basseville, M., Benveniste A., Moustakides G. V., et Rougée A., Optimal Sensor Location for detecting Changes in Dynamical Behavior. *IEEE Transactions on automatic control*, vol. AC-32, no. 12, décembre 1987.
- [Benuzzi et al, 1991] Benuzzi A., Zaldivar J.M., Safety of chemical batch reactors and storage tanks. Eds Kluwer Academic Publishers, Dordrecht and Norwell, Massachusetts, pp. 439, 1991.

-
- [**Bezdek, 1981**] Bezdek J. C. Pattern recognition with fuzzy objective function algorithms, Plenum Press, New York, 1981.
- [**BLS, 1998**] Bureau of Labor Statistics. Occupational injuries and illnesses in the united states by industry. Washington, DC: Government Printing Office, 1998.
- [**Chan et al., 1989**] Chan, M., Aguilar-Martin J., J. Carreté, N., Celsis, P., et Vergnes, J. M., Classification techniques for feature extraction in low resolution tomographic evolutives images :application to cerebral blood flow estimation. In 12th Conf. GRETI, 1989.
- [**CHEM**] Projet CHEM disponible sur : <http://www.chem-dss.org/>
- [**De Luca et al., 1972**] De Luca, A., Termini, S., A definition of a nonprobabilistic entropy in the setting of fuzzy sets theory. Inform. Contr. 20 (4), pp. 301-312, 1972.
- [**Devatine et al., 2003**] Devatine A., Prat L., Cognet P., Cabassud Michel, Gourdon C., Chopard F., Process Intensification, performances evaluation of a new concept « Open Plate Reactor » ECCE 4, 21-25, Granada, Spain, sept. 2003.
- [**Diday, 1973**] Diday, E. The dynamic cluster method in non-hierarchical clustering. J. Comput. Inf. Sci. 2, pp. 61-88, 1973.
- [**Dochain et al., 1996**] Dochain D., Tali- Maamar, Babary J. P., influence of the sensor location on the practical observability of a fixed-bed bioreactor, Proc. 13th IFAC World congress, San Francisco, CA, pp. 491-496, 1996.
- [**Dochain et al., 2001**] Dochain Denis et al., Automatique des Bio procédés. Hermes sciences publications 2001.
- [**Dubes, 1987**] Dubes, R. C.. How many clusters are best ? – an experiment. Pattern Recognition, vol. 20, no. 6, pp. 645-663, novembre 1987.
- [**Dunn, 1973**] Dunn J. C.. A fuzzy Relative of the ISODATA process and its use in detecting compact well-separated clusters, Journal of cybernetics, vol. 3, pp. 32-57, 1973.
- [**Elgue et al., 2004**] Elgue S., Devatine A., Prat L., Cognet P., Cabassud M., Gourdon C., Chopard, F., Dynamic simulation of a novel intensified reactor, Escape 14, Lisbonne, 16-19 mai, 2004.
-

-
- [**Fernandez, 2002**] Fernandez Pellon-Zambrano, R., Desarrollo de algoritmos para la clasificación de secuencias, tesis de maestría, Universidad de las Américas, Puebla (UDLAP), 2002.
- [**Frank et al., 2000**] Frank P. M., Ding S. X., et Marcu T., Model-Based fault diagnosis in technical process. Transactions of the Institution of Measurement and Control, vol. 22 no. 1, pp. 57-101. 2000.
- [**Galvan et al., 1996**] Galvan M., Zaldivar J., Hernández H., Molga E., The use of neural networks for fitting complex kinetic data, Computers and Chemical Engineering, vol. 20, no. 12, pp. 1451-1465, août 1996.
- [**Garcia et al., 1998**] Garcia E., Morant F. “Aplicaciones de los modelos de eventos discretos obtenidos mediante el operador de composición sincronía” XIX Jornadas de Automatica, Madrid., pp. 199-204, 1998.
- [**Garcia et al., 1999**] Garcia E., Morant F., “Diagnostico de fallos modular y control supervisor para un proceso químico de mezclas” XX Jornadas de Automatica, Salamanca, pp. 267-272, 1999.
- [**Gower et al., 1969**] Gower, J. C. and Ross, G. J. S. Minimum spanning trees and single-linkage cluster analysis. *Appl. Stat.* 18, 54-64, 1969.
- [**Haugwitz et al., 2004**] Haugwitz S. and Hagander P., Process Control of an Open Plate Reactor, In proceedings of Reglermöte, May 2004.
- [**Haugwitz et al., 2005**] Haugwitz S. and Hagander P., Temperature control of an Open Plate Reactor, IFAC world Congress, Prague, Czech Republic, July 2005.
- [**Himmelblau et al., 1978**] Himmelblau, D. M. Fault Detection and Diagnosis in Chemical and petrochemical processes. Amsterdam: Elsevier press, 1978.
- [**HYSYS, 1995**] Manuel *Reference* Hyprotech, version 1.0, 1995
- [**HYSYS, 2002a**] Manuel *Customization Guide* version 3.1, Hyprotech 2002, disponible sur: <http://itll.colorado.edu/HYSYS/Doc/HYSYS/CustomGd.pdf>
- [**HYSYS, 2002b**] Manuel *Get Started* version 3.1, Hyprotech 2002 disponible sur: <http://itll.colorado.edu/HYSYS/Doc/HYSYS/GetStart.pdf>
-

-
- [**Isermann, 1984**] Isermann R. Process fault diagnosis based on modelling and estimation methods- A survey. *Automatica*, vol 20, pp. 387-404, 1984.
- [**Jain et al., 1988**] Jain A. K. et Dubes, R. C., *Algorithms for clustering Data* Prentice- Hall advanced references series. Prentice-Hall, Inc., Upper Saddle River, N. J., 1988.
- [**Jain et Dubes,1988**] Jain, A. K. and Dubes, R. C. *Algorithms for Clustering Data*. Prentice-Hall advanced reference series. Prentice-Hall, Inc., Upper Saddle River, NJ., 1988.
- [**Kempowsky, 2004a**] Kempowsky T. « Surveillance de procédés à base de méthodes de classification : Conception d'un outil d'aide pour la détection et le diagnostic des défaillances ». Thèse de doctorat, Institut National des Sciences Appliquées, Toulouse, France, décembre 2004.
- [**Kempowsky, 2004b**] Kempowsky T., SALSA user's manual. Rapport LAAS-CNRS no. 04160, 2004.
- [**King, 1967**] King, B. Step-wise clustering procedures, *Journal of the American Statistical Association*, vol. 69, pp. 86-101, 1967.
- [**Kninchin, 1957**] Kninchin A. I., *Mathematical Foundations of Information Theory*, Dover publications, New York, 1957.
- [**Koodmen, 1989**] Koodmen F., *LAMDA : Logiciel d'Analyse de Données Multidimensionnelles par Apprentissage et Reconnaissances de Formes*. Manuel de reference et d'utilisation. LAAS-CNRS, Toulouse. 1989.
- [**Kramer et al., 1993**] Kramer M. A., et Mah R. S. H., Model-Based Monitoring. In: Rippin D., Hale J., et Davis J. (Eds), *Proceedings of the second International Conference on "foundations of computer-aided process operations" (FOCAPO)*, pp. 45-68,1993.
- [**Kumar et al., 1978**] Kumar S., Seinfeld J. H., Optimal location of measurements in tubular reactors, *Chemical Engineering Science*, vol. 33, pp. 1507-1516, 1978.
- [**Lafortune, 1991**] S. Lafortune et E. Chen, "A relational algebraic approach to the representation and analysis of discrete-event systems" in *proceedings of the American Control Conference*, Boston, MA, pp. 2893-2898, 1991.
- [**Laser, 2000**] Laser, M.. Recent safety and environmental legislation. *Trans IchemE*, vol. 78 (B), pp. 419-422, 2000.
-

-
- [Lees, 1996] Lees, F. P. Les prevention in process industries : hazard identification, assesment and control. London:Butterworth-Heinemann, 1996.
- [Lu et al., 1978] Lu, S. Y. et Fu, K. S. A sentence-to sentence clustering procedure for pattern analysis. *IEEE . IEEE Transactions on systems, man and cybernetics*, vol. 8, pp. 381-389, 1978.
- [Mamdani, 1977] Mamdani E., “application of fuzzy logic to approximate reasoning using linguistic system.” *Fuzzy sets and systems*, vol. 26, pp. 1182-1191, 1977.
- [Mao et al., 1996] Mao, J., et Jain, A. K., A self-organizing network for hyperellipsoidal clustering (HEC). *IEEE Transactions on Neural Networks*, vol. 7, pp.16-29, 1996.
- [Maquin et al., 2000] Maquin Didier et al. Diagnostic des systèmes linéaires. Hermes Science Europe, 2000.
- [McQueen, 1967] McQueen, J. Some methods for classification and analyis of multivariate observations. In proceedings of the Fifth Berkeley Symposium on Mathematical Statistics and Probability, pp. 281-297, 1967.
- [Murtagh, 1984] Murtagh, F., A survey of recent advances in hierarchical clustering algorithms which use cluster centers. *Comput. J.*, vol. 26, pp. 354-359, 1984.
- [NSCI, 1999] National Safety Council Injury facts, Chicago: National Safety Council. 1999.
- [Orantes, 2003] Orantes A., Le Lann M.V., Jempowsky T., Aguilar-Martin J., Detección de fallas de un proceso químico complejo por abstracción de información. 10th International Congress on Computer Science Research (CICC’03), Oaxtepec, Mexico, pp. 314-322, octobre, 2003.
- [Orantes, 2004] Orantes A., Le Lann M. V., Aguilar-Martin J., Classification as an aid tool for the selection of sensors used for fault detection and isolation. Rapport LAAS N°04527, octobre 2004.
- [Papadimitriou, 2003] Papadimitrou, C., Optimal sensor placement methodology for parametric identification of structural systems. *Journal of sound and vibration*, vol. 278, pp. 923-947, 2003.
- [Piera et al., 1989] Piera N., Desroches P., Aguilar-Martin J., *LAMDA : An incremental conceptual clustering method*. Rapport du LAAS No. 89420. Laboratoire d’analyse et d’architecture des systèmes du CNRS, décembre 1989.
-

-
- [**Piera-carreté, 1990**] Piera-carreté, N., Desroches, P., et Aguilar-Martin, J. Variation points in pattern recognition. *Pattern Recognition Letters*, vol. 11, pp. 519-524, 1990.
- [**Qureshi et al., 1980**] Qureshi, Z. H., Ng T.S., Goodwin G. C., Optimum experimental design for identification of distributed parameter systems, *International Journal of Control*, vol. 31, pp. 21-29, 1980.
- [**Rakoto-Ravalontsalama et al., 1995**] Rakoto-Ravalontsalama N., Aguilar-Martin J. *Supervision de processus à l'aide du système expert G2*. 1995.
- [**Ramadge, 1987**] Ramadge P. J. et Wonham W. M., "Modular Feedback logic for discrete event systems," *SIAM Journal of Control and Optimization*, vol. 25, no. 5, pp. 1202-1218, mai 1987.
- [**Sampath et al., 1994**] Sampath M., Sengupta R., Lafortune S., Sinnamohideen K. and Teneketzis D., "Failure diagnosis using discrete-event models" *IEEE Transactions on Control Systems*, vol. 4, no. 2, pp. 105-124, mars 1994.
- [**Sampath et al., 1995**] Sampath M., R. Sengupta, S. Lafortune, K. Sinnamohideen and D. Teneketzis, "Dioagnosability of discrete-event systems" *IEEE Transactions on Automatic Control*, vol. 40, no. 9, pp. 1555-1575, 1995.
- [**Sargousse et al., 1999**] Sargousse A., Le Lann J.M., Joulia X., Jourda L., 1999, DISCo : un nouvel environnement orienté objet, MOSIM'99, 6-8 octobre, Annecy, France.
- [**Shannon, 1948**] Shannon C. E., A mathematical theory of communication, *Bell System Technical Journal*, vol. 27, pp. 379-423 and 623-656, july and October, 1948.
- [**Silviu, 1977**] Silviu Guiasu. *Information Theory with Applications*, McGraw-Hill, 1977.
- [**Sneath, 1973**] Sneath, P. H. A. et Sokal, R. R., *Numerical Taxonomy*. Freeman, London, UK, 1973.
- [**Symon, 1977**] Symon, M. J. Clustering criterion and multi-variate normal mixture. *Biometrics* 77, pp. 35-43, 1977.
- [**Takagi, 1985**] Takagi T., Sugeno M., Fuzzy identification of systems and its application to modelling and control. *IEEE Transactions on systems, man and cybernetics*, vol. 15, no. 1, pp. 116-132, 1985.
-

-
- [**Toscano, 2005**] Toscano R., Commande et diagnostic des systèmes dynamiques: Modélisation, analyse, commande par PID et par retour d'état, diagnostic. Ellipses édition Marketing, S. A., 2005.
- [**Ubrich et al., 1999**] Ubrich O., Srinivasan B. Bonvin D., Stoessel F., Optimisation of a semi-batch reaction system under safety constraints. 5th European Control Conference, 1999.
- [**Ubrich et al., 2001**] Ubrich O., Srinivasan B. Bonvin D., Stoessel F., Optimal Feed Profile for a second order reaction in a semi-batch reactor under safety constraints. Experimental study. Lausanne, Switzerland, 2001.
- [**Venkatasubramanian et al., 2003**] Venkatasubramanian V., Rengaswamy R. Yin K., Kavuri S., A review of process fault detection and diagnosis Part I: Quantitative model-based methods, Computer and Chemical Engineering, 2003.
- [**Waissman-Vilanova, 2000**] Waissman-Vilanova J., *Construction d'un modèle comportemental pour la supervision de procédés : application à une station de traitement des eaux*. Thèse de doctorat, Institut National Polytechnique de Toulouse, France, 2000.
- [**Walter et al.,]** Walter E. and Pronzato L., Identification of Parametric Models from Experimental Data, Masson 1997.
- [**Ward, 1963**] Ward, J. H. Jr. Hierarchical grouping to optimise an objective function. Journal of the American Statistical Association, vol. 58, 236-244, 1963.
- [**Watanabe, 1985**] Watanabe, K., Sasaki M., et Himmelblau, "Determination of optimal measuring sites for fault detection of non-linear systems, Int. J. Syst. Sci., vol. 16, pp. 1345-1364, Nov. 1985.
- [**Winter, 1992**] Winter, "Computer Aided Process Engineering: The Evolution Continues", Chemical Engineering Progress, pp. 77, Feb. 1992.
- [**Winter, 1996**] Winter, "The State of Computing in the CPI 1996: What's going on in the classes of software most important to chemical engineers", 1996.
- [**Wouwer et al., 2000**] Wouwer A. V, Point N., Porteman E., Remy M., An approach to the selection of optimal sensor locations in distributed parameter systems, Journal of process control, vol. 10, pp. 291-300, 2000.
-

- [Zadeh, 1965]** Zadeh, L. Fuzzy sets. *Information Control*, vol. 8, pp. 338-353, 1965.
- [Zadeh, 1978]** Zadeh L. A., "Fuzzy sets a basis for a theory of possibility." *Fuzzy sets and systems*, vol. 1, pp. 3-28, 1978.
- [Zahn, 1971]** Zahn, C. T.. Graph-Theoretical methods for detecting and describing gestalt clusters. *IEEE Trans. Comput. C-20* (Apr.), pp. 68-86., 1971.
- [Zamproga et al., 2004]** Zamproga E., Barolo M., et Seborg D. E., Optimal selection of soft sensor inputs for batch distillation columns using principal component analysis. *Journal of process control*, vol. 15, pp. 39-52, april 2004.
-

ANNEXE A.

A. EXEMPLE DE CALCUL DE LA METHODE LAMDA

Pour mieux comprendre la méthode LAMDA, on donne dans ce qui suit un exemple de classification en mode d'apprentissage supervisé (avec classes « professeur ») utilisant des descripteurs quantitatifs. A titre d'exemple, on a choisi la fonction « Binomiale » pour calculer les MADs et le produit comme opérateur logique pour le calcul du GAD.

Tout d'abord, présentons le contexte de la classification (voir tableau A-1) :

Nombre de descripteurs = 3

Nombre d'éléments déjà classés = 5

Nombre de classes « professeur » = 2 (trois éléments dans la classe 1, deux éléments dans la classe 2).

Tableau A-1 Valeurs brutes des descripteurs d_i pour chaque élément

Éléments	d_1	d_2	d_3	Classes professeur
X^1	0	2	16	1
X^2	3	2	5	1
X^3	10	1	20	2
X^4	4	6	10	2
X^5	5	6	0	1

La première étape consiste à normaliser les éléments par rapport aux valeurs maximale et minimale de chaque descripteur suivant la relation :

$$\hat{x}_j = \frac{x_j - x_{min}}{x_{max} - x_{min}} \quad (A.1)$$

Les valeurs maximale et minimale pour chaque descripteur sont :

	d_1	d_2	d_3
Maximum (x_{max})	10	6	20
Minimum (x_{min})	0	1	0

Le tableau A-2 montre les valeurs normalisées des descripteurs pour chaque élément:

Tableau A-2 Valeurs normalisées des descripteurs pour chaque élément

Éléments normalisés	d ₁	d ₂	d ₃	Classes professeur
X ¹	0	0.2	0.8	C ₁
X ²	0.3	0.2	0.25	C ₁
X ³	1	0	1	C ₂
X ⁴	0.4	1	0.5	C ₂
X ⁵	0.5	1	0	C ₁

L'étape suivante est le calcul des paramètres (les ρ 's) des classes « professeur » et de la classe *NIC*, en évaluant la moyenne des valeurs normalisées :

Pour la classe 1 :

$$\rho_1(d_1/C_1) = (0 + 0.3 + 0.5) / 3 = 0.267$$

$$\rho_2(d_2/C_1) = (0.2 + 0.2 + 1) / 3 = 0.467$$

$$\rho_3(d_3/C_1) = (0.8 + 0.25 + 0) / 3 = 0.35$$

Les paramètres de la classe 1 sont donc : **C1=[0.267, 0.467, 0.35]**

Pour la classe 2 :

$$\rho_1(d_1/C_2) = 0.7$$

$$\rho_2(d_2/C_2) = 0.5$$

$$\rho_3(d_3/C_2) = 0.75$$

Les paramètres de la classe 2 sont : **C2=[0.7, 0.5, 0.75]**

Pour la classe NIC ou classe zéro :

En utilisant la fonction *Binomiale* comme fonction d'appartenance, les paramètres de la classe zéro sont : **C₀=[0.5, 0.5, 0.5]**

L'objectif est de classer un nouvel élément X^6 sur la base des paramètres des classes déjà existantes.

Soit cet élément X^6 avec les valeurs normalisées suivantes :

	d ₁	d ₂	d ₃
X ⁶	0.1	0.1	0.9

On calcule tout d'abord les degrés d'appartenance marginale (MAD) par rapport à chaque descripteur :

$$\mu(x_j/C_i) = \rho_{i,j}^{x_j} (1 - \rho_{i,j})^{(1-x_j)}$$

Où $\rho_{i,j}$ est le paramètre j de la classe i , x_j est la valeur du descripteur j de ce nouvel élément, et $\mu(x_j/C_i)$ est le degré d'appartenance du descripteur j du nouvel élément à la classe i existante.

Pour la classe vide (NIC), dans ce cas, le degré d'appartenance sera toujours de $0.5 \forall x_j$:

$$\mu(x_j/C_0) = (0.5)^{x_j} (0.5)^{(1-x_j)} = 0.5$$

Le résultat des calculs des MAD's par rapport à l'élément X^6 sont les trois vecteurs suivants :

$$\mu(X^6/C_1) = [0.66, 0.52, 0.37]$$

$$\mu(X^6/C_2) = [0.33, 0.5, 0.67]$$

$$\mu(X^6/C_0) = [0.5, 0.5, 0.5]$$

On calcule ensuite les degrés d'appartenance globale (GAD) de l'élément X^6 en utilisant le connectif « produit » :

$$GAD(X/C_i) = \prod_{j=1}^{nd} \mu(x_j/C_i)$$

$$GAD(X/C_i) = \mu(x_1/C_i) \mu(x_2/C_i) \mu(x_3/C_i)$$

Où nd est le nombre de descripteurs. On obtient ainsi les trois degrés d'appartenance globale :

$$GAD(X^6/C_1) = 0.130$$

$$GAD(X^6/C_2) = 0.110$$

$$GAD(X^6/C_0) = 0.125$$

L'élément X^6 appartient à la classe dont le GAD est maximal, c'est-à-dire, à la classe C_1 :

$$GAD(X^6/C_1) = 0.130$$

Les appartenances calculées par l'application de différents connectifs et la décision d'attribution à la classe d'appartenance maximale sont données ci-dessous .:

Tableau A-3 Résultat de décision

		Appartenances			DÉCISION
		NIC	C1	C2	
Connectifs	Produit	0.125	0.13	0.11	C1
	Somme probabiliste	0.75	0.899	0.889	C1
	Minimum	0.5	0.372	0.327	NIC
	Maximum	0.5	0.663	0.672	C2

À cause de l'incorporation du nouvel élément à la classe 1, l'actualisation des paramètres associés aux descripteurs quantitatifs de cette classe est faite par la formule itérative de la moyenne :

$$\rho_{i,j} = \rho_{i,j} + \frac{(x_j - \rho_{i,j})}{N+1} \quad (\text{A.2})$$

où $\rho_{i,j}$ est la nouvelle valeur du paramètre du descripteur j de la classe i et N le nombre d'éléments dans la classe.

$$\rho(x_1/C_1) = 0.267 + \frac{1}{N+1}(0.1 - 0.267) = 0.184$$

$$\rho(x_2/C_1) = 0.284$$

$$\rho(x_3/C_1) = 0.625$$

Alors, la représentation de la classe C_1 actualisée est : $C_1 = [0.18, 0.28, 0.63]$

Par ailleurs, si l'élément X^6 avait appartenu à la classe NIC , dans le cas où le minimum aurait été choisi comme connectif, une nouvelle classe aurait été créée en modifiant les paramètres de la classe NIC selon la relation suivante :

$$\rho_{i,j} = 0.5 + \frac{(x_j - 0.5)}{N_0 + 1} \quad (\text{A.3})$$

où N_0 représente le nombre initial d'éléments dans la classe zéro.

Dans notre exemple, la nouvelle classe créée C_3 aurait eu comme paramètres :

$$\rho(x_1/C_3) = 0.5 + \frac{1}{2}(0.1 - 0.5) = 0.3$$

$$\rho_{3,2} = \rho(x_2/C_3) = 0.3$$

$$\rho_{3,3} = \rho(x_3/C_3) = 0.7$$

La représentation de la nouvelle classe C_3 serait alors : $C_1 = [0.3, 0.3, 0.7]$

ANNEXE B.

B. PROGRAMMATION SOUS HYSYS

Le programme de récupération des mesures sur HYSYS et de génération des défaillances sous *Visual Basic* est le suivant :

```
' L'exécution du macro commence toujours en « Sub Main »
' -----
Sub Main
' Acquière une interface pour le fichier de simulation (case) active
  Dim simcase As Object
  Set simcase = ActiveCase
  If simcase Is Nothing Then
    ' Si fichier de simulation n'est pas active, préviens à
    ' l'utilisateur
    MsgBox "No HYSYS case is open."
    ' et termine...
    End
  End If
' -----
  '
  Open "C:\SALSA05v0-1\DataTr\d_out.dat" For Output As #2
  Close #2
  Simcase.Solver.Integrator.Reset ' réinitialise l'intégrateur de
HYSYS

Do
' -----
  'Commence la modélisation dynamique sur HYSYS
  ' d'abord il termine la simulation, après il continue l'exécution du
  ' programme
  Simcase.Solver.Integrator.RunFor(1,"minutes")
' -----
' -----
' Récupération de toutes les variables sur HYSYS :
  ' item(1)= réacteur
  ' item(2)= colonne de distillation
  ' item(3)= Contrôleur de température du réacteur
  ' item(4)= Contrôleur de niveau du réacteur

  ' Récupération de la pression du réacteur
  Var1 = Simcase.Flowsheet.operations.Item(1).VesselPressure
  ' Récupération de la température du réacteur
  Var2 = Simcase.Flowsheet.operations.Item(1).VesselTemperature
  ' Récupération Du niveau du réacteur
  Var3 = Simcase.Flowsheet.operations.Item(1).liquidPercentLevel
  ' Récupération de la variable de sortie du Contrôleur TC
  Var4 = Simcase.Flowsheet.operations.Item(3).OpValue
  ' Récupération de la variable de sortie du Contrôleur LC
  Var5 = Simcase.Flowsheet.operations.Item(4).OpValue
  ' Récupération de la température de sortie Du débit utilité
```

```

' dans le réacteur
Var6 = Simcase.Flowsheet.operations.Item(1).UtilityOutletTemperature
' Récupération des températures sur les plateaux de la colonne
Set var77 = Simcase.Flowsheet.Operations.Item(2).ColumnFlowsheet
Temp_Profile = var77.Temperatures
' Récupération du débit d'énergie « RebDuty » sur le bouilleur
Var17 = Simcase.Flowsheet.energystreams.Item(2).HeatFlow.Getvalue("KJ/h")
Set var88= Simcase.Flowsheet.Operations.Item(2).ColumnFlowsheet.operations
' Récupération de la variable de sortie du Contrôleur PC dans la
' colonne
var18 = var88.Item(4).OP
' Récupération de la pression du condenseur dans la colonne
var19 = var88.Item(1).VesselPressure
Set var99 = var77.operations 'pfd COL1
' Récupération de la température de sortie d'utilité dans le
' condenseur dans la colonne
var20 = var99.Item(4).ControlValve.OutletTemperature
' Récupération de la température sur le débit de glycol
var21 = simcase.Flowsheet.materialstreams("Glycol").TemperatureValue
' -----
' assignation des températures de la colonne à des autres variables
var77_Temp = Temp_Profile(1)
var88_Temp = Temp_Profile(2)
var99_Temp = Temp_Profile(3)
var10_Temp = Temp_Profile(4)
var11_Temp = Temp_Profile(5)
var12_Temp = Temp_Profile(6)
var13_Temp = Temp_Profile(7)
var14_Temp = Temp_Profile(8)
var15_Temp = Temp_Profile(9)
var16_Temp = Temp_Profile(10)

'Définir les variables avec formats du type : « ##.### »
var11 = Format$(var1,"#.000")
var22 = Format$(var2,"#.000")
var33 = Format$(var3,"#.000")
var44 = Format$(var4,"#.000")
var55 = Format$(var5,"#.000")
var66 = Format$(var6,"#.000")

var77 = Format$(var77_Temp,"#.000")
var88 = Format$(var88_Temp,"#.000")
var99 = Format$(var99_Temp,"#.000")
var1010 = Format$(var10_Temp,"#.000")
var1111= Format$(var11_Temp,"#.000")
var1212 = Format$(var12_Temp,"#.000")
var1313 = Format$(var13_Temp,"#.000")
var1414 = Format$(var14_Temp,"#.000")
var1515 = Format$(var15_Temp,"#.000")
var1616= Format$(var16_Temp,"#.000")

```

```

var1717 = Format$(var17,"#.000")
var1818 = Format$(var18,"#.000")
var1919= Format$(var19,"#.000")
var2020 = Format$(var20,"#.000")
var2121 = Format$(var21,"#.000")
' -----
' COMMUNICATION AVEC L'OUTIL SALSA

' Si existe un fichier appelé: « d_out.dat » créé par l'outil SALSA, alors
' engendre un nouveau fichier « d_in.dat » en contenant les valeurs des
' descripteurs à classer, et supprime le fichier « d_out.dat »
' S'il n'existe pas, il l'attend jusqu'à sa création

F$ = Dir$("C:\SALSA05v0-1\DataTr\d_out.dat")
If F$ <> "d_out.dat" Then
    Do
        F$ = Dir$("C:\SALSA05v0-1\DataTr\d_out.dat")
    Loop Until F$ = "d_out.dat"
    Kill("C:\SALSA05v0-1\DataTr\d_out.dat")
    Open "C:\SALSA05v0-1\DataTr\d_in.dat" For Output As #1
    Print #1, var11;" "; var44;" "; var55;" "; var66;" "; var1616;
    " "; var1818;" "; var2020;
    Close #1
Else
    Kill("C:\SALSA05v0-1\DataTr\d_out.dat")
    Open "C:\SALSA05v0-1\DataTr\d_in.dat" For Output As #1
    Print #1, var11;" ";var44;" ";var55;" ";var66;" "; var1616;" ";
    var1818;" "; var2020;
    Close #1
End If
' -----
' Acquérir le temps actuel de simulation
Tiempo = Simcase.Solver.Integrator.GetTime("minutes")
' -----
' GÉNÉRATION DE DÉFAILLANCES SIMULTANÉES

Select Case Tiempo
' Opération normale du processus
    Case 0 To 20
' Défaillance du débit d'oxyde et de température de refroidissement dans
le ' réacteur
        Case 20 To 120
            oxy=
simcase.Flowsheet.materialstreams.item(0).MolarFlow.getvalue
("kgmole/h")
            oxy2 = oxy - 0.0604
            Simcase.Flowsheet.materialstreams.item(0).MolarFlow.SetValue
oxy2, "kgmole/h"

```



```

        Cool_Inlet=
Simcase.Flowsheet.operations.Item(1).UtilityInletTemperature
    'Rampa Cool 13
        Cool_Inlet2 = Cool_Inlet - 0.02

        Simcase.Flowsheet.operations.Item(1).UtilityInletTemperature.setvalu
e Cool_Inlet2, "C"

' Application constante des défaillances
    Case 120 To 250
        'Simcase.Flowsheet.materialstreams.item(0).MolarFlow.SetValue
oxy2, "kgmole/h"

' Return à l'opération normale du processus
    Case 250 To 400
        Simcase.Flowsheet.materialstreams.item(0).MolarFlow.SetValue
68.04, "kgmole/h"

        Simcase.Flowsheet.operations.Item(1).UtilityInletTemperature.setvalu
e 15, "C"

    Case Else
        End      ' Termine Programme

End Select      ' Termine CASE
'-----
Loop Until False
End Sub

```

ANNEXE C.

C. PROPRIETES DE L'ENTROPIE

Dans ce qui suit, nous donnons les propriétés de l'entropie.

Soient $(p_1, p_2, \dots, p_n)$ et $(q_1, q_2, \dots, q_n)$ deux lois de probabilité, alors : $\sum_{i=1}^n p_i \log \frac{q_i}{p_i} \leq 0$

En effet $\forall x > 0$ on a : $\ln x \leq x-1$ (voir la Figure C-1) :

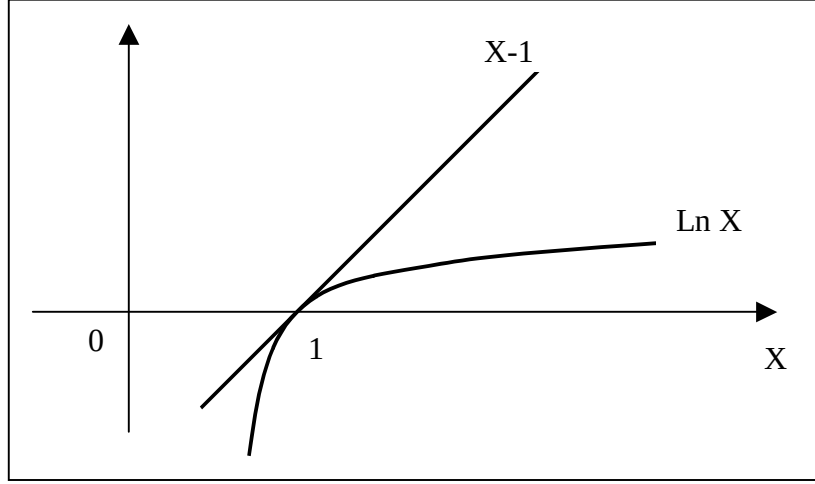


Figure C-1 La fonction logarithme

D'où $\ln \frac{q_i}{p_i} \leq \frac{q_i}{p_i} - 1$, soit $\log \frac{q_i}{p_i} \leq \frac{1}{\ln a} \left(\frac{q_i}{p_i} - 1 \right)$

Donc $\sum_{i=1}^n p_i \log \frac{q_i}{p_i} \leq \frac{1}{\ln a} \sum_{i=1}^n p_i \left(\frac{q_i}{p_i} - 1 \right) = \frac{1}{\ln a} \left(\sum_{i=1}^n q_i - \sum_{i=1}^n p_i \right) = \frac{1}{\ln a} (1-1) = 0$

Propriété 1

L'Entropie d'une variable aléatoire X à n valeurs possibles est maximale et vaut $\log(n)$ lorsque la loi de X est uniforme.

Il suffit d'appliquer le lemme précédent avec $q_1 = q_2 = \dots = q_n = \frac{1}{n}$, ainsi

$$-\sum_{i=1}^n p_i \log p_i \leq -\sum_{i=1}^n p_i \log \frac{1}{n} = -\sum_{i=1}^n \frac{1}{n} \log \frac{1}{n} = -\log n \quad (\text{C.1})$$

L'incertitude sur X est la plus grande si toutes les valeurs possibles ont la même probabilité de se réaliser.

Propriété 2

L'Entropie augmente lorsque le nombre de valeurs possibles augmente.

En effet, soit X à valeurs possibles $\{x_1, x_2, \dots, x_n\}$ de loi $(p_1, p_2, \dots, p_n)$. Supposons que la valeur x_k de probabilité p_k soit fragmentée en deux valeurs y_k et z_k de probabilités α_k, β_k , avec $\alpha_k + \beta_k = p_k$, $\alpha_k \neq 0$ et $\beta_k \neq 0$. Alors, l'entropie de la nouvelle variable aléatoire X' ainsi obtenue s'écrit :

$$H(X') = H(X) + p_k \log p_k - \alpha_k \log \alpha_k - \beta_k \log \beta_k$$

d'où

$$H(X') - H(X) = (\alpha_k + \beta_k) \log p_k - \alpha_k \log \alpha_k - \beta_k \log \beta_k \quad (\text{C.2})$$

$$H(X') - H(X) = \alpha_k \log p_k + \beta_k \log p_k - \alpha_k \log \alpha_k - \beta_k \log \beta_k \quad (\text{C.3})$$

Or, la fonction logarithme étant strictement croissante, on a :

$$\log p_k > \log \alpha_k \quad \text{et} \quad \log p_k > \log \beta_k$$

$$\text{soit } H(X') - H(X) > 0, \text{ c'est à dire : } H(X') > H(X)$$

Propriété 3

L'Entropie est une fonction convexe de $(p_1, p_2, \dots, p_n)$, en effet

$$H(X) = - \sum_{i=1}^n p_i \log p_i = \sum_{i=1}^n \text{gof}_i(p_1, p_2, \dots, p_n) \quad (\text{C.4})$$

Avec f_i l'application projection sur l'axe i :

$$\begin{array}{ll} f_i : [0,1]^n \rightarrow [0,1] & \text{et} \quad g : [0,1] \rightarrow \mathbb{R}^+ \\ (p_1, p_2, \dots, p_n) \rightarrow p_i & p \rightarrow -p \log p \end{array}$$

g est $\cap$ convexe car

$$g'(p) = -\log p - \frac{1}{\ln 2} \frac{p}{p} = -\frac{1}{\ln 2} - \log p \quad (\text{C.5})$$

et

$$g''(p) = -\frac{1}{p} < 0 \quad (\text{C.6})$$

Comme f_i est une forme linéaire, gof_i est convexe $\cap$ et $H(X)$ est convexe $\cap$ car somme de fonctions convexes.

ANNEXE D.

D.DEFAUTS SIMULTANES

Dans ce qui suit, nous donnons les résultats obtenus lors de la reconnaissance de divers scénarii de défauts simultanés, à partir du modèle de comportement avec 7 descripteurs (voir le Tableau 4-N).

Le scénario utilisé pour ce test est la superposition de deux défauts :

- a) Hausse de la température du liquide de refroidissement du réacteur : la valeur de la pente d'incrément de température est de $0.035\text{ }^{\circ}\text{C/min}$ faisant passer la température de sa valeur nominale ($15\text{ }^{\circ}\text{C}$) jusqu'à la valeur finale de la défaillance ($22\text{ }^{\circ}\text{C}$).
- b) Hausse de la vitesse de réaction : ce défaut est engendré en modifiant le facteur de fréquence de la loi d'Arrhenius. Le facteur de fréquence est modifié depuis sa valeur nominale ($1.70\text{E}13$) et passe à $2.00\text{E}+13$.

Scénario 1 : Dans ce scénario le défaut : hausse de la vitesse de réaction, est créé pendant l'application du défaut : hausse de la température du liquide de refroidissement du réacteur, et supprimé après avoir éliminé ce dernier défaut (voir la Figure D-1).

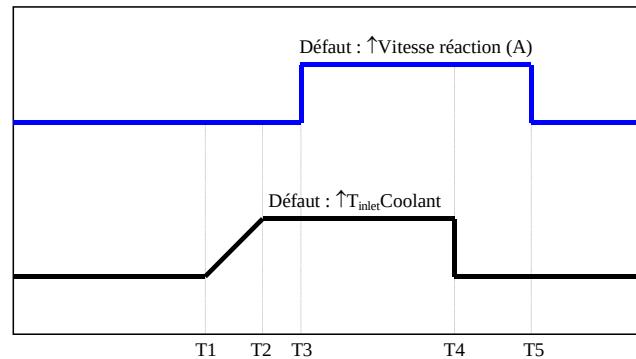


Figure D-1 Scénario 1 : défauts simultanés

Les résultats de l'identification des états fonctionnels du procédé sont rapportés sur la Figure D-2. On peut constater qu'il y a création de 6 états fonctionnels et de l'état zéro ou de non-reconnaissance. Dans ce cas, l'état 10 correspond à l'état normal du procédé (voir le tableau 4-N donnant les états). L'état 11 correspond à l'état H-FREQV (hausse de la vitesse de réaction par modification du facteur de fréquence), tandis que l'état 14 est une alarme de hausse de température du fluide de refroidissement du réacteur. Les autres états (états transitoires) sont dus au comportement oscillatoire des signaux.

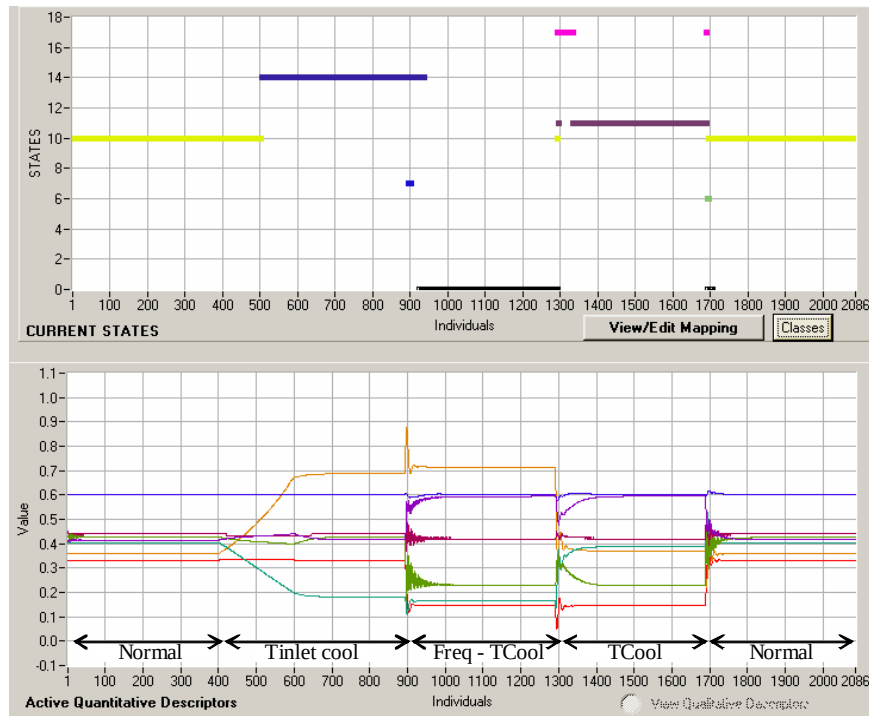


Figure D-2 Identification d'états de défauts simultanés : scénario 1

L'état zéro représente les défauts simultanés et, il n'est pas reconnu, puisqu'il n'a pas été considéré pendant la construction du modèle de comportement. Pour le reconnaître, il faut faire un apprentissage supervisé (*supervised learning*) pour affecter la classe NIC (*Non Information Classe*) à une classe représentant ces défauts simultanés.

Finalement, après cet apprentissage on peut faire la reconnaissance des états, et on peut constater la bonne reconnaissance de l'état numéro 18 (c'est un nouvel état créé) pour représenter la nouvelle défaillance (Figure D-3). Nous pouvons constater dans le Tableau 4-N que l'état 18 n'existait pas dans le modèle construit préalablement.

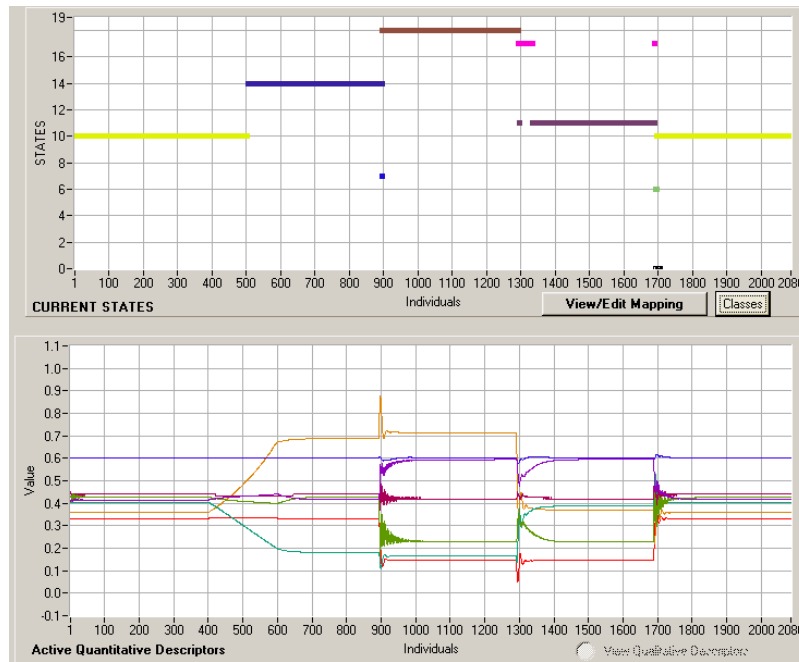


Figure D-3 Scénario 1 : Reconnaissance de défauts simultanés après l'apprentissage supervisé

Scénario 2 : Dans ce scénario le défaut : hausse de la vitesse de réaction est créé et supprimé pendant l'application du défaut : hausse de la température du liquide de refroidissement du réacteur (voir la Figure D-4) :

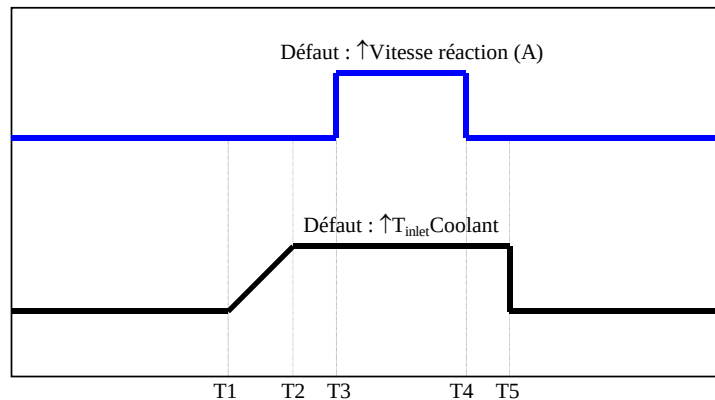


Figure D-4 Scénario 2 : défauts simultanés

La Figure D-5 illustre les résultats obtenus en phase de reconnaissance d'états fonctionnels. On peut vérifier qu'il y a création de 3 états et l'état zéro (état de non-reconnaissance). L'état 10 représente l'opération normale du procédé, l'état 14 est une alarme de hausse de température du fluide de refroidissement du réacteur et l'état 7 correspond au défaut de hausse de la température du fluide de refroidissement du réacteur, mais il est apparu à cause

du comportement oscillatoire des signaux et il n'est composé que d'un seul élément (voir le modèle de comportement donné dans le Tableau 4-N).

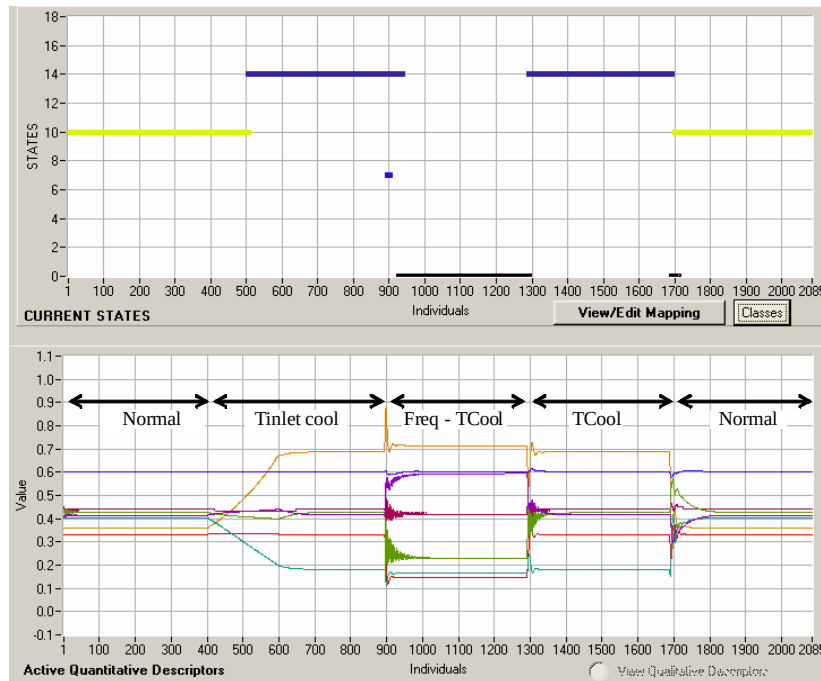


Figure D-5 Identification d'états de défauts simultanés : scénario 2

L'état zéro représente les défauts simultanés et il n'est pas reconnu. Pour le reconnaître, on a fait un nouvel apprentissage supervisé (*supervised learning*) pour affecter la classe non-reconnue à une classe représentant les défauts simultanés. La Figure D-6 illustre les résultats obtenus lors de la reconnaissance des états à partir de l'apprentissage supervisé, et on peut constater la bonne reconnaissance de l'état numéro 18 (c'est un nouvel état créé) pour représenter le nouveau défaut.

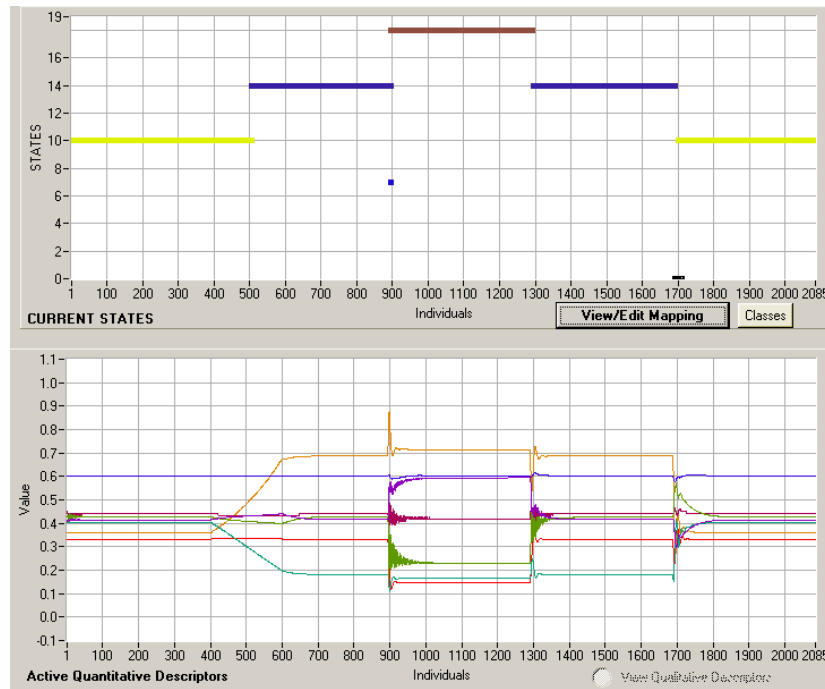


Figure D-6 Scénario 2 : reconnaissance de défauts simultanés après l'apprentissage supervisé

Scénario 3 : Dans ce scénario le défaut : hausse de la vitesse de réaction, est créé plus tôt et supprimé pendant l'application de l'autre défaut : hausse de la température du liquide de refroidissement du réacteur (voir la Figure D-7).

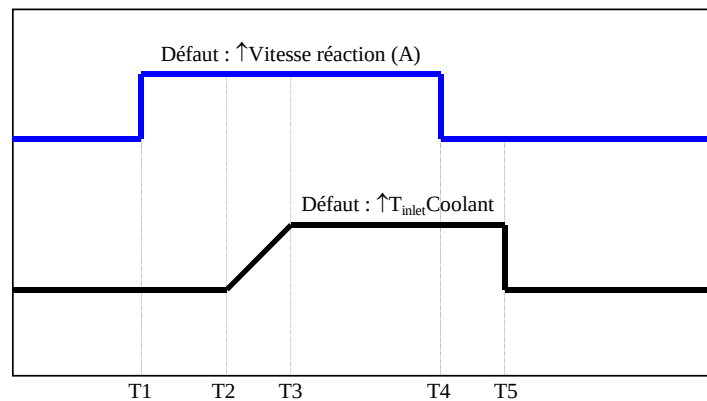


Figure D-7 Scénario 3 : défauts simultanés

La Figure D-8, permet de vérifier qu'il y a création de 4 états et de l'état zéro (état de non-reconnaissance). L'état 10 représente toujours l'opération normale du procédé, l'état 14 est une alarme de hausse de température du fluide de refroidissement du réacteur, l'état 11 correspond à la hausse de la vitesse de réaction (hausse du facteur de fréquence).

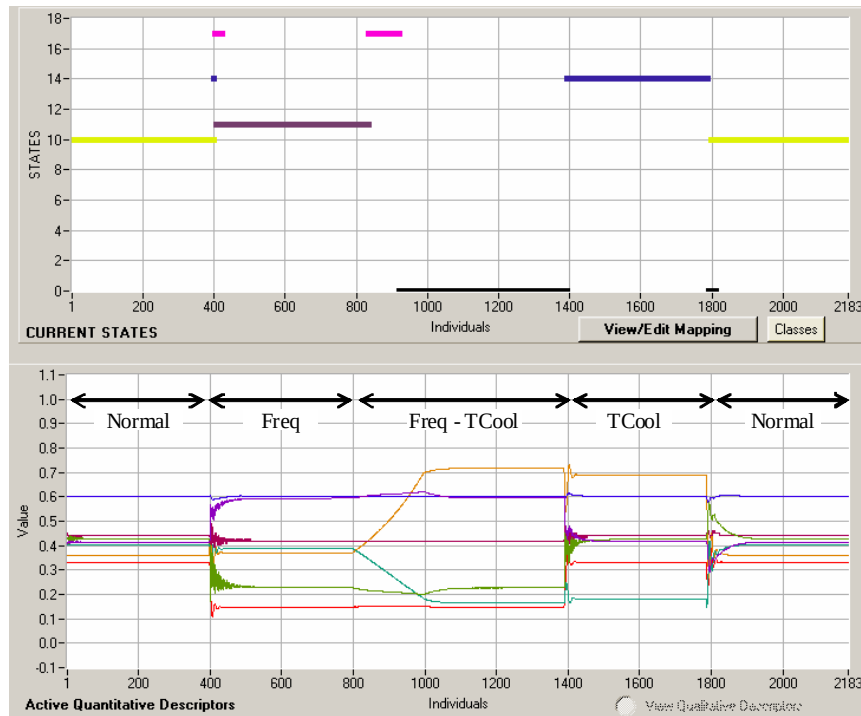


Figure D-8 Identification d'états de défauts simultanés : scénario 3

Comme précédemment, on peut constater qu'après apprentissage supervisé, la simultanéité des défauts est bien reconnue comme le montre la Figure D-9 où l'état numéro 18 (nouvel état créé) représente la nouvelle défaillance.

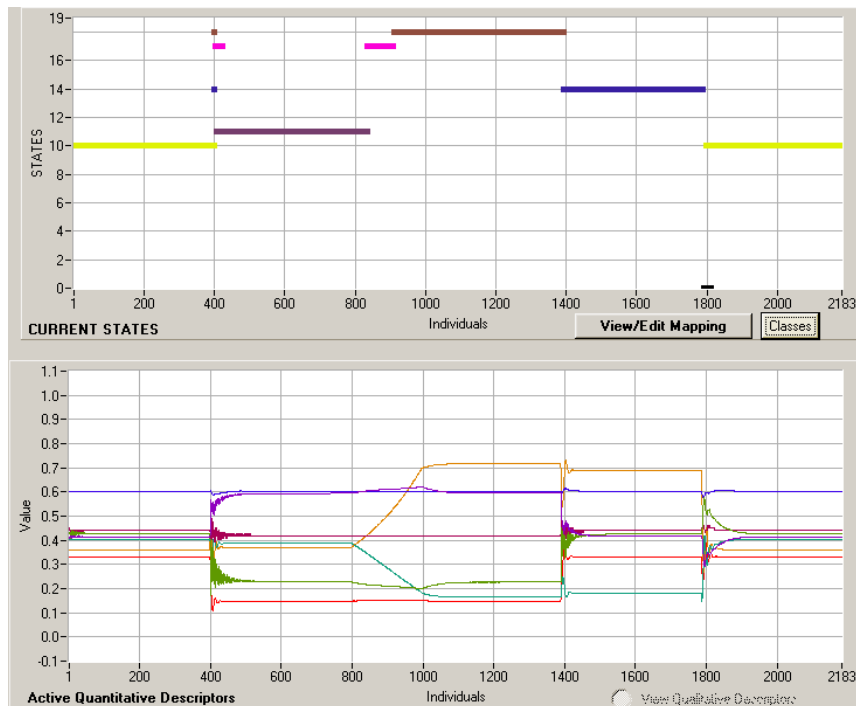


Figure D-9 Scénario 3 : reconnaissance de défauts simultanés après l'apprentissage supervisé

METHODOLOGIE POUR LE PLACEMENT DES CAPTEURS A BASE DE METHODES DE CLASSIFICATION EN VUE DU DIAGNOSTIC

Résumé

Les travaux présentés se situent dans le domaine de l'aide à la décision pour la surveillance et le diagnostic de systèmes complexes tels que les procédés chimiques. Notre travail a permis de concevoir une méthodologie permettant de placer les capteurs les plus pertinents sur un processus en vue de son diagnostic à partir des données d'historiques. Cette méthodologie est basée sur l'association de méthodes de mesure de la quantité d'information (entropie de Shannon) délivrée par des signaux issus d'un système et de la classification de données. A partir des données d'évolution temporelle des capteurs (constituant l'ensemble de tous les possibles) suite à des scénarii de défaillance (par exemple simulés sur un simulateur dynamique), il est possible d'identifier les capteurs les plus pertinents et d'obtenir un modèle de comportement du processus à un niveau d'abstraction tel qu'il soit utilisable pour le diagnostic. Une procédure pour une adaptation du modèle dans le cas de reconnaissance défauts inconnus a été aussi proposée. Des tests de faisabilité sur des cas concrets industriels ont été effectués en simulation tout d'abord en utilisant un simulateur dynamique de procédés mondialement distribué dans l'industrie : *HYSYS* puis sur un nouveau réacteur chimique développé par le Laboratoire de Génie Chimique de Toulouse. Ces travaux rentrent dans le cadre d'un projet financé par l'Institut pour une Culture de Sécurité Industrielle (ICSI).

MOTS CLES : Placement de capteurs, mesure de l'information, détection et localisation de défaillances, diagnostic, classification, apprentissage.

SENSOR PLACEMENT METHODOLOGY FOR FAULT DETECTION BASED ON CLASSIFICATION METHODS

Summary

The presented study is in the field of the decision-making aid for the monitoring and the diagnosis of complex systems such as chemical processes. Our work permitted to design a methodology allowing placing the most relevant sensors on a process for its diagnosis from the analysis of historical data. This methodology is based on the association of methods used for measurement of information quantity (Shannon's entropy) delivered by signals coming from a system and for the classification of data. From time evolution data of sensors (constituting the set of all possible sensors) following scenarios of failure (for example simulated on a dynamic simulator), it is possible to identify the most relevant sensors in the set of the possible sensors and to obtain a model of the process with a level of abstraction such as it is usable for the diagnosis. A procedure for the model adaptation in the case of recognition of unknown failures has also been proposed. Tests of feasibility on concrete industrial cases have been carried out in simulation by using first of all a dynamic process simulator universally distributed in industry: *HYSYS* then on a new chemical reactor developed by the Laboratoire de Génie Chimique of Toulouse. These works are part of a project supported by the "Institut pour une Culture de Sécurité Industrielle" (ICSI).

KEYWORDS: sensor location, information measurement, fault detection and isolation, diagnosis, classification, learning.