

Table des matières

Liste des figures	ix
Liste des tableaux	xi
Introduction	1
PART I. TRAITEMENT DES EAUX USÉES : UN POINT DE VUE DE L'INGÉNIERIE	5
1. Introduction aux procédés biologiques de traitement des eaux	7
1.1. Les procédés de traitement des eaux usées	9
1.1.1. L'eau et la pollution	9
1.1.2. Les indicateurs de qualité de l'eau	11
1.1.3. La législation : quelques éléments	12
1.1.4. Le traitement des eaux usées	15
1.2. Le traitement biologique des eaux usées	18
1.2.1. Les principaux polluants	19
1.2.2. Les micro-organismes épurateurs	20
1.2.3. Les processus métaboliques	22
1.2.4. Les traitements biologiques	24
1.3. L'automatique des bioprocédés	26
1.3.1. Les problèmes spécifiques de l'Automatique des bioprocédés	27
1.3.2. Modélisation et identification des bioprocédés	28
1.3.3. La commande des bioprocédés	31
1.4. Conclusions	36
2. Modélisation dynamique de procédés biologiques	39
2.1. Description des procédés biologiques	42
2.1.1. Les micro-organismes et leur utilisation	42
2.1.2. Les types de bioréacteurs	43
2.1.3. Les trois modes de fonctionnement	43
2.2. Les modèles de bilans-matière	45
2.2.1. Les approches microscopique et macroscopique de modélisation	45
2.2.2. Le taux spécifique de croissance	46
2.2.3. Equations d'état	49
2.3. Influence du transfert de matière	50
2.4. Quelques modèles représentatifs	51
2.4.1. Fermentations continues	51
2.4.2. Fermentations semi-continues	55
2.4.3. Procédés à boues activées	56

2.5.	Modélisation du bioprocédé de traitement des eaux usées par lagunage aérée	59
2.5.1.	Traitement biologique des effluents de papeterie par lagunage aérée.....	60
2.5.2.	Modélisation des cinétiques de croissance et de dégradation	62
2.5.3.	Equations d'évolution du procédé.....	65
2.5.4.	Variables de commande du procédé.....	67
2.6.	Conclusions	69

PART II. IDENTIFICATION - COMMANDE FLOUES ET APPLICATION AU BIOPROCÉDÉ DE TRAITEMENT DES EAUX USÉES 71

3.	Identification de modèles flous Takagi-Sugeno (TS) à partir de données.....	73
3.1.	Modélisation floue de systèmes	76
3.1.1.	Structure générale et différents types de modèles flous	76
3.1.2.	Le modèle flou de type Takagi-Sugeno	84
3.2.	Identification (Construction) de modèles flous.....	88
3.2.1.	Structure du modèle flou	89
3.2.2.	Le problème d'identification et ses solutions.....	90
3.2.3.	L'identification basée sur des données entrée-sortie.....	91
3.3.	Méthodes de coalescence floue (clustering flou)	93
3.3.1.	Notation et Concepts de base	93
3.3.2.	Groupeage <i>c</i> -moyennes floues (algorithme FCM).....	98
3.3.3.	Groupeage avec matrice de covariance floue (algorithme GK).....	101
3.3.4.	Groupeage avec prototypes linéaires (algorithme FCRM)	103
3.3.5.	Groupeage avec régression ellipsoïdale (algorithme FER).....	105
3.3.6.	Groupeage robuste en présence de bruit (algorithme RoFER)	110
3.4.	Construction de modèles TS à partir de données : méthodologie générale	118
3.4.1.	Validation du nombre de clusters.....	122
3.4.2.	Génération des fonctions d'appartenance des antécédents	124
3.4.3.	Obtention des paramètres des conséquents	126
3.4.4.	Validation numérique du modèle flou.....	128
3.4.5.	Outil logiciel pour la modélisation-identification floue de type TS.....	129
3.4.6.	Exemples	130
3.5.	Conclusions	140
4.	Commande floue TS basée sur modèle	143
4.1.	Stabilité et stabilisation à partir de modèles flous.....	146
4.1.1.	Modeles flous dynamiques de type Takagi Sugeno (TS).....	146
4.1.2.	Stabilité des modèles flous	148
4.1.3.	Stabilisation des modèles flous TS.....	150
4.2.	Commande PDC sous optimale par retour d'état.....	153
4.2.1.	Commande LQR pour des modèles affines à temps discret.....	154
4.2.2.	Commande PDC avec compensation de gain statique	158
4.2.3.	Commande PDC avec ajout d'intégrateurs	160
4.2.4.	Choix des matrices de pondération Q et R.....	163
4.3.	Conclusions	165

5. Identification et Commande floues TS du bioréacteur	167
5.1. Description entrée-sortie du bioprocédé de traitement des eaux usées.....	169
5.1.1. Le modèle de bilan matière	169
5.1.2. Les paramètres du modèle.....	170
5.1.3. Les objectifs de conduite du procédé	171
5.1.4. Les contraintes sur les actionneurs.....	172
5.2. Modélisation et identification floue du bioprocédé à partir des données.....	172
5.2.1. Elaboration du jeu des données entrée-sortie	173
5.2.2. Utilisation du clustering flou Gustafson-Kessel et RoFER.....	176
5.3. Commande floue de type TS du bioprocédé basée sur le modèle flou	188
5.4. Conclusions	195
Conclusions et perspectives	197
Références bibliographiques	201

Liste des figures

Figure 1.1. Filière d'épuration des eaux résiduaires (d'après).	16
Figure 1.2. Les micro-organismes et leurs besoins nutritifs et énergétiques (D'après).	22
Figure 1.3. Représentation du système de commande d'un bioprocédé (d'après).	32
Figure 2.1. Les différents modes de fonctionnement des réacteurs biologiques.	43
Figure 2.2. Loi de Monod pour $\mu_{max} = 0,5 \text{ h}^{-1}$ et plusieurs valeurs du paramètre K_S .	47
Figure 2.3. Loi d'Andrews pour $\mu_{max} = 0,5 \text{ h}^{-1}$, $K_S = 0,5 \text{ g/L}$ et plusieurs valeurs du K_I .	48
Figure 2.4. Cycle d'élimination des pollutions azotées.	53
Figure 2.5. Principales étapes de la digestion anaérobie.	54
Figure 2.6. Schéma typique d'un procédé d'épuration par boues activées.	57
Figure 2.7. Pilote expérimental d'un procédé à boues activées (Université de los Andes).	58
Figure 2.8. Schéma synoptique de l'installation du procédé pilote.	62
Figure 2.9. Les interactions microbiennes : compétition et activation/inhibition compétitive.	65
Figure 2.10. Phénomènes d'interactions microbiennes et mortalité bactérienne endogène.	67
Figure 2.11. Alimentation du bioréacteur dans le cas multi-variable.	68
Figure 3.1. Formes typiques représentatives des fonctions d'appartenance.	78
Figure 3.2. Représentation interne d'un système flou.	80
Figure 3.3. Identification par clustering flou (adaptation d'après).	92
Figure 3.4. Interprétation géométrique de la matrice de covariance floue (d'après).	102
Figure 3.5. Illustration des fonctions robustes w_i et ρ_i associées à chacun des clusters.	113
Figure 3.6. Vue générale de la méthodologie d'identification non linéaire basée sur l'approche de clustering flou.	119
Figure 3.7. (a) Illustration du mécanisme de projection et (b) de l'erreur de décomposition au moment de la recombinaison (intersection floue) dans l'espace de l'antécédent.	125
Figure 3.8. Approximation des données projetées par une fonction d'appartenance paramétrique.	126
Figure 3.9. Interprétation géométrique de l'orientation du cluster à partir de la matrice de covariance floue (cas de l'algorithme GK).	127
Figure 3.10. Fonction $y(x) = \sin(x)$.	130
Figure 3.11. Fonction $y(x) = \sin(x)$ avec bruit et points aberrants (cercles).	131
Figure 3.12. Approximation TS de la fonction $y(x) = \sin(x)$ - algorithmes GK et RoFER.	132
Figure 3.13. Approximation TS de la fonction $y(x) = \sin(x) + \xi$ - algorithmes GK et RoFER.	133
Figure 3.14. Fonction vallée de Rosenbrock.	135
Figure 3.15. Fonction vallée de Rosenbrock avec bruit et points aberrants (cercles).	135
Figure 3.16. Approximation TS de la fonction vallée de Rosenbrock.	136
Figure 3.17. Approximation TS de la fonction vallée de Rosenbrock avec bruit et points aberrants.	138
Figure 4.1. Représentation du concept de compensation parallèle distribuée (PDC).	145
Figure 4.2. Schéma général de la commande du type régulateur PDC pour des sous modèles affines.	157

Figure 4.3. Schéma général de la commande du type régulateur PDC avec pré compensateur	160
Figure 4.4. Schéma général de la commande PDC avec action intégrale pour des sous modèles affines	163
Figure 5.1. Schéma de principe entrée-sortie du bioprocédé de traitement des eaux usées...	169
Figure 5.2. Signaux d'entrée $u_1(t)$ et $u_2(t)$ pour l'identification du bioprocédé	173
Figure 5.3. Répartition des données dans l'espace des entrées du procédé.	174
Figure 5.4. Répartition des données dans l'espace des sorties du procédé	175
Figure 5.5. Evolution des sorties du procédé – données d'identification	175
Figure 5.6. Signaux d'entrée $u_1(t)$ et $u_2(t)$ pour le premier test de validation du bioprocédé	177
Figure 5.7. Premier test de validation du modèle flou TS du bioprocédé avec données d'opération en mode "mixte" (batch/continu). Sortie du procédé (ligne en trait plein), sortie du modèle avec clustering GK (ligne pleine) et sortie du modèle avec clustering RoFER (ligne pointillée).	178
Figure 5.8. Deuxième test de validation du modèle flou TS du bioprocédé avec données d'opération en mode continue. Sortie du procédé (ligne en trait plein), sortie du modèle avec groupement GK (ligne pleine) et sortie du modèle avec groupement RoFER (ligne pointillée).	179
Figure 5.9. Partition $X(k)$ avec GK données d'opération en mode continue.....	181
Figure 5.10. Partition $S_x(k)$ avec GK données d'opération en mode continue	181
Figure 5.11. Partition $S_e(k)$ avec GK données d'opération en mode continue	182
Figure 5.12. Partition $X(k)$ avec RoFER données d'opération en mode continue	182
Figure 5.13. Partition $S_x(k)$ avec RoFER données d'opération en mode continue.....	183
Figure 5.14. Partition $S_e(k)$ avec RoFER données d'opération en mode continue.....	183
Figure 5.15. Evolution des concentrations $X(k)$ et $S_x(k)$ suivant le scénario de contrôle désiré.	193
Figure 5.16. Evolution des variables de commande $u_1(k)$ et $u_2(k)$ du bioprocédé.	194

Liste des tableaux

Tableau 1.1. Caractéristiques générales des rejets	13
Tableau 1.2. Valeurs limites des rejets urbains	13
Tableau 1.3. Valeurs limites des rejets industriels	14
Tableau 1.4. Normes sur les rejets (Art. 72 Décret 1594 de 1984, Colombie)	15
Tableau 2.1. Caractéristiques d'effluents d'usines de papier et de pâte à papier	60
Tableau 2.2. Valeurs des variables d'environnement	61
Tableau 3.1. Principales t-normes et t-conormes	81
Tableau 3.2. Règles du modèle flou TS pour la fonction $y(x) = \sinusc(x)$	132
Tableau 3.3. Centres des clusters du modèle flou $y(x) = \sinusc(x)$ – algorithmes GK et RoFER	132
Tableau 3.4. Règles du modèle flou TS pour la fonction $y(x) = \sinusc(x) + \xi$	133
Tableau 3.5. Centres des clusters du modèle flou $y(x) = \sinusc(x) + \xi$ – algorithmes GK et RoFER	133
Tableau 3.6. Evaluation de la qualité de l'approximation de la fonction $y(x) = \sinusc(x)$	134
Tableau 3.7. Règles du modèle flou TS pour la fonction vallée de Rosenbrock – algorithme GK	137
Tableau 3.8. Règles du modèle flou TS pour la fonction vallée de Rosenbrock – algorithme RoFER	137
Tableau 3.9. Centres des clusters du modèle flou pour la fonction vallée de Rosenbrock sans bruit	137
Tableau 3.10. Evaluation de la qualité de l'approximation de la fonction $z(x, y) = rosenbrock(x, y)$	137
Tableau 3.11. Règles du modèle flou TS pour la fonction vallée de Rosenbrock avec bruit et points aberrants– algorithme GK	139
Tableau 3.12. Règles du modèle flou TS pour la fonction vallée de Rosenbrock avec bruit et points aberrants– algorithme RoFER	139
Tableau 3.13. Centres des clusters du modèle flou pour la fonction vallée de Rosenbrock avec bruit et points aberrants	139
Tableau 3.14. Evaluation de la qualité de l'approximation de la fonction $z(x, y) = rosenbrock(x, y) + \xi$	139
Tableau 5.1. Paramètres stoechiométriques et cinétiques du bioprocédé	171
Tableau 5.2. Bornes des variables dans le système d'alimentation du bioprocédé.....	172
Tableau 5.3. Performance numérique de la validation du modèle flou TS pour le bioréacteur – résultats avec l'algorithme GK.....	180
Tableau 5.4. Performance numérique de la validation du modèle flou TS pour le bioréacteur – résultats avec l'algorithme RoFER	180
Tableau 5.5. Paramètres des conséquents du modèle flou TS pour $X(k+1)$ - algorithme GK	184
Tableau 5.6. Centres des clusters du modèle flou TS pour $X(k+1)$ - algorithme GK.....	184
Tableau 5.7. Paramètres des conséquents du modèle flou TS pour $X(k+1)$ - algorithme RoFER.....	185
Tableau 5.8. Centres des clusters du modèle flou TS pour $X(k+1)$ - algorithme RoFER...	185

Tableau 5.9. Paramètres des conséquents du modèle flou TS pour $S_x(k+1)$ - algorithme GK	186
Tableau 5.10. Centres des clusters du modèle flou TS pour $S_x(k+1)$ - algorithme GK	186
Tableau 5.11. Paramètres des conséquents du modèle flou TS pour $S_x(k+1)$ - algorithme RoFER.....	186
Tableau 5.12. Centres des clusters du modèle flou TS pour $S_x(k+1)$ - algorithme RoFER	186
Tableau 5.13. Paramètres des conséquents du modèle flou TS pour $S_e(k+1)$ - algorithme GK	187
Tableau 5.14. Centres des clusters du modèle flou TS pour $S_e(k+1)$ - algorithme GK	187
Tableau 5.15. Paramètres des conséquents du modèle flou TS pour $S_e(k+1)$ - algorithme RoFER.....	188
Tableau 5.16. Centres des clusters du modèle flou TS pour $S_e(k+1)$ - algorithme RoFER	188
Tableau 5.17. Critères des performances en boucle fermée du système de commande flou de type PDC basé sur le modèle dynamique flou affine de type Takagi-Sugeno.....	194

Introduction

L'eau est à l'origine de la vie. Sans eau, aucun organisme, qu'il soit végétal ou animal, simple ou complexe, ne peut vivre. Historiquement, les premières civilisations se sont fondées au bord des mers, lacs et rivières. Les eaux utilisées pour les besoins domestiques, l'agriculture et les activités de production étaient rejetées et elles atteignaient les corps d'eau sans aucun type de traitement. Les rivières arrivaient à se débarrasser de leurs déchets par auto-épuration; et l'eau était suffisante par rapport aux besoins humains quotidiens.

L'équilibre naturel autrefois existant a beaucoup été modifié. L'homme a apporté au cycle de l'eau une perturbation quantitative par ses prélèvements et sa consommation et qualitative par la pollution qu'il engendre. C'est ainsi qu'au fil du temps l'industrialisation et la croissance accélérée des villes ainsi que les besoins d'une population en augmentation constante ont entraîné la nécessité de création d'unités d'épuration pour traiter les volumes rejetés et la concentration en substances polluantes. Ces substances peuvent avoir des origines polluantes diverses : physique (présence des matières en suspension), chimique, organique et bactériologique. Parmi la gamme des différentes formes de pollution, les déchets biodégradables constituent une des sources les plus fréquemment trouvées dans la pratique.

Quand la pollution organique est non toxique, elle peut être dégradée par le milieu naturel - si la masse d'eau est suffisante - grâce au phénomène d'auto-épuration, due aux organismes vivant dans le milieu aquatique : bactéries, protozoaires, algues, qui permettent à l'eau de retrouver sa qualité première. Cependant, quand les rejets de matières organiques sont trop concentrés, la capacité naturelle d'auto-épuration est saturée et la pollution persiste, en affectant l'équilibre écologique. Pour faire face à ce problème, l'homme a mis en place des réseaux d'égout qui collectent les eaux usées et les acheminent vers les stations d'épuration, où le traitement des effluents se fait à l'aide de techniques efficaces, afin de minimiser l'effet polluant avant que l'eau soit rejetée dans le milieu naturel.

L'épuration d'un effluent pollué peut comporter des traitements de nature assez variée : biologiques, chimiques et/ou physico-chimiques. La nature des problèmes d'environnement liés aux activités humaines est une affaire de société qui demande un effort croissant et qui doit être entrepris à l'échelle globale pour mieux comprendre la complexité et maîtriser les phénomènes inhérents à la dépollution. Dans les domaines des bioprocédés (ou les procédés de traitement des eaux sont un cas particulier), il est nécessaire d'avoir une participation pluridisciplinaire afin de couvrir le très large spectre de connaissances impliquées pour parvenir à des procédés de traitement optimisés [STE98]. En effet, pour mieux comprendre les différents acteurs impliqués, du point de vue biotechnologique, l'activité dans le domaine peut grosso modo être regroupée en trois domaines d'activité principaux, évidemment complémentaires [DOC01] :

- *La microbiologie et le génie génétique*, qui ont pour objectifs de développer des microorganismes permettant la production de nouveaux produits et de choisir les meilleures souches de microorganismes de manière à obtenir certains produits désirés ou certaines qualités de produit.

- *Le génie des bioprocédés*, qui s'occupe de choisir les meilleurs modes de fonctionnement ou de développer des procédés et/ou réacteurs qui permettent d'améliorer le rendement et/ou la productivité des bioprocédés.
- *L'automatique des bioprocédés*, qui quant à elle, a pour but d'augmenter le rendement et/ou la productivité en développant des méthodes de surveillance et de commande automatisées permettant l'optimisation en temps réel du fonctionnement des bioprocédés.

Dans les domaines des biotechnologies et du traitement des eaux en particulier, l'amélioration de la performance dans les industries de production ainsi que les nouvelles normes législatives ont rendu nécessaire la modélisation, l'identification et la commande en temps réel des procédés de traitement biologique [QUE00]. La complexité des mécanismes mis en jeu et le fonctionnement au quotidien de tels procédés ont souligné le besoin de mesurer, d'analyser et de contrôler certaines concentrations et variables caractéristiques des effluents. Les cinétiques non-linéaires, les paramètres variant dans le temps, l'absence de mesures fiables et directement accessibles, les fortes variations des conditions opératoires imposent le développement et l'utilisation de techniques avancées de l'Automatique.

Les travaux présentés dans ce manuscrit concernent l'exploration d'outils alternatifs de l'automatique basés sur la logique floue, pour la modélisation et la commande d'un bioprocédé de traitement des eaux usées. Ces travaux ont été développés dans le cadre d'une thèse en cotutelle entre l'Université Paul Sabatier - Laboratoire LAAS-CNRS en France et l'Université de los Andes - Laboratoire GIAP en Colombie. Afin d'apporter des éléments de références par rapport nos travaux, nous nous permettons de présenter en quelques lignes quelques axes de recherche qui sont étudiés dans les établissements scientifiques impliqués.

En ce qui concerne le *Laboratoire LAAS¹-CNRS²* et plus particulièrement le groupe DISCO³, les principaux travaux de recherche portent sur le développement des outils de diagnostic, supervision et conduite pour les systèmes dynamiques complexes, pouvant inclure l'homme. Le groupe met l'accent sur le caractère qualitatif et quantitatif, sans ignorer les aspects continus ou discrets des systèmes. De ce fait, les systèmes hybrides et l'interface entre signaux continus et leur interprétation sous forme plus abstraite à événements discrets sont un des axes de recherche. La nature qualitative des connaissances et l'incertitude entachant les données font appel à des formalismes qualitatifs et symboliques trouvant leurs origines en Intelligence Artificielle comme le raisonnement qualitatif et la modélisation floue et neuronale, de même qu'à des méthodes de classification, d'apprentissage et de reconnaissance des formes. En ce qui concerne l'automatique des bioprocédés, le groupe a traité différents problèmes liés à l'estimation, l'observation, la commande, la classification, la supervision et le diagnostic de certains procédés biologiques et chimiques [ROU92] [BEN96] [DAH96] [ROU01] [KEM04] [ATI05] [HER06] [DAH06].

Pour sa part, le *Laboratoire GIAP⁴* est dédié au développement et aux applications des méthodes et techniques avancées dans plusieurs domaines : modélisation et commande des systèmes, robotique, potabilisation et traitement des eaux, technologies de production, infographie et réalité virtuelle. En ce qui concerne le *groupe d'automatique et d'informatique*, les travaux de recherche développés portent sur la modélisation de systèmes hybrides,

¹ LAAS : *Laboratoire d'Analyse et d'Architecture des Systèmes*.

² CNRS : *Centre National de la Recherche Scientifique*.

³ DISCO : *Diagnostic, Supervision et CONduite qualitatifs*.

⁴ GIAP : *Groupe d'Informatique et d'Automatisation pour la Production*.

l'identification paramétrique des systèmes et la commande des procédés (par des diverses techniques de commande : adaptative, prédictive, robuste, floue et hybride). Le groupe est orienté aussi vers des applications pour l'industrie. Notre travail au sein des Départements de Génie Electrique et Civil-Environnemental a été focalisé sur l'exploration et l'application de techniques alternatives de l'automatique au domaine des procédés biotechnologiques [GAU97] [GIR98] [ROD02] [ORO03] [DUQ03] [MOJ05].

En effet, c'est dans ce cadre que se situent notre motivation et notre intérêt pour la modélisation et la commande des procédés biologiques de traitements des eaux usées. Nous espérons contribuer à l'amélioration du fonctionnement d'une station d'épuration des effluents de papeterie par lagunage aéré en développant des techniques de modélisation et de commande floue de type Takagi-Sugeno (TS). Au niveau de l'identification, nous avons utilisé des méthodes floues de regroupement des données entrée-sortie pour construire un modèle dynamique non linéaire du procédé. Le type d'approche suivi, peut être considéré comme une approche de type « boîte grise » qui représente le système dynamique non linéaire comme un modèle à base de règles du type « *Si-Alors* », concaténant un ensemble de sous modèles localement linéaires sous la forme d'auto-régression non-linéaire (NARX⁵) dans l'espace d'état. De plus, nous avons développé une version graphique de la boîte à outils FMID⁶ pour la modélisation floue des systèmes non linéaires. A partir de cette représentation, nous développons une commande floue TS basée sur le modèle, en utilisant la philosophie de la commande de type compensation parallèle distribuée (PDC⁷). Nous nous sommes intéressés, en particulier, à la commande sous-optimale linéaire quadratique par retour d'état basée sur des modèles TS affines en temps discret.

Pour cela, ce mémoire est structuré en cinq chapitres, dont une première partie concernant le traitement des eaux usées et la seconde, sur la modélisation, l'identification et la commande floues avec application sur le bioprocédé de traitement des eaux résiduaires.

Dans la première partie, constituée des deux premiers chapitres, nous développons de manière non exhaustive, des aspects sur l'ingénierie des traitements des eaux usées, en mettant l'accent sur la modélisation dynamique des procédés biologiques.

Le *premier chapitre* présente une introduction aux procédés de traitement des eaux. Nous positionnons le problème de pollution de l'eau et nous présentons les concepts de base sur les indicateurs de qualité de l'eau et les procédés de traitement. Ensuite nous présentons le traitement biologique des eaux usées, ainsi que la problématique générale de l'automatique des bioprocédés.

Le *deuxième chapitre* aborde la modélisation dynamique des procédés biologiques utilisés dans le domaine du traitement des eaux résiduaires. Nous donnons d'abord un aperçu global en décrivant le rôle des microorganismes, les types de réacteurs et les modes de fonctionnement des bioprocédés. Ensuite, au niveau de la modélisation, nous nous focalisons sur l'approche des bilans de matière des composants principaux pour caractériser la dynamique de tels procédés. Ensuite nous présentons quelques modèles représentatifs des procédés biologiques. Finalement, nous décrivons le procédé considéré au niveau de l'application, qui est un bioréacteur aérobie multivariable en mode continu pour le traitement

⁵ NARX : *Nonlinear AutoRegressive with eXogenous input*.

⁶ FMID : *Fuzzy Model Identification Toolbox for Matlab®*, développé par le Prof. Robert Babuška, Delft University of Technology, The Netherlands.

⁷ PDC : *Parallel Distributed Compensation*.

des eaux usées issues d'une industrie papetière. Le modèle considéré est une version du simulateur utilisé lors du projet européen FAMIMO⁸ pour lequel nous avons proposée l'ajout d'un terme de mortalité endogène au niveau de la dynamique bactérienne.

La deuxième partie de ce mémoire, comportant les chapitres 3 à 5, concerne la modélisation, l'identification et la commande de systèmes complexes en utilisant l'approche floue avec application au bioprocédé de traitement des eaux usées de notre étude.

Dans le *troisième chapitre*, nous traitons du problème d'identification (construction) des modèles flous de type Takagi-Sugeno (TS) à partir de données entrée-sortie. Nous introduisons d'abord la modélisation floue de systèmes, en nous focalisant particulièrement sur le modèle de type TS. En utilisant ce formalisme, nous représentons le comportement non linéaire d'un système par une composition de règles du type « *Si-Alors* », concaténant un ensemble de sous-modèles localement linéaires. Pour la construction de tels modèles nous abordons d'une part, différentes méthodes de classification floue (*clustering*) et, d'autre part, nous proposons l'application d'une méthode d'« agglomération compétitive », robuste en présence de bruit : il s'agit des algorithmes FER⁹ et RoFER¹⁰. Les méthodes considérées appartiennent aux méthodes de clustering basées sur la minimisation d'une fonction objectif. Enfin après avoir considéré la méthodologie générale pour la construction de modèles flous TS à partir de données, nous présentons quelques exemples illustratifs d'application.

Le *quatrième chapitre* concerne la commande des systèmes non linéaires basée sur les modèles flous de type Takagi-Sugeno, en utilisant la philosophie de commande du type compensation parallèle distribuée (PDC). Nous commençons par un aperçu général sur la synthèse (stabilisation) des contrôleurs flous basée sur un modèle flou homogène TS du système. Ensuite nous développons spécifiquement une commande sous-optimale linéaire quadratique par retour d'état basée sur un modèle affine TS en temps discret. Pour cela, nous proposons trois lois de commande qui tiennent compte du terme indépendant caractéristique du modèle affine. Les lois de commande sont de type régulateur, régulateur avec compensation de gain statique et régulateur avec action intégrale, ce qui permet l'annulation des erreurs statiques.

Le *cinquième chapitre* est consacré à l'application des méthodes floues d'identification et de commande au bioprocédé de traitement des eaux usées. Nous reprenons d'abord le modèle du bioréacteur introduit au chapitre 2. Ensuite, nous développons le modèle Takagi-Sugeno pour le bioréacteur en appliquant les techniques floues de construction basées sur des données expérimentales entrée-sortie considérées dans le chapitre 3. En fin, en appliquant la philosophie de commande du type PDC basée sur modèle présentée dans le chapitre 4, nous faisons la synthèse de la loi de commande la mieux adaptée au bioprocédé. Nous nous sommes intéressés, en particulier, à la commande sous-optimale linéaire quadratique par retour d'état pour les modèles TS affines en temps discret, afin de réguler les concentrations de biomasse et de substrat xénobiotique du bioprocédé. Les résultats obtenus en simulation soulignent les caractéristiques d'applicabilité, de performance et certaines limitations pour l'utilisation avec succès de ces techniques sur ce système non linéaire multivariable.

Enfin, nous terminons ce mémoire par une conclusion générale du travail développé et par quelques-unes des perspectives qui nous semble les plus intéressantes à étudier.

⁸ Waste-water Benchmark, in Fuzzy Algorithms for the Control of Multiple-Input, Multiple Output Processes (FAMIMO) project, funded by the European Commission (Esprit LTR 21911).

⁹ FER : Fuzzy Ellipsoidal Regression.

¹⁰ RoFER : Robust Fuzzy Ellipsoidal Regression.

**PART I. TRAITEMENT DES EAUX USÉES :
UN POINT DE VUE DE L'INGÉNIERIE**

Chapitre 1

Introduction aux procédés biologiques de traitement des eaux

Dans ce chapitre, nous présentons une introduction aux procédés de traitement des eaux. Nous positionnons le problème de pollution de l'eau et nous présentons les concepts de base des indicateurs de qualité de l'eau et les procédés de traitement. Ensuite nous présentons le traitement biologique des eaux usées, ainsi que la problématique générale de l'automatique des bioprocédés.

Sommaire

1. Introduction aux procédés biologiques de traitement des eaux.....	7
1.1. Les procédés de traitement des eaux usées	9
1.1.1. L'eau et la pollution.....	9
1.1.2. Les indicateurs de qualité de l'eau	11
1.1.3. La législation : quelques éléments.....	12
1.1.4. Le traitement des eaux usées.....	15
1.2. Le traitement biologique des eaux usées.....	18
1.2.1. Les principaux polluants.....	19
1.2.2. Les micro-organismes épurateurs.....	20
1.2.3. Les processus métaboliques.....	22
1.2.4. Les traitements biologiques	24
1.3. L'automatique des bioprocédés	26
1.3.1. Les problèmes spécifiques de l'Automatique des bioprocédés.....	27
1.3.2. Modélisation et identification des bioprocédés.....	28
1.3.3. La commande des bioprocédés	31
1.4. Conclusions	36

La forme de vie dite « moderne » de l'homme a beaucoup modifié l'équilibre écologique naturel autrefois existant sur la planète. Les problèmes d'environnement liés à la concentration des populations et aux activités humaines, que ce soit au niveau urbain, agricole ou industriel, deviennent de plus en plus importants. La pollution générée par l'homme affecte de plus en plus le cycle de l'eau et des traitements artificiels doivent souvent être appliqués pour compléter les cycles naturels d'auto-épuration. Ces traitements sont en place à l'heure actuelle sur les stations d'épuration.

Le vaste domaine du traitement des eaux usées met en jeu plusieurs disciplines qui se matérialisent par exemple au niveau de la conception, la mise au point et la conduite des stations d'épuration. Cette pluridisciplinarité est constituée de la microbiologie et du génie génétique, en passant par le génie civil, jusqu'à l'automatique. Nous tenterons donc, de présenter dans ce chapitre, quelques concepts fondamentaux liés au problème d'assainissement biologique de l'eau, ainsi que des éléments de base sur l'automatique des bioprocédés, afin de mettre à disposition pour ceux qui ne sont pas forcément spécialistes, des notions de base dans ce domaine multidisciplinaire. Ces éléments nous permettront de fixer un cadre de référence pour mieux comprendre les études développées au cours de cette thèse. Une grande partie de ce chapitre est basée sur les documents [STE98] [HAD99] [QUE00] [DOC01]. Pour une étude plus détaillée et plus spécifique des divers aspects liés au domaine du traitement de l'eau et des bioprocédés, nous suggérons les ouvrages suivants : [HEN02] aborde en profondeur les principes fondamentaux des procédés biologiques et chimiques les plus utilisées dans le traitement de l'eau ; [ORO03] détaille les caractéristiques et le dimensionnement des procédés biologiques du point de vue de l'ingénierie ; [MET03] traite en détail les étapes de la chaîne de traitement et met l'accent sur les aspects de la réutilisation de l'eau et des boues ; [BAS90] et [DOC01] développent et appliquent des outils avancés de l'automatique au domaine des bioprocédés.

1.1. Les procédés de traitement des eaux usées

Dans cette section nous présentons des éléments de base concernant les procédés de traitement des eaux usées. En particulier nous insisterons sur quelques points particuliers : l'eau et la pollution, les indicateurs de qualité de l'eau, la législation en France et le traitement des eaux usées.

1.1.1. L'eau et la pollution

La pollution se définit comme l'introduction dans un milieu naturel de substances provoquant sa dégradation. La pollution des ressources en eau au niveau des stations d'épuration provient de diverses sources, notamment les formes relatives aux activités humaines [HAD99] [STE98] [QUE00] :

- *La pollution domestique et urbaine* : les eaux usées urbaines sont rejetées par les installations collectives (hôpitaux, écoles, commerces,...) et comportent les eaux ménagères (détergents, graisses, ...) et les eaux vannes (eaux sanitaires : matière organique et azotée, germes et matières fécales, ...). Les eaux résiduaires urbaines (ERU) peuvent être considérées comme la plus importante industrie en termes de masse de matériaux bruts à traiter. A titre d'exemple, tant en France qu'en Colombie la consommation moyenne en eau est généralement estimée à 150 litres par jour et par habitant. Dans la communauté européenne il est produit quotidiennement un volume proche à 40 millions de m³ d'eaux usées.

- *La pollution industrielle* : le degré et la nature de la pollution générée par des rejets industriels varient suivant la spécificité de chaque activité industrielle. Certains rejets troublent la transparence et l'oxygénation de l'eau ; ils peuvent avoir un effet nocif sur les organismes vivants et nuire au pouvoir d'auto-épuration de l'eau. Ils peuvent causer aussi l'accumulation de certains éléments dans la chaîne alimentaire (métaux, pesticides, radioactivité,...). Les eaux résiduaires industrielles (ERI) représentent une part importante des rejets arrivant aux stations d'épuration. A titre d'exemple, les deux tiers des industriels redevables des Agences de l'Eau en France (ceux qui génèrent le plus de pollution) sont raccordés aux stations d'épuration des collectivités territoriales. Ils produisent 10% de la charge polluante industrielle brute, ce qui équivaut à un quart de la pollution domestique. Cet apport pose de sérieux problèmes aux exploitants de stations d'épuration urbaines, tant au niveau des capacités que des performances de traitement. En effet, les effluents industriels toxiques¹¹ constituent un danger permanent pour les stations de dépollution biologique. De plus, il faut bien noter que, selon le ministère de l'Environnement Français, 30% des rejets industriels s'échappent encore dans la nature sans aucun traitement !
- *La pollution agricole* : ce type de pollution s'intensifie depuis que l'agriculture est entrée dans un stade d'industrialisation. Les pollutions d'origine agricole englobent à la fois celles qui ont trait aux cultures (pesticides et engrais) et à l'élevage (lisiers et purins). Néanmoins, le problème de la pollution agricole est un peu différent, dans la mesure où cette source de pollution n'arrive qu'indirectement à la station. C'est le cas en particulier des engrais et pesticides qui passent d'abord à travers les milieux naturels (nappes phréatiques, rivières...). C'est aussi le cas des déchets solides issus des industries agro-alimentaires et des concentrations des élevages qui entraînent un excédent de déjections animales (lisiers de porc, fientes des volailles...) par rapport à la capacité d'absorption des terres agricoles ; celles-ci, sous l'effet du ruissellement de l'eau et de l'infiltration dans le sous-sol, enrichissent les cours d'eau et les nappes souterraines en dérivés azotés et constituent aussi une source de pollution bactériologique.
- *La pollution d'origine naturelle et l'eau de pluie* : La teneur de l'eau en substances indésirables n'est pas toujours le fait de l'activité humaine. Certains phénomènes naturels peuvent y contribuer (contact de l'eau avec les gisements minéraux, ruissellement des eaux de pluie, irrptions volcaniques,...). En ce qui concerne l'eau de pluie, bien que longtemps considérée comme propre, l'eau d'origine pluviale est en fait relativement polluée. L'origine de cette pollution peut provenir des gaz ou solides en suspension rejetés dans l'atmosphère par les véhicules, les usines ou les centrales thermiques. Ces polluants (oxyde de carbone, dioxyde de soufre, poussière) sont envoyés vers le sol à la moindre averse. Lorsqu'elle ruisselle, l'eau de pluie a un second effet nocif : elle transporte les hydrocarbures, les papiers, les plastiques et les débris végétaux accumulés sur la terre et les toitures. De plus, cette pollution est déversée sur de courtes périodes et peut atteindre des valeurs très élevées ce qui provoquent un effet de choc sur le milieu biologique.

En ne parlant que de la pollution de l'eau, ce bilan est loin d'être exhaustif puisqu'il faudrait lui rajouter tous les déchets solides, constitués d'ordures ménagères, des déchets ménagers encombrants (mobilier, cuisinières, réfrigérateurs,...), des déchets automobiles (carcasses, batteries, huiles et pneus usagés), des déchets provenant de l'entretien des espaces verts urbains, des déchets d'assainissement des eaux usées (boues), des déchets inertes (les 2/3 des déchets solides industriels), et enfin des déchets produits ou recyclés dans l'agriculture et les industries agro-alimentaires.

¹¹ Ceux qui contiennent une forte proportion de métaux lourds ou de molécules organiques toxiques.

1.1.2. Les indicateurs de qualité de l'eau

Les eaux usées sont des liquides de composition hétérogène, chargés de matières minérales ou organiques pouvant être en suspension ou en solution, et dont certaines peuvent avoir un caractère toxique. L'élaboration et la définition de paramètres qualitatifs de la pollution ont conduit à établir des mesures quantitatives de la pollution [HAD99] [QUE00] :

- *Matières en suspension* (MES) : quantité (en mg/L) de particules solides, de nature minérale ou organique, véhiculées par les eaux usées. La mesure des MES est effectuée par filtration ou bien par centrifugation d'un échantillon, après séchage à 110 °C.
- *Demande chimique en oxygène* (DCO) : permet de mesurer la consommation d'oxygène (en mgO₂/L) dans les conditions d'une réaction d'oxydation complète. C'est une mesure de la pollution organique. On utilise pour cela un oxydant puissant (le bichromate de potassium) dans un milieu acide (acide sulfurique). La quantité d'oxygène cédée par l'oxydant au cours de la réaction constitue la demande chimique en oxygène. Celle-ci est d'autant plus importante que l'eau est polluée. Cette méthode qui permet un résultat rapide, est utilisée surtout pour les eaux très polluées. On sépare généralement la DCO particulaire de la DCO soluble pour différencier les matières en suspension (partie organique) des matières organiques solubilisées.
- *Demande biochimique en oxygène* (DBO) : représente la consommation d'oxygène (en mgO₂/L) résultant de la métabolisation de la pollution organique biodégradable par les micro-organismes. L'échantillon prélevé est dilué par une solution contenant les substances nutritives appropriées et une semence d'un grand nombre de populations bactériennes. La procédure d'analyses biochimique dépendant des prévisions de la charge polluée, normalement plusieurs tests sont effectués avec différentes dilutions, afin de couvrir la gamme complète des concentrations possibles. Le mélange est laissé en incubation dans un flacon étanche, isolé de l'air, pour faciliter les mesures d'oxygène utilisé durant les réactions de métabolisme. Un test supplémentaire est réalisé sur la culture de micro-organismes avec une dilution d'eau, afin d'établir les corrections liées à la fraction d'oxygène consommée par la semence biologique. Les résultats de ce test dépendent de la température ainsi que de la période d'incubation. Si la température standard est fixée à 20 °C, la durée d'oxydation peut varier. La Demande Biochimique en Oxygène (DBO) est souvent mesurée après 5 jours, la notation utilisée dans ce cas est généralement DBO₅. Au cours de cette période, l'oxydation n'est pas totale (de 60 à 70% de la réaction complète), alors qu'elle est estimée entre 90% et 95% après une incubation de 20 jours. La DBO a été reconnue depuis très longtemps comme la principale mesure de la pollution organique.
- *Azote global* (NGI) : quantité totale d'azote (en mgN/L) correspondant à l'azote organique (N_{org}) et ammoniacal (ion ammonium, NH₄⁺) et aux formes minérales oxydées de l'azote : nitrates (NO₃⁻) et nitrites (NO₂⁻). L'analyse de l'ammoniac est réalisée sous un pH élevé par la technique de minéralisation (chauffage et condensation) et un test de colorimétrie. Le test Kjeldahl consiste à faire subir à un échantillon, un processus de digestion où l'azote organique est transformé en ammoniac. Par conséquent, l'azote Kjeldahl (NTK) représente l'azote organique et ammoniacal. Les formes oxydées (nitrates et nitrites) sont mesurées par colorimétrie.

- *Phosphore total* (P_T) : quantité (en mgP/L) correspondant à la somme du phosphore contenu dans les orthophosphates, les polyphosphates et le phosphate organique. Le phosphore qui pollue les eaux est en majeure partie sous forme de phosphates (PO_4^{3-}). Typiquement ce composé est déterminé directement par addition d'une substance chimique qui forme un complexe coloré avec le phosphate.

On pourrait y rajouter des mesures plus spécifiques concernant la présence de toxiques d'origine minérale (mercure, cadmium, plomb, arsenic...) ou organique (composés aromatiques tels que le phénol, PCP...). On trouvera aussi les mesures du Carbone Organique Total (COT), autre mesure de la quantité de matière organique, des Matières Volatiles en Suspension (MVS) qui représentent la partie organique (donc biodégradable) des MES, ou encore des Matières Oxydables (MO). Cette dernière est définie comme:

$$MO = \frac{2DBO_5 + DCO}{3}$$

Cette mesure est particulièrement utilisée par les Agences de l'Eau pour établir les quantités de matières organiques présentes dans un effluent.

1.1.3. La législation : quelques éléments

Les eaux résiduelles sont devenues un des problèmes environnementaux les plus critiques, si nous considérons que l'accroissement démographique de la majorité des centres urbains moyens et grands, est important dans plusieurs pays¹² et régions du monde. Cette situation se reflète dans l'augmentation des décharges de type domestique et productif, en détériorant chaque fois plus l'état de la qualité de l'eau comme ressource. La situation est plus critique lorsqu'on altère les conditions de qualité de l'eau requises pour l'approvisionnement d'activités spécifiques (domestique, industriel, agricole, d'élevage, etc..) et la vie aquatique. Afin de faire face à cette problématique environnementale, plusieurs organismes et gouvernements se sont engagés à ce propos en développant un cadre légal ainsi que des programmes spécifiques pour le traitement des eaux en mettant en commun la problématique et les solutions auprès de l'ensemble des collectivités et des acteurs impliqués. Malgré le fait qu'en pratique on constate des différences propres à la culture et aux intérêts de chaque pays, à la base, l'esprit et les principes de ces normes restent les mêmes : favoriser la protection de l'environnement en gardant un équilibre avec le développement. A titre illustratif et en prenant en compte les pays impliqués dans le cas de la cotutelle de thèse, nous présentons quelques éléments sur la législation en France et en Colombie.

La législation en France

La législation française sur la pollution des eaux, les conditions de rejet et leur traitement repose en grande partie sur la loi sur l'eau n°92-3 du 3 janvier 1992 et les décrets du 29 mars 1992 et du 3 juin 1994 [STE98] [QUE00]. Les arrêtés prévus par ces décrets ont permis à la France de transposer en droit interne les directives européennes "eaux résiduaires urbaines" du 21 mai 1991. Ils imposent aux communes, sur l'ensemble du territoire français, l'élaboration et la mise en œuvre d'un programme d'assainissement avant le 31 décembre 2005, prenant en compte la collecte et le traitement biologique des eaux résiduaires urbaines.

¹² En Colombie le phénomène est probablement plus accentué étant donné la situation actuelle (2007) socio-économique et d'ordre public du pays.

Les conditions de rejets sont fixées par les arrêtés du 22 novembre 1994 pour les effluents urbains et des 1^{er} mars 1993 et 25 avril 1995 pour les effluents industriels. Ces arrêtés précisent en particulier les caractéristiques physico-chimiques des rejets, avec, en particulier, les valeurs limites des critères de pollution (MES, DCO, DBO₅, NGI et P_T) fixées en concentrations et en rendements. Les caractéristiques générales des rejets sont détaillées dans le Tableau 1.1. Les valeurs limites sont résumées dans le Tableau 1.2. pour les rejets urbains et dans le Tableau 1.3. pour les rejets industriels.

Caractéristique	Effluents urbains	Effluents industriels
pH	6 < . < 8,5	5,5 < . < 8,5
Température	< 25°	< 30°
Couleur		< 100 mg Pt/L

Tableau 1.1. Caractéristiques générales des rejets

Remarque

R 1.1 Pour les rejets dans des écosystèmes et milieux aquatiques sensibles à l'eutrophisation, la législation est plus contraignante, tout au moins en ce qui concerne l'azote et le phosphore.

Ainsi, l'application de la Directive Européenne et de la loi sur l'eau de 1992 nécessite non seulement une extension des stations au traitement de l'azote et du phosphore, mais également une fiabilisation de ces traitements, c'est à dire un respect continu des niveaux de rejets. A titre illustratif, en Midi-Pyrénées¹³ par exemple, la situation été relativement alarmante au début du XXI siècle puisque environ 70% des pollutions azotées et 80% des pollutions phosphorées ne sont pas traitées dans les stations d'épuration domestiques.

Paramètre	Pollution journalière kg/L	Valeur limite moyenne/24h mg/L	Rendement minimal %
MES	toutes charges	35	90
DCO	toutes charges	125	75
DBO ₅	120 à 600 > 600	25	70 80
NGI	600 à 6000 > 6000	15 10	70 70
P _T	600 à 6000 > 6000	2 1	80 80

Tableau 1.2. Valeurs limites des rejets urbains

¹³ La région de Midi-Pyrénées avec ses 45348 m² représente 8,3% de la superficie française. Elle est constituée de 8 départements ; parmi eux la Haute-Garonne avec sa capitale Toulouse.

Paramètre	Flux journalier autorisé kg/L	Valeur limite	
		moyenne/24h mg/L	moyenne/mois mg/L
MES	≤ 15 > 15	100 35	
DCO	≤ 100 > 100	300 125	
DBO ₅	≤ 30 > 30	100 30	
NGI	≤ 50	—	30
P _T	≤ 15	—	10

Tableau 1.3. Valeurs limites des rejets industriels

La législation en Colombie

Pour la Colombie, les évaluations reportent que les centres urbains dans le pays recueillent autour de 170 m³/s d'eau laquelle est perdue entre 40% et 50 %. Entre 70% à 80% des eaux consommées retournent à l'environnement sous forme d'eaux résiduelles. On estime qu'en Colombie on apporte quotidiennement environ 700 tonnes de charge organique du secteur domestique urbain aux corps d'eau. L'inventaire des systèmes de traitement d'eaux résiduelles du Ministère de l'Environnement, reporte que seulement 22% des municipalités du pays effectuent un traitement des eaux résiduelles et beaucoup présentent des disfonctionnement. On reporte que les départements qui ont une plus grande couverture d'usines de traitement d'eaux résiduelles (PTAR) sont Cundinamarca (38 PTAR), Antioquia (26 PTAR), Cesar (14 PTAR), Valle del Cauca (14 PTAR) et Tolima (13 PTAR).

En ce qui concerne les normes sur les déversements d'eaux résiduelles et les aspects institutionnels, elles sont fondées sur des politiques nationales et des normes spécifiques mentionnées depuis les années 70. On souligne principalement le Code des Ressources Naturelles (décret de loi 2811 de 1974), le décret 1594 de 1984 et le Règlement Technique d'Eau Potable et d'Assainissement (RAS¹⁴), comme principales normes de règlement environnemental et sanitaire [GUIA02].

Les communes doivent donc s'occuper des normes de déversement établies (cf. Tableau 1.4), et de certains paramètres d'intérêt sanitaire (e.g. Article 72 Décret 1594 de 1984).

¹⁴ RAS : Reglamento Técnico de Agua Potable y Saneamiento (RAS), Ministerio de Desarrollo Económico, Colombia, 2000

Paramètre	Installation déjà existante	Nouvelle Installation
pH	5 - 9	5-9
Température	< 40°	< 40°
Matières grasses et huiles	Efficacité > 80%	Efficacité > 80%
Matières en suspension	Efficacité > 50%	Efficacité > 80%
DBO ₅		
Niveau urbain	Efficacité > 30%	Efficacité > 80%
Niveau industriel	Efficacité > 20%	Efficacité > 80%

Tableau 1.4. Normes sur les rejets (Art. 72 Décret 1594 de 1984, Colombie)

1.1.4. Le traitement des eaux usées

L'eau est le véhicule de transport et de dissémination idéal de nombreux polluants. Les contraintes d'assainissement, de plus en plus strictes, exigent le traitement d'un nombre plus important de polluants (matières organiques, minérales, pathogènes et toxiques). Etant donnée la grande diversité de ces déchets, l'épuration d'un affluent résiduaire comporte plusieurs étapes, chacune spécifique aux caractéristiques particulières des éléments à traiter.

D'un point de vue général, est sans vouloir être exhaustif, compte tenu de la diversité des procédés mis en œuvre selon les cas, l'épuration de l'eau amène toujours avant leur rejet dans le milieu naturel à [QUE00] :

- séparer et éliminer les matières en suspension,
- éliminer la pollution organique, principalement par voie biologique, et, plus récemment les pollutions azotées et phosphorées.

Le schéma fonctionnel de la filière d'épuration des eaux résiduaires est indiqué sur la Figure 1.1 [HAD99] :

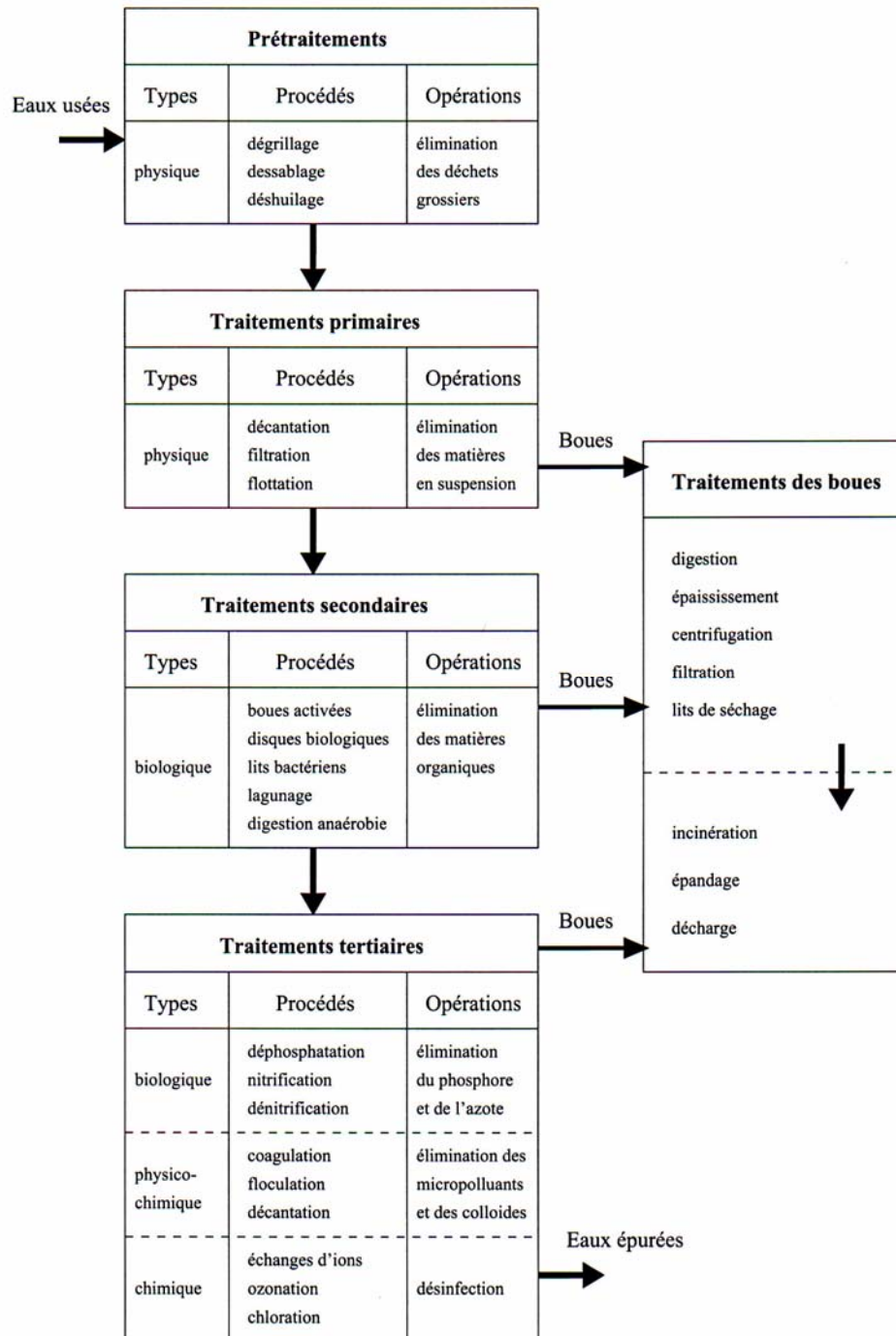


Figure 1.1. Filière d'épuration des eaux résiduaires (d'après [HAD99]).

Globalement, la filière de traitement est constituée de différents modules où les eaux brutes sont soumises à une combinaison ou une succession de différents processus de purification. L'épuration d'un effluent pollué peut comporter cinq phases principales :

- *Le prétraitement* : est un traitement préliminaire comportant un certain nombre d'opérations à caractère mécanique (procédés physiques). Le but est d'extraire de l'eau les gros déchets en suspension ou en flottation (sables, huiles) qui pourraient gêner les traitements subséquents. L'élimination préalable de ces matières permet d'éviter des effets nocifs secondaires (odeurs, colmatage...). Parmi ces méthodes de séparation primaires, les plus courantes sont :

- *le dégrillage*, qui a pour fonction la suppression des déchets les plus grossiers par passage à travers une grille.
- *le dessablage*, qui permet le dépôt du sable et des graviers susceptibles d'endommager les machines de pompage.
- *le déshuilage*, qui favorise, par injection d'air, la flottation des graisses et des hydrocarbures qui sont séparés par raclage de surface.
- *Le traitement primaire* : permet d'éliminer de l'eau plus de la moitié des matières en suspension sous forme de boues dites « boues primaires », recueillies ensuite par pompage de fond. Il fait appel à différents procédés physiques :
 - *La flottation*, visant à séparer les phases solides des phases liquides par la poussée d'Archimède. En flottation naturelle, les floes de faible densité remontent librement à la surface. La flottation assistée s'obtient par injection d'air.
 - *La décantation*, permet aux matières en suspension de se déposer, sous forme de boues, en utilisant la force de gravité pour séparer les particules de densité supérieure à celle du liquide.
 - *La filtration*, est le passage du mélange liquide-solide à travers un milieu poreux (filtre) qui retient les solides et laisse passer les liquides.
- *Le traitement secondaire* : a pour objectif principal l'élimination des matières en suspension et des composants solubles dans l'eau (matières organiques, substances minérales...). Il fait habituellement appel aux procédés biotechnologiques, car les matières présentes dans les eaux usées sont généralement d'origine organique. Au terme du traitement secondaire, l'eau, débarrassée des éléments qui la polluaient est épurée à 90%. Les matières polluantes agglomérées à la suite de ce traitement sont recueillies par décantation sous forme de boues. On utilise typiquement :
 - *Les traitements physico-chimiques*, qui consistent à transformer chimiquement, à l'aide de réactifs, les éléments polluants non touchés par les traitements biologiques (matières non biodégradables).
 - *Les traitements biologiques*, qui sont appliqués aux matières organiques en utilisant des cultures de microorganismes (notamment bactéries) reproduisent le processus de l'auto-épuration naturelle dans des bassins adaptés à ce propos (les bioréacteurs). Plus récemment, le traitement biologique de l'azote a été intégré à cette étape, et de la même manière, le traitement des phosphates commence aussi à y être intégré. Les impuretés sont alors digérées par des êtres vivants microscopiques et transformées en boues. La culture des bactéries se fait soit en milieu aéré (aérobie), soit en absence d'oxygène (anoxie). On distingue aussi les cultures fixées (lits bactériens, disques biologiques), et les cultures libres (lagunage aéré, boues activées). On reviendra dans la section 1.2. plus en détail sur le traitement biologique des eaux usées.
- *Le traitement tertiaire* : est un traitement complémentaire qui a un rôle d'affinage, dans le but soit d'une réutilisation des eaux épurées à des fins agricoles ou industrielles (e.g. refroidissement des turbines), soit de la protection du milieu récepteur pour des usages spécifiques (rejet dans les milieux aquatiques en zone plus ou moins sensible), soit encore de la protection des prises d'eau situées en aval. Différentes méthodes de nature variée peuvent alors être utilisées :
 - *biologiques*, pour l'élimination de l'azote (processus de nitrification-dénitrification) et du phosphore (déphosphatation).

- *physico-chimiques*, pour la précipitation du phosphore (coagulation-décantation) ou l'élimination des dernières matières en suspension (filtration sur lits de sable, tamis métalliques ou charbon actif). L'élimination de l'azote et du phosphore par voie biologique ou chimique évite la prolifération de végétaux dans les corps d'eau¹⁵ (les lacs, les étangs ou les rivières) et protège la vie aquatique.
- *chimiques*, pour la désinfection dans le traitement final des effluents. L'élimination des risques de contamination bactériologique ou virale se fait souvent par *oxydation* en utilisant des agents tels que le chlore et l'ozone. Les procédés chimiques sont utilisés aussi pour l'adoucissement de l'eau et pour la *neutralisation* en agissant sur le pH.
- *radiatifs*, pour les opérations de désinfection de l'eau telle que les rayonnements ultraviolets qui irradiant les cellules vivantes indésirables permettant d'éliminer les risques de contamination due aux bactéries et virus. Suivant la quantité d'énergie UV reçue, elles sont soit stérilisées (effet bactériostatique) soit détruites (effet bactéricide).
- *Le traitement des boues* : a pour but le traitement et le conditionnement des boues résiduelles extraites du décanteur, en permettant de réduire leur volume en éliminant l'eau (épaississement, déshydratation) et les cas échéant de les stabiliser (digestion, compostage). Pour ces boues, quatre destinations sont typiquement possibles :
 - *l'épandage agricole*, qui représente une valorisation de ce sous-produit fertilisant.
 - *l'élaboration de compost*, par incorporation de paille ou de sciure.
 - *l'incinération*, pour quelques grosses unités ou lorsqu'une installation locale existe déjà pour les ordures ménagères.
 - *la mise en décharge*, solution qui devrait être progressivement abandonnée dans les années à venir.

Remarque

R 1.2 Compte tenu de l'évolution des directives, le traitement des pollutions dues aux nitrates et phosphates ainsi que les différentes opérations de traitement biologique sont passé peu à peu du traitement tertiaire vers le traitement secondaire. Mais ceci n'est qu'affaire de présentation et ne change rien dans les principes [QUE00] .

1.2. Le traitement biologique des eaux usées

Le traitement biologique des effluents dans des installations appropriées est un moyen efficace pour répondre aux problèmes environnementaux liés aux activités humaines et compléter l'activité des micro-organismes dans les écosystèmes, en imitant les cycles naturels d'auto-épuration. Les procédés de traitement biologique de l'eau sont particulièrement adaptés à l'épuration d'eaux polluées essentiellement par de la matière organique facilement biodégradable et, dans tous les cas, exemptes de composés toxiques à des concentrations importants. Ces procédés sont donc particulièrement adaptés à l'épuration des eaux résiduelles urbaines. Les eaux industrielles nécessitent généralement des traitements spécifiques. Elles

¹⁵ L'apport constant de substances nutritives (nitrates et phosphates) au corps d'eau facilite la prolifération démesurée d'algues et autres végétaux aquatiques qui favorisent l'appauvrissement du milieu en oxygène, phénomène connu sous le terme d'*eutrophisation*.

peuvent parfois rejoindre la station d'épuration, au prix toutefois d'un traitement physico-chimique préalable, car la présence de toxique détruirait la flore bactérienne [QUE00].

Le principe général d'un procédé biologique, ou bioprocédé, est d'utiliser les propriétés naturelles d'organismes vivants afin de produire ou d'éliminer certaines substances chimiques ou biochimiques, dans des conditions optimales de fonctionnement. En effet, certains micro-organismes ont de grandes facultés de transformation métabolique et de décomposition des matières biodégradables. Ils constituent par leur multiplication rapide et leur action biochimique, des agents épurateurs très efficaces.

Dans cette section nous présentons en quelques lignes une vue panoramique du traitement biologique des eaux usées. En particulier nous insisterons sur quelques points particuliers : les substrats polluants, les micro-organismes épurateurs, les processus métaboliques et les traitements biologiques [HAD99] [MET03].

1.2.1. Les principaux polluants

Les phénomènes de pollution des eaux se traduisent par des effets particuliers liés aux spécificités écologiques propres aux milieux aquatiques. En effet, l'eau peut dissoudre, souvent avec facilité, de nombreuses substances chimiques et biologiques. Par conséquent, tout polluant peut être véhiculé fort loin de la source de contamination. La problématique des déchets présents dans l'eau peut être abordée de plusieurs façons, qui donnent chacune lieu à une classification différente. Ainsi, les impuretés peuvent être identifiées suivant qu'elles soient vivantes ou inertes, minérales ou organiques, solides ou dissoutes. D'autres techniques de classification sont basées sur leur dimension, leurs degrés de toxicité,... Parmi les principaux polluants on peut distinguer les suivants :

- *Les matières organiques* : constituent, de loin, la première cause de pollution des ressources en eaux. Ces matières organiques (déjections animales et humaines, graisses,...) sont notamment issues des effluents domestiques, mais également des rejets industriels (industries agro-alimentaire, en particulier). La pollution organique peut être absorbée par le milieu récepteur tant que la limite d'auto-épuration n'est pas atteinte.
- *Les éléments minéraux* : regroupent essentiellement les produits azotés ainsi que les produits phosphorés. Ces matières proviennent principalement des activités agricoles. La pollution minérale des eaux peut provoquer le dérèglement de la croissance végétale ou des troubles physiologiques chez les animaux.
- *Les métaux lourds* : les plus fréquemment rencontrés mais qui sont aussi les plus dangereux sont le mercure, le cuivre, le cadmium, le chrome, le plomb et le zinc. Ils ont la particularité de s'accumuler dans les organismes vivants ainsi que dans la chaîne trophique. La pollution radioactive peut avoir des effets cancérigènes et mutagènes sur les peuplements aquatiques.
- *Les matières pathogènes* : sont constituées de virus et bactéries entraînant souvent une inhibition des mécanismes biologiques. La pollution micro-biologique se développe conjointement à la pollution organique, par une prolifération des germes d'origine humaine ou animale dont certains sont éminemment pathogènes.
- *Les substances toxiques* : sont des composés chimiques de synthèse, issus des activités industrielles et agricoles. Les conséquences souvent dramatiques de la pollution chimique sur les écosystèmes, varient suivant la concentration de composés dans les rejets.

- *Les hydrocarbures* : provenant des industries pétrolières et des transports, ces composés chimiques sont des substances peu solubles dans l'eau et sont difficilement biodégradables. Leur densité inférieure à l'eau les fait surnager et leur vitesse de propagation dans le sol est 5 à 7 fois supérieure à celle de l'eau. Ils constituent un redoutable danger pour les nappes phréatiques.

Une autre classification très importante est fondée sur le pouvoir de dégradation des déchets polluants. On distingue ainsi deux classes principales :

- *Les matières biodégradables* : affectées par les activités biologiques des micro-organismes, ces substances sont soumises aux divers processus biochimiques de conversion. Cette fraction biodégradable peut être structurée en deux groupes :
 - *Les matières aisément dégradables*, composées des substances solubles, ces matières ont la caractéristique de pouvoir être directement absorbées par les bactéries.
 - *Les matières lentement dégradables*, composées des substrats particuliers formés par un mélange de substances organiques solides, colloïdales et solubles. Ces matières sont soumises à certains processus intermédiaires avant d'être absorbées par les populations bactériennes.
- *Les matières non-biodégradables* : ces substances inertes ne subissent aucun phénomène biologique de transformation. Ces matières sont soit présentes dans les eaux résiduaires, comme les métaux lourds, soit issues des phénomènes de mortalité des micro-organismes au cours des processus biologiques d'épuration. Les composants non-biodégradables solubles peuvent traverser la station d'épuration sans être modifiés mais les matières inertes en suspension peuvent être éliminées par des mécanismes de décantation.

La structure chimique des polluants permet de distinguer deux types de composés :

- *Les matières organiques* : elles sont constituées d'un grand nombre de composés qui ont la particularité commune de posséder au moins un atome de carbone, d'où leur nom de substance carbonées. Ces atomes de carbone sont oxydés biologiquement par les micro-organismes pour fournir l'énergie nécessaire à leur croissance.
- *Les matières inorganiques* : sont des substances ne contenant pas de carbone. La fraction minérale des eaux résiduaires représente principalement les produits azotés et phosphorés.

1.2.2. Les micro-organismes épurateurs

Le monde vivant est classé en trois catégories principales : les végétaux, les animaux et les protistes (organismes eucaryotes unicellulaires), qui se distinguent des deux autres règnes par la structure relativement simple et la multiplication rapide de leurs individus. Ces micro-organismes sont composés essentiellement des bactéries, des algues et des protozoaires (prédateurs des bactéries). Certains de ces populations microbiologiques ont la faculté de dégrader les substances polluantes présentes dans les eaux résiduaires pour les convertir en eau, en dioxyde de carbone et en matières minérales dont l'effet polluant est moins nuisible pour les milieux récepteurs. Ces micro-organismes sont à la base de l'épuration biologique qui est le procédé le plus utilisé pour restaurer la qualité de l'eau en la débarrassant de ses principales impuretés pourvu qu'elles soient plus ou moins biodégradables et ne contiennent pas de toxiques qui font l'objet d'un traitement particulier (épuration physico-chimique).

Parmi tous les individus du monde protiste, trois populations jouent un rôle fondamental dans le traitement :

- *Les bactéries* : unicellulaires, ces micro-organismes possèdent la structure interne la plus simple de toutes les espèces vivantes. Les tests effectués sur différentes populations ont montré que les bactéries sont composées de 80% d'eau et 20% de matière sèche dont 90% est organique. Elles croissent et se multiplient en général par scissiparité (scission binaire). Ces cellules représentent la plus importante population de la communauté microbienne dans tous les procédés biologiques, avec souvent des concentrations qui dépassent 10^6 bactéries/mL.
- *Les protozoaires* : de structure plus complexe que celle des bactéries, la distinction de protozoaires est plus simple. Certains groupes de protozoaires sont de redoutables prédateurs pour les bactéries. Ils ont la faculté de se déplacer et sont classifiés suivant leur mode de mouvement (nageurs, rampants,...). Ces organismes peuvent jouer un rôle important au cours du processus d'épuration par leur abondance et leurs interactions avec les bactéries épuratrices (compétition et prédation).
- *Les algues* : ce sont des organismes photosynthétiques unicellulaires ou multicellulaires formant une population hétérogène. Les algues sont indésirables dans les sources d'eau car elles affectent leur goût et leur odeur. Dans le traitement, on les retrouve dans deux types de procédés uniquement : les lits bactériens ainsi que les bassins de lagunage, mais ce n'est que dans ces derniers qu'elles jouent un rôle bénéfique dans l'épuration.

Une grande partie du poids sec des micro-organismes est constitué en général, du carbone (50%), d'oxygène (20%), d'azote (10 à 15%), d'hydrogène (8 à 10%) et du phosphore (1 à 3%). Afin d'accomplir ses fonctions biologiques, cette faune microscopique se nourrit des matières organiques qui contiennent ces éléments nutritifs vitaux pour son développement. La consommation des polluants permet aux populations bactériennes de :

- fournir la matière nécessaire pour la synthèse cytoplasmique dont le composé fondamental est le carbone ;
- assurer leurs besoins nutritifs pour leur croissance, l'azote étant nécessaire pour la synthèse des protéines par exemple ;
- servir comme source d'énergie pour la croissance cellulaire et les réactions biosynthétiques, le phosphore joue un rôle très important dans le transfert d'énergie ;
- servir comme accepteur final d'électrons libérés au cours des réactions d'oxydation, l'oxygène joue le rôle d'agent oxydant dans le processus de respiration aérobie.

Selon la nature de leurs besoins nutritifs, les micro-organismes peuvent être classés en plusieurs catégories. Sans vouloir rentrer dans les détails d'autres classifications plus élaborées, les catégories les plus couramment utilisées sont [HAD99] :

- *Forme chimique du carbone nécessaire* :
 - *les autotrophes*, utilisent le dioxyde de carbone (CO_2) et le carbonate (HCO_3^-) comme unique source de carbone pour la synthèse de leurs biomolécules (les plantes par exemple).
 - *les hétérotrophes*, utilisent le carbone sous forme organique relativement complexe pour la synthèse de nouvelles cellules vivantes (les animaux par exemple).

- *Source d'énergie :*
 - les *phototrophes*, utilisent la lumière solaire comme source d'énergie (les plantes par exemple).
 - les *chimiotrophes*, obtiennent leur énergie à partir des réactions d'oxydo-réduction (les animaux par exemple).
- *Accepteur final d'électrons :*
 - *aérobies*, se servent des molécules d'oxygène comme accepteurs d'électrons. Les micro-organismes aérobies se développent donc en présence d'oxygène (les animaux par exemple).
 - *anaérobies*, ont la faculté d'utiliser d'autres molécules que l'oxygène comme accepteurs d'électrons (dioxyde de carbone, sulfate,...). Les micro-organismes anaérobies peuvent se passer d'oxygène pour se développer (certaines souches de bactéries comme les acidogènes et les méthanogènes par exemple).
 - *facultatifs*, utilisent comme agents oxydants, l'oxygène ou d'autres composants chimiques comme le nitrate (conditions d'anoxie). Cependant, leur croissance est plus efficace sous les conditions d'aérobiose (les plantes par exemple).

La Figure 1.2 présente un résumé des populations biologiques et leurs besoins nutritifs et énergétiques :

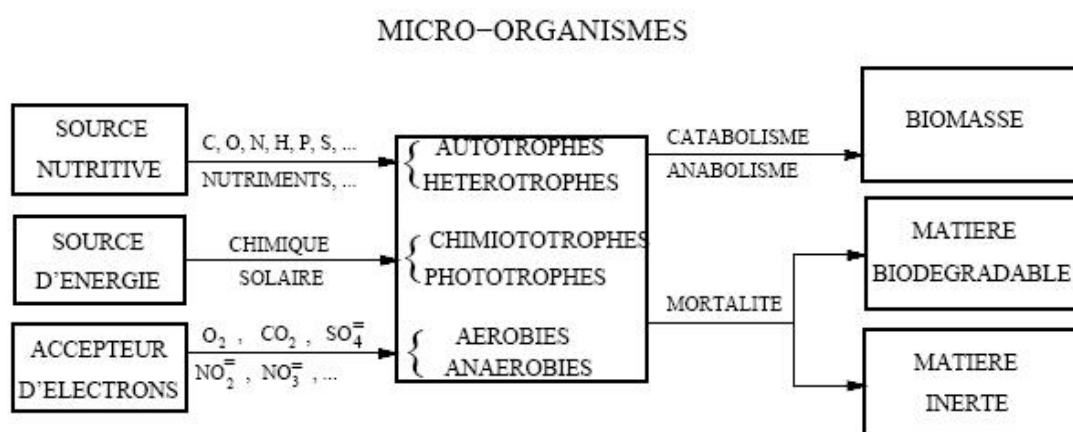


Figure 1.2. Les micro-organismes et leurs besoins nutritifs et énergétiques (D'après [HAD99]).

1.2.3. Les processus métaboliques

Les composants chimiques consommés par les micro-organismes sont soumis à de nombreuses réactions biochimiques qui font partie d'un des deux mécanismes métaboliques fondamentaux pour le développement des bactéries :

- *le catabolisme* : représente l'ensemble des réactions d'oxydation et de dégradation enzymatique. C'est une activité exothermique qui libère l'énergie inhérente à la structure complexe des molécules organiques et minérales, et qui est stockée par les micro-organismes.
- *l'anabolisme* : représente l'ensemble des réactions de réduction et de synthèse enzymatique. C'est une activité endothermique qui utilise l'énergie libérée par les processus catabolique pour développer la taille et la structure chimique des composants organiques.

Parmi le grand nombre de processus biochimiques mis en jeu au cours du traitement biologique des polluants par les différentes populations bactériennes, on peut distinguer principalement les activités suivantes :

- *Oxydation* : c'est une réaction qui implique une perte d'électrons suivie d'une production d'énergie. Une partie des matières absorbées par les micro-organismes est utilisée pour fournir l'énergie nécessaire afin d'accomplir leurs fonctions biologiques. Selon la nature de l'accepteur final d'électrons, le processus d'oxydation peut s'effectuer sous les conditions d'aérobiose (présence d'oxygène), d'anaérobiose (absence d'oxygène) ou d'anoxie (présence de nitrate).
- *Digestion* : la caractéristique de la digestion appelée aussi *fermentation* est que ce processus ne nécessite pas un accepteur d'électrons externe. C'est un mécanisme anaérobie de production d'énergie qui n'implique pas de chaîne de transport d'électrons. La fermentation est provoquée par des bactéries anaérobies capables de décomposer la matière organique en acides et alcools et de donner du méthane (CH_4) et du gaz carbonique (CO_2).
- *Nitrification* : c'est une transformation chimique de l'azote organique en nitrate (NO_3^-) par des organismes dits nitrifiants (*Nitrosomas* et *Nitrobacter*). La nitrification a lieu en trois étapes dans les conditions d'aérobiose : l'*ammonification* (conversion de l'azote organique en ammoniac (NH_4^+)), la *nitritation* (oxydation de l'ion ammoniac en nitrite (NO_2^-)) et la *nitratation* (oxydation du nitrite en nitrate).
- *Dénitrification* : c'est un processus de conversion du nitrate, effectué par les hétérotrophes facultatifs sous les conditions d'anoxie. La dénitrification peut avoir lieu selon deux activités biologiques différentes : l'*assimilation*, où le nitrate est réduit en ammoniac qui peut servir comme source d'azote pour la synthèse cellulaire et la *dissimilation*, qui joue un rôle plus important dans l'élimination totale du nitrate.
- *Floculation* : dans les cultures en suspension où les micro-organismes flottent librement dans les eaux à traiter, ceux-ci ont tendance à s'agglutiner sous forme de petits amas appelés bioflocs. Certaines substances biodégradables (particules, colloïdes, grosses molécules) ne sont pas directement absorbées par les bactéries. La floculation permet aux micro-organismes d'améliorer leurs caractéristiques d'adsorption des aliments sur leur membrane cellulaire.
- *Adsorption* : certains composés organiques comme les substances biodégradables particulières ne peuvent pas être directement absorbés par les bactéries. Ces matières sont d'abord adsorbées par les micro-organismes et stockées à leur surface avant de subir un processus de conversion qui génère des substances simples aisément dégradables.
- *Hydrolyse* : c'est une série de réactions enzymatiques extra-cellulaires appliquées aux substances adsorbées et qui ont lieu à la surface des micro-organismes. Les molécules organiques complexes sont converties en molécules plus simples qui peuvent diffuser à travers la membrane cellulaire.
- *Mortalité* : la population biologique est soumise à divers phénomènes de mortalité : la prédation (protozoaires), le processus de lyse (dissolution) et la respiration endogène où une partie des composants cellulaires est oxydée pour satisfaire les besoins en énergie nécessaires pour maintenir en vie les autres cellules quand le substrat n'est pas disponible.

Une fraction des produits de mortalité est biodégradable, l'autre partie représente les résidus endogènes inertes.

1.2.4. Les traitements biologiques

Le traitement biologique est largement utilisé dans la dépollution de l'eau, compte tenu des grandes facultés de certains micro-organismes, notamment des bactéries, de transformation métabolique et de décomposition de la matière organique biodégradable, lorsque les composés toxiques sont à des concentrations faibles. L'épuration biologique consiste alors à favoriser la prolifération de ces micro-organismes pour utiliser leurs propriétés remarquables dans les conditions les plus adaptées au résultat désiré. Ces conditions reproduisent le processus de l'auto-épuration naturelle dans les bioréacteurs, des bassins adaptés à ce propos. Les impuretés sont alors digérées par les micro-organismes et transformées en boues. Les polluants ingérés par la biomasse sont d'origine organique mais certains éléments minéraux ou inertes sont affectés par les processus métaboliques ou floculés au cours des réactions biochimiques [MET03].

L'élimination de la production polluante conduit toujours, en fonction des caractéristiques physico-chimiques des rejets et du degré d'épuration souhaité, à la conception d'une chaîne de traitement constituée d'une succession d'opérations unitaires ou de stades de traitement entre lesquels il existe généralement des interactions. On peut cependant, dans de nombreux cas, simplifier le fonctionnement global du cycle d'épuration de l'eau (de son prélèvement à son renvoi dans le milieu naturel, en passant par son utilisation, les traitements physico-chimiques et biologiques, et la gestion des boues produites) en plusieurs sous-systèmes traités indépendamment. C'est en particulier le cas pour les stations de traitement biologique, dans lesquelles les traitements de la pollution peuvent être, selon les cas, considérés de manière globale ou découpée [QUE00].

Historiquement, le premier objectif des installations de traitement des eaux usées était l'élimination de la pollution carbonée. La nitrification, mieux connue et mieux maîtrisée que la dénitrification, a ensuite été ajoutée à ces installations. Les procédés de dénitrification sont apparus plus tardivement et se sont intégrés aux systèmes nitrifiants préexistants. L'étape de traitement du phosphore, quant à elle, commence aussi à y être intégrée dans les procédés modernes [STE98]. Cette chronologie historique explique le nombre important de réacteurs en cascade pouvant être observés dans les configurations de traitement tertiaire existantes à l'heure actuelle.

Parmi les différentes possibilités de classement de bioprocédés de traitement des eaux usées, nous voulons faire appel à ceux qui correspondent au trois critères suivants : l'emplacement des populations bactériennes, les conditions de l'environnement et les types de procédés biologiques.

Suivant l'emplacement des bactéries épuratrices dans le bioréacteur, on distingue deux dispositifs [EDE97]:

- *Le procédé à cultures libres* : Les micro-organismes sont maintenus en suspension dans le mélange à épurer. La biomasse entre ainsi constamment en contact avec les polluants. Ces dispositifs ont l'avantage d'avoir un traitement plus homogène et une meilleure maîtrise des facteurs d'épuration (apport d'eau résiduaire et de masse bactérienne) comparés aux procédés à culture fixée.

- *Le procédé à cultures fixées* : Les micro-organismes sont fixés sur des supports. Le contact entre les eaux à traiter et les cellules épuratrices est assuré soit par arrosage des supports avec l'eau usée (lits bactériennes), soit par rotation des supports dans le mélange pollué (disques biologiques). Ces procédés permettent d'obtenir des concentrations en biomasse plus importantes et donc des traitements intensifs, mais présentent des risques de colmatage ou d'émanation d'odeurs. Une purge des boues accumulées dans l'ouvrage est cependant régulièrement nécessaire pour prévenir le colmatage du filtre. Néanmoins, les procédés à biomasse fixée gagnent en compacité par rapport aux procédés à boues activées et ne nécessitent pas la présence d'un décanteur en sortie.

Suivant les conditions de l'environnement des cellules dans l'unité de dépollution, on distingue deux modes de traitement [HAD99] :

- *Le traitement aérobie* : Ce type de traitement fait appel aux bactéries aérobies qui se développent en présence d'oxygène. La dégradation des polluants est effectuée par des réactions d'oxydation dans un milieu aéré.
- *Le traitement anaérobie* : Ce traitement s'effectue en condition d'anaérobiose, c'est-à-dire, en absence d'oxygène. Les bactéries anaérobies assurent la décomposition métabolique des composés biodégradables par des processus de fermentation.

Parmi l'ensemble des procédés biologiques utilisés dans le traitement des eaux usées, on peut citer les principaux procédés suivants :

- *Les lits bactériens* : Ce procédé aérobie à cultures fixées consiste à faire supporter les micro-organismes par des matériaux poreux. L'effluent est distribué par aspersion en surface et l'oxygénation est apportée par ventilation naturelle de bas en haut. L'affluent arrive par la partie supérieure alors que l'effluent est évacué par le fond afin de ne pas perturber la fonction aérobie. De ce fait, ce système présente un inconvénient majeur, dans le sens qu'il nécessite un dispositif de relevage. La biomasse se développe à la surface du support. Lorsqu'elle devient trop importante, la pellicule bactérienne se détache naturellement ; elle doit alors être séparée de l'effluent par décantation.
- *Les disques biologiques* : Dans ce procédé, les micro-organismes sont fixés sur des disques à demi immergés et tournent lentement (quelques tours par minute) autour d'un axe horizontal. La biomasse est ainsi alternativement mouillée par les eaux résiduaire et aérée par l'air ambiant. Cette technique présente l'avantage d'être peu coûteuse en énergie mais peut entraîner l'émanation d'odeurs.
- *Les boues activées* : Ce système comprend deux compartiments principaux. Le premier est le *bassin d'aération* où ont lieu les activités biologiques de transformation des polluants biodégradables par l'intermédiaire des micro-organismes en suspension. Le bassin est associé à un décanteur-clarificateur qui permet de recycler les boues en tête du bassin d'une part, et de récupérer l'effluent traité d'autre part. Outre les matières organiques assimilés par les hétérotrophes, principaux constituants des boues activées, les composés azotés peuvent aussi être oxydés par des phénomènes de nitrification-dénitrification. Les bactéries floculantes utilisées dans ce système, ont la faculté de transformer les éléments ingérés en matière corpusculaire. Les flocs formés dans le bassin d'aération sont alors conduits vers un second compartiment appelé *décanteur secondaire* où a lieu la séparation des solides de la phase liquide par décantation.

De nombreuses variantes du procédé par boues activées sont disponibles. Par exemple, selon le cas, le bassin d'aération peut être constitué de plusieurs bassins en série privilégiant chacun le traitement d'une pollution spécifique (organique, nitrification, dénitrification, déphosphatation), ou d'un seul bassin permettant de réaliser les différentes réactions biologiques simultanément, par exemple par alternance de phases aérées et non aérées intégrant même au sein d'un unique bassin des cultures aérobies et anoxies.

- *Le lagunage* : Il s'agit d'un étang ou un système de lagunes mettant en œuvre une culture mixte algo-bactérienne. Le lagunage est un procédé de traitement extensif qui repose sur le principe de la dégradation en eau libre de la pollution organique. Il s'agit d'une boue activée sans recyclage de boue, avec ou sans décanteur final [EDE97]. L'aération peut être naturelle ou artificielle, lorsque l'apport d'oxygène est assuré par des aérateurs externes. Suivant la profondeur du bassin, on peut distinguer différents régimes de fonctionnement.
 - En zone peu profonde, le traitement s'effectue dans des conditions d'aérobiose. Les deux populations vivent en symbiose. Les besoins en oxygène des bactéries sont principalement assurés par l'activité photosynthétique des algues exposées à la lumière. De leur côté, les végétaux profitent du gaz carbonique ainsi que des nutriments inorganiques produits au cours des réactions métaboliques des cellules vivantes.
 - Dans le cas de lagunes plus profondes, en plus de la zone supérieure aérobie, on peut distinguer une région intermédiaire où la disponibilité de l'oxygène dépend de la lumière solaire. Le traitement a lieu dans des conditions d'aérobiose le jour, et en anaérobiose durant la nuit. Les dépôts de boues au fond des bassins suffisamment profonds forment une couche anaérobie où ont lieu des processus de fermentation.
- *La digestion anaérobie* : Ce procédé repose sur les réactions de fermentation qui ont lieu dans des conditions d'anaérobiose et sont effectuées par une biomasse anaérobie complexe. De manière générale, on distingue trois étapes dans ce processus. Dans la phase de liquéfaction, des bactéries hydrolytiques solubilisent les matières complexes. Dans la deuxième étape, une population acétogène transforme les composés organiques en acides, qui serviront ultérieurement de nourriture à des bactéries méthanogènes. Au cours de l'étape finale les acides organiques sont transformés en méthane et gaz carbonique. Le métabolisme anaérobie libérant très peu d'énergie comparé au processus aérobie, provoque une faible production de boues. Nonobstant, l'épuration anaérobie est plus lente que l'épuration aérobie.

1.3. L'automatique des bioprocédés

Comme on vient de le voir dans les sections 1.1 et 1.2, les procédés de traitement des eaux usées et en particulier les procédés biologiques sont des systèmes complexes qui font appel à la croissance des micro-organismes par la consommation de substrats ou de nutriments en imitant les cycles d'auto-épuration trouvés dans la nature. Au niveau des stations d'épuration, cette croissance n'est possible qu'en présence de conditions « environnementales » favorables. Par conditions environnementales, on entend les conditions physico-chimiques (pH, température, agitation, aération,...) nécessaires à une bonne activité des micro-organismes. Pour parvenir à ces conditions et aux exigences techniques, économiques et environnementales actuelles, il est nécessaire de faire appel à des systèmes de mesure, de commande, de supervision et de surveillance. Ces problématiques sont traitées dans le domaine de l'Automatique.

C'est dans le cadre de cette approche que nous abordons dans cette section quelques aspects généraux concernant l'automatique des bioprocédés, en faisant référence aux travaux de [DOC01] et [QUE00]. Cela nous permettra ensuite de mieux comprendre les problèmes spécifiques, ainsi que la contribution que l'automatique peut apporter aux procédés de traitement biologique des eaux usées.

1.3.1. Les problèmes spécifiques de l'Automatique des bioprocédés

Le domaine des procédés biotechnologiques est assez vaste, en effet il englobe plusieurs types de procédés comme : les fermentations pour la production de substances chimiques à haute valeur ajoutée (e.g. pharmaceutiques), la production en grande quantité d'aliments et de boissons (e.g. yaourt, fromage, bière) ainsi que les traitements biologiques des résidus, qu'ils soient solides (compostage), liquides (e.g. boues activées) ou gazeux (biofiltres). L'utilisation industrielle des procédés biotechnologiques a considérablement augmentée ces dernières décennies pour diverses raisons (normes législatives plus sévères dans les industries de traitement, amélioration de l'efficacité et de la qualité des procédés, etc.). Les problèmes posés par cette industrialisation sont globalement les mêmes que ceux rencontrés dans n'importe quelle industrie des procédés et on rencontre, dans le domaine des bioprocédés, la quasi-totalité des problématiques abordées en Automatique. Ainsi, les besoins en systèmes de commande, de supervision et de surveillance des procédés sont nécessaires afin d'en optimiser le fonctionnement, d'établir une meilleure conduite ou encore de détecter des dysfonctionnements. Néanmoins, dans la pratique industrielle, très peu d'installations sont munies de tels systèmes. Deux raisons principales expliquent cet état de fait : *la complexité de ces procédés* et *l'absence de capteurs* dans la plupart des cas, à la fois peu coûteux et fiables pour des applications en ligne.

Les procédés biologiques sont des procédés complexes, due à la nature très variée des phénomènes multiples qui se déroulent en faisant intervenir des organismes vivants. Du point de vue de l'Automatique, ces procédés sont caractérisés de manière générale comme des systèmes multivariables, non-linéaires et non-stationnaires [ROU01]. De fait, la modélisation de ces systèmes se heurte à deux difficultés principales : d'une part, un manque de reproductibilité des expérimentations et une imprécision des mesures qui entraînent à la fois des difficultés relatives au choix de la structure des modèles, mais également des difficultés liées aux procédures d'identification; d'autre part, on rencontre également des difficultés lors de la phase de validation de ces modèles dont les jeux de paramètres peuvent précisément avoir évolué au cours du temps [DOC01]. Ces variations peuvent être la conséquence de changements métaboliques de la biomasse ou encore de modifications génétiques non prévisibles et inobservables d'un point de vue macroscopique.

Étant donné que les organismes vivants sont le cœur de ces procédés, il en résulte que la modélisation mathématique de ces procédés est une tâche difficile. En effet les modèles sont centraux au développement de systèmes de commande, donc la commande en ligne devient elle aussi complexe.

Une seconde difficulté majeure est l'absence quasi systématique de capteurs permettant d'accéder aux mesures nécessaires pour connaître le fonctionnement interne des procédés biologiques [VAN98]. C'est bien là que se trouve la principale limitation retardant l'automatisation complète des processus biotechnologiques, et de manière plus cruciale encore, des procédés de traitement des eaux usées. La plupart des variables clés de ces systèmes (concentrations en biomasses, en substrats et en produits) ne peuvent être mesurées

qu'à l'aide d'analyses de laboratoire, dont la durée, les coûts et surtout le mode opératoire limitent la fréquence et l'automatisation des mesures. C'est ainsi que la plupart des stratégies de commande existantes se limitent bien souvent à de la commande de type indirecte des procédés de fermentation par des boucles de régulation des variables environnementales telles que la concentration en oxygène dissous, la température, les débits de liquide et de gaz le pH, etc. [DOC01]. Cependant, le problème des capteurs reçoit une attention de plus en plus importante au cours des ces années. Malgré cela, même si un effort important a été fait au niveau des centres de recherche pour développer et fiabiliser des capteurs, le transfert de ces derniers vers l'industrie et les collectivités territoriales reste naissant et doit encore se développer dans les années à venir.

Concernant les capteurs, quelles que soient les solutions retenues, on voit toujours apparaître deux types de systèmes : les *capteurs in situ*, directement implantés dans les réacteurs biologiques et les *capteurs en ligne*, implantés sur une boucle spéciale de prise d'échantillons. En général, il apparaît diverses solutions qui tentent de résoudre le problème de l'instrumentation :

- Intégration de capteurs existant dans une boucle de mesure, l'idée conductrice est dans ce cas d'automatiser des capteurs hors ligne, de manière à obtenir, sans intervention humaine, les mesures sur site des variables intervenant directement dans les modèles bilans-matières. Ceci se fait à partir d'un échantillonneur implanté directement dans le bioréacteur, ou plus généralement dans une boucle de mesure.
- Couplage de mesures indirectes et de modèles mathématiques pour reconstruire au travers d'observateurs les variables du procédé (souvent appelés *capteurs logiciels* ou virtuels). Dans la plupart des cas, c'est un observateur de type filtre de Kalman étendu qui permet de reconstruire les variables à contrôler.
- Développement de capteurs spécifiques, appelés *biocapteurs*, qui sont des dispositifs qui permettent, par l'intermédiaire d'une couche de reconnaissance ionique ou moléculaire, de transformer l'information "biologique" en un signal électrique analogique. Leur utilisation reste cependant limitée, due au degré de spécialisation des mesures et aux limitations de transfert au niveau industriel.

1.3.2. Modélisation et identification des bioprocédés

La notion de modèle dynamique joue un rôle central en Automatique. C'est en effet sur la base de la connaissance de l'évolution au cours du temps du procédé que s'effectue toute la conception, l'analyse et la mise en œuvre des méthodes de commande et de surveillance. Dans le cadre des bioprocédés, la manière la plus courante de déterminer les modèles qui permettent de caractériser la dynamique du procédé est de considérer des équations différentielles de bilans de matière (et éventuellement d'énergie) des composants principaux du procédé [DOC01]. Ce type d'approche est basé sur des lois physiques et/ou principes fondamentaux (approche mécanistique ou phénoménologique) qui mettent en évidence le fonctionnement interne du système, au moins du point de vue macroscopique. C'est une approche de modélisation où les paramètres du modèle sont liés à la description physique des conversions biochimiques. Ces modèles sont très appropriés pour des objectifs de conception, de dimensionnement et d'optimisation de bioprocédés [TEB98].

Un des aspects importants des modèles de bilan matière est qu'ils sont constitués de deux types de termes représentant respectivement la conversion et le transport. La conversion englobe la cinétique à laquelle se déroulent les diverses réactions biochimiques du procédé et

les rendements de conversion des divers substrats en biomasse et produits. La dynamique de transport regroupe le transit de la matière au sein du procédé sous forme solide, liquide ou gazeuse, et les phénomènes de transfert entre phases. Concernant la modélisation dynamique de procédés biologiques nous abordons, dans le chapitre 2, une approche dite « boîte blanche ».

D'autres approches de modélisation existent, certaines d'entre elles plus orientées vers la modélisation de la *connaissance*, problématique qui peut compléter la tradition des sciences dites objectives, lesquelles se préoccupent essentiellement de la modélisation de l'univers physique. Ces techniques ont montrée leur applicabilité comme des méthodes alternatives pour la modélisation de systèmes dynamiques non-linéaires [NAR90] [BAB98a]. Ce sont des structures mathématiques flexibles, qui exhibent comme caractéristiques, la capacité d'approximation des comportements non-linéaires et l'*apprentissage* à partir des données. Les problèmes de représentation et d'utilisation des connaissances sont au centre d'une discipline scientifique relativement nouvelle et en tout cas controversée : l'intelligence artificielle. Cette discipline a eu un impact limité, jusqu'à une date récente, sur les applications industrielles, parce qu'elle a mis l'accent, de façon pratiquement exclusive, sur le traitement symbolique de la connaissance, par opposition à la modélisation numérique utilisée traditionnellement dans les sciences de l'ingénieur. Plus récemment, on a assisté à un retour du numérique dans ces problèmes d'intelligence artificielle, avec les réseaux de neurones artificiels¹⁶ et la logique floue¹⁷. Alors que les réseaux de neurones proposent une approche implicite de type « boîte noire » de la représentation des connaissances, faisant abstraction de la représentation interne du système très analogue à la démarche d'identification classique des systèmes en Automatique, la logique floue est plus proche de l'intelligence artificielle symbolique, qui met en avant la notion de raisonnement, et où les connaissances sont codées explicitement. Néanmoins, la logique floue permet de faire le lien entre la modélisation numérique et la modélisation symbolique. Le comportement du système se matérialise par une représentation à base de règles du type « *Si-Alors* » qui apportent une description linguistique à la fois *transparente* et *interprétable* pour l'utilisateur du modèle. Cette représentation peut être construite en intégrant la connaissance sur le comportement du système, donnée par des experts, mais aussi à partir des données entrée-sortie du système, ce qui correspond à une construction plus systématique du modèle. Nous développerons cette approche plus en détail dans le chapitre 3, en nous focalisant sur la modélisation et l'identification floue de type Takagi-Sugeno [TAK85].

La théorie des ensembles flous a également donné naissance à un traitement original de l'incertitude, fondée sur l'idée de définir des ensembles pouvant contenir des éléments de façon partielle (et pourtant avec un relatif « degré d'appartenance »), permettant de formaliser le traitement de l'ignorance partielle. Les ensembles flous ont également eu un impact sur les techniques de classification automatique et la commande non-linéaire et ont donc contribué à un certains renouvellement des approches existantes de l'aide à la décision. Du point de vue de l'identification des systèmes et plus particulièrement en ce qui concerne la modélisation des systèmes non-linéaires, ces outils ont permis le développement d'une approche dite « boîte grise » qui permet d'un côté d'introduire dans la modélisation la connaissance *a priori* sur le système à identifier (connaissance de l'expert sous la forme de règles, choix d'entrées/sous-modèles, descriptions de régions de validité,...) et d'autre part, plus récemment, d'incorporer des données entrée-sortie pour la construction du modèle.

¹⁶ En anglais : *ANN*, *Artificial Neural Networks*.

¹⁷ En anglais : *Fuzzy Logic*.

Dans le domaine des bioprocédés ces méthodes de modélisation trouvent leurs origines en Intelligence Artificielle. On peut citer à titre d'exemple le cas de la modélisation *hybride* (semi-mécanistique) qui combine une partie mécanistique (bilan de matière) et un modèle neuronal [PSI92] ou flou [BAB96a] [VANL02] qui a pour but de modéliser la cinétique mal connue à partir des données expérimentales. Dans ce dernier article, l'auteur propose un modèle flou qui prédit le taux de conversion d'obtention de la Pénicilline-G. Le modèle décrit qualitativement la dépendance de la cinétique de l'enzyme par rapport aux concentrations de composants qui interviennent dans la conversion. Le modèle flou qui modélise la cinétique, a été incorporé dans un modèle « boîte blanche » à base de bilan matière. D'autre part on peut citer aussi l'application des techniques *neurofloues* qui visent à intégrer les capacités d'apprentissage à partir des données qui distinguent les réseaux de neurones avec la capacité de description linguistique du comportement qui caractérisent les modèles flous. Comme exemple, [TAY00] propose une représentation neurofloue pour des systèmes anaérobies de traitement des eaux usées basé sur le modèle ANFIS¹⁸ [JAN93]. Le modèle conceptuel proposé est utilisé pour la prédiction de la réponse du système (production volumétrique de Méthane, DCO et concentration d'AGV¹⁹) aux variations des charges organique, hydraulique et d'alcalinité ; [BEL99] présente le développement d'un modèle pour la prédiction de la DCO de l'effluent dans une station d'épuration. Pour cela il utilise des réseaux hétérogènes neurofloues basés sur des variables opérationnelles de l'installation industrielle tels que le débit, la DCO et le TSS²⁰ de l'affluent. Dans le même ordre d'idée, [TOM99] propose l'obtention d'un modèle neuroflou de simulation de la DCO pour un procédé de boues activées.

En ce qui concerne l'application directe de la logique floue à la problématique de la modélisation des bioprocédés, plusieurs travaux se sont déroulés depuis deux décennies (cf. [HOR02] et références incluses). Dans les premières applications de la théorie des ensembles flous aux procédés biologiques, l'approche floue est vue comme une méthode pour l'automatisation des bioprocédés où la modélisation linguistique de l'expertise des opérateurs joue un rôle important dans leur opération et commande. Ainsi, [FIL85] présente la modélisation floue pour la simulation d'une fermentation alcoolique basée sur des règles de production ; [DOH85] propose l'étude d'un modèle flou d'un fermenteur qui utilise des descriptions linguistiques du comportement ; [CZO89] considère la modélisation d'un contrôleur flou appliqué sur des procédés biologiques ; [HOR95] applique des méthodes statistiques et des procédures d'identification floue pour la caractérisation des phases de cultures microbiennes ; [EST95] présente la modélisation floue d'un bioprocédé de digestion anaérobie en utilisant une description semiquantitative du procédé ; [POL02] réalise une étude comparative avec des résultats expérimentaux concernant la modélisation floue d'un digesteur anaérobie. L'incorporation des coefficients flous sur le taux de croissance permet de prendre en considération l'influence de la température et du pH sur la dynamique du système. D'autres approches plus récentes concernant la modélisation floue des bioprocédés à partir de *données* se sont développées au cours des dernières années. Ces approches se sont focalisées sur les comportements *entrée-sortie* du système, en particulier sous la forme de modèle flou de type Takagi-Sugeno (TS). On peut citer à titre d'exemple quelques applications : [GEO01] présente la modélisation floue TS, sous forme d'entrée-sortie et d'espace d'état, d'un procédé biotechnologique de fermentation en mode batch pour la production de gomme de xanthane²¹.

¹⁸ ANFIS : *Adaptive-Network-based Fuzzy Inference Systems*.

¹⁹ AGV : *Acides Gras Volatils*.

²⁰ TSS : *Total Suspended Solids*.

²¹ Fabriqué par la bactérie *Xanthomonas Campestris*, le xanthane est le plus connu des polysaccharides pouvant être produits par voie microbienne. Cette substance est ajoutée à de nombreux aliments comme stabilisateur.

Le modèle permet la description des dynamiques non-linéaires de la croissance de biomasse, de la consommation de substrat et de l'accumulation de produit ; [RAG01] considère le problème de définition de la structure et de la détermination des paramètres d'un modèle flou TS appliqué à la prédiction dynamique de l'absorption UV comme mesure indirecte de la DCO dans une station d'épuration des eaux usées ; [GRI05] propose l'utilisation des techniques de clustering flou, à partir des données entrée-sortie, pour l'identification d'un modèle dynamique TS de type NARX qui représente le comportement non-linéaire des concentrations de biomasse et des substrats dans un bioprocédé de traitement des eaux usées.

1.3.3. La commande des bioprocédés

La commande des procédés a pour objectif l'amélioration (voir l'optimisation) de la performance, de la qualité et de l'efficacité économique des procédés tout en garantissant leur sûreté générale et celle des opérateurs, ainsi que le respect de la législation et de l'environnement. Dans le contexte des bioprocédés, caractérisés par nature par des dynamiques non-linéaires et complexes, cet objectif principal peut se traduire par le développement et l'implantation de réalisations matérielles et/ou informatiques, afin de maintenir le procédé biologique dans des conditions opératoires stables et, si possible, optimales en dépit de perturbations internes/externes pouvant affecter son bon fonctionnement [DOC01]. Les perturbations qui agissent sur ces systèmes sont dues à la non stationnarité des phénomènes biologiques (variations internes) ou à l'action sur le système d'entrées inconnues (non mesurées) telles que les concentrations des différents substrats, le pH, la présence de toxiques, la température, les biais sur les capteurs et/ou actionneurs, etc. [STE98]. Bien que l'utilisation de stratégies de commande s'avère nécessaire pour optimiser le fonctionnement des procédés biologiques, le contrôle, dans le sens conventionnel employé en Sciences pour l'Ingénieur, ne s'applique encore aujourd'hui que difficilement. Comme nous l'avons déjà souligné, deux types de difficultés expliquent largement ce manque d'application généralisée de systèmes de commande sophistiqués aux bioprocédés : la première liée à la modélisation de la dynamique des bioprocédés et la seconde liée à l'instrumentation et à la difficulté de mesurer en ligne de manière fiable toutes les variables-clés du procédé.

L'utilisation d'un ordinateur pour contrôler et commander un procédé biologique est représentée schématiquement sur la Figure 1.3. L'élément central de ce schéma est le procédé (représenté ici par un réacteur). Sur ce procédé, on effectue une série de mesures, soit dans le milieu liquide (typiquement, des mesures de pH, de température, ou de concentrations), soit dans le milieu gazeux (sous forme de mesures de débits gazeux). Ces mesures sont envoyées à un ordinateur qui peut en particulier les traiter dans un objectif de commande (qui pourrait être aussi de surveillance) du procédé. Le régulateur (ici intégré dans l'ordinateur) permet de calculer une action de commande à appliquer au bioprocédé sur la base des mesures (physiques et logicielles), de la connaissance de la dynamique du procédé et des objectifs de commande choisis. Dans la situation décrite, cette action de commande est représentée schématiquement par l'ouverture d'une vanne permettant de moduler le débit de substrat d'alimentation dans le réacteur. Sa valeur est la sortie d'un algorithme de commande qui utilise les informations disponibles sur le procédé. Ces informations regroupent d'une part, l'état du procédé jusqu'à l'instant considéré (i.e. les mesures) et, d'autre part, des connaissances disponibles *a priori* (par exemple sous la forme d'un modèle du type « bilan-matière ») et relatives à la dynamique du procédé biologique et aux interactions mutuelles des différentes variables du procédé. Nous retrouvons dans le schéma la notion de *boucle fermée*. Cette boucle est constituée de l'ensemble procédé-capteurs-ordinateur-actionneurs.

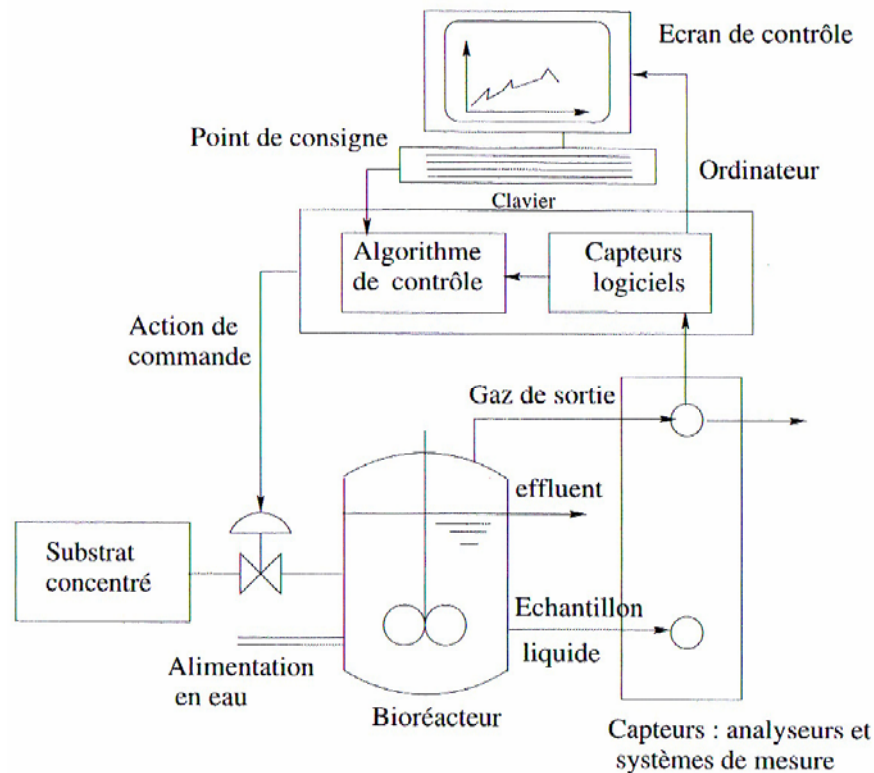


Figure 1.3. Représentation du système de commande d'un bioprocédé (d'après [DOC01]).

Bien que les principes de base de la commande des bioprocédés restent les mêmes, ils existent des différences en fonction du secteur d'application. Ainsi, dans le cas des procédés de dépollution, il est important de noter que réguler la concentration de substrat revient à contrôler le degré de pollution rejetée dans l'environnement alors que dans d'autres procédés biologiques (i.e., à finalité agro-alimentaire, pharmaceutique,...), réguler le substrat permet généralement d'optimiser la production de biomasse et/ou de produit en limitant les inhibitions [STE98]. En ce qui concerne les procédés biologiques de dépollution des eaux résiduelles, selon les connaissances disponibles sur le procédé et les objectifs de commande spécifiés par l'utilisateur, nous allons pouvoir développer et mettre en œuvre des algorithmes de commande plus ou moins complexes. On peut alors distinguer deux types de stratégie de commande : une approche « directe », qui consiste à minimiser la pollution en sortie du procédé en régulant la concentration de l'effluent et une approche « indirecte », qui consiste en le maintien de conditions opératoires permettant d'assurer un fonctionnement nominal du réacteur sans réguler la concentration du polluant en sortie. Dans ce cas, la connaissance des spécialistes en biotechnologie et des experts des procédés nous assure simplement que, certaines variables étant régulées à des niveaux donnés, le fonctionnement du procédé sera satisfaisant.

Remarque

R 1.3 L'approche indirecte concerne par exemple la régulation de la concentration en oxygène dissous à l'intérieur d'un réacteur aérobique dans lequel cette concentration doit être maintenue suffisamment haute pour garantir une activité bactérienne satisfaisante. Il est évident que le maintien de ce taux minimum sans autre considération sur les variables comme le pH, l'agitation ou encore la température ne garantit en rien une dépollution effective [STE98].

En général, la méthode de synthèse des lois de commande des bioprocédés dépend de nombreux facteurs. Tout d'abord, ce choix dépend des objectifs mêmes de commande et de la stratégie utilisée (directe ou indirecte). Plus les spécifications de commande vont être précises, plus l'approche analytique va être justifiée. Bien évidemment, on peut ensuite envisager diverses combinaisons « précision des spécifications/complexité de la stratégie de commande ». On peut ainsi imaginer des spécifications de commande très strictes mais très faciles à vérifier, dans la mesure où la variable à commander est directement mesurée par un moyen très complexe. A l'inverse, un problème de régulation d'une variable environnementale facilement mesurable peut être résolu par un régulateur très simple basé sur un modèle linéaire très succinct.

Le choix du type de loi de commande va ensuite dépendre de la connaissance disponible *a priori* sur le procédé. De ce point de vue, il convient de rappeler ici quelques règles élémentaires pouvant orienter l'utilisateur dans son choix [DOC01] :

- lorsqu'un modèle analytique suffisamment précis - quelle qu'en soit sa nature - est disponible, alors il faut l'utiliser ;
- si un modèle analytique n'est pas disponible mais que l'expertise disponible sur le fonctionnement du procédé est importante, alors les approches heuristiques (logique floue, systèmes à base de connaissance, etc.) peuvent être utilisées ;
- lorsqu'un modèle analytique s'avère trop long à obtenir, que peu d'expertise est disponible, mais qu'en revanche une grande quantité de données est disponible, alors il peut être intéressant d'utiliser une approche statistique ou neuronale/floue permettant en particulier d'obtenir des modèles de comportement non linéaires (cf. par exemple [MON94], [GRI05]).

On constate dans les faits que les boucles de commande implantées sur les procédés au stade industriel sont typiquement des régulateurs de type PID²² complètement découplées [QUE00] et pourtant mono-variables. Les algorithmes de commande classique ont été appliqués aux bioprocédés pour la régulation autour d'un point de fonctionnement ou pour le suivi d'une trajectoire préétablie dans des conditions opératoires bien spécifiques [VAN98a]. Cependant, l'efficacité de ces algorithmes est souvent limitée en pratique, compte tenu de la diversité des phénomènes mis en jeu dans les bioprocédés, ce qui nécessite de contrôler à la fois plusieurs variables et de tenir compte des variations des conditions opératoires et du comportement des micro-organismes. Cela a conduit à la recherche de solutions de contrôle plus avancées et plus performantes. A la lecture de la bibliographie sur les techniques de contrôle des bioprocédés, on peut retrouver plusieurs approches de commande selon les caractéristiques les plus importantes du système à commander, la connaissance du procédé et les objectifs de commande définis par l'utilisateur. Sans vouloir être exhaustifs, nous essayons de faire un tour d'horizon concernant les techniques le plus souvent rencontrés dans la littérature : *la commande adaptative*, consistent à réajuster certains paramètres intervenant dans le calcul de la commande en fonction de la dynamique du processus pour maintenir les performances du système lorsque les paramètres du modèle varient. Elle est de loin la technique la plus largement développée au niveau de la commande pour les bioprocédés [BAS90] [DOC91] [ROU92] [ROU94] [BAS95] [BEN96] [DAH06] ; *la commande robuste*, orientée vers la conception de correcteurs à paramètres fixes capables d'assurer ces propriétés en présence de perturbations et d'incertitudes paramétriques du modèle valable dans un domaine de fonctionnement [QUE97] [JUL97] [PAI97] [GOM02] [RAP02] ; *la commande*

²² PID : Contrôleur *Proportionnel - Intégral - Dérivé*.

optimale, intéressée à trouver à partir d'un modèle et parmi les commandes admissibles, celle qui permet à la fois : de vérifier des conditions initiales et finales données, de satisfaire diverses contraintes imposées et d'optimiser un critère mathématique choisi [OHN76] [DUV88] [TSO95] [ASE96] [DON01] ; la *commande prédictive*, qui se base sur l'utilisation d'un modèle dynamique du système pour anticiper son comportement futur [MOR92] [BEN96] [ZHU00] [RAM05] ; la *commande neuronale*, qui s'avère intéressante pour la commande des systèmes en s'appuyant sur des modèles non-linéaires d'entrée-sortie obtenus à partir des données [MUR95] [EMM96] [PAT97] [YU98] [WEN98] [OBE99] ; la *commande floue*, technique appropriée tant pour formaliser la connaissance heuristique exprimée par des règles, que pour intégrer cette connaissance sur le comportement des procédés avec l'information obtenue à partir des données numériques. Nous allons faire une étude bibliographique plus approfondie de cette méthode de commande en nous focalisant sur le domaine d'application qui nous intéresse, dans les paragraphes qui suivent.

Remarque

R 1.4 Les techniques de commande que l'on vient de décrire peuvent être complémentaires en visant l'exploitation des caractéristiques les plus avantageuses de chacune. On peut citer à titre illustratif la commande optimale adaptative [VAN98a] ou bien la commande prédictive non-linéaire basée sur des modèles neuronal ou flou du procédé [TEB98] [VEN99]. Avec une compréhension croissante du comportement biochimique ainsi que les progrès au niveau des technologies de mesure, on peut prévoir l'amélioration et l'apparition de nouvelles approches de commande des bioprocédés dans la décennie à venir.

En ce qui concerne l'application de la logique floue à la commande des procédés biologiques [HON00] [HOR02], il faut souligner que dans un premier temps, la commande floue consiste à représenter sous la forme de règles linguistiques la connaissance et le savoir faire des opérateurs experts pour assurer le bon fonctionnement du procédé. Le contrôleur flou détermine directement la sortie de commande à partir de la base de connaissance experte et des données en ligne. On peut ainsi citer plusieurs exemples : [NAK85] présente la commande floue du taux d'alimentation de sucre dans un procédé de fermentation; des approches similaires ont été appliquées pour la commande des concentrations de glucose et d'éthanol en cultures semi-continues [PAR93] [SHI94] et pour le pilotage de fermentations de levures [SII95]. Dans le cas particulier des procédés biologiques de dépollution, les premiers systèmes utilisant l'expertise humaine ont été développés sur la base de l'approche floue vers la fin des années 1970 [DOC01]. En utilisant le profil d'oxygène dissous, un système d'inférence floue a par exemple permis d'estimer la consommation du substrat dans un procédé à boues activées et de donner des conseils aux opérateurs humains sur le changement des consignes des régulateurs classiques de l'installation [FLA80]. La même année, vingt règles floues ont été proposées pour la régulation des débits d'un procédé de boues activées [TON80]. Plus récemment, un système de commande flou des solides en suspension pour le même type de procédé a été développé [TSA96]. Dans le même ordre d'idées, l'approche floue a été utilisée avec succès pour le contrôle des perturbations dans un processus de traitement des eaux résiduaires [MUL97], ainsi que pour le contrôle de l'abatement de la demande chimique en oxygène dans une station d'épuration urbaine [FU98]. A la même période, [EST97] propose l'analyse et la commande par logique floue d'un procédé de digestion anaérobie en lit fluidisé pour la dépollution des vinasses de vin. Dernièrement, la commande floue pour maintenir la concentration de MES constante dans une procédé de dépollution des eaux usées a été proposé [TRA03]. L'implantation industrielle de la

commande dite « intelligente²³ » dans les automates programmables a été aussi décrite depuis peu [MAN98]. On peut retrouver aussi l'application de la logique floue pour le réglage des contrôleurs traditionnels (e.g. de type PID) qui opèrent au niveau du procédé biologique lui-même [ROD99] ou bien comme un contrôle supervisé [KIT94] [VON94] ou encore comme un outil pour l'identification de l'état physiologique du procédé (différentes phases des cultures microbiennes) avant d'appliquer la stratégie de commande elle-même [KON89] [KIS91] [KON92]. On assiste à un déplacement progressif de la conception des contrôleurs flous vers l'établissement des procédures de synthèses plus méthodiques qui se rapprochent beaucoup plus des techniques traditionnelles de l'Automatique, en s'appuyant sur des modèles. A titre illustratif, [GEO96] présente une approche pour la synthèse de la commande des procédés biologiques en mode batch, basée sur un modèle qui combine une description déterministe avec un modèle flou, dans lequel le modèle flou sert à représenter le comportement entrée-sortie du procédé. Egalement, [POL01] propose la synthèse d'observateurs pour plusieurs concentrations et l'alcalinité dans un réacteur de digestion anaérobie en utilisant un modèle flou de type Takagi-Sugeno. En effet, bien que le formalisme flou soit capable d'intégrer l'information heuristique, il nous semble percevoir en pratique une orientation vers des techniques qui s'appuient sur des modèles mathématiques non-linéaires, focalisés sur le comportement entrée-sortie des procédés. C'est justement cette capacité à intégrer la connaissance experte avec la puissance de la commande basée sur des modèles, qui nous a motivé pour explorer cette dernière approche dans notre travail. En particulier, nous développerons une proposition pour la commande floue sous-optimale basée sur des modèles flous de type Takagi-Sugeno dans le chapitre 4.

Jusqu'à présent, nous avons peu abordé la question de la supervision des bioprocédés, terrain sur lequel les techniques issues de l'Intelligence Artificielle (plus spécifiquement la logique floue et le raisonnement qualitatif) jouent un rôle majeur ces dernières décennies, comme des outils d'aide à la décision. La *supervision*²⁴, représente la surveillance²⁵ (détermination de l'état) d'un système physique et la prise de décisions appropriées en vue de maintenir son opération face à des défaillances²⁶. Les données provenant du système sont utilisées par la *surveillance* pour représenter l'état de fonctionnement puis détecter les évolutions de comportements anormaux. Le *diagnostic* identifie la cause, la localisation et l'occurrence de ces évolutions, ensuite l'*aide à la décision* intervient pour proposer des actions correctives. Cette approche d'intégration des outils de surveillance, de diagnostic et d'aide à la décision au niveau de la supervision [TRA97] est pertinente dans les procédés biologiques en complément de systèmes primaires de mesure et de commande, compte tenu de la difficulté de conduite. Pour le cas spécifique des bioprocédés de traitement des eaux usées, le niveau de supervision est important dans la mesure où ces procédés sont en général aussi plus complexes que les procédés « classiques » (i.e., ceux utilisés dans les industries agro-alimentaires ou dans la pharmacie) puisqu'il s'agit généralement d'écosystèmes et non pas de cultures pures. D'autre part, ils sont soumis aussi à un nombre beaucoup plus significatif de perturbations à la fois externes et internes.

²³ Nous préférons plus modestement de parler ici de systèmes de commande issus de techniques de l'Intelligence Artificielle.

²⁴ D'après la définition qui a été discuté au sein du *SAFEPROCESS Technical Committee* de l'*International Federation of Automatic Control*.

²⁵ La *surveillance* est une tâche continue et en temps réel, de détermination de l'état d'un système. Elle consiste en l'enregistrement de l'information ainsi qu'en la reconnaissance et l'indication de comportements anormaux.

²⁶ Une *défaillance* ou *panne* est une interruption permanente de la capacité du système à remplir une fonction requise dans des conditions d'opération spécifiées.

Cette approche intégrée de la supervision est un domaine de recherche actif, mais en même temps très vaste, nous allons simplement citer ici quelques axes de travaux et certains développements qui pourraient être appliquées aux procédés biologiques : la détection des défaillances, le diagnostic, la reconfiguration du processus et la maintenance [ISE97] [COL00]. Nous suggérons au lecteur intéressé à se reporter à l'ouvrage [DOC01] (en particulier au chapitre 8 [STE01] et les références incluses, consacré aux outils d'aide au diagnostic et détection de pannes), pour un traitement approfondi sur le sujet dans le cas des procédés biologiques de dépollution.

1.4. Conclusions

Le traitement des eaux usées est un domaine vaste, qui englobe plusieurs disciplines, comme la microbiologie, le génie des procédés et l'Automatique. Ce premier chapitre nous a permis de présenter les notions de base des procédés biologiques de dépollution des eaux résiduaires ainsi que le positionnement de l'Automatique dans ce contexte, avant d'aborder les problèmes de modélisation et de commande floues représentant les principaux objectifs de ce travail de thèse.

La pollution des ressources en eau dans les stations d'épuration provient de différentes origines, notamment des activités humaines au niveau domestique, agricole et industrielle. Nous avons commencé par une description des diverses sources de pollution ainsi que les indicateurs de qualité de l'eau. Après avoir présenté un tour d'horizon sur la législation correspondante en France, en concordance avec une normativité mondiale à chaque fois plus contraignante, nous avons fait une description globale d'une filière de traitement en mentionnant les différentes étapes d'une chaîne de dépollution des eaux résiduaires. En effet, étant donnée la grande diversité de déchets, les eaux brutes sont soumises à une combinaison ou une succession de différents processus de purification, de nature variée : physique, chimique, biologique,... Nous avons détaillé, plus particulièrement, le traitement biologique, sur lequel porte spécifiquement notre étude, compte tenu de l'efficacité montré pour ce type de traitement pour restaurer la qualité de l'eau en dégradant principalement les matières organiques. Nous avons décrit les principaux polluants, les micro-organismes épurateurs, leurs besoins nutritifs et les processus métaboliques les plus importants pour le développement des populations bactériennes. Nous avons aussi décrit différents modes et procédés spécifiques de traitement biologique. Par la suite, le lagunage faisant intervenir des bactéries aérobies sera la base d'une première étude de modélisation en utilisant l'approche de bilans de matières dans le chapitre suivant. Une deuxième étude comportant la modélisation et la commande, en utilisant l'approche floue, fera l'objet d'application au chapitre 5.

Le fait de faire intervenir des organismes vivants dans les procédés biologiques de traitement des eaux usées rend ce type de procédés très intéressante du point de vue scientifique et d'ingénierie, mais en même temps très complexe étant donnée la nature très diverse des interactions et des phénomènes multiples mis en jeu. De manière générale, ces procédés sont caractérisés par des systèmes multivariables, non-linéaires et non-stationnaires. Malgré la difficulté liée à la modélisation de la dynamique des bioprocédés et le manque actuel d'une instrumentation de mesure en ligne permettant de connaître le fonctionnement interne des ces procédés, nous sommes vraiment convaincus que les besoins croissants en systèmes de supervision et de commande demanderont une intervention encore plus poussée de l'Automatique dans les années à venir, afin de garantir un rejet faible et stable des effluents dans le milieu naturel avec une gestion efficace des installations. Nous considérons également que l'approche floue, qui a déjà montrée son applicabilité pour l'automatisation des

bioprocédés où la modélisation linguistique de l'expertise des opérateurs qui a joué un rôle important dans leur conduite, peut contribuer aussi au niveau de la modélisation et la commande, en permettant une intégration de la connaissance experte avec l'incorporation des données du procédé et l'utilisation des modèles pour le contrôle. En effet, nous assistons à un changement d'orientation d'une conception heuristique de ces systèmes vers une autre plus systématique que nous développerons par la suite, dans le cadre de la modélisation - avec une démarche d'identification très similaire à celle classique des systèmes en Automatique - à partir des données entrée-sortie du comportement du procédé dans le chapitre 3 et pour la commande floue basée sur de modèles Takagi-Sugeno, dans le chapitre 4.

Dans le chapitre 2, nous allons donc aborder la modélisation dynamique des bioprocédés, afin d'établir un modèle de connaissance initial qui va nous servir de référence pour comparer la performance des modèles flous comme modèles de description de comportement non-linéaire entrée-sortie du système.

Chapitre 2

Modélisation dynamique de procédés biologiques

Dans ce chapitre, nous faisons une description des procédés biologiques utilisés dans le domaine de traitement des eaux résiduaires. Nous commençons d'abord par un aperçu global en décrivant le rôle des microorganismes, les types de réacteurs et les modes de fonctionnement des bioprocédés. Ensuite, au niveau de la modélisation, nous nous focalisons sur l'approche de modélisation macroscopique qui considère les bilans de matière des composants principaux pour caractériser la dynamique de tels procédés. Ensuite nous présentons quelques modèles représentatifs : des fermentations continues, semi-continues et de procédés à boues activées. Finalement, nous décrivons le procédé considéré au niveau de l'application, qui est un bioréacteur aérobie multivariable en mode continu pour le traitement des eaux usées issues d'une industrie papetière. Le modèle considéré est une version du simulateur utilisé lors du projet européen FAMIMO²⁷ pour lequel nous avons proposée l'ajout d'un terme de mortalité endogène au niveau de la dynamique bactérienne.

²⁷ Waste-water Benchmark, in Fuzzy Algorithms for the Control of Multiple-Input, Multiple Output Processes (FAMIMO) project, funded by the European Commission (Esprit LTR 21911).

Sommaire

2. Modélisation dynamique de procédés biologiques.....	39
2.1. Description des procédés biologiques.....	42
2.1.1. Les micro-organismes et leur utilisation.....	42
2.1.2. Les types de bioréacteurs.....	43
2.1.3. Les trois modes de fonctionnement.....	43
2.2. Les modèles de bilans-matière	45
2.2.1. Les approches microscopique et macroscopique de modélisation.....	45
2.2.2. Le taux spécifique de croissance.....	46
2.2.3. Equations d'état.....	49
2.3. Influence du transfert de matière	50
2.4. Quelques modèles représentatifs	51
2.4.1. Fermentations continues	51
2.4.2. Fermentations semi-continues.....	55
2.4.3. Procédés à boues activées	56
2.5. Modélisation du bioprocédé de traitement des eaux usées par lagunage aérée.....	59
2.5.1. Traitement biologique des effluents de papeterie par lagunage aérée.....	60
2.5.2. Modélisation des cinétiques de croissance et de dégradation.....	62
2.5.3. Equations d'évolution du procédé.....	65
2.5.4. Variables de commande du procédé.....	67
2.6. Conclusions	69

Fidèle à la procédure traditionnelle en Automatique, une première étape vers des actions hiérarchiques supérieures (e.g. commande, supervision, optimisation...) est forcément l'établissement d'un *modèle*. Bien que dans le domaine scientifique on utilise souvent ce terme, il convient de donner une définition afin d'établir sa portée, plus spécifiquement dans le cas de procédés biochimiques :

Un *modèle* est un système physique, mathématique ou logique représentant d'une façon simplifiée et relativement abstraite, les structures essentielles d'une réalité, en vue de la décrire, de l'expliquer ou de la prévoir.

Il est possible d'affirmer que la modélisation des procédés biologiques sur lesquels se basent les traitements des eaux résiduaires, est un exercice délicat, compte tenu de la complexité des différents phénomènes mise en jeu. En effet, il n'existe pas de lois caractérisant l'évolution des micro-organismes [BER01]. D'autre part, la qualité du modèle, ainsi que sa structure devront avant tout correspondre à l'objectif pour lequel le modèle a été construit. En effet, un modèle peut être développé dans des buts très différents, qu'il faudra clairement identifier dès le départ. Ainsi, le modèle pourra-t-il être utilisé pour :

- reproduire un comportement observé,
- expliquer un comportement observé,
- prédire l'évolution d'un système,
- aider à comprendre les mécanismes du système étudié,
- estimer des variables qui ne sont pas mesurées,
- estimer des paramètres du procédé,
- agir sur le système pour le gouverner et contrôler ses variables,
- détecter une anomalie dans le fonctionnement du procédé,
- etc.

La construction de modèles biochimiques devient une tâche complexe qui demande normalement des compétences spécialisés. Dans la plupart des cas, le comportement au niveau microscopique, les interactions entre l'environnement et les métabolismes cellulaires ne sont pas connus en détail. Il est difficile de tenir compte de tous les facteurs qui peuvent influencer l'évolution des micro-organismes pour élaborer des modèles pour ces procédés [ROU01]. Dans le cadre de nos travaux, nous nous intéressons dans ce chapitre à définir les principes de base de la modélisation de procédés biochimiques en utilisant une description au niveau macroscopique de l'ensemble des réactions biologiques et chimiques, en mettant en évidence le fonctionnement interne des bioprocédés. Après avoir fait une description sur les bioréacteurs et leurs modes de fonctionnement, nous considérerons les modèles de bilans de matière ainsi que les cinétiques des réactions. Finalement, nous décrirons le procédé considéré au niveau de l'application, qui est un bioréacteur aérobie en mode continu pour le traitement des eaux usées issues d'une industrie papetière. Comme nous l'avons indiqué précédemment, ce modèle de connaissance nous servira de référence pour comparer la performance du modèle flou représentant le comportement non-linéaire entrée-sortie du système. La modélisation floue du bioprocédé sera une partie du chapitre 5.

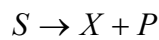
Nous invitons le lecteur à se référer à [PAV94] pour une réflexion plus approfondie sur la modélisation en biologie et en écologie et aux ouvrages [EDE97] et [ORO03] pour une description de l'énergétique, du métabolisme et de la dynamique des populations microbiennes. Une grande partie de ce chapitre est basée sur les documents [STE98] [QUE00] [BER01] et [BEN96].

2.1. Description des procédés biologiques

Dans cette section nous donnons d'abord un aperçu global sur les procédés biologiques, en décrivant les éléments de base, les types de réacteurs et les modes de fonctionnement.

2.1.1. Les micro-organismes et leur utilisation

La fermentation²⁸ microbienne est un procédé qui consiste à faire croître une population de micro-organismes (bactéries, levures, moisissures...) aux dépens de certains éléments nutritifs (nutriments) sous des conditions environnementales (température, pH, agitation, aération...) favorables. Elle correspond à la transformation de substances (substrats carbonés généralement) en produits résultant de l'activité métabolique des cellules [BER01]. De façon schématique, tout processus dans lequel interviennent des micro-organismes peut être décrit par la réaction générale suivante :



Les composés principaux de la réaction sont ainsi :

- *les substrats*, notés S , qui sont les éléments nutritifs nécessaires pour la croissance des micro-organismes, ou bien qui sont des précurseurs d'un composé à produire. Ces substrats contiennent le plus souvent une source de carbone (glucose, éthanol,..) et parfois d'azote (NO_3 , NH_4 ,...), de phosphore (PO_4 ,...), etc. ;
- *la population microbienne* encore appelée *biomasse*²⁹, notée X ;
- *les produits*, notés P , à finalité agro-alimentaire (huiles, fromages, bières, vins,...), chimique (solvants, enzymes, acides aminés,...), pharmaceutique (antibiotiques, hormones, vitamines,...) ou pour la production d'énergie (éthanol, biogaz,...), etc.

A ces composés principaux s'ajoutent les sels minéraux et les vitamines, qui, bien qu'apparaissant rarement dans les modèles, n'en sont pas moins indispensables à la croissance. En effet, la population microbienne, une foisensemencée dans un milieu de culture, ne peut s'y développer correctement que si elle y trouve des aliments et des conditions d'environnement favorables. De plus, la vitesse de croissance des micro-organismes est proportionnelle à la vitesse du processus interne le plus lent. Il s'agit donc d'analyser et de résoudre les diverses étapes limitantes pouvant survenir [STE98] :

- *au niveau physiologique* : synthèses des enzymes, transferts des substrats de l'extérieur vers l'intérieur des cellules, des produits de l'intérieur vers l'extérieur des cellules, inhibitions par excès de substrats, par accumulations intracellulaires de métabolites, par les produits excrétés dans le milieu de culture,...
- *au niveau de l'environnement des micro-organismes* : pH, température, apports d'oxygène et de substrats compatibles avec les exigences de la croissance, présence de gaz carbonique,...

Chaque type de micro-organisme possède des caractéristiques liées à son patrimoine génétique et à ses systèmes de régulation interne. Ainsi, la fermentation peut avoir différentes utilisations :

²⁸ Transformation de la matière organique sous l'action des micro-organismes.

²⁹ Masse totale des organismes vivant présente dans un espace ou un volume déterminés.

- *croissance microbienne* : l'objectif premier est la croissance du micro-organisme. C'est le cas des fermentations visant à produire la levure de boulanger.
- *production métabolique* : l'objectif est de favoriser la synthèse d'un métabolite par la cellule (éthanol, pénicilline,...).
- *consommation de substrat* : dans ce cas, c'est la dégradation du substrat qui est privilégiée. On trouvera essentiellement dans cette catégorie les procédés de dépollution (traitement biologique des eaux usées, dégradation de polluants spécifiques,...).
- *étude phénoménologique* : ici, la fermentation a pour but l'étude du micro-organisme. Cela peut être par exemple, pour mieux comprendre comment le micro-organisme se développe dans le milieu naturel.

La majorité des procédés biotechnologiques développés à l'échelle industrielle utilisent des cultures microbiennes composées d'une seule espèce de micro-organisme pour la synthèse éventuelle d'un produit bien déterminé (culture pure). Dans certains cas cependant, plusieurs espèces peuvent être amenées à croître en parallèle (culture mixte), mais cela ne peut être possible que si elles ne sont pas trop compétitives. Les procédés biologiques de traitement des eaux usées se caractérisent par la présence de cultures mixtes.

2.1.2. Les types de bioréacteurs

Du point de vue de la modélisation mathématique, les réacteurs biologiques peuvent se décomposer en deux grandes classes [BAI86] :

- les *réacteurs infiniment mélangés* (en anglais *stirred tank reactors*, notés STR), pour lesquels le milieu réactionnel est homogène et la réaction est décrite par des équations différentielles ordinaires ;
- les *réacteurs à gradient spatial de concentration*, tels que lits fixes, lits fluidisés, air lifts, etc., pour lesquels la réaction est décrite par des équations aux dérivées partielles.

Nous nous intéressons, dans ce chapitre, uniquement à la première classe, les réacteurs infiniment mélangés et le lecteur intéressé par la seconde classe pourra se référer à [DOC94].

2.1.3. Les trois modes de fonctionnement

Dans la pratique, les modes opératoires se caractérisent par les échanges liquides, c'est-à-dire par la façon dont les réacteurs biologiques sont alimentés en substrat [BER01]. Nous distinguons trois modes principaux (Cf. Figure 2.1).

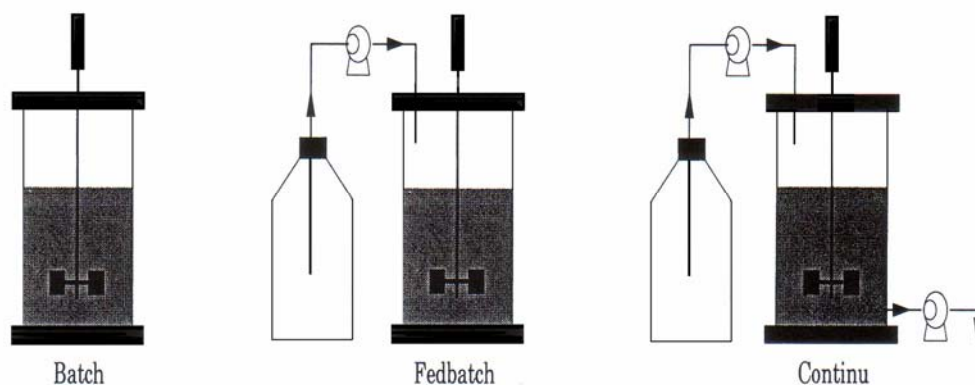


Figure 2.1. Les différents modes de fonctionnement des réacteurs biologiques

Le mode discontinu (ou batch)

Dans ce mode de fonctionnement la totalité des éléments nutritifs nécessaires à la croissance biologique est introduite lors du démarrage de la réaction. Aucun apport ni prélèvement (excepté bien sûr pour quelques mesures hors lignes éventuellement) n'est par la suite réalisé et la réaction se déroule à volume constant. Les seules actions possibles de l'opérateur ne concernent que les variables d'environnement (pH, température, vitesse d'agitation, aération,...). Peu de moyens sont ainsi nécessaires à sa mise en œuvre, ce qui en fait son attrait du point de vue industriel (e.g. secteurs chimique et pharmaceutique). Le second avantage est de garantir l'axénicité des cultures, car les risques de contamination de la culture sont minimales. Le revers de la médaille est le peu de moyens qui permettent d'agir sur le fermenteur pour optimiser l'utilisation des micro-organismes. Il souffre cependant de deux inconvénients majeurs:

- sa durée requise, car à chaque nouveau remplissage du réacteur, avant que la réaction ne démarre, on observe une période de latence liée au temps d'adaptation des micro-organismes, ce qui entraîne la génération de temps morts lors de l'opération ;
- son inadéquation lorsque la réaction générale est sensible à une inhibition par le substrat (les quantités introduites sont en effet importantes car devant couvrir l'ensemble des besoins des micro-organismes sur la totalité du fonctionnement) ou par un produit (toute production est en effet conservée dans le réacteur jusqu'à l'achèvement attendu de la réaction), ce qui se traduit par des durées de traitement allongées et des limitations de la charge initiale admissible.

Pour le cas des procédés biologiques de dépollution des eaux usées, ce mode de fonctionnement exigera la présence en amont d'un dispositif de stockage des affluents. Par ailleurs, de par les contraintes imposées par le traitement des effluents (celui-ci doit être continu), ce mode de fonctionnement n'est que très rarement utilisé en dépollution.

Le mode semi-continu (ou fedbatch)

Tout en nécessitant un dispositif de stockage des affluents, ce mode de fonctionnement se distingue du précédent par un apport des différents éléments nutritifs au fur et à mesure des besoins constatés des micro-organismes. La variation du volume du milieu réactionnel est donc une fonction directe de l'état d'avancement de la réaction. Ce mode permet essentiellement de lever les problèmes d'inhibition associés au mode précédent, et de fonctionner à des taux spécifiques de croissance proches de leur valeur maximale. A partir d'un volume initial préalablementensemencé, le réacteur est alimenté par un débit augmentant typiquement de façon exponentielle, nécessitant un contrôle en boucle fermée. C'est d'ailleurs ce dernier point qui a fortement limité l'utilisation du *fedbatch* en milieu industriel.

Ce mode de fonctionnement est essentiellement utilisé pour des productions issues de procédés déjà anciens : production de levure, de vinaigre, de bière,... Il est par ailleurs, tout comme le précédent, également préconisé lorsque la récupération des produits est réalisée en discontinu (accumulation intracellulaire par exemple) ou que l'on ne peut se permettre de relarguer des matières toxiques résiduelles (cas du fonctionnement en continu).

Le mode continu

Caractérisé par un volume réactionnel constant, il est soumis à un soutirage de milieu réactionnel égal au flux d'alimentation en matière nutritive. Les procédés continus

fonctionnent en régime permanent, en maintenant, pour des conditions d'alimentation fixées, le système dans un état stationnaire, tout en évitant tout phénomène inhibiteur grâce à l'effet de dilution dû à l'alimentation. En agro-alimentaire, pharmacie ou autres industries productrices, ils autorisent de plus des productions importantes dans des réacteurs de taille réduite. Dans le domaine de la dépollution biologique des eaux résiduaires, il s'agit là du mode de fonctionnement préférentiel.

Avant de présenter le principe de la modélisation par bilans-matières, il convient de dire qu'en pratique il est possible de trouver des réacteurs avec une opération qui est en fait une combinaison au cours du temps des différents modes de fonctionnement précédents³⁰. L'idée est de récupérer la biomasse par sédimentation ou le produit du surnageant entre deux séquences de production (ou de traitement).

2.2. Les modèles de bilans-matière

Contrairement à la physique où il existe des lois connues depuis des siècles (loi d'Ohm, loi des gaz parfaits, relation fondamentale de la dynamique, principes de la thermodynamique,...), la plupart des modèles en biologie reposent sur des lois empiriques. La construction de modèles biochimiques devient une tâche délicate, car il n'existe pas de lois caractérisant l'évolution des micro-organismes. Dans le cadre des bioprocédés, la manière la plus courante de déterminer les modèles qui permettent de caractériser la dynamique du procédé est de considérer des équations différentielles de bilans de matière des composants principaux du procédé [DOC01]. Ce type d'approche est basé sur des lois physiques et/ou principes fondamentaux (approche mécanistique ou phénoménologique) qui mettent en évidence le fonctionnement interne du système, au moins du point de vue macroscopique. Nous allons explorer cette approche par la suite.

2.2.1. Les approches microscopique et macroscopique de modélisation

Comme nous l'avons souligné, les phénomènes mis en jeu dans les procédés biologiques possèdent une dimension particulière par rapport aux exemples habituels d'applications : ils relèvent du vivant. Il est ainsi logique de mettre le micro-organisme au centre de toute analyse. C'est la démarche du biologiste ou du biochimiste qui cherche à expliquer les modifications du milieu réactionnel à partir des lois de croissance et de fonctionnement de la population microbienne. Deux approches essentielles permettent d'appréhender les phénomènes intervenants lors d'une réaction biologique [STE98] :

- *approche microscopique* pour laquelle chaque cellule est considérée comme un réacteur élémentaire. Cette méthodologie se heurte cependant au manque de capteurs. Il n'est en effet que peu envisageable de pouvoir un jour mesurer en ligne les concentrations des composants intracellulaires ou encore la composition d'une paroi cellulaire,
- *approche macroscopique* au cours de laquelle le bilan global de la réaction est modélisé. Cette approche possède le mérite d'affronter des phénomènes plus facilement accessibles même si, à l'heure actuelle, toute l'information n'est pas physiquement disponible. Néanmoins, il est dans ce cas possible d'appliquer certaines techniques de filtrage pour reconstituer l'état global de la fermentation à partir des mesures indirectes (Cf. par exemple l'application du filtre de Kalman et ses dérivés).

³⁰ C'est le cas des *réacteurs batch séquentiels* (Sequencing Batch Reactors, *SBR*).

Les modèles de bilans de matière appartiennent à la deuxième approche. Ces types de modèles permettent de décrire le comportement dynamique des différents composants de la réaction biologique en exprimant par des équations différentielles la variation de la quantité d'un composé comme égale à la somme de ce qui est produit ou apporté, diminué de ce qui est consommé ou soutiré. Un des aspects importants des modèles de bilan matières est qu'ils sont constitués de deux types de termes représentant respectivement la conversion et le transport. La conversion englobe la cinétique à laquelle se déroulent les diverses réactions biochimiques du procédé et les rendements de conversion des divers substrats en biomasse et produits. Pour sa part, la dynamique de transport regroupe le transit de la matière au sein du procédé sous forme solide, liquide ou gazeuse, et les phénomènes de transfert entre phases. Dans ce chapitre nous nous intéressons au développement des expressions qui nous permettront de modéliser les phénomènes globaux de bioconversion.

Bien que la recherche des expressions mathématiques décrivant le fonctionnement du processus demeure encore aujourd'hui une phase capitale et une source d'étude remarquable, une distinction importante doit toutefois être faite entre les modèles de connaissance et les modèles d'action [STE98]. En effet, la structure et la complexité d'un modèle ne peuvent être définies indépendamment de l'objectif désiré. Il faut préciser le but d'utilisation auquel on le destine. Alors que les modèles de connaissance se proposent de décrire des phénomènes extrêmement complexes en incluant un certain nombre d'étages de conversions enzymatiques - et par là même un grand nombre de paramètres - les modèles d'action sont essentiellement élaborés en vue de la commande. Ces derniers, souvent obtenus par simplification des modèles de connaissance, font intervenir deux aspects de modélisation :

- détermination des cinétiques de conversion au travers des vitesses de réaction des variables macroscopiques,
- détermination des équations d'état du système par l'intermédiaire des bilans de matière relatifs aux principaux composés intervenants dans la réaction.

2.2.2. Le taux spécifique de croissance

Dans tout procédé biologique, la vitesse de croissance de la masse cellulaire - v_X - est proportionnelle à la variation de la concentration en microorganismes - X - à tout instant. Nous obtenons :

$$v_X = \frac{1}{V} \frac{d(XV)}{dt} \quad (2.1)$$

Le taux spécifique de croissance, notée μ , est définie comme étant la vitesse de croissance ramenée à l'unité de masse de la biomasse :

$$\mu = \frac{1}{XV} \frac{d(XV)}{dt} \quad (2.2)$$

Ce paramètre est un des plus importants pour la description de l'évolution de la population microbienne. Malheureusement, d'un point de vue général, le taux de croissance dépend fortement de plusieurs facteurs tels que les conditions opératoires (température, pH...) et le milieu réactionnel (concentrations en composés carbonés, azotés, phosphores, en produits formés, en sels minéraux, en oxygène...). La modélisation de ces différentes influences au sein d'une expression mathématique unique est très complexe, voire impossible. La méthodologie usuellement employée consiste donc à fixer les variables d'environnement à des valeurs déterminées expérimentalement comme étant optimales. Ne demeurent alors que les influences des substrats, produits et éventuellement de la biomasse.

Historiquement l'expression la plus couramment utilisée est le modèle empirique de Monod [MON42], qui lui a permis de décrire la croissance bactérienne de la loi introduite au début du XX^e siècle par Michaëlis-Menten pour une cinétique enzymatique :

$$\mu(t) = \mu_{\max} \frac{S(t)}{K_S + S(t)} \quad (2.3)$$

avec $S(t)$: concentration en substrat (g/L)³¹,
 $\mu_{\max}$: taux spécifique de croissance maximal (h⁻¹),
 K_S : constante de demi-saturation (g/L), correspond à l'affinité des micro-organismes pour le substrat limitant.

Cette expression décrit que le taux spécifique de croissance μ , en présence d'une quantité donnée de substrat S , commence par augmenter rapidement lorsqu'on augmente S , puis devient finalement indépendante de S par un effet de saturation. Egalement l'expression permet de décrire le phénomène de limitation de la croissance par manque de substrat et l'arrêt complet lorsque le substrat n'est plus disponible. Le comportement de μ en fonction de la concentration de substrat en suivant la loi de Monod est représenté dans la Figure 2.2 pour plusieurs valeurs du paramètre $K_S \in [0,1 ; 5]$ [STE98].

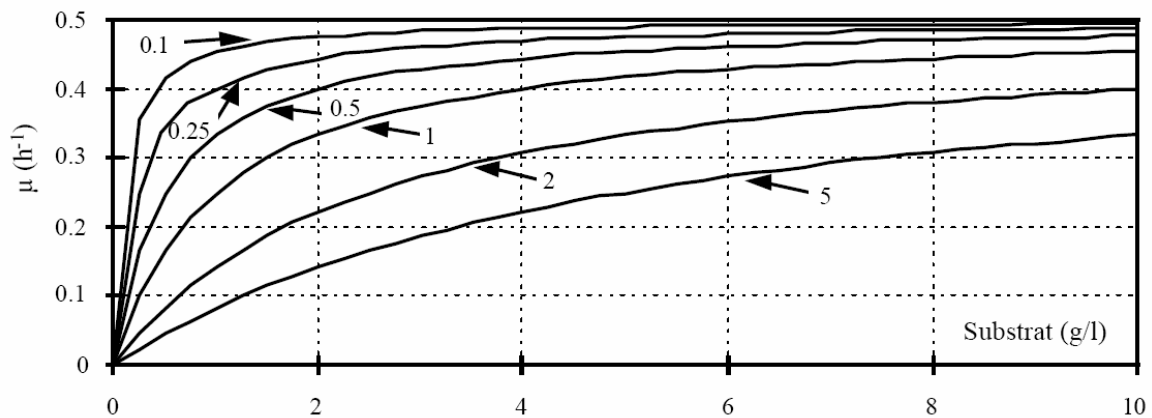


Figure 2.2. Loi de Monod pour $\mu_{\max} = 0,5 \text{ h}^{-1}$ et plusieurs valeurs du paramètre K_S

Toutefois, le modèle décrit par l'expression précédente possède l'inconvénient de ne traduire que la limitation en substrat. Pour prendre en compte les effets inhibiteurs des hautes concentrations, ces phénomènes sont généralement modélisés par l'expression de Haldane, introduite dans le cas des réactions enzymatiques et reprise par Andrews [AND68] dans le cas de réactions biologiques, en particulier dans la digestion anaérobie :

$$\mu = \mu_{\max} \frac{S}{K_S + S + S^2 / K_I} \quad (2.4)$$

où K_I est ici la constante d'inhibition (g/L).

Le comportement du taux spécifique de croissance μ en fonction de la concentration de substrat en suivant la loi d'Andrews est représenté dans la Figure 2.3 pour K_S constant et plusieurs valeurs du paramètre $K_I \in [0,1 ; 100]$ [STE98].

³¹ Afin d'éviter toute confusion au niveau de la notation, nous utiliserons toute au long de ce document « L » pour dénoter « litre » comme l'unité de mesure de volume.

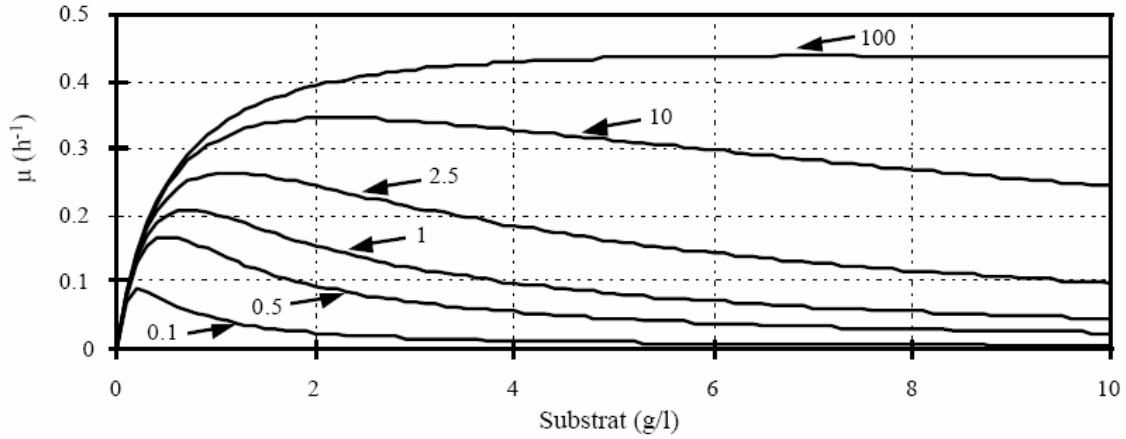


Figure 2.3. Loi d'Andrews pour $\mu_{\max} = 0,5 \text{ h}^{-1}$, $K_S = 0,5 \text{ g}_S/\text{L}$ et plusieurs valeurs du paramètre K_I .

Il convient de noter que de nombreuses autres relations algébriques ont été établies pour décrire ces phénomènes de limitation et/ou d'inhibition, mais que leur utilisation reste marginale. De la même manière, certains modèles prennent en compte l'influence de la concentration en micro-organismes [CON59], en cométabolite, de la température, du pH,... A ce propos, une liste de plus de 50 expressions analytiques peut être trouvée dans [DOC86].

Le cas de l'oxygène est un peu particulier. En effet, dans le cas de procédés fonctionnant en aérobie, l'oxygène correspond à un cosubstrat de la réaction et peut ainsi être traité comme tel, c'est-à-dire intervenir sous la forme d'un terme de type Monod dans l'expression du taux de croissance, conduisant ainsi à l'expression suivante:

$$\mu = \mu_{\max} \left(\frac{S}{K_S + S} \right) \left(\frac{O_2}{K_{O_2} + O_2} \right) \quad (2.5)$$

avec O_2 : concentration en oxygène dissous (g_{O_2}/L),

K_{O_2} : constante de demi-saturation pour l'oxygène (g_{O_2}/L).

Cependant, cette expression est souvent omise sous l'hypothèse que le réacteur est suffisamment aéré et que K_{O_2} est très petit par rapport à la concentration en oxygène dissous présente dans le réacteur en fonctionnement normal.

Dans certains cas, la présence du cosubstrat peut aussi se manifester de manière compétitive. Considérons à cet effet une bactérie telle que *pseudomonas putida* qui peut, sous certaines conditions, utiliser simultanément le phénol (S_1) et le glucose (S_2) comme sources d'énergie. Le taux spécifique de croissance peut alors s'exprimer comme la somme de deux taux spécifiques de croissance sur chacun des substrats :

$$\mu = \mu_1(S_1 + S_2) + \mu_2(S_1 + S_2) \quad (2.6)$$

Le phénomène d'activation/inhibition compétitive dans le mélange bisubstrat peut être décrit par une expression de la forme [WAN96] :

$$\mu_i(S_i + S_j) = \frac{\mu_{\max,i} S_i}{K_{S,i} + S_i + S_i^2 / K_{I,i} + K_{3,i} S_j + K_{4,i} S_i S_j}, \quad i, j = 1, 2, \quad i \neq j \quad (2.7)$$

En l'absence du substrat S_j , $\mu_i(S_i, 0)$ se ramène à l'expression d'Andrews précédemment décrite. Par un choix judicieux de $K_{3,i}$, $K_{4,i}$, $i = 1, 2$, on peut alors aisément décrire les phénomènes d'inhibition compétitive ou non, croisée ou partielle.

2.2.3. Equations d'état

Quel que soit le mode de fonctionnement (batch, fed batch, continu), le comportement dynamique des différents composants de la réaction biologique (biomasse, substrat, produits) découle directement de l'expression des bilans de matières. L'équation générale d'évolution de chacun de ces éléments sur un intervalle de temps déterminé est régie par :

$$\left[\begin{array}{c} \text{Variation de la quantité} \\ \text{dans le réacteur} \end{array} \right] = \left[\begin{array}{c} \text{Quantité} \\ \text{apportée} \end{array} \right] - \left[\begin{array}{c} \text{Quantité} \\ \text{soutirée} \end{array} \right] + \left[\begin{array}{c} \text{Quantité} \\ \text{produite} \end{array} \right] - \left[\begin{array}{c} \text{Quantité} \\ \text{consommée} \end{array} \right]$$

Dans le cas général d'un réacteur infiniment mélangé³² à alimentation et soutirage permanent et à volume variable, les équations de fonctionnement décrivant la croissance d'une population de micro-organismes sur un simple substrat est alors représentée par les équations différentielles ordinaires suivantes³³ [STE98] :

<i>Variation</i>	<i>Quantité apportée</i>		<i>Quantité soutirée</i>		<i>Quantité produite</i>		<i>Quantité consommée</i>
$\frac{d(VX)}{dt} =$	0	-	$Q_{out} X$	+	$v_X V$	-	0
$\frac{d(VS)}{dt} =$	$Q_{in} S_{in}$	-	$Q_{out} S$	+	0	-	$v_S V$
$\frac{dV}{dt} =$	Q_{in}	-	Q_{out}	+	0	-	0

avec S_{in} : concentration de substrat dans la solution d'alimentation (gs/L),
 Q_{in} : débit d'alimentation en substrat (L/h),
 Q_{out} : débit de soutirage (L/h),
 V : volume réactionnel (L),
 X : concentration en micro-organismes (gx/L),
 v_X : vitesse de croissance de la biomasse (gx/Lh),
 v_S : vitesse de consommation du substrat (gs/Lh).

Le couplage entre croissance de la biomasse et consommation du substrat est généralement décrit par la relation algébrique suivante :

$$v_S = \frac{\mu}{Y_{X/S}} \quad (2.8)$$

avec $Y_{X/S}$ le rendement de conversion substrat/biomasse (gx/g_S).

³² *Stirred Tank Reactor*, noté STR, signifie que le milieu réactionnel est homogène.

³³ En raison de notre domaine d'application, nous sommes intéressés par les évolutions de la biomasse et du substrat. Toutefois, un raisonnement analogue permettrait d'obtenir l'équation d'évolution du produit formé.

En se ramenant aux concentrations des composés, nous obtenons:

$$\left\{ \begin{array}{l} \frac{dX}{dt} = \mu X - \frac{Q_{out}}{V} X \end{array} \right. \quad (2.9)$$

$$\left\{ \begin{array}{l} \frac{dS}{dt} = -\frac{\mu}{Y_{X/S}} X + \frac{Q_{in}}{V} (S_{in} - S) \end{array} \right. \quad (2.10)$$

$$\left\{ \begin{array}{l} \frac{dV}{dt} = Q_{in} - Q_{out} \end{array} \right. \quad (2.11)$$

Il est possible de retrouver les équations de fonctionnement des trois principaux modes opératoires des bioprocédés :

- *pour un procédé batch* : $Q_{in} = Q_{out} = 0$, la totalité des éléments nutritifs est introduite dès le lancement de la fermentation et seules les conditions d'environnement du milieu fermentaire subissent une action de l'opérateur (pH, température, vitesse d'agitation, débit d'aération,...) ;
- *pour un procédé semi-continu ou fed batch* : $Q_{in} \neq 0$ et $Q_{out} = 0$, les éléments nutritifs sont ici apportés au fur et à mesure des besoins des micro-organismes par action sur le débit d'alimentation ;
- *pour un procédé continu* : $Q_{in} = Q_{out} \neq 0$, dans ce cas, le procédé se déroule à volume constant, les apports continus d'éléments nutritifs étant balancés par un soutirage en tout point identique.

2.3. Influence du transfert de matière

Comme certains procédés de dépollution fonctionnent en aérobie, nous devons également établir l'expression de la variation de la concentration en oxygène dissous. Dans les réacteurs conventionnels, cette concentration est assurée par l'utilisation des agitateurs. L'agitation possède 4 fonctions principales :

1. assurer l'homogénéité du milieu de culture à travers un macro-mélange,
2. promouvoir le transfert de masse interfacial à travers un micro-mélange,
3. promouvoir le transfert de chaleur,
4. assurer une aire interfaciale importante pour un bon transfert de matière entre gaz et liquide [DEM86].

Malgré la présence des agitateurs au niveau industriel, les réacteurs de grande échelle ne sont généralement pas parfaitement mélangés. Les rendements de conversion et la croissance des micro-organismes peuvent donc diminuer par rapport à des réacteurs de plus petites dimensions. Des règles - parfois empiriques - sont toutefois disponibles pour effectuer ces changements d'échelle et les simulations dynamiques peuvent être bénéfiques pour aider au dimensionnement des installations. Malheureusement, les débits liquides et gazeux internes au réacteur étant extrêmement complexes, un modèle mathématique rigoureux du transport de l'oxygène est très difficile à développer [STE98].

En ce qui concerne la modélisation de la dynamique de l'oxygène, elle est obtenue de façon identique au cas précédent, à partir de l'équation générale de bilan de masse de l'oxygène dissous :

$$\frac{d(VO_2)}{dt} = -q_{O_2} XV + K_1 a V (O_2^* - O_2) - Q_{out} O_2 \quad (2.12)$$

avec O_2^* : concentration de saturation³⁴ en oxygène dissous (g_{O_2}/L),
 q_{O_2} : vitesse spécifique de consommation d'oxygène ($g_{O_2}/g_X/h$),
 $K_1 a$: coefficient global de transfert de matière (h^{-1}).

Pour sa part, la vitesse spécifique de consommation d'oxygène est donnée par l'expression :

$$q_{O_2} = \frac{\mu}{Y_{X/O}} + m_{O_2} \quad (2.13)$$

avec $Y_{X/O}$: rendement de conversion (g_X/g_{O_2}),
 m_{O_2} : coefficient de maintenance ($g_{O_2}/g_X/h$).

Le coefficient de transfert de matière $K_1 a$ dépend fortement des conditions de fonctionnement et en particulier de l'agitation, de la pression, du débit d'aération, du milieu de culture.... Par exemple, dans le cas d'une agitation assurée par des turbines, une masse de travail considérable a été fournie par le passé pour modéliser le transfert de matière de l'oxygène de la phase gazeuse à la phase liquide. L'expression du $K_1 a$ est alors donnée par une fonction de la puissance dissipée par unité de volume en présence de gaz et de la vitesse ascensionnelle du gaz, dans laquelle interviennent plusieurs constantes liées à la géométrie du réacteur et à celle de l'agitateur. En résumé, la détermination du paramètre $K_1 a$ doit être faite en fonction des conditions opératoires. C'est d'ailleurs ce problème qui fait que la mesure de l'oxygène a rarement été utilisée à des buts de contrôle, bien qu'elle soit très facile à obtenir.

2.4. Quelques modèles représentatifs

Sans vouloir faire une liste exhaustive, nous présentons à titre illustratif, quelques modèles représentatifs des procédés biologiques de dépollution basés sur l'approche de modélisation de bilan de matières [QUE00].

2.4.1. Fermentations continues

Reprenons le modèle (2.9) - (2.11) décrivant la croissance d'une population de micro-organismes X sur un substrat limitant S . On obtient directement, dans le cas de la fermentation continue, i.e., à volume V constant, et en rajoutant l'équation relative à un métabolite P , le modèle classique suivant :

$$\left\{ \begin{array}{l} \frac{dX}{dt} = \mu X - DX \end{array} \right. \quad (2.14)$$

$$\left\{ \begin{array}{l} \frac{dS}{dt} = -\frac{\mu}{Y_{X/S}} X + D(S_{in} - S) \end{array} \right. \quad (2.15)$$

$$\left\{ \begin{array}{l} \frac{dP}{dt} = \frac{\mu}{Y_{X/P}} X - DP \end{array} \right. \quad (2.16)$$

³⁴ Concentration qui dépend du milieu de culture (sels minéraux) et de la température essentiellement.

dans lequel $D=Q_{in}/V=Q_{out}/V$ représente le *taux de dilution* (h^{-1}), $Y_{X/P}$ le rendement de conversion (g de biomasse/g de produit) et μ est décrit par (2.3). L'équation correspondant à la production du métabolite est généralement omise lorsque cette variable n'a pas d'influence sur la croissance (μ ne dépend pas de P) et qu'elle n'est pas mesurée (donc non utilisée au niveau d'observateurs et/ou de contrôleurs).

D'autres phénomènes cellulaires tels que la *respiration endogène* et le besoin de *maintenance* peuvent être pris en compte au niveau de la modélisation [EDE97] :

- La *respiration endogène*, correspond à l'utilisation différée des réserves par combustion de substrats endogènes, par opposition aux autres substrats qui sont exogènes. Le phénomène commence par la combustion d'abord des réserves, puis de certains constituants cellulaires. Bien qu'il s'agisse d'une respiration endogène, il y a deux conceptions à son propos, selon qu'on considère l'ensemble de la consommation d'oxygène associée à la consommation des réserves et à la lyse³⁵ cellulaire, ou seulement cette dernière.
- On appelle énergie de maintien celle qui est nécessaire pour la resynthèse de certains composants cellulaires. Les besoins de *maintenance* des cellules correspondent à une consommation de substrat externe non accompagné de croissance cellulaire, et fournissant à la cellule le minimum d'énergie dont elle a besoin pour remplacer ses déchets et conserver son intégrité.

Ces deux phénomènes sont représentés sous la forme de termes supplémentaires dans les équations précédents : b représente le taux de mortalité bactérienne (h^{-1}) associé à la respiration endogène et m le coefficient de maintenance (que l'on retrouve systématiquement sur l'équation d'évolution de l'oxygène (2.13)) de la manière suivante:

$$\left\{ \begin{array}{l} \frac{dX}{dt} = \mu X - DX - bX \end{array} \right. \quad (2.17)$$

$$\left\{ \begin{array}{l} \frac{dS}{dt} = -\frac{\mu}{Y_{X/S}} X + D(S_{in} - S) - mX \end{array} \right. \quad (2.18)$$

$$\left\{ \begin{array}{l} \frac{dP}{dt} = \frac{\mu}{Y_{X/P}} X - DP \end{array} \right. \quad (2.19)$$

On peut étendre directement le type de modèle (2.17) - (2.19) à des réactions plus complexes comportant plusieurs populations de micro-organismes couplées ou non, et/ou plusieurs réactions enchaînées. Ainsi par exemple, nous présenterons en détail le modèle d'un traitement biologique des effluents de papeterie par lagunage aéré comportant une croissance sur mélange bi-substrat avec activation/inhibition compétitive dans la section 2.5.

Nitrification

Dans une station de traitement biologique de l'azote des effluents urbains, le cycle de transformation de l'azote est le même que dans la nature (Cf. Figure 2.4). La première partie du traitement de la pollution azotée est donc la nitrification, qui consiste à transformer par voie biologique l'azote ammoniacal (NH_4^+) en en nitrate (NO_3^-), en présence d'oxygène [PRO90] [QUE98]. Cette réaction comprend deux étapes. La première, la nitritation met en

³⁵ Dissolution ou destruction des cellules par rupture de leur membrane libérant le contenu intérieur.

jeu des bactéries autotrophes de type *Nitrosomonas* (X_{NS}), qui transforment l'azote ammoniacal en nitrite (NO_2^-). La deuxième étape concerne, quant à elle, la conversion du nitrite en nitrate par les bactéries autotrophes *Nitrobacter* (X_{NB}).

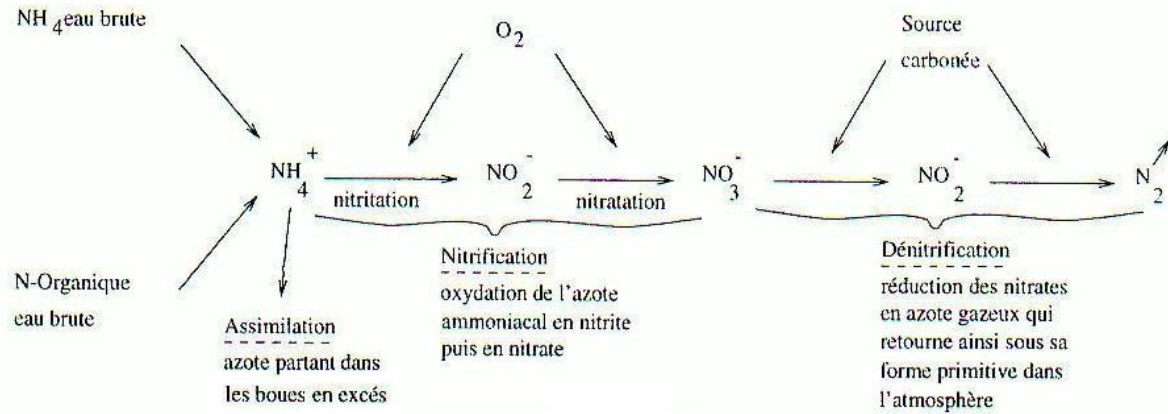


Figure 2.4. Cycle d'élimination des pollutions azotées.

En réacteur complètement mélangé et alimenté continûment en azote ammoniacal NH_4^+ , la nitrification peut être modélisée par les équations différentielles issues des bilans de matières relatifs aux cinq variables principales de la réaction, à savoir, les deux populations bactériennes, l'azote ammoniacal, le nitrite et le nitrate :

$$\left\{ \begin{array}{l} \frac{dX_{NS}}{dt} = \mu_{NS} X_{NS} - DX_{NS} \quad (2.20) \\ \frac{dX_{NB}}{dt} = \mu_{NB} X_{NB} - DX_{NB} \quad (2.21) \\ \frac{d\text{NH}_4^+}{dt} = -\frac{\mu_{NS}}{Y_{NS}} X_{NS} + D(\text{NH}_4^+ - \text{NH}_4^+) \quad (2.22) \\ \frac{d\text{NO}_2^-}{dt} = \frac{\mu_{NS}}{Y_{NS}} X_{NS} - \frac{\mu_{NB}}{Y_{NB}} X_{NB} - D\text{NO}_2^- \quad (2.23) \\ \frac{d\text{NO}_3^-}{dt} = \frac{\mu_{NB}}{Y_{NB}} X_{NB} - D\text{NO}_3^- \quad (2.24) \end{array} \right.$$

Y_{NS} et Y_{NB} représentent les rendements des deux étapes du procédé. Les taux de croissance μ_{NS} et μ_{NB} sont décrits par des termes de Monod relatifs aux substrats limitants des deux populations, NH_4^+ et NO_2^- respectivement :

$$\mu_{NS} = \mu_{\max_{NS}} \frac{\text{NH}_4^+}{K_{\text{NH}_4^+} + \text{NH}_4^+} \quad (2.25)$$

$$\mu_{NB} = \mu_{\max_{NB}} \frac{\text{NO}_2^-}{K_{\text{NO}_2^-} + \text{NO}_2^-} \quad (2.26)$$

Cette représentation largement acceptée dans la littérature sous-entend que les deux étapes de la nitrification sont indépendantes (pas de compétition entre les micro-organismes),

le co-métabolite de la première devenant le substrat de la seconde. Il convient aussi de noter que le modèle (2.20) - (2.24) n'explicite pas la présence de l'oxygène, bien que celle-ci soit indispensable à la croissance des deux populations autotrophes, partant du principe qu'il est apporté en quantité suffisante pour ne pas limiter le phénomène d'oxydation biologique.

Digestion anaérobie

La digestion anaérobie est le procédé de consommation de déchets solides composés de matière organique par voie biologique. La matière organique complexe (macromolécules) est transformée en biogaz (méthane et gaz carbonique) par une séquence de réactions suivant quatre étapes principales : l'hydrolyse, l'acidogénèse, l'acétogénèse et la méthanogénèse [DEN88]. Ces étapes sont schématisées sur la Figure 2.5. L'hydrolyse consiste en la transformation de matière organique particulaire lentement biodégradable X_s en matière organique solubilisée facilement biodégradable S_s . Ces molécules peuvent servir de source de carbone aux biomasses hétérotrophes intervenant en particulier lors de la dénitrification. Dans le cas contraire, elles sont transformées en différents acides organiques, ou acides gras volatils (AGV) par les deux étapes d'acidogénèse et d'acétogénèse. Sachant que l'on ne mesure généralement que l'ensemble de ces acides, et que c'est l'acétate qui est dominant, certains modèles court-circuitent l'acétogénèse pour ramener la réaction en un processus en trois étapes [CHA97] [MOL86]. Enfin, la dernière étape, réalisée par les bactéries méthanogènes anaérobies strictes transforment les AGV en méthane et en gaz carbonique. Le méthane peut alors être récupéré et valorisé comme source d'énergie.

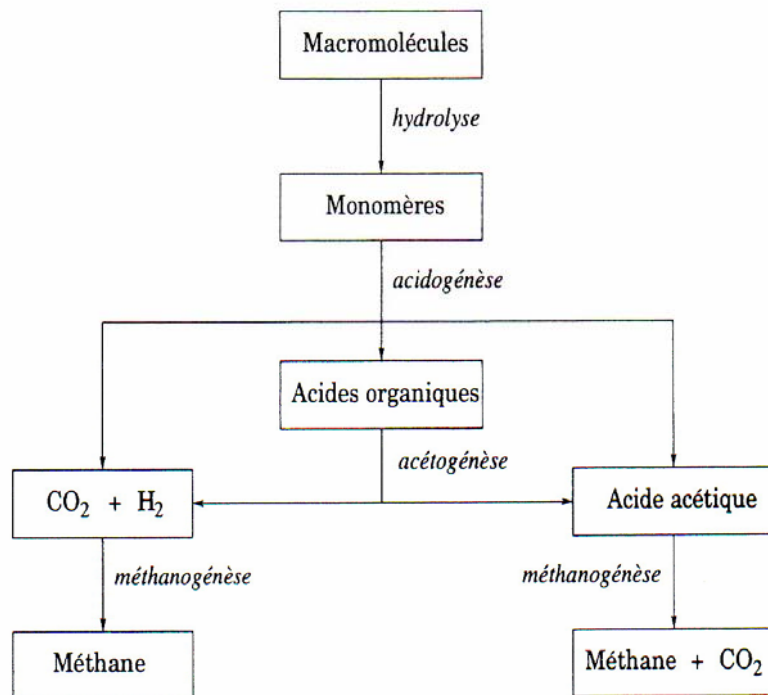


Figure 2.5. Principales étapes de la digestion anaérobie

En considérant donc la digestion anaérobie comme un processus en trois étapes, comportant trois populations de micro-organismes, les bactéries hydrolytiques X_h , les bactéries acidogènes X_a et les bactéries méthanogènes X_m , et leurs substrats respectifs, le déchet organique solide X_s , la matière organique solubilisée S_s et les acides gras volatils A , le modèle suivant est obtenu pour décrire la croissance microbienne en réacteur alimenté en continu par les déchets solides $X_{s,in}$ [GER00] :

$$\left\{ \begin{array}{l} \frac{dX_h}{dt} = \mu_h X_h - b_h X_h - X_h D \quad (2.27) \\ \frac{dX_a}{dt} = \mu_a X_a - b_a X_a - X_a D \quad (2.28) \\ \frac{dX_m}{dt} = \mu_m X_m - b_m X_m - X_m D \quad (2.29) \\ \frac{dX_s}{dt} = -\frac{1}{Y_h} \mu_h X_h + D(X_{s_{in}} - X_s) \quad (2.30) \\ \frac{dS_s}{dt} = \left(\frac{1-Y_h}{Y_h} \mu_h X_h \right) (1-f_{X_I}) - \frac{1}{Y_a} \mu_a X_a + D(S_{s_{in}} - S_s) \quad (2.31) \\ \frac{dA}{dt} = \left(\frac{1-Y_a}{Y_a} \mu_a X_a \right) (1-f_{S_I}) - \frac{1}{Y_m} \mu_m X_m + D(A_{in} - A) \quad (2.32) \\ \frac{dX_I}{dt} = f_{X_I} \left(\frac{1-Y_h}{Y_h} \mu_h X_h \right) + D(X_{I_{in}} - X_I) \quad (2.33) \\ \frac{dS_I}{dt} = f_{S_I} \left(\frac{1-Y_a}{Y_a} \mu_a X_a \right) + D(S_{I_{in}} - S_I) \quad (2.34) \end{array} \right.$$

X_I et S_I représentent respectivement les concentrations en matière inerte particulaire et solubilisée. Les taux de croissance sont modélisés par la loi de Monod relative aux substrats respectifs de chacune des populations. L'équation de production de méthane peut éventuellement être rajoutée si celui-ci est mesuré et présente donc un intérêt au niveau des étapes d'observation et/ou de commande.

Des modèles plus complexes peuvent être proposés [GIR86] [DEN88] [KIE97], qui prennent en compte davantage de variables, mais cela pose beaucoup plus de problèmes pour déterminer les valeurs numériques des nombreux paramètres associés.

2.4.2. Fermentations semi-continues

L'extension du modèle (2.9) - (2.11) au cas des fermentations semi-continues est obtenue directement en considérant $Q_{out} = 0$. Cette stratégie est particulièrement adaptée au cas de polluants toxiques qui ne doivent pas se retrouver dans les effluents, et qui ne permettent pas d'obtenir des productivités élevées en cultures discontinues [QUE00].

A titre illustratif, considérons le cas de la biodégradation du phénol. C'est un polluant toxique contenu dans les eaux usées de nombreuses industries chimiques, pétrochimiques et agrochimiques. Sa dégradation par des bactéries l'acceptant comme seule source de carbone et d'énergie est fortement inhibée, même à de très faibles concentrations. L'étude de la croissance de *Ralstonia eutropha* [LEO99] a permis d'établir un modèle bilans-matières dans lequel le taux de croissance est modélisé par l'expression de Haldane (2.4) et est couplé à la vitesse spécifique de dégradation du phénol. La croissance aérobie peut ainsi être décrite par les équations suivantes :

$$\left\{ \begin{array}{l} \frac{dX}{dt} = \mu X - \frac{Q_{in}}{V} X \\ \frac{dS}{dt} = -\frac{\mu}{Y_{x/s}} X + \frac{Q_{in}}{V} (S_{in} - S) \\ \frac{dP}{dt} = v_p X - \frac{Q_{in}}{V} P \\ \frac{dO_2}{dt} = K_l a (O_2^* - O_2) - q_{O_2} X - \frac{Q_{in}}{V} O_2 \\ \frac{dV}{dt} = Q_{in} \end{array} \right. \quad \begin{array}{l} (2.35) \\ (2.36) \\ (2.37) \\ (2.38) \\ (2.39) \end{array}$$

P est un co-métabolite de la croissance, l'acide 2-hydroxymuconique semialdéhyde (noté 2-hms). L'accumulation de 2-hms se traduit par une coloration jaune de plus en plus intense, corrélée avec le taux de croissance :

$$v_p = \alpha_0 + \alpha_1 \mu \quad (2.40)$$

où α_1 représente un terme de rendement, alors que α_0 , qui n'a pas de réelle signification physique, a été introduit pour assurer l'ajustement des données [LEO99].

Concernant la dynamique de l'oxygène, comme nous l'avons dit dans un paragraphe précédent, elle n'a d'intérêt que dans la mesure où l'on veut se servir de l'oxygène comme source d'information de la réaction. La vitesse spécifique est donnée par l'expression (2.13) et le coefficient de transfert de matière $K_l a$ doit être déterminé en fonction des conditions opératoires [QUE00a].

2.4.3. Procédés à boues activées

Les procédés à boues activées sont très largement utilisés pour le traitement biologique des eaux usées. Traditionnellement, ils sont composés d'un réacteur biologique et d'un décanteur/clarificateur schématisés sur la Figure 2.6. L'eau polluée provenant d'une source externe circule dans le bassin d'aération³⁶ dans lequel la biomasse bactérienne dégrade la matière organique. Les micro-organismes s'agglomèrent en floes et produisent les boues. La sortie du bassin d'aération comportant la liqueur mixte est ensuite envoyée dans le décanteur où la séparation de l'eau épurée et des floes bactériens est faite par gravité.

Dans l'aérateur se produit le processus d'épuration de l'eau proprement dit. En pratique l'aération est réalisée en utilisant des agitateurs superficiels ou diffuseurs, chargés d'apporter l'oxygène nécessaire pour la réaction. Le bassin d'aération est alors considéré hypothétiquement comme infiniment mélangé. La composition est supposée uniforme dans tout le bassin (le gradient de concentration est nul), autant en concentration de biomasse que de substrat rémanent et d'oxygène dissous. Une fraction des boues décantées est recyclée vers l'aérateur pour maintenir sa capacité d'épuration, l'autre est éliminée. La recirculation permet de réinjecter une partie de la biomasse qui est plus concentrée et en même temps, de découpler

³⁶ Le bassin d'aération est en fait, selon le cas, constitué de plusieurs bassins en série privilégiant chacun le traitement d'une pollution spécifique (organique, nitrification, dénitrification, phosphatation), ou d'un seul bassin permettant de réaliser les différentes réactions biologiques en même temps [QUE00].

le temps de séjour hydraulique du temps de séjour cellulaire (âge des boues). L'effluent épuré est ainsi récupéré en sortie du décanteur.

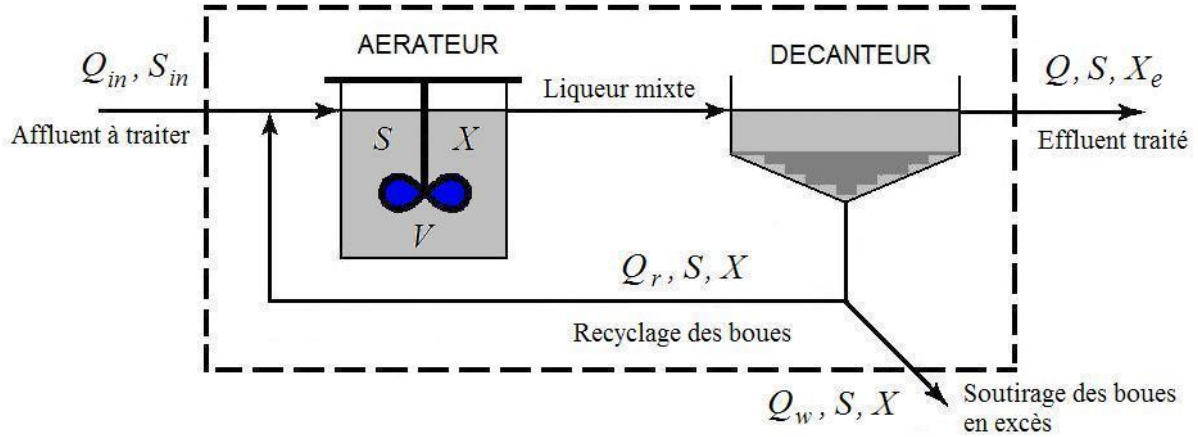


Figure 2.6. Schéma typique d'un procédé d'épuration par boues activées

Pour ce procédé le substrat affluent a une concentration S_{in} (g/L) et un débit $Q_{in}=Q$ (L/h). On doit ajouter au substrat affluent le débit de recirculation Q_r (L/h) provenant du décanteur avec une concentration S (g/L). Le débit de purge est représenté par Q_w (L/h). En appliquant les bilans-matière pour le procédé, on arrive aux expressions présentées ensuite.

Pour le substrat, dans le bassin d'aération,

$$V \frac{dS}{dt} = QS_{in} + Q_r S - (Q + Q_r - Q_w) S - Q_w S - v_s XV \quad (2.41)$$

où $v_s XV$ représente la consommation totale de substrat par des micro-organismes (g/h). Il faut remarquer que comme le bassin d'aération est supposé infiniment mélangé, le substrat a la même concentration tant dans le bassin que dans la sortie et dans la recirculation (en considérant qu'il n'y a pas de réaction dans le décanteur). Après simplification, l'expression (2.41) peut s'exprimer également par :

$$\frac{dS}{dt} = D(S_{in} - S) - \mu X \quad (2.42)$$

dans laquelle $D=Q/V$ représente le taux de dilution (h^{-1}).

Pour la biomasse, le bilan de matière en considérant le procédé complet (réacteur et décanteur) afin de prendre en compte les solides effluents X_e est donné par :

$$V \frac{dX}{dt} = -(Q - Q_w) X_e - Q_w X + \mu XV \quad (2.43)$$

expression dans laquelle μXV représente la production de biomasse par la dégradation du substrat (g/h).

En ce qui concerne l'applicabilité, les procédés à boues activées sont souvent utilisés à une échelle pilote pour caractériser le comportement des micro-organismes pour l'épuration d'un affluent spécifique, ainsi que pour la sélection et l'optimisation des expériences du traitement, avant de passer à une échelle industrielle. La Figure 2.7 montre une unité pilote typique d'un procédé à boues activées implémenté au cours de la formation de Doctorat au

Centre d'Informatique et d'Automatisation pour la Production de l'Université de los Andes (Colombie). Le pilote est utilisé pour la caractérisation expérimentale et la mise au point d'un procédé de traitement des effluents d'une industrie locale de boissons gazeuses. Le montage physique a permis d'intégrer dans une seule structure le bassin d'aération (à gauche) et le décanteur (à droite). L'aération est assurée par une batterie d'aérateurs externes qui injectent de l'oxygène dans le réacteur. La couleur plus claire de l'effluent sur la partie supérieure du bassin de décantation met en évidence le bon fonctionnement du procédé.

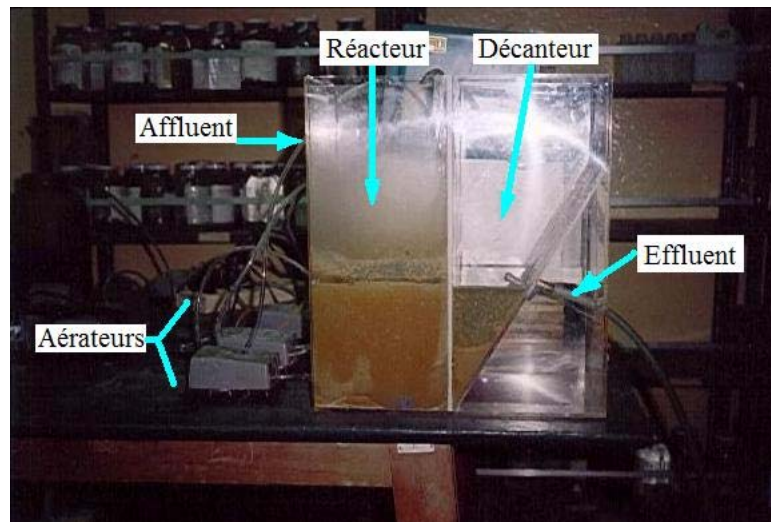


Figure 2.7. Pilote expérimental d'un procédé à boues activées implémenté à l'Université de los Andes

Le schéma de base pour les procédés à boues activées montré dans la Figure 2.6 peut être complété par d'autres réacteurs en série (permettant de favoriser les différentes populations de micro-organismes en optimisant les conditions opératoires de chaque réacteur), par un décanteur primaire et par des boucles de recirculation interne. D'autre part, afin de traiter les pollutions composées d'azote organique, le réacteur biologique peut prendre en compte à la fois des conditions aérobies (pour la nitrification), et anoxies (pour la dénitrification en particulier). Ces opérations peuvent être réalisées soit dans deux bassins placés successivement, l'un aéré, l'autre non, soit dans un seul bassin autorisant des périodes alternées d'aération et de non-aération [JUL97] [SPE98]. Mis à part la modélisation des étapes de traitement des déchets (nitrification, dénitrification, digestion...), un élément clé de ces procédés est la modélisation du décanteur afin de gérer au mieux la recirculation des boues [QUE00b] [TRA03].

L'intérêt des procédés à boues activées a donné lieu à des nombreuses études où la modélisation dynamique de ces procédés a beaucoup avancée. A ce propos, l'IAWQ³⁷ a développé un modèle qui a été validé sur de nombreux cas pratiques, servant de référence dans le domaine du traitement combiné des matières organiques biodégradables (lentement et rapidement) et de l'azote dans les effluents urbains [HEN87]. Dans sa deuxième version, l'IAWQ a également élargie ce modèle pour tenir compte du traitement biologique du phosphore [HEN94]. Le modèle se présente sous forme matricielle, les lignes correspondant aux différentes réactions biologiques et les colonnes aux variables mises en jeu dans ces réactions. De très nombreux composants interviennent dans les schémas réactionnels et, au final, un "grand" modèle est obtenu. Ainsi, le modèle n° 2 de l'IAWQ fait intervenir 19 variables d'états et 65 paramètres et, même si les auteurs de ce modèle proposent des valeurs

³⁷ IAWQ : *International Association for Water Quality*.

pour ces paramètres, ils doivent souvent être recalculés pour modéliser le traitement d'un effluent spécifique. Ceci étant, même si ce modèle ne constitue pas une réponse définitive au problème de la modélisation des procédés biologiques de dépollution, il se révèle particulièrement utile pour répondre à plusieurs besoins, telles que la sélection des expériences pour la mise au point de procédés ainsi que l'optimisation, le diagnostic et l'assistance au dimensionnement de procédés à l'échelle industrielle [STE98].

2.5. Modélisation du bioprocédé de traitement des eaux usées par lagunage aérée

Dans cette section nous décrivons en détail le procédé considéré dans le cadre de notre étude. Il s'agit d'un réacteur biologique aérobie multivariable en mode continu pour le traitement des eaux usées issues d'une industrie papetière. Nous avons repris les travaux développés par [BEN96], notamment l'étude de caractérisation du procédé qui a nécessité le montage d'une unité pilote afin de reproduire les conditions opératoires d'une lagune aérée pour le traitement de ces effluents. Le modèle mathématique ainsi obtenu a été utilisé comme prototype pour le développement et le test de performance de différents algorithmes d'identification et de commande dans le cadre du projet européen FAMIMO³⁸ [ROU99]. Le modèle que nous développons est une version améliorée du simulateur utilisé lors du projet FAMIMO, en effet nous avons proposée l'ajout d'un terme de mortalité endogène au niveau de la dynamique bactérienne [GRI05].

Une caractéristique particulière du procédé considéré - à différence de la plupart des applications de synthèse de produits exploitées en bioindustrie qui utilisent des populations microbiennes composées d'une seule espèce de micro-organismes - est la présence d'une culture mixte, comme c'est le cas de la plupart des procédés d'épuration devant traiter des effluents septiques situés en environnements naturels. Les micro-organismes s'y développent généralement en présence de faibles concentrations en divers éléments nutritifs servant à leurs fonctions métaboliques comme source carbonée ou azotée. Ces écosystèmes naturels sont hétérogènes et les conditions de croissance peuvent changer constamment au cours du temps du fait de compétitions permanentes entre les différentes espèces de bactéries pour la digestion de substrats multiples. Ainsi, en traitement des eaux usées, les cultures mettant en jeu un seul type de micro-organisme et limitées par un seul type de substrat sont très rares. Ces cultures sont en fait caractérisées par une grande diversité de micro-organismes ayant chacun une affinité différente pour de multiples substrats. De ce fait, la modélisation et la maîtrise des systèmes multi-organismes sont plus difficiles, motivant les recherches fondamentales et appliquées [BEN96].

Alors qu'en culture pure toute apparition de bactéries étrangères au procédé est considérée comme agent contaminant, paradoxalement, les interactions compétitives au sein de cultures mixtes conduisent, dans la plupart des cas, à des systèmes plus stables et plus flexibles. Il a été démontré aussi pour certaines applications que de meilleurs rendements de production pouvaient être obtenus en choisissant des cultures mixtes [LIN82]. Les avancées réalisées en matière d'écologie microbienne contribuent à une meilleure compréhension de la stabilité et des interactions au sein des cultures mixtes. Le message le plus intéressant résultant de ces travaux est que les performances obtenues par les cultures mixtes sont plus qu'une simple addition des potentiels d'activité de chaque micro-organisme. En fait, les mécanismes qui

³⁸ Waste-water Benchmark, in Fuzzy Algorithms for the Control of Multiple-Input, Multiple Output Processes (FAMIMO) project, funded by the European Commission (Esprit LTR 21911).

gouvernent la vie d'une population mixte résultent également d'un large spectre d'interactions microbiennes, notamment, la compétition et la sélection naturelle de micro-organismes sous limitation de substrats nutritifs.

2.5.1. Traitement biologique des effluents de papeterie par lagunage aérée

Principe du lagunage aéré

Le lagunage naturel est certainement la plus ancienne parmi les techniques d'épuration d'eaux usées. Elle consiste à déverser l'effluent dans un bassin aéré naturellement (photosynthèse et aération de surface) où séjourne une population mixte destinée à digérer la matière polluante (cf. section 1.2.4). Afin d'intensifier le traitement et de le rendre plus performant, la lagune est équipée d'aérateurs permettant une meilleure oxygénation des eaux. Cette technique, dite lagunage aéré, est de tradition pour un certain nombre d'industries (secteurs de l'agro-alimentaire, de la santé, du papier, de la chimie fine, etc.). En comparaison avec les procédés dits à boues activées ou à cultures fixées, le lagunage aéré possède l'avantage considérable d'être économiquement plus intéressant, notamment sur les plans investissement et fonctionnement. Il existe différentes configurations d'installation de lagunage aéré suivant le type d'effluents à épurer (bassin unique ou association de bassins).

Caractéristiques des effluents de papeterie

La production industrielle de pâtes à papier génère une quantité variée de résidus gazeux émis dans l'air ambiant et des effluents chargés en divers résidus sous forme liquide et solide. Les impuretés drainées par les eaux usées proviennent de la matière première (bois) et des procédés chimiques utilisés pour la fabrication de la pâte à papier (produits réactifs). Après différents prétraitements physico-chimiques (cuisson, blanchiment), l'effluent industriel est traité dans une station d'épuration. Trois exemples de composition d'effluent d'industrie papetière (Compagnie Saint Gobain, France) sont répertoriés sur le tableau suivant [ROL95] :

Processus	Sulfite d'ammonium	Kraft	Fibres recyclées
Produit	Pâte Fluff	Papier Kraft	Carton ondulé
Production (t/j)	400	1 200	420
Débit de l'effluent (m ³ /j)	24000	60000	4300
Charge en DCO (t/j)	80	25	15
Charge en DBO ₅ (t/j)	27	10	7
Principaux constituants	lignosulfonate	lignine	amidon
	acides résiniques	hémicellulose	cellulose
	cellulose, hémicellulose	méthanol	acides gras volatiles
	acides gras volatiles		

Tableau 2.1. Caractéristiques d'effluents d'usines de papier et de pâte à papier

Bien que la composition de la plupart des effluents d'industries papetières soit assez complexe (Cf. Tableau 2.1), il est cependant possible de définir deux catégories de fraction de

la DBO₅ en fonction de leur niveau de forte/faible biodégradabilité. Cette distinction correspond à la vitesse rapide/lente à laquelle ces fractions sont dégradées. Dans les deux cas, les composés sont biodégradables. Le premier type de substrat, dit *énergétique*, correspond à de la cellulose. Il est caractérisé par une forte biodégradabilité et constitue la principale source énergétique de croissance des micro-organismes. Le second type de substrat, dit *xénobiotique* (synthèse issue de l'activité humaine), est caractérisé par une faible biodégradabilité. Principalement constitué d'un mélange de dérivés ligneux et de lignosulfonates, il est considéré comme le principal substrat polluant à dégrader par voie microbienne.

Le principe du lagunage aéré peut être appliqué efficacement à ce type d'effluent si le comportement hydrodynamique général de la lagune est correctement adapté aux cinétiques de la réaction biologique. Dans ce contexte, l'étude de l'optimisation de la biodégradation des résidus carbonés peut être menée en utilisant le concept des bioréacteurs alimentés en mode continu. L'amélioration des performances de biodégradation dépend alors du degré d'adéquation existant entre le type de polluant à traiter, les micro-organismes et les caractéristiques du milieu environnemental (température, pression d'oxygène dissous, pH).

Bioréacteur continu pour la dégradation de composés xénobiotiques

L'étude de caractérisation du procédé a nécessité le montage d'une unité pilote permettant de reproduire le plus fidèlement possible les conditions opératoires et la configuration géométrique d'une lagune aérée pour le traitement des eaux usées d'industrie papetière. L'équipe de génie microbiologique du département de Génie Biochimique et Alimentaire de l'INSA³⁹ a mis en œuvre à l'échelle du laboratoire un bioréacteur contenant une population mixte naturelle (prélevée sur site) et alimenté par un effluent comportant deux substrats carbonés. Ce type d'effluent, dit *bi-substrat*, est composé de cellulose et de lignosulfonate. Le pilote est constitué d'un bioréacteur homogène alimenté en mode continu par une solution contenant un mélange bi-substrat (Cf. Figure 2.8). L'équipement de conduite du bioréacteur est constitué d'une unité Sétric, interface réalisant l'acquisition et/ou la régulation des variables d'environnement. La température est mesurée par une thermosonde à semi-conducteur, le pH par une sonde pH et l'oxygène dissous par une sonde O₂. Un analyseur automatisé pour la mesure en ligne de la concentration en substrat énergétique et de la concentration en biomasse est associé au bioréacteur. L'aération du réacteur est assurée par un dispositif composé d'un débitmètre massique et d'un filtre stérilisable [BEN96].

La culture microbienne mixte prélevée sur la lagune aérée permet la dégradation biologique dans des conditions opératoires déterminées par les valeurs des variables d'environnement du Tableau 2.2 :

Variable	Valeur
Température (°C)	30
pH	7
Volume (L)	1,5
Agitation (tpm)	[400 ;800]
Aération (vvm) ⁴⁰	[0,5 ;1]

Tableau 2.2. Valeurs des variables d'environnement

³⁹ INSA : Institut National des Sciences Appliquées de Toulouse.

⁴⁰ vvm : volume d'air/volume de liquide/minute.

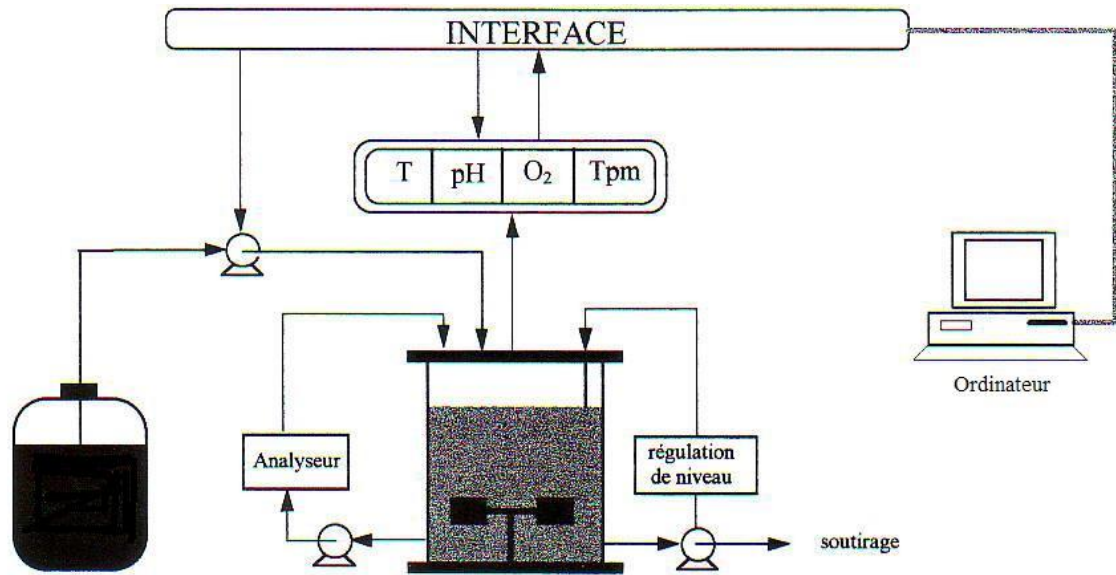


Figure 2.8. Schéma synoptique de l'installation du procédé pilote

2.5.2. Modélisation des cinétiques de croissance et de dégradation

Le comportement des populations mixtes dans un mélange multi-substrat dépend des affinités de chaque micro-organisme vis-à-vis de chaque substrat. Les cinétiques d'utilisation des substrats sont gouvernées par les concentrations et les types de substrat présents dans le milieu mais également par les interactions existant entre les différentes espèces microbiennes. En ce qui concerne le procédé d'épuration des effluents de papeterie étudié, il est indispensable d'intégrer dans la modélisation les phénomènes de compétition et d'activation/inhibition microbiennes.

Compétition microbienne pour le mélange bi-substrat

La compétition reflète un conflit entre les différentes espèces de micro-organismes pour la consommation des éléments nutritifs et leur prolifération dans l'espace vital disponible. Le phénomène de compétition est lié à une limitation de croissance par le substrat. La modélisation du taux spécifique de croissance de la population mixte et des taux spécifiques de dégradation en substrats énergétique et xénobiotique est une étape d'autant plus complexe que la croissance s'effectue dans un mélange multi-substrat. La théorie de Monod, aboutissant à des modèles donnés par l'équation (2.3), a été étendue pour inclure les cas où plusieurs substrats en concentration limitante sont présents lors de la croissance d'un seul type de micro-organisme. Trois hypothèses différentes ont été proposées pour décrire l'effet d'une culture mixte sur le taux spécifique de croissance.

1. Certains types de culture traduisent une limitation de croissance par un seul des deux substrats suivant un point de fonctionnement précis : la limitation est imposée par le substrat dont la vitesse de dégradation est la plus faible [RYD72] [SYK73]. Ainsi, pour une croissance en présence de deux substrats $S_1(t)$ et $S_2(t)$, le taux de croissance cellulaire $\mu_X(t)$ résultant est modélisé par :

$$\mu_X(t) = \begin{cases} \mu_{S_1}(t) & \text{si } \mu_{S_1}(t) < \mu_{S_2}(t) \\ \mu_{S_2}(t) & \text{si } \mu_{S_1}(t) > \mu_{S_2}(t) \end{cases} \quad (2.44)$$

où $\mu_{S_1}(t)$ et $\mu_{S_2}(t)$ sont décrits par le modèle de Monod (2.3).

2. Megee et coll. [MEG72] proposent, au contraire, un modèle bi-substrat dans lequel les deux substrats affectent simultanément le taux spécifique de croissance. Le modèle correspondant est donné par la relation :

$$\mu_x(t) = \mu_{s1}(t) \cdot \mu_{s2}(t) \quad (2.45)$$

ce qui implique une croissance nulle en l'absence d'un des deux substrats.

3. Enfin, dans d'autres applications [YOO77] [GOT82], le taux spécifique de croissance de la biomasse est défini comme étant la somme des deux taux de croissance liés à chaque type de substrat :

$$\mu_x(t) = \mu_{s1}(t) + \mu_{s2}(t) \quad (2.46)$$

Ceci traduit le fait qu'en l'absence d'un des deux substrats limitants, par exemple $S_1(t)$, la biomasse continue à croître selon $\mu_{s2}(t)$, c'est à dire en utilisant le second substrat $S_2(t)$.

En étendant ce dernier modèle au cas des cultures mixtes, il est montré que la coexistence de diverses espèces en régime permanent dans une culture continue est possible lorsque plus d'un substrat est disponible. Ceci est vrai même si ces organismes ont une préférence marquée pour un de ces différents substrats. La diversité des sources de substrats dans un système réel est donc un facteur clé pour le maintien d'une population hétérogène. Lorsque l'évolution de cette population est suivie par la concentration en biomasse totale $X(t)$, le taux spécifique de croissance global (ou apparent) s'exprime alors par la somme des taux de croissance sur chaque substrat. Il semble que ce soit ce dernier type de mécanisme qui gère la croissance de la biomasse pour le procédé d'épuration des effluents de type bi-substrats [BEN96] ; la croissance de la population mixte en présence des substrats xénobiotique $S_x(t)$ et énergétique $S_e(t)$ ⁴¹ est donc du type :

$$\mu_x(t) = \mu_{sx}(t) + \mu_{se}(t) \quad (2.47)$$

Cependant, ce modèle ne décrit pas complètement les différentes situations de croissance des micro-organismes. En effet, les premières observations relatives aux cinétiques de dégradation des substrats de l'effluent de papeterie ont montré que, en présence de fortes concentrations en substrat énergétique $S_e(t)$, la vitesse de dégradation du substrat xénobiotique $S_x(t)$ est plus lente que celle qui devrait théoriquement être obtenue par une loi de Monod classique (équation (2.3)). Pour prendre en compte ces effets croisés liés à la présence, dans l'effluent, de substrats de nature très différente, Yoon et coll. [YOO77] ont fait appel au modèle de Monod généralisé. Cette théorie permet d'intégrer l'effet inhibiteur ou activateur de la présence d'un substrat sur la dégradation des autres substrats présents dans le milieu de culture microbienne.

Phénomène d'activation/inhibition compétitive dans le mélange bi-substrat

En culture pure, un excès de substrat se traduit par une inhibition de la croissance microbienne. Ce type d'interaction peut être représenté par le modèle de Haldane (Cf.

⁴¹ La notation originale utilisée par [BEN96] a été volontairement modifiée afin de la rendre plus conforme avec les pratiques internationales. Ainsi, les concentrations des substrats seront identifiés par S_i (où l'indice $i=\{e,x\}$ représente soit le substrat énergétique, soit le substrat xénobiotique) et la concentration de biomasse par X .

équation (2.4)) où l'inhibition est proportionnelle au carré du substrat limitant. Cependant, lorsque le milieu de culture comporte plusieurs substrats, ce modèle nécessite une modification permettant de tenir compte des effets croisés de chaque substrat sur l'assimilation des autres. Ainsi, dans le mélange bi-substrat, les effets croisés d'activation ou inhibition des substrats énergétique $S_e(t)$ et xénobiotique $S_x(t)$ sur l'évolution des taux spécifiques de croissance $\mu_{sx}(t)$ et $\mu_{se}(t)$ peut être correctement représentée par le modèle de Monod généralisé de la forme [YOO77] :

$$\mu_{sx}(t) = \mu_{sx}^{\max} \frac{S_x(t)}{K_{sx} + S_x(t) + a_{se}S_e(t)} \quad (2.48)$$

$$\mu_{se}(t) = \mu_{se}^{\max} \frac{S_e(t)}{K_{se} + S_e(t) + a_{sx}S_x(t)} \quad (2.49)$$

où $\mu_{sx}^{\max}$ et $\mu_{se}^{\max}$ correspondent aux taux spécifiques de croissance maximums. Les paramètres de Michaelis-Menten, K_{sx} et K_{se} , permettent de tenir compte de l'effet de limitation de croissance par chaque substrat. La constante a_{se} permet de modéliser l'effet inhibiteur du substrat énergétique $S_e(t)$ sur la consommation du substrat xénobiotique $S_x(t)$ (et réciproquement pour la constante a_{sx}). Théoriquement, ces deux paramètres doivent vérifier la relation $a_{se} = 1/a_{sx}$. Cependant, a_{sx} et a_{se} doivent être considérés comme des paramètres indépendants à ajuster en fonction des données expérimentales [YOO77]. Pour l'étude développée par [BEN96], la consommation préférentielle du substrat énergétique par les micro-organismes se traduit par une faible valeur de a_{sx} et par conséquent une forte valeur de a_{se} . Le terme *d'inhibition* est donc relatif à l'effet du substrat énergétique sur la dégradation du substrat xénobiotique.

En conclusion, en tenant compte de la structure de $\mu_{sx}(t)$ et $\mu_{se}(t)$ ainsi définie, nous obtenons l'expression du taux spécifique de croissance $\mu_X(t)$ de la population mixte :

$$\begin{aligned} \mu_X(t) &= \mu_{sx}(t) + \mu_{se}(t) \\ &= \mu_{sx}^{\max} \frac{S_x(t)}{K_{sx} + S_x(t) + a_{se}S_e(t)} + \mu_{se}^{\max} \frac{S_e(t)}{K_{se} + S_e(t) + a_{sx}S_x(t)} \end{aligned} \quad (2.50)$$

La Figure 2.9 permet la visualisation des phénomènes de compétition et d'activation/inhibition compétitive de la culture mixte en milieu bi-substrat. Il s'agit d'une simulation du procédé à taux de dilution nul (culture discontinue ou *batch*). Dans la première partie de cette courbe (zone 1), nous pouvons remarquer que la dégradation du substrat xénobiotique $S_x(t)$ est presque complètement inhibée par la seule présence du substrat énergétique $S_e(t)$. Lorsque le substrat $S_e(t)$ est pratiquement consommé (zone 2), la croissance de la population mixte s'effectue plus lentement (pente p_{sx} plus faible que p_{se}) mais en dégradant le substrat xénobiotique $S_x(t)$.

Une conséquence du phénomène d'inhibition compétitive est qu'il serait possible d'augmenter la dégradation du substrat xénobiotique en l'absence de substrat énergétique. En dehors du fait que la présence du substrat énergétique dans l'effluent de la lagune aérée est imposée par la nature même du procédé de fabrication de la pâte à papier, celle-ci est

indispensable pour assurer une croissance suffisante des micro-organismes. En effet, pour qu'un procédé alimenté en mode continu soit stable, il faut qu'il y ait au moins un substrat qui puisse supporter la croissance des micro-organismes. Ce rôle de support est joué par le substrat énergétique qui, comme nous venons de le voir, est le substrat qui favorise le mieux la croissance de la population mixte à l'intérieur de la lagune. Ainsi, la conversion du substrat énergétique en biomasse permet de maintenir un taux spécifique de croissance au moins supérieur au taux de dilution du réacteur sans quoi celui-ci serait inexorablement lavé⁴². La diminution du taux de dilution de la lagune, et par conséquent de la charge de pollution, dans le but d'améliorer les vitesses de dégradation des deux substrats, n'est pas envisageable pour une station d'épuration dont l'objectif est d'intensifier le traitement des eaux usées.

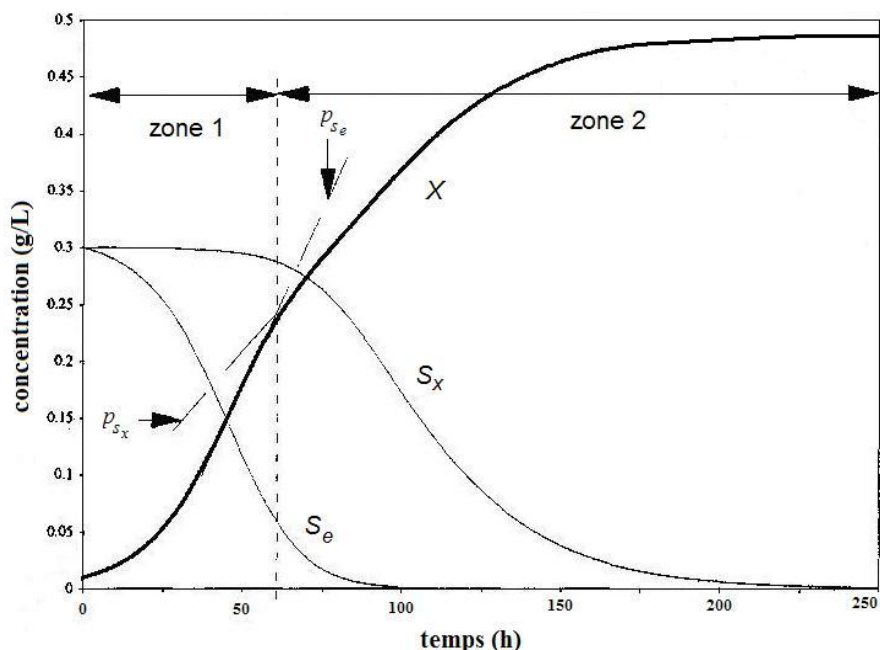


Figure 2.9. Les interactions microbiennes : compétition et activation/inhibition compétitive

2.5.3. Equations d'évolution du procédé

A partir des équations de bilans matière établis pour chaque élément macroscopique de la réaction biologique, les dynamiques du procédé de traitement des eaux résiduaires décrivant la croissance de la biomasse microbienne sur l'effluent bi-polluant considéré (substrats xénobiotique et énergétique), dans le cas d'un réacteur homogène à alimentation et soutirage permanent (donc de volume constant) sont modélisées par le système d'équations différentielles suivant :

$$\left\{ \begin{array}{l} \frac{dX(t)}{dt} = \mu_X(t)X(t) - D(t)X(t) \end{array} \right. \quad (2.51)$$

$$\left\{ \begin{array}{l} \frac{dS_x(t)}{dt} = -v_{Sx}(t)X(t) + D(t)S_{x-in}(t) - D(t)S_x(t) \end{array} \right. \quad (2.52)$$

$$\left\{ \begin{array}{l} \frac{dS_e(t)}{dt} = -v_{Se}(t)X(t) + D(t)S_{e-in}(t) - D(t)S_e(t) \end{array} \right. \quad (2.53)$$

⁴² Lavage : phénomène indésirable de perte de la biomasse qui se produit dans un bioréacteur lorsque le taux de dilution est supérieur au taux de croissance maximum.

où $X(t)$, $S_x(t)$ et $S_e(t)$ représentent respectivement les concentrations en (g/L) de biomasse, de substrat xénobiotique et de substrat énergétique. $\mu_x(t)$ représente la vitesse spécifique de croissance de la biomasse. $v_{sx}(t)$ et $v_{se}(t)$ représentent respectivement les vitesses spécifiques de dégradation du substrat xénobiotique et du substrat énergétique. Le taux de dilution $D(t)$ correspond au rapport entre le débit d'alimentation et le volume (constant) du bioréacteur. $S_{x-in}(t)$ et $S_{e-in}(t)$ sont respectivement les concentrations en substrat xénobiotique et énergétique du milieu d'alimentation. Généralement, les vitesses de dégradation en substrats sont exprimées en fonction des rendements de conversions du substrat xénobiotique en biomasse $Y_{X/Sx}$ et du substrat énergétique en biomasse $Y_{X/Se}$ par :

$$v_{sx}(t) = \frac{\mu_{sx}(t)}{Y_{X/Sx}} \quad (2.54)$$

$$v_{se}(t) = \frac{\mu_{se}(t)}{Y_{X/Se}} \quad (2.55)$$

Dans ces expressions, $\mu_{sx}(t)$ (respectivement $\mu_{se}(t)$) représente le taux spécifique de croissance des micro-organismes provenant de la conversion du substrat xénobiotique (respectivement énergétique) en biomasse. Ces paramètres dépendent des conditions opératoires (pH, température, aération) ainsi que des concentrations en substrats $S_x(t)$ et $S_e(t)$ (cf. équations (2.48) et (2.49)). La complexité de modélisation de ces variables est néanmoins réduite par le fait que leur expression est établie pour des conditions opératoires optimales fixes.

Nous proposons d'améliorer le modèle du procédé biologique utilisé par [BEN96] et qui a servi de benchmark dans le projet européen FAMIMO, en considérant la mortalité bactérienne due à la respiration endogène (Cf. section 2.4.1). Ce phénomène peut être représenté sous la forme d'un terme supplémentaire b (h^{-1}) dans l'équation (2.51) qui décrit la dynamique du comportement de la population bactérienne dans le bioréacteur :

$$\frac{dX(t)}{dt} = \mu_x(t)X(t) - D(t)X(t) - bX(t) \quad (2.56)$$

En reprenant les équations (2.52) et (2.53) pour la dynamique de dégradation des substrats xénobiotique et énergétique ainsi que l'équation (2.56) pour l'évolution de la biomasse, nous visualisons sur la Figure 2.10 les phénomènes de compétition, d'activation/inhibition compétitive et de mortalité bactérienne endogène.

Il s'agit de nouveau d'une simulation du procédé en mode d'opération discontinu et par conséquence à taux de dilution nul. Après la dégradation consécutive des substrats énergétique et xénobiotique avec une augmentation de la biomasse due à la biotransformation de matières organiques (zones 1 et 2), on remarque la réduction de la concentration de la biomasse bactérienne une fois que les substrats sont presque épuisés (zone 3). En effet, sans nutriments les bactéries doivent faire appel à une consommation des substances de réserves endogènes et puis à certains constituants cellulaires. S'il s'agissait du mode continu de fonctionnement il faudrait apporter en permanence une quantité minimale de nutriments, afin de promouvoir la bonne performance du bioprocédé tout en évitant les phénomènes de lavage du bioréacteur (quand le soutirage est supérieur à l'apport pour la croissance de la biomasse) et/ou de mortalité totale de la biomasse (quand l'apport de nutriments ne suffit pas à

compenser la mortalité microbienne). Nous avons retenu, pour notre étude de modélisation et de commande floues du bioprocédé, le modèle plus complet qui prend en compte ce phénomène de mortalité bactérienne endogène.

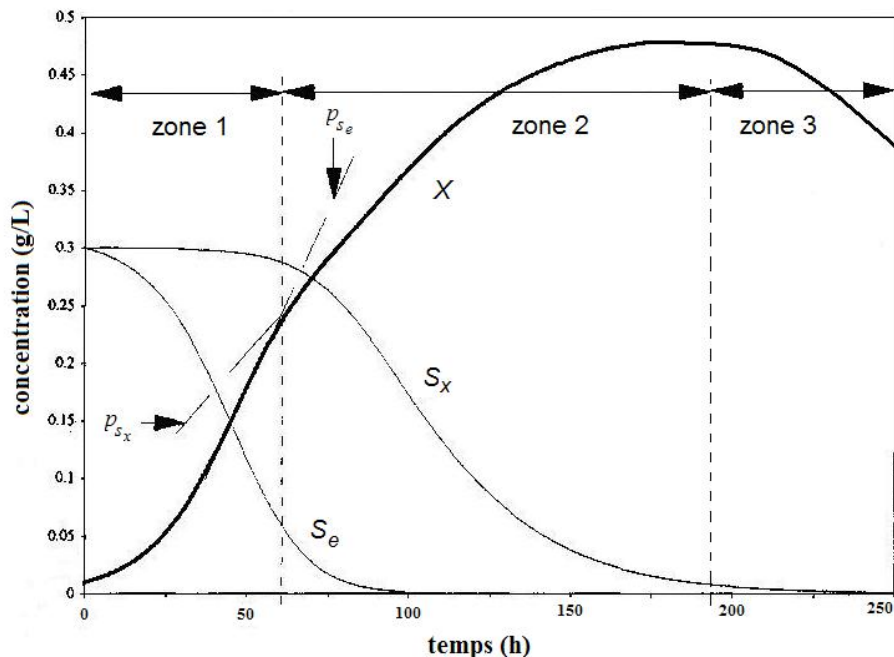


Figure 2.10. Phénomènes d'interactions microbiennes et mortalité bactérienne endogène

2.5.4. Variables de commande du procédé

En technologie d'épuration d'eaux usées, le premier objectif est évidemment de dégrader biologiquement un maximum de polluant tout en assurant de bonnes conditions de croissance des micro-organismes. La détermination des matières en suspension (MES) en sortie du réacteur permet d'estimer grossièrement la biomasse bactérienne, la fraction minérale étant relativement constante pour un type d'effluent. La surveillance des MES en sortie d'un réacteur permet d'évaluer l'excès de biomasse évacuée ou de détecter un éventuel lavage du réacteur dû par exemple à une surcharge accidentelle (toxicité, taux de dilution supérieur au taux de croissance maximum des micro-organismes). Pour le procédé de traitement des effluents de papeterie, l'excès de biomasse est considéré comme un facteur de pollution pour le milieu récepteur. Le double objectif de conduite de ce type de système est donc de maximiser la dégradation de substrat xénobiotique tout en contrôlant la concentration en biomasse à la sortie du procédé biologique.

Pour notre cas d'étude, l'entrée du système est définie en fonction des objectifs de commande mono-variable ou multi-variable que l'on désire se fixer. Dans le cas mono-variable la grandeur de commande correspond au taux de dilution $D(t)$ du réacteur. Nous nous sommes intéressés au cas multi-variable, où il est nécessaire de disposer d'une autre commande pour mener à bien la stratégie de conduite. Généralement, on prend la concentration en substrats xénobiotique et énergétique du milieu d'alimentation $S_{x-in}(t)$ et $S_{e-in}(t)$. En pratique, il suffit de dissocier l'alimentation classique en deux alimentations, l'une apportant de l'eau et l'autre le milieu contenant les deux sources de carbone aux concentrations maximales $S_{x-in}^{max}(t)$ et $S_{e-in}^{max}(t)$. Les relations qui lient $D(t)$, $S_{x-in}(t)$ et $S_{e-in}(t)$ aux variables de commande auxiliaires sont données par :

$$\left\{ \begin{array}{l} D(t) = u_1(t) + u_2(t) \\ S_{x-in}(t) = \frac{u_2(t)}{u_1(t) + u_2(t)} S_{x-in}^{\max} \\ S_{e-in}(t) = \frac{u_2(t)}{u_1(t) + u_2(t)} S_{e-in}^{\max} \end{array} \right. \quad \begin{array}{l} (2.57) \\ (2.58) \\ (2.59) \end{array}$$

Dans cette expression nous avons deux variables de commande : $u_1(t)$ qui correspond au taux de dilution en eau et $u_2(t)$ qui correspond au taux de dilution en substrats. Au niveau laboratoire, le bidon d'alimentation initial est remplacé par deux bidons, l'un contenant une solution d'eau et l'autre contenant une solution de substrat suivant le schéma d'alimentation représenté sur la Figure 2.11.

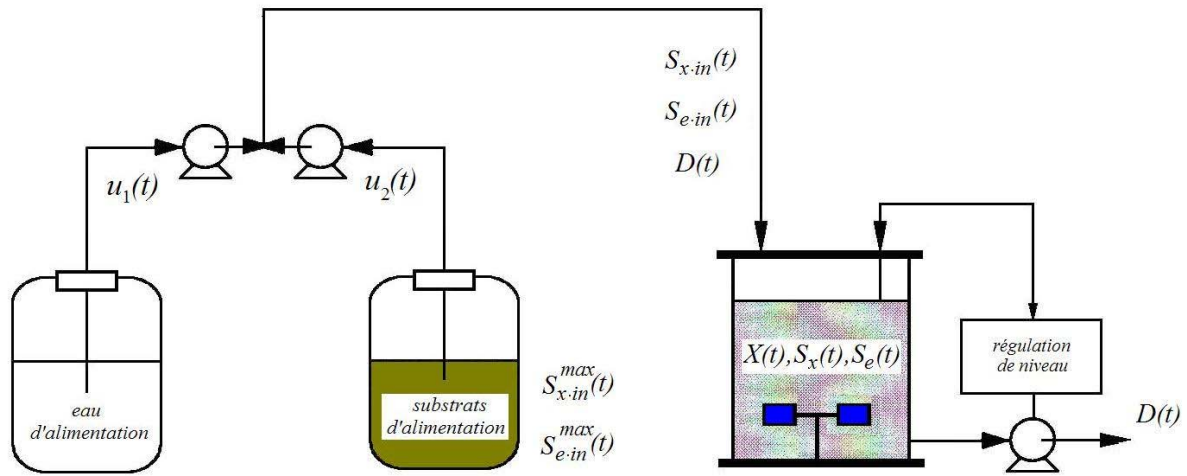


Figure 2.11. Alimentation du bioréacteur dans le cas multi-variable

En combinant les expressions de la dynamique du bioprocédé (2.56), (2.52) et (2.53) avec l'expression (2.50) du taux spécifique de croissance $\mu_X(t)$ de la population mixte et les équations (2.57) - (2.59) concernant les variables de commande $u_1(t)$ et $u_2(t)$, nous obtenons le modèle complet du procédé biologique de dépollution des eaux usées considéré :

$$\frac{dX(t)}{dt} = \left\{ \mu_{Sx}^{\max} \frac{S_x(t)}{K_{Sx} + S_x(t) + a_{Se} S_e(t)} + \mu_{Se}^{\max} \frac{S_e(t)}{K_{Se} + S_e(t) + a_{Sx} S_x(t)} \right\} X(t) - (u_1(t) + u_2(t))X(t) - bX(t) \quad (2.60)$$

$$\frac{dS_x(t)}{dt} = -\frac{\mu_{Sx}^{\max}}{Y_{X/Sx}} \frac{S_x(t)}{K_{Sx} + S_x(t) + a_{Se} S_e(t)} X(t) - u_1(t)S_x(t) + (S_{x-in}^{\max} - S_{x-in}(t))u_2(t) \quad (2.61)$$

$$\frac{dS_e(t)}{dt} = -\frac{\mu_{Se}^{\max}}{Y_{X/Se}} \frac{S_e(t)}{K_{Se} + S_e(t) + a_{Sx} S_x(t)} X(t) - u_1(t)S_e(t) + (S_{e-in}^{\max} - S_{e-in}(t))u_2(t) \quad (2.62)$$

Nous reprendrons ce modèle dans le chapitre 5, qui concerne l'application de la modélisation floue de type Takagi-Sugeno (TS) ainsi que la commande floue TS basée sur le modèle flou du bioprocédé.

2.6. Conclusions

Comme nous l'avons souligné, les phénomènes mis en jeu dans les procédés biologiques possèdent une dimension particulière par rapport aux exemples habituels de modélisation car ils relèvent du vivant. Le développement et la disponibilité d'un modèle mathématique décrivant les différents phénomènes biochimiques sur lesquels se basent les procédés biologiques d'épuration est une tâche délicate nécessitant des connaissances issues de domaines divers allant de la microbiologie aux mathématiques et à l'informatique. Il est bien connu que les processus métaboliques qui ont lieu dans les unités biologiques de dépollution, sont très nombreux et complexes. En effet, contrairement à la physique où il existe des lois connues depuis des siècles, la plupart des modèles en biologie reposent sur des lois empiriques. Toutefois, les recherches menées dans ce domaine ont cependant permis de mettre en évidence certains mécanismes biochimiques, du moins les plus influents. Loin d'être exhaustif, ce chapitre avait simplement pour objectif de présenter une introduction à la modélisation dynamique des procédés biologiques de traitement des eaux résiduaires en utilisant l'approche des bilans de matière.

Nous avons vu dans ce chapitre que, quel que soit le procédé biologique et son mode de fonctionnement (batch, fed-batch ou continu), le principe de modélisation, basé sur les bilans de matière, reste le même. Cette approche met en évidence, au moins du point de vue macroscopique (phénomènes globaux de bioconversion), le fonctionnement interne du système au cours duquel le bilan global des réactions est modélisé. Le comportement dynamique du bioprocédé est ainsi représenté par des équations différentielles non linéaires qui expriment la variation au cours du temps de la quantité de matière des principaux composants du procédé (biomasse, substrats, produits). En général deux types de termes sont présents dans ces modèles, représentant respectivement la conversion et le transport. La conversion, qui correspond à la cinétique à laquelle se déroulent les diverses réactions biochimiques conduisant à la conversion des divers substrats en biomasse (et produits) et la dynamique de transport, qui regroupe le transit de la matière au sein du procédé. En utilisant l'approche de modélisation basée sur des bilans de matière, nous avons décrit quelques procédés biologiques représentatifs tels que les fermentations continues (e.g. croissance d'une population des micro-organismes dégradant un substrat limitant, nitrification, digestion anaérobie), les fermentations semi-continues et les procédés à boues activées. La complexité de ces modèles dépend essentiellement du nombre de réactions internes considérées et de la structure choisie pour décrire le taux de croissance microbienne.

Finalement, nous avons décrit le procédé considéré au niveau de l'application dans notre étude, qui est un bioréacteur aérobie multivariable en mode continu pour le traitement des eaux usées issues d'une industrie papetière. Ce procédé est alimenté par un mélange bi-polluant (substrats énergétique et xénobiotique), dans lequel les micro-organismes exhibent des interactions telles que la compétition et l'activation/inhibition compétitive. Le modèle que nous avons retenu, développé sous Matlab®, est une version améliorée de celui utilisé lors du projet européen FAMIMO, pour lequel nous avons ajouté un terme de mortalité endogène bactérienne. Comme nous l'avons indiqué précédemment, ce modèle de connaissance nous servira de base pour comparer en simulation la performance d'un modèle flou pour représenter le comportement dynamique non-linéaire entrée-sortie du système.

Nous abordons maintenant la deuxième partie des travaux, concernant la modélisation, l'identification et la commande floues de type Takagi-Sugeno. Nous développerons l'application de ces techniques sur notre bioprocédé d'épuration dans le chapitre 5.

**PART II. IDENTIFICATION - COMMANDE FLOUES
ET APPLICATION AU BIOPROCÉDÉ DE
TRAITEMENT DES EAUX USÉES**

Chapitre 3

Identification de modèles flous Takagi-Sugeno (TS) à partir de données

Dans ce chapitre, nous traitons du problème d'identification (construction) des modèles flous à partir de données entrée-sortie. Nous introduisons d'abord la modélisation floue de systèmes, en nous focalisant particulièrement sur le modèle de type Takagi-Sugeno. En utilisant ce formalisme, nous représentons le comportement non linéaire d'un système par une composition de règles du type « *Si-Alors* », concaténant un ensemble de sous-modèles localement linéaires. Pour la construction de tels modèles nous abordons d'une part, différentes méthodes de classification floue et, d'autre part, nous proposons l'application d'une méthodologie d'« agglomération compétitive », robuste en présence de bruit : il s'agit des algorithmes FER⁴³ et RoFER⁴⁴. Les méthodes considérées appartiennent aux méthodes de coalescence floue (*clustering*) basées sur la minimisation d'une fonction objectif. Enfin après avoir considéré la méthodologie générale pour la construction de modèles flous Takagi-Sugeno à partir de données, nous présentons quelques exemples illustratifs d'application.

⁴³ FER : *Fuzzy Ellipsoidal Regression*.

⁴⁴ RoFER : *Robust Fuzzy Ellipsoidal Regression*.

Sommaire

3. Identification de modèles flous Takagi-Sugeno (TS) à partir de données.....	73
3.1. Modélisation floue de systèmes.....	76
3.1.1. Structure générale et différents types de modèles flous	76
3.1.2. Le modèle flou de type Takagi-Sugeno.....	84
3.2. Identification (Construction) de modèles flous	88
3.2.1. Structure du modèle flou.....	89
3.2.2. Le problème d'identification et ses solutions.....	90
3.2.3. L'identification basée sur des données entrée-sortie.....	91
3.3. Méthodes de coalescence floue (clustering flou).....	93
3.3.1. Notation et Concepts de base.....	93
3.3.2. Groupage <i>c</i> -moyennes floues (algorithme FCM).....	98
3.3.3. Groupage avec matrice de covariance floue (algorithme GK).....	101
3.3.4. Groupage avec prototypes linéaires (algorithme FCRM).....	103
3.3.5. Groupage avec régression ellipsoïdale (algorithme FER)	105
3.3.6. Groupage robuste en présence de bruit (algorithme RoFER).....	110
3.4. Construction de modèles TS à partir de données : méthodologie générale.....	118
3.4.1. Validation du nombre de clusters	122
3.4.2. Génération des fonctions d'appartenance des antécédents	124
3.4.3. Obtention des paramètres des conséquents.....	126
3.4.4. Validation numérique du modèle flou	128
3.4.5. Outil logiciel pour la modélisation-identification floue de type TS.....	129
3.4.6. Exemples.....	130
3.5. Conclusions	140

Le développement de modèles mathématiques des systèmes est un sujet central dans plusieurs disciplines des sciences et de l'ingénierie. Traditionnellement, la modélisation est vue comme la double conjonction entre la compréhension de la nature et du comportement d'un système ainsi que le traitement mathématique approprié qui conduise à l'obtention d'un modèle utilisable. Néanmoins, le besoin d'une forte compréhension des éléments physiques de base - comme c'est le cas dans l'approche de type boîte blanche - constitue une grande restriction au niveau pratique quand on est confronté aux systèmes complexes et/ou pauvrement compris. Il est à remarquer que les procédés chimiques et biologiques font souvent partie de ces systèmes. D'un autre côté, dans l'approche de type boîte noire, il est supposé que le procédé étudié peut être représenté en utilisant une structure générale qui approxime le comportement du système. Dans ce cas, le problème de modélisation consiste alors à proposer la structure appropriée pour l'approximateur et d'estimer les paramètres du modèle, en utilisant habituellement des données entrée-sortie représentatives du comportement, afin de capturer correctement les dynamiques et la non linéarité du système.

Le fait que les humains puissent souvent gérer différents types de situations complexes en utilisant des informations dont beaucoup sont subjectives et imprécises a stimulé la recherche de paradigmes alternatifs de modélisation et de commande. Ainsi, le concept de la modélisation dite « floue » a trouvé ses origines dans la théorie des ensembles flous proposée en 1965 par Zadeh [ZAD65], comme une manière de traiter l'incertitude, fondée sur l'idée de définir des ensembles pouvant contenir des éléments de façon graduelle. Cette théorie a introduit une manière de formaliser les méthodes humaines de raisonnement en utilisant des bases de règles et variables linguistiques pour la représentation de connaissances [ZAD73]. Comme nous l'avons déjà souligné dans le chapitre 1, bien que la commande floue ait été une des grandes applications de la logique floue, la plupart des applications développées dans les années 80-90 avaient pour cadre une démarche « basée connaissance » reposant sur l'expertise d'un opérateur pour un problème de commande donné, de complexité limitée (contrôleurs dits de type Mamdani). Lorsqu'on a voulu passer à des problèmes plus complexes, il était difficile d'écrire (même pour un expert) des bases de règles volumineuses et l'approche basée connaissance n'était plus appropriée. Pour pallier la croissance des règles, des approches tels que la fusion des variables et la structuration de la commande ont été proposées [LAC03].

Dans le but de faire face à un problème de commande complexe, on peut tirer profit de différents types de connaissance [BOR03] :

- de l'expertise sur certains points limités,
- de modèles locaux à validité restreinte,
- de nombreuses données entrées-sorties sur le procédé.

Pour prendre en compte toutes ces connaissances, l'approche basée connaissance/expertise, qui consistait à écrire directement un contrôleur, n'était plus adaptée et l'on s'est tourné vers une approche plus générale, rejoignant la démarche classique de l'automatique qui est une démarche « basée modèle », grâce notamment aux idées de modélisation floue de Sugeno [SUG88]. Cette démarche amène à manipuler des modèles dits flous pour faire la synthèse de commandes dites floues - même si l'aspect flou est de moins en moins présent dans ces nouveaux formalismes, rejoignant des formalismes utilisés par ailleurs, hors de la « communauté du flou », par exemple ceux de commandes supervisées, de systèmes linéaires variant avec le temps, de représentations polytopiques, etc. Comme on le verra, cette approche de conception de systèmes flous basée modèle est en général plus systématique, du point de vue des performances numériques, que celle basée

connaissance/expertise, malgré le fait que plusieurs problèmes théoriques restent encore ouverts [SAL05].

Dans ce chapitre nous présentons les éléments de base concernant l'identification de modèles flous à partir des données entrée-sortie du procédé. Nous nous intéressons particulièrement aux modèles de type Takagi-Sugeno et à leur construction en utilisant des techniques de coalescence floue (clustering). Une grande partie de ce chapitre est basée sur les ouvrages [BAB98a] [FOU03a] [FOU03b]. Pour une étude plus vaste concernant les divers aspects liés au domaine de la logique floue appliquée à la modélisation et à l'identification des systèmes, nous suggérons les documents suivants : [YAG94] qui présente une introduction complète et bien développée au sujet de la modélisation et de la commande floues classiques ainsi que ses applications ; [MUR97] qui aborde tout un panorama de paradigmes existants, basées sur l'approche multi-modèle, pour leur application en identification, modélisation et commande non linéaires ; [HEL97] qui traite différentes approches pour l'identification de modèles flous tels que les méthodes de clustering flou, les réseaux de neurones et les algorithmes génétiques ; [NGU98] qui considère divers aspects sur l'utilisation de la logique floue pour la commande de systèmes (techniques d'apprentissage et de réglage) ainsi que des méthodes de synthèse et des relations avec la commande linéaire.

3.1. Modélisation floue de systèmes

Les modèles flous peuvent être considérés comme des modèles logiques qui utilisent des règles du type « Si...Alors... » pour établir des relations qualitatives entre les variables du modèle. Ils reposent sur une dualité linguistique/numérique dans laquelle les ensembles flous servent d'interfaces entre les variables qualitatives (linguistiques) impliquées dans les règles et les valeurs numériques présentes dans les données entrées/sorties du modèle. La nature des modèles flous, basée sur des règles, permet l'utilisation d'information sous la forme des expressions du langage naturel, ce qui facilite la formalisation de la connaissance des experts ainsi que la transparence et l'interprétabilité du modèle. De plus, les modèles flous ont une structure mathématique flexible capable d'approximer un grand nombre de systèmes non linéaires complexes avec une bonne précision [KOS94] [WAN94] [ZEN95]. Les algorithmes d'apprentissage sont aussi utilisés pour la modélisation floue, ils peuvent même être combinés avec des techniques conventionnelles de régression [TAK85] [LIN94] [BAB98a].

Dans cette section nous introduisons d'abord d'une manière très succincte la structure générale ainsi que différents types de modèles flous. Ensuite nous nous focalisons en particulier sur une description du modèle flou de type Takagi-Sugeno. Enfin nous abordons le but et les différentes approches pour l'identification de modèles flous en mettant l'accent sur l'approche de construction de modèles à partir de données.

3.1.1. Structure générale et différents types de modèles flous

En général les systèmes flous s'appuient sur une représentation de la connaissance sous forme de règles « Si-Alors » qui permettent de représenter les relations entre les variables d'entrée et de sortie dont l'expression générique est de la forme :

Si antécédent **Alors** conséquent

Dans un premier temps et afin de faciliter l'interprétation, on peut considérer l'*antécédent* (*prémisse*) comme une description linguistique qui indique les *conditions* de validité du phénomène représenté. Pour sa part, le *conséquent* (*conclusion*) représente le *comportement* associé aux conditions de validité décrites par l'antécédent.

Considérons à titre illustratif la règle suivante :

Si concentration de polluant est élevée **Alors** temps de dégradation est long

Les règles floues établissent des relations logiques entre les variables du système en associant valeurs *qualitatives* d'une variable (la concentration de polluant est *élevée*) avec celles d'une autre variable (le temps de dégradation est *long*). Les valeurs qualitatives ont typiquement une interprétation linguistique, elles sont nommées *termes linguistiques* (étiquettes). La signification des termes linguistiques par rapport aux variables d'entrée/sortie numériques (concentration de polluant, temps de dégradation) est définie par des ensembles flous appropriés. Dans ce sens, les ensembles flous, ou plus précisément leurs fonctions d'appartenance, fournissent une interface entre les variables numériques d'entrée/sortie et les valeurs linguistiques qualitatives dans les règles.

Selon la structure particulière de la proposition conséquent, on peut distinguer trois types de modèles flous basés sur des règles [BAB98a] :

- *Modèle flou linguistique* (ou modèle Mamdani), dans lequel l'antécédent et le conséquent sont tout les deux des propositions floues qui utilisent des variables linguistiques [ZAD73] [MAM77].
- *Modèle flou relationnel*, qui peut être considéré comme une généralisation du modèle linguistique dans lequel il est possible d'associer une proposition antécédent spécifique avec plusieurs propositions conséquents différentes via une relation floue. Cette relation floue représente des associations entre les ensembles flous individuels définis dans les domaines d'entrée/sortie du modèle [PED84] [YI93].
- *Modèle flou Takagi-Sugeno* (TS), dans lequel le conséquent utilise des variables numériques plutôt que des variables linguistiques, sous la forme d'une constante, d'un polynôme ou de manière plus générale d'une fonction ou d'une équation différentielle dépendant des variables associées à la proposition antécédent [TAK85].

Dans le cadre des applications de la logique floue en commande des systèmes, les types de modèles les plus souvent utilisés sont ceux de Mamdani et de Takagi-Sugeno (noté TS). Les exemples suivants correspondent à ces deux types de règles :

Si polluant est élevée et biomasse est faible **Alors** temps de dégradation est long

Si polluant est élevée et biomasse est faible **Alors** $t = 30P + 50B + 10$

La première règle est celle d'un modèle de type Mamdani utilisant les variables linguistiques associées aux concentrations de polluant, biomasse et le temps de dégradation tandis que la seconde est celle d'un modèle de type Takagi-Sugeno où la variable t du conséquent (temps de dégradation) est numérique sous la forme d'une fonction des variables associées à l'antécédent. Ici P et B représentent respectivement les valeurs numériques des concentrations de polluant et de biomasse.

Dans un cadre plus formel, la représentation d'un modèle flou demande souvent l'utilisation de plusieurs règles et comporte diverses opérations mathématiques pour le traitement d'informations symboliques, hors de la portée de ce document. Dans la suite nous abordons d'une manière succincte les notions de base de la logique floue appliquée à la modélisation et à la commande. Nous supposons que le lecteur a une connaissance préalable dans ce domaine. Nous recommandons aux personnes intéressées de se référer, entre autres, aux ouvrages [DUB80] [BOU95] [YAG94] [BAB98a] pour un traitement plus approfondi.

Afin d'établir quelques éléments de base de la logique floue nécessaires dans le cadre de nos travaux, reprenons une règle linguistique R_i de la forme générale suivante :

$$R_i : \quad \text{Si } \mathbf{x} \text{ est } A_i \text{ Alors } \mathbf{y} \text{ est } B_i, \quad i = 1, \dots, r \quad (3.1)$$

où r dénote le nombre de règles du modèle. $\mathbf{x} \in X \subset \mathbb{R}^p$ est la variable de l'*antécédent*, qui représente l'entrée du système et $\mathbf{y} \in Y \subset \mathbb{R}^q$ est la variable du *conséquent*, qui représente la sortie du système flou. Habituellement la dimension de l'espace du conséquent est prise comme $q = 1$. X et Y correspondent respectivement aux domaines (*univers de discours*) des variables d'entrée et de sortie. A_i et B_i sont des *termes linguistiques* qui correspondent aux valeurs qualitatives associées aux variables de base $\mathbf{x}$ et $\mathbf{y}$. Ces termes linguistiques sont décrits par des ensembles flous définis par des *fonctions d'appartenance* μ , qui établissent une correspondance de chacun des univers de discours vers l'intervalle $[0, 1]$, ainsi :

$$\mu_{A_i}(\mathbf{x}) : X \rightarrow [0, 1], \quad \mu_{B_i}(\mathbf{y}) : Y \rightarrow [0, 1] \quad (3.2)$$

La théorie des ensembles flous [ZAD65] permet une appartenance partielle d'un élément à un ensemble. Prenons à titre d'exemple $\mu_{A_i}(\mathbf{x})$, qui symbolise la valeur de la fonction d'appartenance (nommé le *degré d'appartenance*) de $\mathbf{x}$ à l'ensemble caractérisé par A_i . Si le degré d'appartenance est égal à un, alors l'élément $\mathbf{x}$ appartient complètement à l'ensemble. Si celui-ci est égal à zéro, alors l'élément $\mathbf{x}$ n'appartient pas à cet ensemble. Si le degré d'appartenance se trouve entre zéro et un, alors $\mathbf{x}$ appartient partiellement à l'ensemble flou. Le degré d'appartenance permet alors une gradation du degré de constatation des faits.

Les ensembles flous A_i définissent des régions floue dans l'espace de l'antécédent pour lequel la prémisse du conséquent est valable. Le niveau de précision avec lequel peut être représenté un système donné par un modèle flou est associé avec la *granularité* du traitement de l'information, qui dépend entre autres, du nombre de termes linguistiques ainsi que de la forme particulière et du degré de chevauchement (*overlapping*) des fonctions d'appartenance utilisées dans le modèle.

Typiquement plusieurs formes de fonctions d'appartenance sont utilisées dans les applications. La Figure 3.1 montre un aperçu des fonctions d'appartenance les plus usuelles :

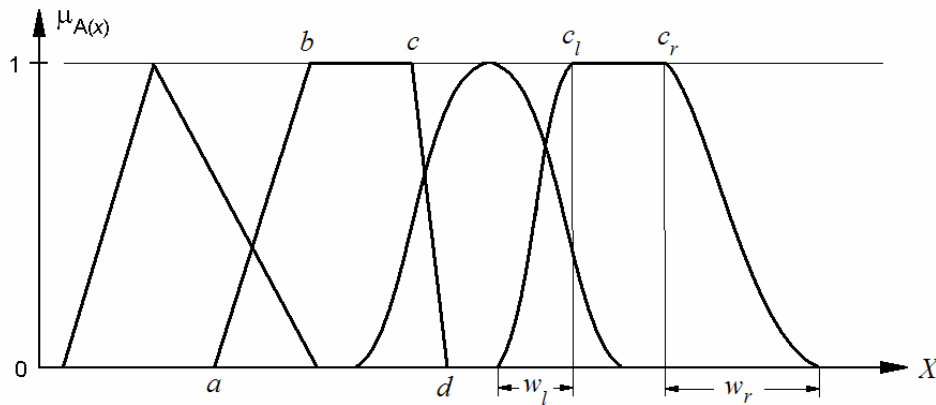


Figure 3.1. Formes typiques représentatives des fonctions d'appartenance

Les fonctions d'appartenance peuvent être définies par le développeur du modèle (l'expert) en utilisant la connaissance préalable ou bien l'expérimentation. Cette dernière approche est typique de la commande floue basée sur la connaissance [DRI93]. Dans ce cas,

les fonctions d'appartenance sont conçues de telle manière qu'elles représentent la signification des termes linguistiques dans le contexte donné. Quand des données entrée-sortie du procédé sont disponibles, il est alors possible d'appliquer des méthodes pour l'obtention ou l'adaptation des fonctions d'appartenance. Nous donnons un aperçu de cette approche dans la section 3.2.

Dans les univers de discours continus, les ensembles flous sont définis analytiquement par des fonctions d'appartenance. Afin de faciliter les calculs, elles sont souvent paramétrées [BAB98a]. Voyons deux cas assez généraux :

Fonction d'appartenance trapézoïdale :

$$\mu(x; a, b, c, d) = \max\left(0, \min\left(\frac{x-a}{b-a}, 1, \frac{d-x}{d-c}\right)\right) \quad (3.3)$$

où a, b, c, d sont les coordonnées des sommets du trapèze (voir Figure 3.1). Il faut remarquer que quand $b = c$ on obtient une fonction d'appartenance triangulaire.

Fonction d'appartenance exponentielle par morceaux :

$$\mu(x; c_l, c_r, w_l, w_r) = \begin{cases} \exp\left(-\left(\frac{x-c_l}{2w_l}\right)^2\right) & \text{si } x < c_l \\ \exp\left(-\left(\frac{x-c_r}{2w_r}\right)^2\right) & \text{si } x > c_r \\ 1 & \text{les autres cas} \end{cases} \quad (3.4)$$

où les paramètres $\{w_l, w_r, c_l \text{ et } c_r\}$ correspondent aux points définis sur la Figure 3.1. Pour les conditions $c_l = c_r$ et $w_l = w_r$, on obtient la fonction d'appartenance Gaussienne.

En revenant sur la règle linguistique générale décrite par l'expression (3.1), les systèmes flous basés sur des règles du type « Si...Alors » ont des antécédents et des conséquents qui sont spécifiés de manière symbolique. Dans le cadre de la modélisation et la commande de systèmes, l'exploitation d'une telle connaissance nécessite en général la mise en place d'interfaces numérique/symbolique (N/S) et symbolique/numérique (S/N). Ces dernières sont en effet des passerelles indispensables à l'établissement d'un lien entre l'ensemble de règles (*base de règles*) qui interfacent le système flou et le procédé, sur lequel seules des mesures et des actions numériques sont envisageables. De manière classique, le fonctionnement interne des systèmes flous repose sur une structure, représentée sur la Figure 3.2 [BAB98a] [GAL01], qui inclut :

- la *fuzzification* des variables d'entrée, avec éventuellement un prétraitement⁴⁵ de l'information ;
- l'*inférence* à partir d'une *base de connaissance*,
- la *défuzzification*, avec éventuellement un post-traitement d'information.

⁴⁵ Dans la plupart des applications d'ingénierie, les entrées et les sorties sont des valeurs numériques. Dans ce cas, les éventuels prétraitement et post-traitement d'information peuvent inclure des opérations telles que le filtrage dynamique et la normalisation/dénormalisation des données.

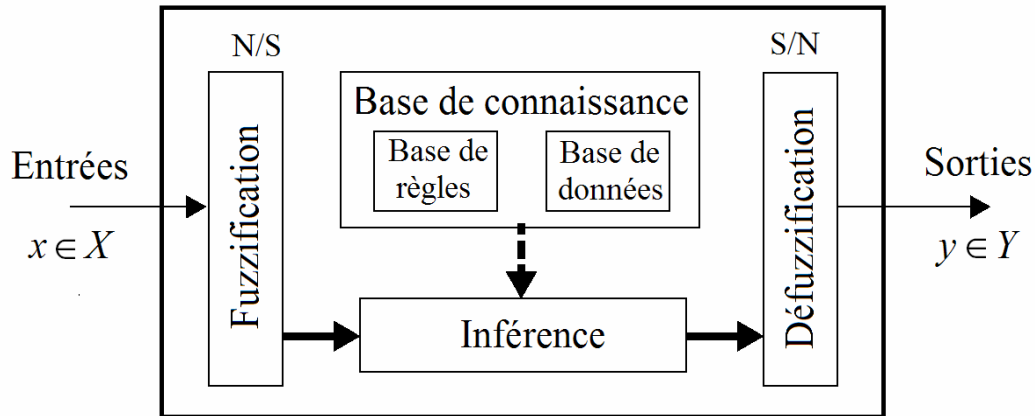


Figure 3.2. Représentation interne d'un système flou

L'étape de *fuzzification* consiste à transformer les entrées numériques disponibles en parties floues. Celles-ci alimentent alors le mécanisme d'*inférence* qui à partir d'une valeur d'entrée et selon la connaissance fournie par la *basse de connaissance*, détermine la valeur correspondante de la sortie. Cette base de connaissance est composée par la *base de règles* et par la *base de données*, qui stocke les fonctions d'appartenance associées aux termes linguistiques employés dans le système flou. Cela constitue le fondement du « *raisonnement approché* » du système, car la combinaison des entrées avec les règles floues permet de tirer des conclusions. Enfin, la *défuzzification* est l'opération inverse de la fuzzification, en convertissant les parties floues relatives aux sorties du mécanisme d'inférence en sorties numériques.

Remarque

R 3.1 Bien que nous ayons associé - à titre illustratif - une interprétation de conversion numérique/symbolique pour l'étape de fuzzification et inversement pour celle de défuzzification, il doit être souligné que des études approfondies [FOU95] ont montré qu'il s'agit d'un cas particulier. En effet, en considérant la nature numérique ou symbolique des données entrant et sortant du bloc d'inférence, il est possible de mettre en évidence différentes implantations possibles sous la forme d'une typologie des systèmes flous. Ainsi, l'implantation montrée sur la Figure 3.2 est dite linguistique, ce sont les blocs de fuzzification et de défuzzification qui réalisent les interfaces N/S et S/N. Elle permet de mener à bien l'inférence directement au niveau symbolique et autorise ainsi l'exploitation de règles linguistiques pondérées [GAL01].

Prémises de l'antécédent et composition de la base de règles

La proposition de l'antécédent peut contenir des ensembles flous définis directement dans le domaine du vecteur $\mathbf{x} \in X \subset \mathcal{R}^p$, comme dans le cas de la règle (3.1), ce qui correspond aux ensembles flous *multidimensionnels*. Néanmoins, plus habituellement les règles sont représentées sous sa forme *décomposée*, dans laquelle l'antécédent est défini comme une combinaison de propositions floues simples, définies sur les composants individuels x_i de $\mathbf{x}$. Les opérateurs logiques de conjonction (intersection), disjonction (union) et négation (complément) peuvent être utilisés afin de construire une proposition floue composée, par exemple :

$$R_i : \text{ Si } x_1 \text{ est } A_{i1} \text{ et } x_2 \text{ est } A_{i2} \text{ ou } x_3 \text{ n'est pas } A_{i3} \text{ Alors } y \text{ est } B_i \quad (3.5)$$

dans laquelle des ensembles flous *unidimensionnels* sont définis pour chaque composant du vecteur *antécédent*. Le degré d'accomplissement β_i (*degree of fulfillment*) de la règle est alors calculé en employant les opérateurs d'intersection, d'union et de complément appropriés, par exemple :

$$\beta_i = \mu_{A_{i1}}(x_1) \wedge \mu_{A_{i2}}(x_2) \vee (1 - \mu_{A_{i3}}(x_3)) \quad (3.6)$$

où l'opérateur *minimum* ($\wedge$) représente dans ce cas la conjonction (**et**), l'opérateur *maximum* ($\vee$) représente la disjonction (**ou**), et $1 - \mu$ est le complément (négation : **n'est pas**). Néanmoins, la généralisation des opérations d'union et d'intersection ensembliste classique aux sous-ensembles flous n'est pas définie de manière unique. Il existe, en fait, une multitude d'opérateurs de disjonction et de conjonction utilisés dans les propositions floues. Ces opérateurs sont regroupés en deux familles : les normes triangulaires, notées *t-normes*, qui définissent les opérateurs d'intersection ou de conjonction et les conormes triangulaires, notées *t-conormes*, qui définissent les opérateurs d'union ou de disjonction. De manière générale ces opérateurs sont des fonctions définies dans l'espace $[0, 1] \times [0, 1] \rightarrow [0, 1]$, mais elles ont des propriétés distinctes qui vont influencer le type de raisonnement approché qui va en découler [BOU95]. En prenant comme les opérands $x, y \in [0, 1]$, le Tableau 3.1 regroupe quelques définitions de *t-normes* ($t(x, y)$) et *t-conormes*⁴⁶ ($s(x, y)$) duales utilisées fréquemment [BAB98a] [TRA03]. Pour plus de détails sur les différents opérateurs le lecteur peut consulter, par exemple, [KLI95] [LEE90a] [LEE90b].

t-norme : $t(x, y)$ (intersection floue)	t-conorme : $s(x, y)$ (union floue)	Nom
$\min(x, y)$	$\max(x, y)$	min et max de Zadeh
$x \cdot y$	$x + y - x \cdot y$	Produit et somme algébrique (Bandler)
$\max(0, x + y - 1)$	$\min(1, x + y)$	Produit et somme bornés (Lukasiewicz)

Tableau 3.1. Principales t-normes et t-conormes

Les systèmes d'inférence floue (SIF) sont composés par une collection de plusieurs règles. Dans le cas des SIF avec une structure Mamdani conventionnelle (modèle linguistique), la forme la plus commune de ces règles est la forme dite *conjonctive*⁴⁷ donnée par :

$$R_i : \text{ Si } x_1 \text{ est } A_{i1} \text{ et } x_2 \text{ est } A_{i2} \text{ et } \dots \text{ et } x_p \text{ est } A_{ip} \text{ Alors } y \text{ est } B_i \quad (3.7)$$

Pour l'*i*-ème règle R_i du système, les p variables (x_1 à x_p) de l'antécédent sont associées aux ensembles flous (A_{i1} à A_{ip}) et le conséquent y est associé à un autre ensemble flou B_i . Les termes linguistiques représentés par A_i et B_i sont définis par des ensembles flous

⁴⁶ Dans la littérature anglaise les *t-conormes* sont souvent appelées aussi *s-normes*.

⁴⁷ La base de règles est supposée tacitement relié par l'opérateur **ou** (union) : R_1 ou R_2 ou ... ou R_r .

caractérisés par des fonctions d'appartenance $\mu_{A_i}(\mathbf{x}) : \mathfrak{R}^p \rightarrow [0, 1]$ et $\mu_{B_i}(y) : \mathfrak{R} \rightarrow [0, 1]$. Dans ce cas on peut voir le système flou comme une représentation qualitative qui décrit le comportement du système en utilisant une description linguistique.

Inférence et défuzzification dans le modèle linguistique

L'*inférence* dans les systèmes flous basés sur des règles, est le processus consistant à obtenir un ensemble flou de sortie étant donné une collection de règles et les entrées du système. Le mécanisme d'inférence dans le modèle linguistique est basé sur la *règle compositionnelle d'inférence* [ZAD73]. L'idée de base consiste à dériver un ensemble flou de sortie qui représente une conclusion globale issue de l'évaluation de l'ensemble des règles (*agrégation*) et du constat d'un certain fait d'entrée. Par simplicité, considérons ici une seule règle avec antécédent et conséquent scalaires de la forme "**Si** x est A **Alors** y est B " ainsi que le fait (x est A'), qui signifie que le sous-ensemble flou d'entrée A' est en un certain degré différent du sous-ensemble A de la prémisse. Le raisonnement mené dans le mécanisme d'inférence est basé sur le *modus ponens* généralisé, illustré de la manière suivante :

$$\begin{array}{ll} \text{Règle :} & \text{Si } x \text{ est } A \text{ Alors } y \text{ est } B \\ \text{Fait :} & x \text{ est } A' \\ \text{Conclusion :} & \frac{}{y \text{ est } B'} \end{array}$$

Le raisonnement approché est représenté dans ce cas par la proposition floue (y est B'), qui met en évidence que la conclusion B' dérivé est "moins certaine" que B , le sous-ensemble flou du conséquent.

Pour le cas des entrées numériques précises (*crisp*) et en utilisant la *t-norme*, le schéma de raisonnement peu être simplifié, aboutissant à celui bien connu dans la littérature appelé le max-min ou inférence de Mamdani [KL195]. Pour le cas de la règle composée définie par l'expression (3.7), l'inférence de Mamdani peut être résumée comme suit [BAB98a] :

Pas 1 : Pour chaque règle R_i du système ($i = 1, \dots, r$), calculer le degré d'accomplissement β_i de l'antécédent :

$$\beta_i = \mu_{A_{i1}}(x_1) \wedge \mu_{A_{i2}}(x_2) \wedge \dots \wedge \mu_{A_{ip}}(x_p) \quad (3.8)$$

Pas 2 : L'ensemble flou de sortie B'_i est dérivé pour chaque règle, en utilisant la *t-norme* du minimum :

$$\mu_{B'_i}(y) = \beta_i \wedge \mu_{B_i}(y), \quad \forall y \in Y \quad (3.9)$$

Pas 3 : L'ensemble flou de sortie agrégé B' est calculé en prenant le maximum (union) des conclusions individuelles B'_i :

$$\mu_{B'}(y) = \max_{i=1,2,\dots,r} \mu_{B'_i}(y), \quad \forall y \in Y \quad (3.10)$$

L'expression (3.32) peut être généralisée en introduisant des opérateurs quelconques de type *t-norme/t-conorme* (voir par exemple [DUB87] [BOU93]). Le sous-ensemble flou de sortie B' peut alors être obtenu par l'expression :

$$\mu_{B'}(y) = \bigwedge_{i=1,2,\dots,r} T(\mu_{A'}(x), I(\mu_{A'}(x), \mu_B(y))), \quad \forall y \in Y \quad (3.11)$$

où $\bigwedge$ représente une t -conorme, T est une t -norme et I représente l'implication floue. Cette définition générale tient compte de l'aspect graduel des sous-ensembles flous : ainsi la conclusion (y est B') est d'autant plus proche de la conclusion de la règle (y est B) que la donnée (x est A') est proche du prémisses de la règle (x est A).

Dans les SIF le traitement des informations fait intervenir les opérateurs logiques **et** et **ou** pour lier les variables floues d'entrée au niveau des conditions de chaque règle. Le terme **Alors**, introduisant la conclusion de chaque règle, est réalisée à partir de l'expression (3.11). L'implication floue (I) et la t -norme de cette expression permettent de former des fonctions d'appartenance partielles (par règle) qui sont ensuite combinés au cours de la phase d'agrégation de règles, pour fournir en sortie, la fonction d'appartenance globale résultante $\mu_{B'}(y)$ qui caractérise le sous-ensemble flou B' . Nous l'avons déjà souligné, il existe plusieurs méthodes de réalisation des opérateurs **et** ou **ou**, dont les noms sont associés au choix des méthodes d'agrégation des règles et d'implication. On utilise généralement l'une des trois méthodes d'inférence suivantes : *max-min* (Mamdani), *max-prod* et *somme-prod*, où le terme *somme* correspond au calcul de la valeur moyenne et *prod* au produit.

Comme on vient de le voir, le résultat de l'inférence floue est le sous-ensemble flou B' . Cet ensemble flou est caractérisé par la fonction d'appartenance globale résultante pour la variable de sortie $y \in Y \subset \mathfrak{R}$ définie sur un univers de discours, disons, $[a, b]$. Comme nous l'avons remarqué dans la Figure 3.2 pour le bloc de fuzzification, ici aussi il faudra envisager l'opération inverse permettant de transformer les grandeurs floues en grandeurs numériques pour les transmettre à l'extérieur du système. La *défuzzification* consiste alors en une transformation qui remplace le sous-ensemble flou de sortie de l'inférence par une valeur numérique unique représentative de cet ensemble. Il existe plusieurs techniques de défuzzification mais la plus utilisée est la méthode par centre de gravité. Il s'agit de déterminer la coordonnée en y du centre de gravité (noté COG⁴⁸) du sous-ensemble flou B' résultant de l'inférence :

$$cog(B') = \frac{\int_a^b y \cdot \mu_{B'}(y) dy}{\int_a^b \mu_{B'}(y) dy} \quad (3.12)$$

Le numérateur correspond au moment et, le dénominateur, à la surface. La défuzzification peut aussi être réalisée par d'autres méthodes, par exemple en prenant la valeur maximale de la fonction d'appartenance résultante ou bien la moyenne des abscisses du maximum (en anglais MOM⁴⁹). A partir d'une interprétation géométrique de cette dernière, la méthode MOM peut être utilisée par exemple dans des applications de prise de décisions pour privilégier la sortie la "plus plausible". La méthode du centre de gravité (COG) est utilisée avec l'inférence max-min de Mamdani, dans la mesure qu'elle fournit une interpolation proportionnelle à la taille des sous-ensembles flous individuels des conséquents. Cela est nécessaire, car la méthode d'inférence de Mamdani n'a pas d'effet interpolateur, et l'utilisation de la méthode MOM dans ce cas-ci, fournit une sortie abrupte avec des changements par morceaux [BAB98a].

⁴⁸ COG : *Center of Gravity*

⁴⁹ MOM : *Mean of Maxima*

3.1.2. Le modèle flou de type Takagi-Sugeno

Un autre type de modèle flou, approprié pour l'approximation d'une classe générale de systèmes non linéaires est celui proposé par Takagi et Sugeno [TAK85]. Ce type de modèle est, comme celui de Mamdani, construit à partir d'une base de règles "Si...Alors...", dans laquelle si la prémisse est toujours exprimée linguistiquement, le conséquent utilise des variables numériques plutôt que des variables linguistiques. Le conséquent peut s'exprimer par exemple, sous la forme d'une constante, d'un polynôme ou de manière plus générale d'une fonction ou d'une équation différentielle dépendant des variables associées à l'antécédent. D'une manière générale, un modèle de type Takagi-Sugeno (TS) est basé sur une collection des règles R_i du type :

$$R_i : \quad \text{Si } \mathbf{x} \text{ est } A_i \quad \text{Alors } y_i = f_i(\mathbf{x}), \quad i = 1, \dots, r \quad (3.13)$$

où R_i dénote la i -ème règle du modèle est r est le nombre de règles que contient la base de règles. $\mathbf{x} \in \mathcal{R}^p$ est la variable d'entrée (*antécédent*) et $y \in \mathcal{R}$ est la variable de sortie (*conséquent*). A_i est le sous-ensemble flou de l'antécédent de l' i -ème règle, définie, dans ce cas, par une fonction d'appartenance (*multivariable*) de la forme :

$$\mu_{A_i}(\mathbf{x}) : \mathcal{R}^p \rightarrow [0, 1] \quad (3.14)$$

Comme dans le modèle linguistique, la proposition de l'antécédent " $\mathbf{x}$ est A_i " est normalement exprimée comme une combinaison logique de propositions simples avec des sous-ensembles flous *unidimensionnels* définis pour les composants individuels du vecteur $\mathbf{x}$, usuellement dans la forme *conjonctive* suivante :

$$R_i : \quad \text{Si } x_1 \text{ est } A_{i1} \text{ et } x_2 \text{ est } A_{i2} \text{ et } \dots \text{ et } x_p \text{ est } A_{ip} \quad \text{Alors } y_i = f_i(\mathbf{x}), \quad i = 1, \dots, r \quad (3.15)$$

Typiquement les fonctions f_i sont choisies comme des fonctions paramétrées appropriées, avec la même structure pour chaque règle où seuls les paramètres varient. Une forme de paramétrisation souvent utilisée est la forme affine, donnée par :

$$y_i = \mathbf{a}_i^T \mathbf{x} + d_i \quad (3.16)$$

où $\mathbf{a}_i \in \mathcal{R}^p$ est un vecteur de paramètres et d_i est un scalaire. Ce modèle est appelé le *modèle affine Takagi-Sugeno*. Les conclusions des règles dans ce modèle sont alors des hyperplans (sous-espaces linéaires p -dimensionnels) dans l'espace $\mathcal{R}^{p+1}$. Ainsi, en modélisation floue des systèmes, l'antécédent de chaque règle définit une région (floue) de validité pour le sous-modèle correspondant du conséquent. Le modèle global est composé par la concaténation des modèles locaux (linéaires) et peut être vue comme une approximation par morceaux d'une surface non linéaire correspondant à la sortie du système.

Un cas particulier de la fonction du conséquent s'obtient quand $d_i = 0$, $i = 1, \dots, r$. Dans ce cas, le modèle est appelé modèle *homogène* Takagi-Sugeno [BAB98a], de la forme :

$$R_i : \quad \text{Si } \mathbf{x} \text{ est } A_i \quad \text{Alors } y_i = \mathbf{a}_i^T \mathbf{x}, \quad i = 1, \dots, r \quad (3.17)$$

Ce modèle a des propriétés d'approximation plus limitées que le modèle affine TS [FAN96] [ROV96]. Cependant, dans le cadre de la commande des systèmes, l'absence du terme scalaire d_i ⁵⁰ (*biais*) facilite la synthèse du contrôleur ainsi que l'analyse de la stabilité

⁵⁰ Terme de déplacement (En anglais : *offset* term ou *bias* term).

basé sur des systèmes homogènes TS parce que le modèle peut être traité comme un modèle quasi-linéaire [TAN92] [WAN95] [ZHA97].

Dans le cas où $\mathbf{a}_i = 0$, pour $i = 1, \dots, r$, les conséquents dans le modèle (3.16) prennent la forme d'une constante. Dans ce cas, le modèle, appelé *singleton*, est obtenu par :

$$R_i : \quad \text{Si } \mathbf{x} \text{ est } A_i \quad \text{Alors } y_i = d_i, \quad i = 1, \dots, r \quad (3.18)$$

Ce modèle peut aussi être vu comme un cas particulier du modèle flou linguistique, dans lequel les sous-ensembles flous des conséquents se réduisent aux singletons.

Inférence dans le modèle Takagi-Sugeno

Avant de pouvoir inférer la sortie, il faut calculer d'abord le degré d'accomplissement $\beta_i(\mathbf{x})$ de l'antécédent. Pour les règles avec des sous-ensembles flous multivariés dans l'antécédent (cf. expressions (3.13) et (3.14)), le degré d'accomplissement est simplement égal au degré d'appartenance de l'entrée multidimensionnelle $\mathbf{x}$; c'est-à-dire $\beta_i = \mu_{A_i}(\mathbf{x})$. Quand des connectives logiques sont utilisées, le degré d'accomplissement de l'antécédent est calculé comme une combinaison des degrés d'appartenance des propositions individuelles en utilisant les opérateurs de la logique floue.

Dans la modélisation Takagi-Sugeno, l'obtention de la sortie du modèle est réalisée à partir d'une combinaison des opérations d'inférence et de défuzzification [GAL01]. La sortie finale se calcule comme la moyenne des sorties correspondants aux règles R_i , pondérées par le degré d'accomplissement normalisé, selon l'expression [TAK85] :

$$y = \frac{\sum_{i=1}^r \beta_i(\mathbf{x}) \cdot y_i}{\sum_{i=1}^r \beta_i(\mathbf{x})} \quad (3.19)$$

En notant λ_i le degré normalisé d'accomplissement conformément à l'expression :

$$\lambda_i = \frac{\beta_i(\mathbf{x})}{\sum_{i=1}^r \beta_i(\mathbf{x})} \quad (3.20)$$

le modèle affine Takagi-Sugeno, avec une structure commune du conséquent, peut être exprimé comme un modèle pseudo-linéaire avec des paramètres dépendants des entrées :

$$y = \left(\sum_{i=1}^r \lambda_i(\mathbf{x}) \mathbf{a}_i^T \right) \mathbf{x} + \sum_{i=1}^r \lambda_i(\mathbf{x}) d_i = \mathbf{a}^T(\mathbf{x}) \mathbf{x} + d(\mathbf{x}) \quad (3.21)$$

Les paramètres $\mathbf{a}(\mathbf{x})$ et $d(\mathbf{x})$ sont des combinaisons linéaires convexes des paramètres des conséquents $\mathbf{a}_i$ et d_i , de la forme :

$$\mathbf{a}(\mathbf{x}) = \sum_{i=1}^r \lambda_i(\mathbf{x}) \mathbf{a}_i, \quad d(\mathbf{x}) = \sum_{i=1}^r \lambda_i(\mathbf{x}) d_i \quad (3.22)$$

De cette façon, le modèle TS peut être vu comme une correspondance (*mapping*) entre l'espace de l'antécédent (entrée) et une région convexe (polytope) dans l'espace des paramètres d'un système quasi-linéaire. Cette propriété a certains avantages qui facilitent l'analyse des systèmes flous TS dans le cadre des systèmes polytopiques [BOY94].

Cependant, cette représentation a aussi des propriétés indésirables d'interpolation qui affectent la capacité d'approximation des systèmes. En effet, une amélioration du mécanisme d'inférence du modèle TS a été proposée [BAB96a]. L'idée centrale consiste, à partir de l'analyse des propriétés du mécanisme d'inférence/interpolation conventionnel du modèle TS, à remplacer l'expression de la moyenne pondérée (3.19) afin d'adoucir l'interpolation.

Dans le cadre de l'identification des systèmes, un système flou approxime typiquement un modèle de régression dynamique non linéaire de la forme :

$$y(k+1) = f_{NL}(\mathbf{x}(k)) \quad (3.23)$$

où f_{NL} représente la fonction non linéaire et le vecteur de régression $\mathbf{x}(k)$ à l'instant k contient une collection des entrées u et des sorties y précédentes du système. Bien que les modèles de Mamdani et de Takagi-Sugeno soient les plus utilisés, dans le contexte des systèmes dynamiques le modèle TS l'est plus fréquemment⁵¹ car il est mieux adapté pour le traitement et l'exploitation numérique des systèmes à entrées/sorties multiples. Les règles du *modèle dynamique TS* prennent alors la forme générale suivante :

$$R_i : \quad \text{Si } \mathbf{z}(k) \text{ est } A_i \quad \text{Alors } y_i(k+1) = \theta_i^T \mathbf{w}(k), \quad i = 1, \dots, r \quad (3.24)$$

où les variables de l'antécédent $\mathbf{z}(k)$ et du conséquent $\mathbf{w}(k)$ sont usuellement choisies à partir du vecteur de régression $\mathbf{x}(k)$. Le modèle a deux jeux de paramètres, les paramètres θ_i des conséquents et les paramètres définissant les fonctions d'appartenance pour les sous-ensembles flous A_i . Ainsi, le modèle flou de type Takagi-Sugeno peut approximer les systèmes non linéaires avec des non linéarités douces et abruptes (à travers la structure des conséquents et la forme des fonctions d'appartenance). Il permet de faire un choix spécifique des variables de l'antécédent et du conséquent (e.g., systèmes Wiener et Hammerstein) [ABO03] et de modéliser des structures dynamiques qui changent en fonction de variables connues (systèmes commutés, modes de défaillance, etc.) [SAL05].

Dans le cadre de nos travaux, nous nous sommes particulièrement intéressés aux modèles dynamiques affines Takagi-Sugeno en temps discret sous la forme Non Linéaire Auto-Régressive avec entrée eXogène (NARX). Ce type de modèle établit une relation entre les valeurs précédents des entrées et sorties avec la sortie prédite $\hat{y}$, de la manière suivante :

$$\hat{y}(k+1) = f_{NL}(y(k), \dots, y(k-n_y+1), u(k), \dots, u(k-n_u+1)) \quad (3.25)$$

où k dénote l'instant d'échantillonnage, n_y et n_u sont des entiers liés à l'ordre (structure) du système. Une fois que la structure appropriée est établie, la fonction f_{NL} peut être approximée en utilisant une régression statique non linéaire, correspondant dans notre cas au modèle flou de type TS. Dans le modèle NARX, le vecteur de régression est composé par un certain nombre d'entrées et de sorties précédentes, $\mathbf{x} = [y(k), \dots, y(k-n_y+1), u(k), \dots, u(k-n_u+1)]^T$ et $\hat{y}(k+1)$ est la sortie prédite. Le cas des systèmes avec retard pur entre l'entrée et la sortie peut être directement incorporé sur le vecteur de régression sous la forme donnée par $\mathbf{x} = [y(k), \dots, y(k-n_y+1), u(k-n_d+1), \dots, u(k-n_d-n_u+2)]^T$, où n_d est une valeur entière du retard en nombre d'échantillons. Par simplicité, nous prenons par la suite $n_d = 1$.

⁵¹ En effet, les modèles flous de type TS sont souvent associés à une recherche de performance numérique alors que ceux de Mamdani sont orientés vers une possible prise en compte des connaissances expertes [GAL01]. Cependant, lorsqu'il s'agit de systèmes complexes avec un grand nombre de variables en entrée et sortie, le nombre de règles devient prohibitif. L'approche TS est alors mieux adaptée pour diminuer le nombre de règles.

Dans ce cas les règles du *modèle dynamique affine TS* prennent la forme :

$$\begin{aligned}
 R_i : \quad & \text{Si } y(k) \text{ est } A_{i1} \text{ et } y(k-1) \text{ est } A_{i2} \text{ et } \dots \text{ et } y(k-n_y+1) \text{ est } A_{in_y} \\
 & \text{et } u(k) \text{ est } B_{i1} \text{ et } u(k-1) \text{ est } B_{i2} \text{ et } \dots \text{ et } u(k-n_u+1) \text{ est } B_{in_u} \\
 \text{Alors } & y_i(k+1) = \sum_{j=1}^{n_y} a_{ij} y(k-j+1) + \sum_{j=1}^{n_u} b_{ij} u(k-j+1) + d_i, \quad i=1, \dots, r
 \end{aligned} \tag{3.26}$$

La sortie globale du modèle est calculée à partir de l'expression :

$$y(k+1) = \sum_{i=1}^r \lambda_i(k) y_i(k+1) \tag{3.27}$$

Dans laquelle $\lambda_i(k) \in [0, 1]$ est le degré normalisé d'accomplissement conformément à l'expression (3.136) et $\sum_{i=1}^r \lambda_i(k) = 1$.

Il faut remarquer que le modèle flou dynamique Takagi-Sugeno représenté par les règles de la forme (3.26) inclut aussi le cas $d_i = 0$ dit homogène [BAB98a] ou linéaire [ZHA97]. Cependant, il convient également de souligner que le terme "linéaire" signifie tout simplement que la sortie individuelle de chaque règle est une combinaison linéaire des régresseurs. En effet, le modèle dynamique représenté par l'expression (3.26) est en réalité non linéaire car les coefficients du modèle dépendent non linéairement des entrées et sorties du système via le degré normalisé d'accomplissement. Le fait que le terme indépendant d_i ne soit pas considéré dans le modèle homogène limite ces capacités d'approximation par rapport au modèle affine, mais d'autre part, facilite énormément la synthèse de la commande ainsi que l'analyse de la stabilité. En effet, certaines méthodes ont été développées pour la synthèse de contrôleurs avec des caractéristiques de performance désiré en boucle fermée [FIL96] [ZHA97] ainsi que pour l'analyse de la stabilité [TAN92] [TAN96].

Dans le cadre de la modélisation de systèmes, le modèle flou dynamique de type Takagi-Sugeno permet aussi la représentation d'autres structures non linéaires entrée-sortie utilisés fréquemment, telles que les formes NOE⁵² et NARMAX⁵³. Cependant, dans une approche d'identification à partir des données entrée-sortie, le problème de régression dans ces structures doit être résolu itérativement car le vecteur de régression ne peut pas être construit directement à partir des données puisque il contient les valeurs précédentes de la sortie du modèle. En principe, les méthodes d'identification par *clustering*, que nous aborderons dans la section 3.3, peuvent être appliquées à ces structures, mais comme ces méthodes de clustering flou sont itératives, cela amène à un problème d'optimisation plutôt complexe [BAB98a]. Par la suite nous nous sommes intéressés aux modèles dynamiques affines TS du type NARX.

D'autre part, dans le modèle dynamique représenté par (3.26)-(3.27) il est supposé que y est un scalaire, c'est-à-dire, le système étudié est un système de type MISO⁵⁴. Avec les modèles entrée-sortie, les systèmes MIMO⁵⁵ peuvent être représentés de deux manières : soit la fonction non linéaire f_{NL} est une fonction vectorielle, soit le système MIMO est décomposé dans un ensemble de systèmes MISO couplés. En modélisation flou l'approche de

⁵² NOE : *Nonlinear Output Error*.

⁵³ NARMAX : *Nonlinear AutoRegressive Moving Average with eXogenous inputs*.

⁵⁴ MISO : *Multiple-Input Single-Output*.

⁵⁵ MIMO : *Multiple-Input Multiple-Output*.

décomposition est la plus adoptée. La raison est qu'il y a plus de flexibilité si chaque sortie est associée à différent type de non linéarité pour la même région de l'espace de l'antécédent [BAB98a]. En décomposant le système MIMO en plusieurs sous-systèmes MISO, le nombre de fonctions d'appartenance et de règles peut être réduit. Dans la suite de ce mémoire nous utilisons cette dernière approche.

En plus de la représentation des structures entrée-sortie les plus fréquentes, les modèles flou dynamiques de type Takagi-Sugeno peuvent aussi représenter les systèmes non linéaires sous la forme d'espace d'état :

$$\begin{aligned} \mathbf{x}(k+1) &= g(\mathbf{x}(k), \mathbf{u}(k)) \\ \mathbf{y}(k) &= f(\mathbf{x}(k)) \end{aligned} \quad (3.28)$$

où g et h sont des fonctions statiques. La fonction g , nommée la fonction de transition d'état, établie la correspondance entre l'état actuel $\mathbf{x}(k)$ et l'entrée $\mathbf{u}(k)$ avec le nouveau état $\mathbf{x}(k+1)$. Quant à elle, la fonction f établie la correspondance entre l'état actuel $\mathbf{x}(k)$ et la sortie du système $\mathbf{y}(k)$. Un exemple d'une représentation basée sur des règles d'un système non linéaire dans l'espace d'état est le modèle dynamique Takagi-Sugeno suivant :

$$R_i : \quad \text{Si } \mathbf{x}(k) \text{ est } A_i \text{ et } \mathbf{u}(k) \text{ est } B_i \quad \text{Alors} \quad \begin{cases} \mathbf{x}(k+1) = \mathbf{A}_i \mathbf{x}(k) + \mathbf{B}_i \mathbf{u}(k) \\ \mathbf{y}(k) = \mathbf{C}_i \mathbf{x}(k) \end{cases}, \quad i = 1, \dots, r \quad (3.29)$$

Le modèle comporte r règles dans lesquelles A_i et B_i représentent des sous-ensembles flous, les matrices $\mathbf{A}_i$, $\mathbf{B}_i$ et $\mathbf{C}_i$ sont des matrices de dimensions appropriés associés à la représentation de la dynamique du système. La sortie globale du modèle est déterminée en utilisant le même principe de calcul (moyenne pondérée) présenté dans l'équation (3.27).

Comme nous l'avons vu, le modèle dynamique flou de type Takagi-Sugeno comporte les paramètres de l'antécédent et du conséquent pour chacune des règles. Dans la section suivante nous donnons un aperçu général sur la construction des modèles flous. Nous nous sommes plus spécialement intéressés à l'identification de modèles dynamiques TS en utilisant des données entrée-sortie du procédé. Nous développerons l'approche d'identification basée sur des méthodes de clustering flou dans la section 3.3.

3.2. Identification (Construction) de modèles flous

Typiquement, la construction de modèles flous peut se faire à partir de deux sources d'information différentes, telles que la connaissance préalable et les données (mesures du procédé). La connaissance préalable peut être plutôt d'une nature qualitative ou heuristique, issue de la connaissance des "experts", i.e., des designers de processus, des opérateurs. Dans ce sens, les modèles flous peuvent être vus comme des *systèmes experts flous* [ZIM87].

D'un autre côté, pour certains procédés, des données sont disponibles sous forme d'enregistrements de l'opération du procédé ou bien il est possible de réaliser des expériences d'identification afin d'obtenir les données appropriés du comportement du système. La construction de modèles flous à partir de données implique des méthodes basées sur la logique floue et le raisonnement approché, mais aussi des idées issues du domaine des

réseaux de neurones, de l'analyses de données et de l'identification conventionnelle de systèmes [HEL97].

L'utilisation des techniques et d'algorithmes pour la construction de modèles flous à partir de données est habituellement appelée *identification floue* [BAB98a]. Deux approches principales pour l'intégration de la connaissance et des données dans un modèle flou peuvent être distinguées :

1. La connaissance experte exprimée sous une forme verbale est traduite dans une collection de règles du type "Si-Alors". De cette façon, une certaine structure du modèle est créée. Les paramètres dans cette structure (fonctions d'appartenance, paramètres des conséquents) peuvent être réglés avec précision en utilisant les données entrée-sortie. Les algorithmes particuliers d'ajustement exploitent le fait qu'au niveau du calcul, le modèle flou peut être vu comme une structure par couches (réseaux), similaire aux réseaux artificiels de neurones, pour laquelle des algorithmes standard d'apprentissage peuvent être appliqués. Cette approche est habituellement nommée *modélisation neurofloue* [JAN93] [PED95].
2. Aucune connaissance antérieure n'est initialement utilisée pour formuler les règles et le modèle flou est construit à partir des données. On s'attend à ce que les règles extraites puissent fournir une interprétation postérieure sur le comportement du système. Un expert peut confronter cette information avec sa propre connaissance, peut modifier les règles ou fournir de nouvelles règles et peut concevoir de nouvelles expériences additionnelles afin d'obtenir plus de données informatives.

Ces techniques peuvent, évidemment, être combinées selon l'application particulière. Dans la suite, nous décrivons les étapes principales pour la sélection de la structure d'un modèle flou ainsi que pour la détermination de paramètres à partir de données entrée-sortie.

3.2.1. Structure du modèle flou

En modélisation floue des systèmes, la sélection de la structure implique habituellement les choix suivants :

- *Variables d'entrée et de sortie.* Dans les systèmes complexes, il n'est pas toujours facile de déterminer quelles variables devraient être employées comme entrées du modèle. Dans le cas des systèmes dynamiques, on doit également estimer l'ordre du système. Pour le modèle entrée-sortie NARX représenté par l'expression (3.25) cela signifie qu'il faut définir le nombre de retards, d'entrée et de sortie, n_u et n_y . Plusieurs sources d'information typiques permettent de faciliter ce choix, tels que la connaissance préalable (e.g., connaissance sur le comportement du système). Parfois, des méthodes de sélection automatique de structure (l'ordre du système) basées sur des données peuvent être utilisées pour comparer différents choix en termes de critères de performance [HE93] [BOM98] [RHO98] [HAD02] [FEI04] [SIN04] [UNR07].
- *Structure des règles.* Ce choix implique le type de modèle (linguistique, singleton, Takagi-Sugeno) et la forme de l'antécédent (fonctions d'appartenance multivariées ou dans la forme composée).
- *Nombre et type de fonctions d'appartenance pour chaque variable.* Ce choix détermine le niveau de détail (granularité) du modèle. De nouveau, le but de la

modélisation ainsi que la connaissance disponible vont influencer cette sélection. Des méthodes automatiques basées sur des données peuvent être utilisés pour ajouter ou supprimer des fonctions d'appartenance dans le modèle.

- *Type de mécanisme d'inférence, opérateurs logiques, méthode de défuzzification.* Ces choix sont restreints par le type de modèle flou choisi (Mamdani ou TS), ainsi que par les opérateurs flous sélectionnés. Pour faciliter l'optimisation des modèles flous basée sur les données (apprentissage), des opérateurs différentiables (produit, somme) sont souvent préférés aux opérateurs standard min et max.

Une fois que la structure est fixée, la performance de la méthode de modélisation peut être réglée avec précision en ajustant les paramètres. Les termes ajustables des modèles linguistiques sont les paramètres des fonctions d'appartenance de l'antécédent et du conséquent (détermination de sa forme et de sa position) ainsi que les règles (détermination de la correspondance entre les régions floues de l'antécédent et du conséquent). Pour leurs parts, les modèles flous Takagi-Sugeno ont des paramètres tant dans les fonctions d'appartenance de l'antécédent que dans les fonctions du conséquent (paramètres $\mathbf{a}_i, \mathbf{b}_i, d_i$ pour le modèle dynamique affine TS (3.26)). Un aperçu général du problème d'identification est donné dans le paragraphe suivant.

3.2.2. Le problème d'identification et ses solutions

Les deux étapes de base en identification des systèmes sont la *détermination de la structure* et l'*estimation des paramètres*. Le choix de la structure du modèle (variables, fonctions d'appartenance, etc...) est très important, car elle détermine la flexibilité du modèle dans l'approximation des systèmes. Un modèle avec une structure riche peut approximer des fonctions plus compliquées, mais, en même temps, il aura des propriétés de généralisation plus mauvaises. Le problème d'estimation de paramètres peut être formulé comme la minimisation d'un critère non linéaire des moindres carrés, de la façon suivante :

$$\{\mathbf{a}_1, \dots, \mathbf{a}_r, \boldsymbol{\theta}_1, \dots, \boldsymbol{\theta}_r\} = \arg \min \sum_{k=1}^n \left(y_k - \sum_{i=1}^r \lambda_i(\mathbf{z}_k, \mathbf{a}_i) \boldsymbol{\theta}_i^T \mathbf{w}_k \right)^2 \quad (3.30)$$

où les variables de l'antécédent $\mathbf{z}$ et du conséquent $\mathbf{w}$ sont choisies à partir du vecteur de régression $\mathbf{x}$, et n correspond au nombre disponible de paires de données $(\mathbf{x}_k, y_k)$, pour $k = 1, \dots, n$. Les paramètres libres sont les vecteurs des paramètres de l'antécédent et du conséquent $(\mathbf{a}_i$ et $\boldsymbol{\theta}_i^T)$. $\lambda_i(\mathbf{z})$ représente le degré normalisé d'accomplissement pour l' i -ème règle, étant donnée une base de r règles. Les techniques d'optimisation généralement utilisées peuvent être divisées en deux principales catégories [SAL05] :

- (1) Méthodes basées sur l'optimisation globale de tous les paramètres, tels que les algorithmes génétiques, les techniques d'apprentissage neuroflou (rétropropagation du gradient et ses variantes), clustering flou dans l'espace produit entrée-sortie, etc.
- (2) Méthodes qui exploitent le fait que l'expression (3.30) est non linéaire en $\mathbf{a}_i$ (du à la paramétrisation non linéaire des fonctions d'appartenance) et qu'elle est linéaire en $\boldsymbol{\theta}_i$. Typiquement, le problème d'estimation linéaire est résolu comme un problème local à chaque itération du problème d'optimisation des paramètres de l'antécédent.

Comme en principe toutes les méthodes de régression non linéaire peuvent être adoptées pour la modélisation floue, alors l'éventail des techniques disponibles est vaste. Néanmoins, la construction des modèles flous implique un compromis entre plusieurs caractéristiques,

telles que la précision, la transparence et l'interprétabilité du modèle [JIN00] [GUI01]. Beaucoup de recherches ont été aussi consacrées aux méthodes de réduction de la complexité de systèmes flous, par exemple en utilisant des mesures de similarité [SET98], transformations orthogonales [YEN99] ou bien fusion de variables ainsi que des méthodes de structuration (hiérarchisation, décentralisation) [LAC03].

Parmi l'éventail des méthodes, nous nous intéressons en particulier, aux méthodes de construction de modèles flous de type Takagi-Sugeno à partir des données entrée-sortie du procédé en utilisant des techniques de *clustering* flou. Le paragraphe suivant donne une vue générale à ce propos.

3.2.3. L'identification basée sur des données entrée-sortie

On se propose donc d'obtenir un modèle flou directement à partir des données numériques issues du système à modéliser ; ces données peuvent être bruitées. L'objectif des entrées appliquées au système (en boucle ouverte ou fermée) est de parcourir l'ensemble de l'espace dans lequel on recherche à modéliser le comportement du système. Le modèle flou obtenu offre une faible capacité d'extrapolation en dehors de cet espace [BOR03].

Parmi la gamme d'approches de modélisation floue à partir des données entrée-sortie, les suivantes sont dignes d'être mentionnées :

- *Modélisation basée sur des prototypes (template-based)*, dans laquelle les domaines des variables de l'antécédent sont simplement partitionnés dans un nombre spécifié de fonctions d'appartenance avec la même forme et distribution dans l'espace.
- *Modélisation basée sur l'ajout progressif de fonctions d'appartenance*, dans laquelle est générée une partition de l'espace avec une base de règles à complexité contrôlée.
- *Modélisation basée sur des techniques de coalescence floue (fuzzy clustering)*, dans laquelle les règles du modèle flou peuvent être extraites à partir de la conformation de groupes de données (classes) dans l'espace produit d'entrée-sortie.

Nous abordons par la suite ces approches de manière plus approfondi, en mettant particulièrement l'accent sur les méthodes de coalescence floue que nous développerons dans la section 3.3.

Modélisation basée sur des prototypes (*template-based*)

Dans cette approche il est supposé que l'expert peut fournir initialement les fonctions d'appartenance et les règles. Néanmoins, pour divers systèmes, une connaissance particulière sur la partition de l'espace d'opération n'est pas disponible. Dans ce cas, les domaines des variables de l'antécédent peuvent être simplement partitionnés dans un nombre spécifié de fonctions d'appartenance avec la même forme et distribution dans l'espace. La base de règles peut être établie afin de couvrir toutes les combinaisons des termes de l'antécédent. La procédure d'identification consiste alors à estimer, à partir des données, les paramètres restants dans le modèle flou. S'il n'y a pas de connaissance disponible sur quelles variables génèrent la non linéarité du système, alors toutes les variables de l'antécédent sont partitionnées uniformément, ce qui conduit à une croissance exponentielle du nombre de règles [KOS91].

Modélisation basée sur l'ajout progressif de fonctions d'appartenance

L'objectif de cette modélisation est de générer automatiquement une partition floue de l'espace d'état ou de l'espace des entrées du processus. A chaque règle floue de type Takagi-Sugeno est associé un modèle linéaire calculé à partir des données mesurées sur le système réel. La précision du modèle peut être accrue par ajouts successifs de fonctions d'appartenance (et donc de règles) dans l'espace de chaque entrée du modèle, par ajout de fonctions d'appartenance unidimensionnelles ou, dans l'espace des entrées, par ajout de fonctions d'appartenance multidimensionnelles. Typiquement ces méthodes utilisent une procédure itérative pour la construction du modèle, en effet, le modèle flou est augmenté jusqu'à atteindre une partition optimale et une base de règles qui satisfassent un certain critère de performance [NAK97a] [NAK97b]. Par exemple, la méthode itérative s'arrête quand l'erreur entre le modèle et le système réel est inférieure à une valeur prédéfinie [BOR99] [BOR03].

Modélisation basée sur des techniques de coalescence floue

Les méthodes d'identification basées sur la coalescence floue (*fuzzy clustering*) sont des méthodes qui ont des liens avec les domaines de l'analyse des données et de la reconnaissance des formes. Dans ces domaines, le concept d'une appartenance partielle est utilisé pour représenter le degré avec lequel un certain objet, représenté comme un vecteur dans un espace de caractéristiques (vecteur caractéristiques), est similaire à un certain objet prototype. Le degré de similitude est habituellement calculé en utilisant une mesure appropriée de distance. Basé sur la similitude, les vecteurs caractéristiques peuvent être regroupés de manière que les vecteurs d'un même groupe (*cluster*) soient aussi semblables que possible, de telle manière que des vecteurs de groupes différents soient aussi dissemblables que possible. La mesure de distance quantifie la distance entre les données, représentées comme des points dans l'espace caractéristique, et les objets prototype.

L'idée de base du *clustering flou* est représentée sur la Figure 3.3a [BAB98a], dans laquelle les données sont regroupées dans deux groupes avec des prototypes v_1 et v_2 , en utilisant une mesure Euclidienne de distance dans un espace caractéristique $\mathbb{R}^2$. La partition des données est représentée dans la *matrice de partition floue*, $U=[\mu_{ik}]$, dont les éléments μ_{ik} sont les degrés d'appartenance des points $[x_k, y_k]$ dans des clusters flous avec prototypes v_i .

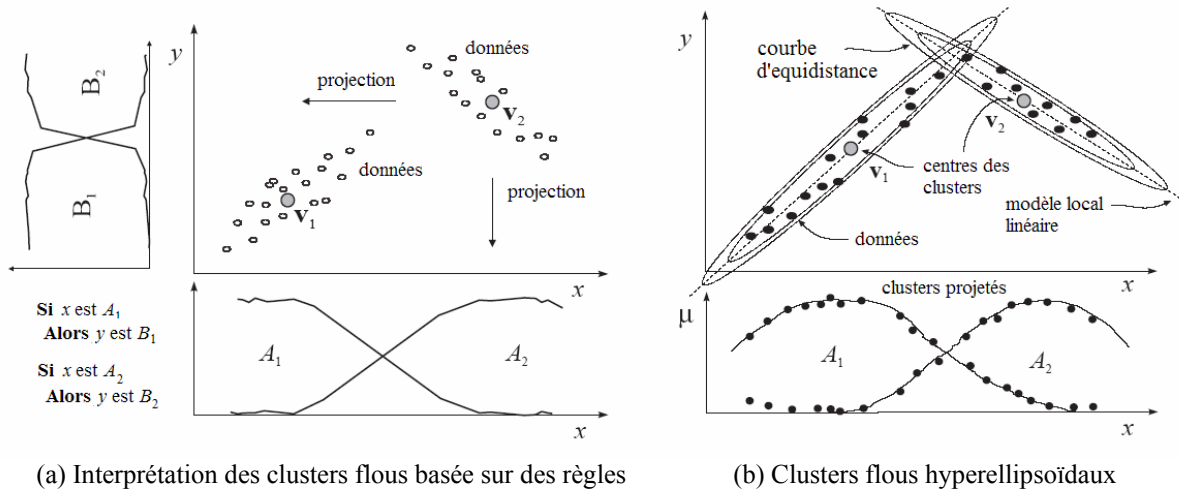


Figure 3.3. Identification par clustering flou (adaptation d'après [BAB98a])

Des règles floues de type "Si-Alors" peuvent être extraites en projetant les clusters sur les axes des coordonnées. La Figure 3.3a montre un ensemble de données avec deux clusters apparents et deux règles floues associées. Le concept de similarité des données par rapport à un prototype permet d'envisager plusieurs formes possibles pour la définition d'une mesure de distance appropriée et le caractère du prototype. Par exemple, les prototypes peuvent être définis comme des sous-espaces linéaires (lignes, plans et hyperplans) [BEZ81] ou bien comme des hyperellipsoïdes avec une mesure adaptative de distance [GUS79] (voir Figure 3.3b). Pour de tels clusters, les fonctions d'appartenance et les paramètres des conséquents d'un modèle de type Takagi-Sugeno peuvent être extraits [BAB95] :

$$\begin{aligned} \text{Si } x \text{ est } A_1 & \quad \text{Alors } y_1 = a_1x + d_1, \\ \text{Si } x \text{ est } A_2 & \quad \text{Alors } y_2 = a_2x + d_2 \end{aligned} \quad (3.31)$$

Chaque cluster est représenté par une règle dans le modèle Takagi-Sugeno. Les fonctions d'appartenance pour les ensembles flous A_1 et A_2 sont générées par la projection point par point de la matrice de partition sur les variables des antécédents. Ces ensembles flous définis point par point peuvent alors être approximés par des fonctions paramétriques appropriées (cf. expressions (3.3) et (3.4)). Les paramètres du conséquent pour chaque règle peuvent être obtenus par l'intermédiaire d'une estimation de type moindres carrés, que nous aborderons plus tard dans les sections suivantes.

3.3. Méthodes de coalescence floue (clustering flou)

Comme nous l'avons déjà indiqué, les modèles flous (en particulier ceux de type Takagi-Sugeno) sont considérés sous l'optique d'une approche de modélisation multi-linéaire, qui essaie de résoudre un problème complexe de modélisation en le décomposant en plusieurs sous-problèmes plus simples. La théorie des ensembles flous offre alors un excellent outil pour représenter l'incertitude associée à la tâche de décomposition, en fournissant des transitions douces entre les sous-modèles linéaires et pour intégrer divers types de connaissance dans un même cadre. Dans ce cas, l'identification floue peut être vue comme une recherche de la décomposition du système, en exploitant le fait qu'en général la complexité des systèmes n'est pas uniforme. Puisqu'il n'est pas possible de prévoir à l'avance qu'il y ait suffisamment de connaissance concernant cette tâche, nous considérons alors des méthodes de génération automatique de la décomposition, principalement à partir des données du système. Dans cette section nous abordons une classe appropriée d'algorithmes de coalescence floue (*clustering*) utilisée à ce propos.

3.3.1. Notation et Concepts de base

Les techniques de classification sont un outil important qui a trouvée des applications variées, telles que l'exploitation des données, la recherche documentaire, la compression de signal, la reconnaissance de la parole et des caractères manuscrits, la segmentation d'image, la modélisation des systèmes, la détection des défaillances et le diagnostic.

D'un point de vue général, les méthodes de classification ont pour but de regrouper les éléments d'un ensemble $Z = \{z_1, z_2, \dots, z_N\}$ en un nombre optimal de classes selon leurs ressemblances. Il existe un grand nombre de typologies de méthodes de classification que l'on peut regrouper en plusieurs catégories. Une première distinction consiste à séparer les

techniques en trois grandes familles, selon le type d'*apprentissage*⁵⁶ sous-jacent, utilisé lors du classement. D'une façon très succincte, ces types d'apprentissage sont :

- L'*apprentissage supervisé*, appelé aussi « apprentissage avec un maître », dans lequel on dispose d'un ensemble de données dont on connaît *a priori* la classe d'appartenance. Pour chaque vecteur d'entrée, la sortie désirée correspondante est connue à tout moment. Le but de la *classification supervisée* consiste alors à ajuster les paramètres du classifieur pour classer le plus correctement les données en utilisant la différence entre sa sortie et la sortie de référence. Un exemple largement connu associé à ce type d'apprentissage est la technique de rétropropagation du gradient, souvent utilisée dans les systèmes neuroflous.
- L'*apprentissage par renforcement*, contrairement à l'apprentissage supervisé, il cherche à faire émerger des *comportements* permettant d'atteindre un objectif, sans autre information qu'un signal scalaire, le renforcement. Dans ce type d'apprentissage, un « agent » analyse en permanence les conséquences de ses actions, ayant tendance à reproduire préférentiellement celles qui, dans les mêmes circonstances, ont conduit à des succès. L'agent est un terme très général qui désigne ici un système sous apprentissage : un programme informatique, un réseau de neurones, un système d'inférence floue, etc. [GLO03] [CER03].
- L'*apprentissage non supervisé*, qui consiste à établir des relations d'entrée/sortie en cherchant la formation de classes dans un ensemble d'observations sur lesquelles on ne dispose pas d'informations quant à leur appartenance aux éventuelles classes. Les méthodes de *classification non supervisée* permettent le regroupement des éléments ayant des propriétés statistiques, géométriques ou linguistiques semblables. Les méthodes de coalescence floue (regroupement ou *clustering*) font partie de ces méthodes, car elles n'utilisent pas des étiquettes préexistantes pour les classes⁵⁷. Elles peuvent être utilisées comme un outil pour obtenir une partition des données quand les transitions entre les classes sont plutôt "douces" que "fortes" [BAB98a].

Une autre distinction peut être faite en considérant le principe de représentation ou de construction du classement. De son côté [BIE99] présente cinq catégories de méthodes : *les méthodes statistiques*, *les méthodes hiérarchiques*, *les méthodes par graphes*, *les méthodes itératives* et *les méthodes connexionnistes*. Pour sa part [BEZ81] présente une distinction en se basant sur l'approche algorithmique des différentes techniques. Ainsi par exemple, les méthodes *hiérarchiques* forment de nouveaux groupes (*clusters*) en redistribuant les appartenances d'un point à la fois, basé sur une certaine mesure convenable de similitude. Dans les méthodes *par graphes*, l'ensemble Z , est considéré comme un ensemble de nœuds. Les poids des liaisons entre paires de nœuds sont basés sur une mesure de similitude entre ces nœuds.

Par ailleurs, la définition des classes et des relations entre celles-ci dépend de la structure de classification utilisée. La plupart des algorithmes de clustering présents dans la littérature appartiennent à une des deux catégories suivantes : clustering hiérarchique et clustering par partition [JAI88]. Les algorithmes de clustering *hiérarchique* considèrent seulement des

⁵⁶ Malgré la complexité qu'implique une définition étendue du concept d'apprentissage, nous considérons ici d'une façon très modeste, une définition assez limitée mais appropriée à nos recherches : on peut considérer l'apprentissage comme le processus par lequel un utilisateur fournit des données (dans notre cas numériques) à un logiciel dans le but que celui s'en serve comme modèle pour pouvoir ensuite les reconnaître.

⁵⁷ Néanmoins, il faut dire que pour la plupart des algorithmes conventionnels de coalescence floue le nombre de classes est établi *a priori*.

voisins locaux et créent une décomposition hiérarchique des données, en générant une séquence de partitions arborescentes. Ces algorithmes peuvent être divisés en méthodes agglomératives (*bottom-up*) et méthodes par division (*top-down*) :

- Les algorithmes *agglomératifs* commencent en considérant un grand nombre de clusters individuels et successivement les fusionnent selon une mesure de distance. Le clustering peut s'arrêter quand toutes les données sont dans un seul groupe ou bien à n'importe quel point intermédiaire désiré. Ces méthodes suivent une stratégie de fusion de clusters de type constructive.
- Les algorithmes *par division* suivent la stratégie opposée. Ils commencent par un groupe qui contient toutes les données et progressivement le groupe est divisé en groupes plus petits jusqu'à atteindre la situation limite dans laquelle chaque donnée constitue un groupe, ou bien avant cette limite, selon ce que l'on désire.

Les algorithmes hiérarchiques sont, en principe, plus efficaces dans le traitement du bruit et des points aberrants (*outliers*). Cependant, ces techniques hiérarchiques souffrent du fait qu'une fois qu'une fusion ou qu'une division est faite, elle ne peut pas être défaite ou raffinée.

Concernant les algorithmes de clustering *par partition*, ces algorithmes diffèrent des techniques hiérarchiques dans le fait qu'ils admettent la relocation dans les classes (dans le processus itératif les données peuvent se déplacer d'un cluster à un autre). Cela permet qu'une partition initiale pauvre, puisse être corrigée à une étape postérieure. Ils diffèrent aussi parce que leurs solutions n'expriment pas nécessairement un rapport hiérarchique entre les classes. Ces algorithmes génèrent une partition unique des données et peuvent incorporer la connaissance de la forme ou de la taille des clusters en utilisant des prototypes et des mesures de distance (mesure de similitude) adéquats dans leur formulation. Le clustering en utilisant une mesure de *similitude* permet la définition des classes décrivant des régions distinctes dans l'espace de représentation. Le terme de similitude est typiquement défini au moyen d'une *norme de distance*. La distance peut être mesurée parmi les vecteurs de données, ou bien comme une distance d'un vecteur de données à un certain objet prototypique du groupe. Les *prototypes* peuvent être des vecteurs de la même dimension que les objets de données, mais ils peuvent également être définis en tant qu'objets géométriques "de plus haut niveau", tels que des sous-espaces ou des fonctions linéaires ou non linéaires. En particulier, la formulation de ces algorithmes est basée sur l'utilisation d'une *fonction objectif* (critère de classification) pour établir itérativement une partition floue appropriée d'un ensemble de données. Des algorithmes d'optimisation non linéaires sont utilisés pour chercher l'extremum local d'une telle fonction. Dans ces méthodes la fonction objectif est différentiable, ce qui constitue une propriété utile pour l'optimisation. Les principaux inconvénients de l'approche par partition sont la difficulté pour estimer le nombre de clusters, ainsi que la sensibilité par rapport à l'initialisation, le bruit et les points aberrants [FRI96].

Dans la suite de ce chapitre nous nous focalisons sur les techniques de clustering flou basées sur une fonction objectif. Nous verrons comment, en utilisant de telles techniques, il est possible de construire des modèles flous de type Takagi-Sugeno à partir des données entrée-sortie du procédé.

Notation

Dans les méthodes de clustering flou, les données sont typiquement des *observations* (mesures) issues d'un certain processus physique. Chaque k -ième observation constitue un

vecteur noté par $\mathbf{z}_k = [z_{1k}, z_{2k}, \dots, z_{nk}]^T$, avec $1 \leq k \leq N$, $\mathbf{z}_k \in \mathfrak{R}^n$, où N représente le nombre des observations et n correspond au nombre de variables mesurées, attributs ou *caractéristiques* dans l'espace de représentation $\mathfrak{R}^n$. Un ensemble de N observations est dénoté par $Z = \{\mathbf{z}_k \mid k = 1, 2, \dots, N\}$ et il est représenté comme une matrice sous la forme :

$$Z = \begin{bmatrix} z_{11} & z_{12} & \cdots & z_{1N} \\ z_{21} & z_{22} & \cdots & z_{2N} \\ \vdots & \vdots & \vdots & \vdots \\ z_{n1} & z_{n2} & \cdots & z_{nN} \end{bmatrix} \quad (3.32)$$

Dans la littérature sur la reconnaissance des formes, les colonnes de cette matrice sont appelées les objets ou patrons et Z est appelée la matrice patronne ou *matrice des données*. Quand le *clustering* est appliqué pour la modélisation et l'identification des systèmes dynamiques, les colonnes de Z contiennent les échantillons des signaux observés au cours du temps, et les lignes sont, par exemple, les variables physiques mesurés dans le système (e.g. température, pH, etc.). Afin de représenter la dynamique du système, il est nécessaire d'inclure typiquement les valeurs passées des variables dans la matrice Z . Par exemple, pour le cas d'un système non linéaire représenté par un modèle NARX de deuxième ordre $y(k+1) = f_{NL}(y(k), y(k-1), u(k), u(k-1))$, les règles R_i d'un modèle dynamique affine Takagi-Sugeno ont la forme suivante :

$$\begin{aligned} R_i : \quad & \text{Si } y(k) \text{ est } A_{i1} \text{ et } y(k-1) \text{ est } A_{i2} \text{ et } u(k) \text{ est } B_{i1} \text{ et } u(k-1) \text{ est } B_{i2} \\ & \text{Alors } y(k+1) = a_{i1}y(k) + a_{i2}y(k-1) + b_{i1}u(k) + b_{i2}u(k-1) + d_i \end{aligned} \quad (3.33)$$

Dans ce cas, la matrice des données prend la forme :

$$Z^T = \begin{bmatrix} y(1) & y(2) & u(1) & u(2) & \vdots & y(3) \\ y(2) & y(3) & u(2) & u(3) & \vdots & y(4) \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ y(N-2) & y(N-1) & u(N-2) & u(N-1) & \vdots & y(N) \end{bmatrix} \quad (3.34)$$

Les algorithmes de clustering flou visent à réaliser une classification d'un ensemble de données en établissant une partition floue des observations en un certain nombre de classes. La notion de *partition* floue, qui s'est avéré d'une grande utilité pour le développement des techniques floues de classification, a été introduite par Ruspini [RUS69]. Au sens de Ruspini (partition floue stricte), une c -partition floue d'un ensemble Z peut être obtenue en définissant des c sous-ensembles flous de Z tel que la somme des degrés d'appartenance pour chaque observation de Z soit unitaire. En fait, on associe à chaque observation $\mathbf{z}_k \in Z$ un vecteur de degrés d'appartenance aux différentes classes. La juxtaposition de ces vecteurs pour l'ensemble des N observations de Z amène à la définition d'une matrice d'appartenance $\mathbf{U}$ (de dimension $c \times N$) où l'élément μ_{ik} représente le coefficient d'appartenance⁵⁸ de l'observation $\mathbf{z}_k$ à la classe c . Cette matrice établit une relation d'ordre floue et traduit l'idée d'une partition floue en c classes.

⁵⁸ Afin de simplifier la notation, nous utilisons dans cette section μ_i au lieu de la forme usuelle μ_{A_i} . De plus, pour faciliter la notation matricielle, μ_{ik} est noté $\mu_i(\mathbf{z}_k)$.

La matrice d'appartenance $\mathbf{U} = [\mu_{ik}]$ est également appelée *matrice de partition floue* et respecte les propriétés suivantes :

$$\mu_{ik} \in [0, 1], \quad 1 \leq i \leq c, \quad 1 \leq k \leq N, \quad (3.35)$$

$$0 < \sum_{k=1}^N \mu_{ik} < N, \quad 1 \leq i \leq c. \quad (3.36)$$

$$\sum_{i=1}^c \mu_{ik} = 1, \quad 1 \leq k \leq N, \quad (3.37)$$

L'expression (3.37) traduit une condition de normalisation d'inspiration probabiliste. Cependant, bien que cette approche de la définition d'une partition floue soit la plus utilisée, différents auteurs [KRI93] [KHO97] présentent des approches où cette contrainte est relaxée. Il s'agit d'approches qui conduisent à une partition dite *possibiliste*⁵⁹, dans laquelle en plus de respecter les expressions (3.35) et (3.36), la contrainte exprimée par (3.37) est remplacée par une autre moins restrictive :

$$\exists i, \mu_{ik} > 0, \quad 1 \leq k \leq N \quad (3.38)$$

Cette expression est dérivée du fait que la contrainte (3.37) ne peut pas être complètement enlevée, afin d'assurer que chaque point soit assigné à au moins un des sous-ensembles flous avec une appartenance supérieur à zéro. Dans la suite nous nous intéressons aux techniques qui conduisent à une partition floue des données définie par l'ensemble des expressions (3.35) à (3.37).

Algorithmes de coalescence floue

Comme nous l'avons indiqué précédemment, les algorithmes de clustering flou optimisent itérativement un critère de classification afin d'établir une partition floue d'un ensemble de données en un certain nombre de classes (clusters). Le groupement des données est fait à partir d'une phase d'"apprentissage" en utilisant une mesure de similitude par l'intermédiaire de techniques de classification. Au niveau de la modélisation, à chaque cluster correspond théoriquement un fonctionnement homogène du système qui peut se présenter sous la forme d'un modèle linéaire local. Le nombre de modèles linéaires locaux est égal au nombre de règles du modèle. Cela permet le recours à la théorie de l'Automatique linéaire pour la définition du contrôleur.

Dans les paragraphes suivants, un ensemble non exhaustif d'algorithmes de coalescence floue est présenté, basés sur la minimisation d'une fonction objectif. Parmi ceux-ci, le plus connu est l'algorithme des *c-moyennes floues*, en anglais *Fuzzy c-Means* (FCM) ainsi qu'une de ses extensions, l'algorithme de *Gustafson-Kessel* (GK). Le premier permet la détection de classes hypersphériques, tandis que le deuxième, détecte des classes hyperellipsoïdales, typiquement mieux adaptées à la géométrie des observations. Dans la gamme des algorithmes de clustering avec prototypes linéaires utiles à nos propos de modélisation floue des systèmes, nous présentons l'algorithme *Fuzzy c-Regression Models* (FCRM). Ensuite, nous introduisons l'algorithme *Fuzzy Ellipsoidal Regression* (FER), utilisé pour favoriser l'interprétabilité du modèle flou résultant. Ces quatre algorithmes de coalescence floue nécessitent la connaissance préalable du nombre de clusters. Pour finir, nous présentons l'algorithme *Fuzzy Robust Ellipsoidal Regression* (RoFER), sur la base d'une approche d'« agglomération

⁵⁹ Le terme "possibiliste" (partition, clustering, etc.) a été introduit par Krishnapuram et Keller [KRI93]. Dans la littérature, les termes "partition floue contrainte" et "partition floue non contrainte" sont utilisés aussi pour dénoter les partitions qui impliquent respectivement les expressions (3.37) et (3.38) [BAB98a].

compétitive », qui non seulement permet d'établir "automatiquement" le nombre de clusters linéaires, mais aussi qui est robuste en présence de bruit et de points aberrants⁶⁰ (*outliers*).

3.3.2. Groupage *c*-moyennes floues (algorithme FCM)

L'algorithme *Fuzzy c-Means* (FCM), issu des travaux de Dunn [DUN74] et amélioré plus tard par Bezdek [BEZ81], constitue une référence parmi les différentes méthodes de coalescence floue basées sur la minimisation de la fonction objectif *c-mean*, de la forme :

$$J_{FCM}(Z; \mathbf{U}, \mathbf{V}) = \sum_{i=1}^c \sum_{k=1}^N (\mu_{ik})^m D_{ik\mathbf{A}}^2 \quad (3.39)$$

où Z est l'ensemble des données,

$\mathbf{U} = [\mu_{ik}]$ est la matrice de partition floue (de dimension $c \times N$),

$\mathbf{V} = [\mathbf{v}_1, \mathbf{v}_2, \dots, \mathbf{v}_c]$ est un vecteur de *prototypes de clusters* (centres) qui doit être déterminé, avec $\mathbf{v}_i \in \mathcal{R}^n$ centre de la i -ème classe, $1 \leq i \leq c$.

$$D_{ik\mathbf{A}}^2 = \|\mathbf{z}_k - \mathbf{v}_i\|_{\mathbf{A}}^2 = (\mathbf{z}_k - \mathbf{v}_i)^T \mathbf{A} (\mathbf{z}_k - \mathbf{v}_i), \quad 1 \leq i \leq c, \quad 1 \leq k \leq N \quad (3.40)$$

est une norme de distance quadratique dans l'espace considéré, qui définit la mesure de distance entre l'observation $\mathbf{z}_k$ et le centre $\mathbf{v}_i$ au sens de la métrique induite par $\mathbf{A}$.

$m \in [1, \infty)$ est un facteur qui désigne le degré de flou⁶¹ de la partition.

Le paramètre m influence directement la forme des clusters dans l'espace des données du système. Quand le facteur m s'approche de la valeur 1, la forme de la fonction d'appartenance pour chaque cluster est presque booléenne ($\mu_{ik} \in \{0,1\}$) et $\mathbf{v}_i$ sont les moyennes ordinaires des clusters. De l'autre côté, plus m est grand, plus la partition est floue. Quand $m \rightarrow \infty$, la partition devient floue au maximum ($\mu_{ik} = 1/c$) et les moyennes des clusters sont toutes égales à la moyenne de Z . Ces propriétés limites sont indépendantes de la méthode d'optimisation sélectionnée [PAL95]. Bien que le choix de m dépend des données [YU04], habituellement ce paramètre est initialisé à une valeur entre 1,5 et 2,5.

Dans l'équation (3.39), la mesure de non-similarité exprimée par le terme $J_{FCM}(Z; \mathbf{U}, \mathbf{V})$ est la somme des carrés des distances entre chaque vecteur de données $\mathbf{z}_k$ et le prototype (centre) du cluster correspondant $\mathbf{v}_i$. L'effet de cette distance est pondéré par le degré d'activation du cluster $(\mu_{ik})^m$ correspondant au vecteur de données $\mathbf{z}_k$. La valeur de la fonction coût $J_{FCM}(Z; \mathbf{U}, \mathbf{V})$ peut être vue comme une mesure de la variance totale de $\mathbf{z}_k$ par rapport aux prototypes $\mathbf{v}_i$.

La minimisation de la fonction objectif $J_{FCM}(Z; \mathbf{U}, \mathbf{V})$ est un problème d'optimisation non linéaire qui peut être résolu par différentes méthodes, parmi lesquelles la plus utilisée est basée sur l'itération de Picard à travers les conditions du premier ordre pour les points stationnaires [BEZ80]. Cette méthode est connue comme l'algorithme *Fuzzy c-Means* (FCM). Les points stationnaires de la fonction objectif (3.39) peuvent être trouvés en ajoutant la contrainte de normalisation (3.37) au critère J_{FCM} au moyen des multiplicateurs de Lagrange :

⁶⁰ *Outlier* : Dans un ensemble d'observations, c'est la donnée qui s'écarte nettement de cet ensemble.

⁶¹ m : *fuzziness weighting exponent*.

$$\bar{J}_{FCM}(Z; \mathbf{U}, \mathbf{V}, \boldsymbol{\lambda}) = \sum_{i=1}^c \sum_{k=1}^N (\mu_{ik})^m D_{ik\mathbf{A}}^2 + \sum_{k=1}^N \lambda_k \left[\sum_{i=1}^c \mu_{ik} - 1 \right] \quad (3.41)$$

où les termes λ_k sont les coefficients multiplicateurs de Lagrange. La minimisation du critère J_{FCM} s'obtient alors en annulant le gradient de $\bar{J}_{FCM}$ ⁶² par rapport $\mathbf{U}, \mathbf{V}$ et $\boldsymbol{\lambda}$. Si $D_{ik\mathbf{A}}^2 > 0, \forall i, k$ et $m > 1$, alors on obtient les relations de mise à jour suivantes :

Pour les coefficients μ_{ik} de la matrice de partition floue :

$$\mu_{ik} = \frac{1}{\sum_{j=1}^c (D_{ik\mathbf{A}}^2 / D_{jk\mathbf{A}}^2)^{1/(m-1)}}, \quad 1 \leq i \leq c, \quad 1 \leq k \leq N. \quad (3.42)$$

Pour les prototypes (centres) $\mathbf{v}_i$ des clusters :

$$\mathbf{v}_i = \frac{\sum_{k=1}^N (\mu_{ik})^m \cdot \mathbf{z}_k}{\sum_{k=1}^N (\mu_{ik})^m}, \quad 1 \leq i \leq c. \quad (3.43)$$

Cette solution satisfait aussi les contraintes (3.35) et (3.36). Il faut remarquer que l'expression (3.43) exprime $\mathbf{v}_i$ comme la moyenne pondérée des données qui appartiennent à un cluster, où le facteur de pondération est associé aux degrés d'appartenance floue. C'est pourquoi l'algorithme est appelé "*c-moyennes*" (*c-means*).

Le processus itératif s'arrête lorsque la partition devient stable, c'est-à-dire lorsqu'elle n'évolue plus significativement, entre deux itérations successives. Ceci s'exprime de manière générale par la vérification de l'expression (3.44) où le terme de gauche traduit une norme matricielle et le coefficient ε définit le seuil de convergence :

$$\|\mathbf{U}^{(t)} - \mathbf{U}^{(t-1)}\| < \varepsilon \quad (3.44)$$

L'expression $\mathbf{U}^{(t)}$ représente la matrice de partition floue à l'itération (t).

L'algorithme général de résolution *fuzzy c-mean* (FCM) peut se formuler ainsi :

Algorithme FCM - *Fuzzy c-means*

Etant donné un ensemble de vecteurs de données Z :

Définir la matrice $\mathbf{A}$ de métrique de la distance $D_{ik\mathbf{A}}^2$

Fixer le nombre de *clusters* $1 < c < N$

Fixer le degré de flou $m > 1$ et la tolérance de fin d'algorithme $\varepsilon > 0$

Initialiser la matrice $\mathbf{U} = [\mu_{ik}]$ de partition floue

Répéter

Pas 1. Calculer les prototypes (centres) des clusters (Eq. (3.43))

Pas 2. Calculer les distances pour chaque cluster (Eq. (3.40))

Pas 3. Mettre à jour la matrice $\mathbf{U}$ de partition floue (Eq. (3.42))

Jusqu'à obtenir la stabilité de la partition (Eq. (3.44))

⁶² Pour la suite du chapitre, remarquons que par rapport la fonctionnelle J , la notation $\bar{J}$ comprend les contraintes d'égalité pour la minimisation.

Remarques

- R 3.2** Une singularité dans l'algorithme FCM se produit au Pas 3 quand $D_{isA} = 0$ pour un certain $\mathbf{z}_k$ et un ou plusieurs prototypes de groupe $\mathbf{v}_s, s \in S \subset \{1, 2, \dots, c\}$. Dans ce cas, le degré d'appartenance dans l'expression (3.42) ne peut pas être calculé. Quand cela arrive, la valeur zéro est assignée à chaque $\mu_{ik}, i \in \bar{S}$ et l'appartenance est distribuée arbitrairement parmi μ_{sj} soumis à la contrainte $\sum_{s \in S} \mu_{sj} = 1, \forall k$ [BAB98a].
- R 3.3** Les équations (3.42), (3.43) sont seulement des conditions nécessaires du premier ordre pour les points stationnaires de la fonctionnelle (3.39). La suffisance de ces conditions et la convergence de l'algorithme a été abordée par Bezdek [BEZ80]. Néanmoins, la minimisation du critère J_{FCM} n'est pas simple à cause des minima locaux⁶³. En effet, l'existence de "points selle" a été démontrée [TUC87] dans plusieurs contre-exemples au supposé théorème de convergence. Les contre-exemples sont expliqués plus en détail dans l'article de Bezdek [BEZ87].
- R 3.4** La procédure d'optimisation alternative utilisée par l'algorithme FCM itère à travers les estimations $\mathbf{U}^{(t-1)} \rightarrow \mathbf{V}^{(t)} \rightarrow \mathbf{U}^{(t)}$ et se termine aussitôt que $\|\mathbf{U}^{(t)} - \mathbf{U}^{(t-1)}\| < \varepsilon$. Un autre choix, consiste à l'initialiser avec $\mathbf{V}^{(0)}$, itérer après à travers $\mathbf{V}^{(t-1)} \rightarrow \mathbf{U}^{(t)} \rightarrow \mathbf{V}^{(t)}$ et finir quand $\|\mathbf{V}^{(t)} - \mathbf{V}^{(t-1)}\| < \varepsilon$. La norme d'erreur dans le critère de terminaison est usuellement choisie comme $\max_{ik} (|\mu_{ik}^{(t)} - \mu_{ik}^{(t-1)}|)$. Des résultats différents peuvent être obtenus avec la même valeur de ε , puisque le critère de terminaison utilisé dans l'algorithme FCM exige que les valeurs de plusieurs paramètres s'approchent les uns des autres. La valeur usuelle pour le seuil de terminaison est $\varepsilon = 0.001$, bien qu'une valeur de $\varepsilon = 0.01$ convienne dans la plupart des cas [BAB98a].

L'algorithme FCM réalise une partition floue d'un ensemble d'observations en un nombre connu *a priori* de classes. Comme on vient de l'indiquer, cet algorithme converge en général vers un minimum local de la fonction objectif. Le résultat obtenu dépend aussi de l'initialisation des prototypes (centres) des classes, ainsi que du choix de la métrique employée et des différents paramètres de l'algorithme (c, m, ε). Le nombre de clusters c est le paramètre le plus important, les autres paramètres ont une influence mineure au niveau de la partition, en comparaison avec les effets du nombre de clusters. Ceci implique habituellement de réitérer plusieurs fois l'algorithme avec des valeurs différentes, puis de sélectionner ultérieurement la meilleure partition au cours d'une phase de validation, sujet qui sera abordé dans le paragraphe 3.4.1.

D'autre part, la forme des clusters dans l'algorithme FCM est déterminée par le choix de la matrice de norme induite $\mathbf{A}$ et par la mesure de distance D_{ikA}^2 (cf. expression (3.40)). Un choix simple consiste à sélectionner $\mathbf{A}$ comme la matrice d'identité, $\mathbf{A} = \mathbf{I} \in \mathcal{R}^{n \times n}$, ce qui correspond à la norme standard Euclidienne :

$$D_{ik}^2 = (\mathbf{z}_k - \mathbf{v}_i)^T (\mathbf{z}_k - \mathbf{v}_i) \quad (3.45)$$

⁶³ Il serait souhaitable que deux points équidistants du centre de la classe à laquelle ils appartiennent aient le même degré d'appartenance à cette classe. Ce n'est pas le cas avec l'algorithme FCM. On peut avoir des points équidistants qui n'ont pas le même degré d'appartenance, cela est dû à l'équation (3.42) qui est fonction de la distance du point à la classe, mais aussi du point à toutes les autres classes. De ce fait des singularités se produisent car deux points bien qu'équidistants à une classe ne le sont pas par rapport aux autres [ATI05].

Notons qu'avec ce choix de la matrice $\mathbf{A}$, l'expression (3.45) mesure la distance en pondérant toutes les observations à par égale par rapport au prototype (centres) des clusters. L'interprétation géométrique de l'algorithme FCM en utilisant une mesure de distance Euclidienne, donne alors des clusters de forme hypersphérique, c'est-à-dire, des clusters dont les surfaces des fonctions d'appartenance constantes (surfaces de niveau) sont des hypersphères dans l'espace de représentation $\mathcal{R}^n$.

Nous abordons par la suite deux algorithmes qui représentent des extensions de l'algorithme FCM. L'algorithme de *Gustafson-Kessel* (GK), qui utilise une mesure adaptative de distance permettant la détection des classes hyperellipsoïdales, et l'algorithme *Fuzzy c-Regression Models* (FCRM), basé sur des prototypes hyperplanaires.

3.3.3. Groupage avec matrice de covariance floue (algorithme GK)

L'extension de l'algorithme FCM proposée par Gustafson et Kessel [GUS79] a l'avantage de tenir compte de la forme propre de chacune des classes. En particulier, il permet la détection des classes hyperellipsoïdales, typiquement mieux adaptées à la géométrie des observations. Pour cela, pour chaque cluster, la mesure de distance est définie à l'aide d'une métrique $\mathbf{A}_i$:

$$D_{ik\mathbf{A}_i}^2 = (\mathbf{z}_k - \mathbf{v}_i)^T \mathbf{A}_i (\mathbf{z}_k - \mathbf{v}_i) \quad (3.46)$$

où la matrice de norme induite $\mathbf{A}_i$ est propre à chaque cluster, avec $1 \leq i \leq c$. Les matrices $\mathbf{A}_i$ sont utilisées comme variables d'optimisation dans la fonctionnelle *c-means*, permettant ainsi à chaque cluster d'adapter la norme de distance à la structure des données.

Soit $\mathbf{A}$ qui représente un c -uplet de matrices de norme induite définies positives $\mathbf{A} = (\mathbf{A}_1, \mathbf{A}_2, \dots, \mathbf{A}_c)$. La fonction objectif de l'algorithme GK est définie par :

$$J_{GK}(\mathbf{Z}; \mathbf{U}, \mathbf{V}, \mathbf{A}) = \sum_{i=1}^c \sum_{k=1}^N (\mu_{ik})^m D_{ik\mathbf{A}_i}^2 \quad (3.47)$$

où $\mathbf{U} \in \mathcal{R}^{c \times N}$, $\mathbf{V} \in \mathcal{R}^{n \times c}$ et $\mathbf{A} \in \mathcal{R}^{n \times n}$ correspondent aux solutions de la minimisation (points stationnaires) de la fonctionnelle J_{GK} . Pour $\mathbf{A}$ fixe, les conditions (3.42) et (3.43) peuvent être appliquées. Néanmoins, la fonction objectif (3.47) ne peut pas être minimisée directement par rapport à $\mathbf{A}_i$ puisque elle est linéaire en $\mathbf{A}_i$. Pour obtenir une solution réalisable, $\mathbf{A}_i$ doit être contrainte d'une certaine façon, pour cela, il suffit d'imposer une contrainte sur le déterminant de $\mathbf{A}_i$ [BAB98a]. En permettant à la matrice $\mathbf{A}_i$ de varier alors que son déterminant reste fixe, cela correspond à optimiser la forme du cluster tandis que son volume reste constant. C'est pourquoi on considère l'algorithme GK comme une méthode de clustering avec une mesure adaptative de la distance. La contrainte sur le volume se traduit par l'expression :

$$|\mathbf{A}_i| = \psi_i \quad (3.48)$$

Nous avons la possibilité d'affecter des volumes différents pour chaque cluster par l'intermédiaire de ψ_i , initialisé par défaut à la valeur 1. En utilisant la méthode des multiplicateurs de Lagrange [GUS79], la minimisation du critère d'optimisation conduit à l'expression de $\mathbf{A}_i$ suivante :

$$\mathbf{A}_i = (\psi_i \cdot \det(\mathbf{F}_i))^{1/n} \mathbf{F}_i^{-1} \quad (3.49)$$

où $\mathbf{F}_i$ est la *matrice de covariance floue* du i -ème cluster, définie par l'expression :

$$\mathbf{F}_i = \frac{\sum_{k=1}^N (\mu_{ik})^m (\mathbf{z}_k - \mathbf{v}_i)(\mathbf{z}_k - \mathbf{v}_i)^T}{\sum_{k=1}^N (\mu_{ik})^m} \quad (3.50)$$

La structure propre de la matrice de covariance floue $\mathbf{F}_i \in \mathfrak{R}^{n \times n}$ fournit une information précieuse sur la forme et l'orientation des clusters hyperellipsoïdaux. En effet, la racine carrée de chacune de ses valeurs propres $\lambda_j, j=1, \dots, n$ représente les longueurs des axes dans l'hyperespace. Les directions des axes sont données par les vecteurs propres ϕ_j . Cela est illustré sur la Figure 3.4 pour un espace de dimension $n = 2$:

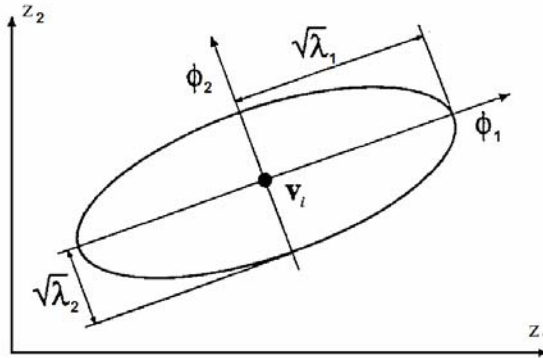


Figure 3.4. Interprétation géométrique de la matrice de covariance floue (d'après [BAB98a])

L'équation $(\mathbf{z} - \mathbf{v})^T \mathbf{F}^{-1} (\mathbf{z} - \mathbf{v}) = 1$ définit un hyperellipsoïde dans $\mathfrak{R}^n$.

La substitution des équations (3.49) et (3.50) dans l'expression (3.46) donne une norme quadratique de distance Mahalanobis généralisée entre le vecteur $\mathbf{z}_k$ et la moyenne (centre) du cluster $\mathbf{v}_i$, où la covariance est pondérée par les degrés d'appartenance en $\mathbf{U}$. La mesure de distance pour l'algorithme GK est définie alors comme :

$$D_{ik\mathbf{A}_i}^2 = (\mathbf{z}_k - \mathbf{v}_i)^T (\psi_i \det(\mathbf{F}_i))^{1/n} \mathbf{F}_i^{-1} (\mathbf{z}_k - \mathbf{v}_i) \quad (3.51)$$

L'algorithme de clustering de *Gustafson-Kessel* (GK) peut se formuler ainsi :

Algorithme GK - Gustafson-Kessel

Etant donné un ensemble de vecteurs de données Z :

Fixer le nombre de *clusters* $1 < c < N$

Fixer le degré de flou $m > 1$, le volume ψ_i des clusters
et la tolérance de fin d'algorithme $\varepsilon > 0$

Initialiser la matrice $\mathbf{U} = [\mu_{ik}]$ de partition floue

Répéter

Pas 1. Calculer les prototypes (centres) des clusters (Eq. (3.43))

Pas 2. Calculer les matrices de covariance de cluster (Eq. (3.50))

Pas 3. Calculer les distances pour chaque cluster (Eq. (3.51))

Pas 4. Mettre à jour la matrice $\mathbf{U}$ de partition floue (Eq. (3.42))

Jusqu'à obtenir la stabilité de la partition (Eq. (3.44))

Remarques

- R 3.5** En utilisant la structure propre de la matrice de covariance floue $\mathbf{F}_i$, il est possible de dériver des sous-espaces linéaires de données. En effet, ceux-ci peuvent être représentés par des hyperellipsoïdes plats, que l'on peut voir comme des hyperplans. Le vecteur propre qui correspond à la valeur propre la plus petite détermine la normale à l'hyperplan et peut aussi être utilisé pour calculer les sous modèles locaux linéaires. Cette approche sera développée dans la section 3.4 [BAB98a].
- R 3.6** Une avantage de l'algorithme GK sur le FCM consiste en ce que l'algorithme GK peut détecter des clusters de formes et d'orientations différentes pour un ensemble de données dans l'espace de représentation. Cependant, il est plus lourd en temps de calcul, car l'inverse et le déterminant de la matrice de covariance floue de cluster doivent être calculés à chaque itération.

Les deux algorithmes décrits jusqu'à présent se caractérisent par une représentation des clusters par l'intermédiaire des points *prototypes* (*centres*), $\mathbf{v}_i \in \mathcal{R}^n$: ce sont des structures géométriques du même "type" que les vecteurs de données. L'algorithme FCM utilise une norme de distance fixe et préfère ainsi fortement les clusters d'une forme géométrique induite par cette norme. L'algorithme GK remédie à cet inconvénient en adaptant localement la norme de distance. Une autre approche conceptuellement différente consiste à définir les prototypes comme des sous-espaces r -dimensionnels linéaires⁶⁴ ou non linéaires dans l'espace de représentation $\mathcal{R}^n$, où $1 \leq r \leq n$. Nous abordons cette approche dans le paragraphe suivant, en présentant l'algorithme FCRM.

3.3.4. Groupage avec prototypes linéaires (algorithme FCRM)

L'algorithme *Fuzzy c-Regression Models* (FCRM), proposé par Hathaway et Bezdek [HAT93] appartient à la gamme des algorithmes de clustering avec prototypes linéaires. Contrairement aux deux méthodes exposées précédemment, les prototypes des clusters ne sont pas des objets géométriques dans l'espace des données, mais ils sont définis explicitement par des relations fonctionnelles en termes d'équations de régression. Cet algorithme estime les paramètres du c modèles de régression conjointement avec une c -partition floue de l'ensemble des données. Les modèles de régression prennent la forme générale donnée par :

$$y_k = f_i(\mathbf{x}_k; \boldsymbol{\theta}_i) \quad (3.52)$$

où les fonctions f_i sont paramétrées par $\boldsymbol{\theta}_i \in \mathcal{R}^{p_i}$. Le degré d'appartenance $\mu_{ik} \in \mathbf{U}$ est interprété comme une pondération qui représente le domaine pour lequel la valeur prédite par le modèle $f_i(\mathbf{x}_k; \boldsymbol{\theta}_i)$ correspond à la sortie y_k . L'erreur de prédiction peut être définie par :

$$E_{ik}(\boldsymbol{\theta}_i) = (y_k - f_i(\mathbf{x}_k; \boldsymbol{\theta}_i))^2 \quad (3.53)$$

La famille des fonctions objectif définie pour les c modèles flous de régression, avec $1 \leq i \leq c$, et le jeu de données $Z = \{(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_N, y_N)\}$, est donnée par l'expression :

$$E(Z; \mathbf{U}, \{\boldsymbol{\theta}_i\}) = \sum_{i=1}^c \sum_{k=1}^N (\mu_{ik})^m E_{ik}(\boldsymbol{\theta}_i) \quad (3.54)$$

⁶⁴ C'est-à-dire : lignes ($r = 1$), plans ($r = 2$) et hyperplans ($2 < r < n$).

Une approche possible pour la minimisation de la fonction objectif (3.54) est la méthode de minimisation de coordonnée groupée [HAT91], dans laquelle la mise à jour de la matrice de partition floue est donnée par l'expression :

$$\mu_{ik} = \frac{1}{\sum_{j=1}^c (E_{ik} / E_{jk})^{2/(m-1)}}, \quad 1 \leq i \leq c, \quad 1 \leq k \leq N \quad (3.55)$$

L'algorithme général FCRM est alors formulé comme :

Algorithme FCRM - Fuzzy c -Regression Models

Etant donné un ensemble de données $Z = \{(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_N, y_N)\}$:

Fixer le nombre de *clusters* $1 < c < N$

Fixer la structure des modèles de régression (Eq. (3.52)) et la mesure de l'erreur de prédiction (Eq. (3.53))

Fixer le degré de flou $m > 1$ et la tolérance de fin d'algorithme $\varepsilon > 0$

Initialiser la matrice $\mathbf{U} = [\mu_{ik}]$ de partition floue

Répéter

Pas 1. Calculer des valeurs pour les paramètres du modèle θ_i qui minimisent globalement l'expression $E(Z; \mathbf{U}, \{\theta_i\})$ (Eq. (3.54))

Pas 2. Mettre à jour la matrice $\mathbf{U}$ de partition floue (Eq. (3.55))

Jusqu'à obtenir la stabilité de la partition (Eq. (3.44))

L'algorithme FCRM présente l'inconvénient que les clusters ne sont pas limités dans leur taille, ce qui génère une tendance à connecter des clusters qui peuvent être bien séparés. Par conséquent, l'algorithme est plutôt sensible à l'initialisation. L'avantage de cet algorithme est qu'il peut aussi ajuster des modèles locaux non linéaires aux données, comme des polynômes, qui sont toujours linéaires en leurs paramètres. Cela conduit à un problème d'estimation linéaire dans le Pas 1 de l'algorithme [BAB98a].

En fait, quand les fonctions de régression f_i dans l'expression (3.52) sont linéaires par rapport aux paramètres θ_i , i.e., $f_i(\mathbf{x}_k; \theta_i) = \mathbf{a}_i^T \mathbf{x} + d_i$, ceux-ci peuvent être obtenus alors comme une solution d'un problème des moindres carrés pondérés où les degrés d'appartenance de la matrice de partition floue $\mathbf{U}$ servent comme facteurs de pondération (cf. paragraphe 3.4.3, concernant l'obtention des paramètres des conséquents du modèle Takagi-Sugeno). En choisissant la matrice et le vecteur de données (entrée et sortie, respectivement) $\mathbf{X} \in \mathbb{R}^{N \times p}$ et $\mathbf{y} \in \mathbb{R}^N$, ainsi que la matrice de pondération $\mathbf{W}_i \in \mathbb{R}^{N \times N}$ tel que :

$$\mathbf{X} = \begin{bmatrix} \mathbf{x}_1^T \\ \mathbf{x}_2^T \\ \vdots \\ \mathbf{x}_N^T \end{bmatrix}, \quad \mathbf{y} = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_N \end{bmatrix}, \quad \mathbf{W}_i = \begin{bmatrix} \mu_{i1} & 0 & \cdots & 0 \\ 0 & \mu_{i2} & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \mu_{iN} \end{bmatrix} \quad (3.56)$$

Alors les paramètres optimaux pour chacun des sous modèles linéaires sont calculés par l'expression :

$$\theta_i = (\mathbf{X}_e^T \mathbf{W}_i \mathbf{X}_e)^{-1} \mathbf{X}_e^T \mathbf{W}_i \mathbf{y} \quad (3.57)$$

3.3.5. Groupage avec régression ellipsoïdale (algorithme FER)

Dans ce paragraphe nous introduisons l'algorithme de groupage avec régression ellipsoïdale, en anglais *Fuzzy Ellipsoidal Regression* (FER⁶⁵) [BAR03] [GRI04], dans lequel la mesure de distance est définie comme une combinaison de la version d'axes parallèles de l'algorithme GK et de l'algorithme FCRM. Dans le cadre de la modélisation des systèmes, l'algorithme proposé est orienté afin de favoriser l'interprétabilité du modèle flou résultant.

Dans l'algorithme GK conventionnel, les clusters hyperellipsoïdaux obtenus permettent de tenir compte de la forme propre à chacune des classes par rapport aux clusters hypersphériques issus de l'utilisation de l'algorithme FCM. Il y a des améliorations possibles en relation avec l'interprétabilité des règles du modèle flou Takagi-Sugeno associé. En effet, une fois que les données sont regroupées au niveau des clusters, il faut procéder à la construction du modèle flou correspondant, en générant les règles constitutives de l'antécédent et du conséquent (voir le paragraphe 3.4.2). Si l'on souhaite, exprimer les règles sous une forme *conjonctive*, une *erreur de décomposition* apparaît lors d'une utilisation du *mécanisme de projection* compte tenu de la rotation naturelle des clusters par rapport aux axes des entrées, à partir des ensembles flous multidimensionnelles. Cette rotation des clusters est tout à fait en concordance avec la définition de la matrice $\mathbf{A}_i$ de norme induite dans l'expression (3.49), car cette matrice fait intervenir la matrice $\mathbf{F}_i$ de covariance floue (3.50), qui est une matrice symétrique avec des coefficients hors de la diagonale principale. Pour pallier cet inconvénient et réduire l'erreur de décomposition, une possibilité consiste à imposer des clusters hyperellipsoïdaux *parallèles* aux axes des entrées dans l'espace de représentation. Cela se traduit par une matrice $\mathbf{A}_i$ diagonale, notée $\mathbf{A}_{id}$ et exprimée par convenance sous la forme suivante :

$$\mathbf{A}_{id} = \begin{bmatrix} \mathbf{A}_{ixd} & 0 \\ 0 & \sigma_i \end{bmatrix} \quad (3.58)$$

dans laquelle la matrice $\mathbf{A}_{ixd}$ et le terme σ_i correspondent respectivement aux composants de l'espace d'entrée et de sortie. Ainsi, la fonction objectif pour l'algorithme GK avec des axes parallèles (GKP) peut se exprimer par :

$$J_{GKP}(Z; \mathbf{U}, \mathbf{V}, \mathbf{A}) = \sum_{i=1}^c \sum_{k=1}^N (\mu_{ik})^m (\mathbf{z}_k - \mathbf{v}_i)^T \mathbf{A}_{id} (\mathbf{z}_k - \mathbf{v}_i) \quad (3.59)$$

ou, en tenant compte des composants d'entrée et sortie, par :

$$J_{GKP}(Z; \mathbf{U}, \mathbf{V}, \mathbf{A}) = \sum_{i=1}^c \sum_{k=1}^N (\mu_{ik})^m \left((\mathbf{x}_k - \mathbf{v}_{ix})^T \mathbf{A}_{ixd} (\mathbf{x}_k - \mathbf{v}_{ix}) + \sigma_i (y_k - v_{iy})^2 \right) \quad (3.60)$$

L'ensemble des données est vu sous la forme $\mathbf{z}_k = [\mathbf{x}_k^T, y_k]^T$, dans laquelle $\mathbf{x}_k$ et y_k sont un vecteur et une constante qui contiennent respectivement les composants d'entrée et de sortie. D'une façon analogue, $\mathbf{v}_i = [\mathbf{v}_{ix}^T, v_{iy}]^T$, où $\mathbf{v}_{ix}$ et v_{iy} sont un vecteur et une constante contenant les composants d'entrée et de sortie du centre de l' i -ème cluster.

⁶⁵ Le nom FER (*Fuzzy Ellipsoidal Regression*) introduit par Grisales et utilisé pour la première fois dans le cadre de cette thèse, fait appel à la double interprétation géométrique de l'algorithme : les clusters hyperellipsoïdaux parallèles dans l'espace d'entrée obtenus avec la mesure de distance de l'algorithme GK et d'autre part, la prise en compte de prototypes hyperplanaires au niveau du calcul des conséquents, d'une manière similaire à l'algorithme FCRM. En résumé, FER = GK Parallèle dans $\mathfrak{R}^p$ + FCRM dans $\mathfrak{R}^{p+1}$.

Dans la fonctionnelle donnée par l'expression (3.60) apparaissent des clusters hyperellipsoïdaux d'axes parallèles par rapport aux variables d'entrée. Afin de générer des sous modèles linéaires, le terme associé à la sortie est modifié, ce qui donne comme fonction objectif pour l'algorithme FER, la forme suivante :

$$J_{FER}(Z; \mathbf{U}, \mathbf{V}, \mathbf{A}, \boldsymbol{\theta}) = \sum_{i=1}^c \sum_{k=1}^N (\mu_{ik})^m \left((\mathbf{x}_k - \mathbf{v}_{ix})^T \mathbf{A}_{ixd} (\mathbf{x}_k - \mathbf{v}_{ix}) + \sigma_i (y_k - f_i(\mathbf{x}_k; \boldsymbol{\theta}_i))^2 \right) \quad (3.61)$$

où $f_i(\mathbf{x}_k; \boldsymbol{\theta}_i) = \mathbf{a}_i^T \mathbf{x} + d_i$ représente chacun des i -ème sous modèles linéaires dans $\mathcal{R}^{p+1}$. Notons que le terme associé à la variable de sortie a la même forme que celui correspondant à l'erreur de prédiction dans l'algorithme FCRM (cf. expression (3.53)). D'autre part, la contrainte (3.48) associée à l'hypervolume du cluster s'exprime comme :

$$\mathbf{A}_{id} = |\mathbf{A}_{ixd}| \sigma_i = \psi_i \quad (3.62)$$

Autrement l'algorithme conduirait à la solution évidente $\mathbf{A}_{ixd} = 0$ et $\sigma_i = 0$. La fonctionnelle est également soumise à la contrainte de partition floue stricte donnée par (3.37).

Dans la suite, afin de simplifier la notation dans la procédure de minimisation, nous récrivons l'expression de la fonctionnelle J_{FER} proposée, de la manière suivante :

$$J_{FER}(Z; \mathbf{U}, \mathbf{V}, \mathbf{A}, \boldsymbol{\theta}) = \sum_{i=1}^c \sum_{k=1}^N (\mu_{ik})^m d_{ik}^2 \quad (3.63)$$

avec

$$d_{ik}^2 = \left([\mathbf{x}_k, y_k]^T - [\mathbf{v}_{ix}, f_i(\mathbf{x}_k; \boldsymbol{\theta}_i)]^T \right)^T \mathbf{A}_{id} \left([\mathbf{x}_k, y_k]^T - [\mathbf{v}_{ix}, f_i(\mathbf{x}_k; \boldsymbol{\theta}_i)]^T \right) \quad (3.64)$$

Pour minimiser la fonctionnelle (3.63) on utilise une procédure itérative similaire à celle décrite pour l'algorithme FCM à travers les conditions du premier ordre. Les points stationnaires de la fonction objectif peuvent être trouvés en ajoutant les contraintes d'égalité (3.37) et (3.62) à J_{FER} au moyen des multiplicateurs de Lagrange, et l'on obtient :

$$\bar{J}_{FER}(Z; \mathbf{U}, \mathbf{P}, \boldsymbol{\lambda}, \boldsymbol{\beta}) = \sum_{i=1}^c \sum_{k=1}^N (\mu_{ik})^m d_{ik}^2 + \sum_{k=1}^N \lambda_k \left[\sum_{i=1}^c \mu_{ik} - 1 \right] + \sum_{i=1}^c \beta_i \left[|\mathbf{A}_{id}| - \psi_i \right] \quad (3.65)$$

où $\mathbf{P} = [\mathbf{V}, \mathbf{A}, \boldsymbol{\theta}]$ représente l'ensemble des matrices associées aux paramètres qui caractérisent les clusters et les termes λ_k , β_i sont les coefficients multiplicateurs de Lagrange. La minimisation du critère J_{FER} s'obtient alors en annulant le gradient de $\bar{J}_{FER}$ par rapport $\mathbf{U}$, $\mathbf{P}$, $\boldsymbol{\lambda}$ et $\boldsymbol{\beta}$. Etant donnée que les colonnes de la matrice de partition floue $\mathbf{U}$ sont indépendants les uns des autres, on peut réduire le problème de minimisation à chacune des observations (colonnes), i.e., à N problèmes d'optimisation indépendants. Ainsi, en annulant initialement le gradient de $\bar{J}_{FER}$ par rapport aux termes correspondants en $\mathbf{U}$ et $\boldsymbol{\lambda}$ on obtient respectivement les expressions :

$$\frac{\partial \bar{J}_{FER}}{\partial \mu_{ik}} = m(\mu_{ik})^{m-1} d_{ik}^2 + \lambda_k = 0, \quad \forall i, k \quad (3.66)$$

$$\frac{\partial \bar{J}_{FER}}{\partial \lambda_k} = \sum_{i=1}^c \mu_{ik} - 1 = 0, \quad \forall k \quad (3.67)$$

L'expression (3.66) permet alors d'exprimer le terme μ_{ik} de la manière suivante :

$$\mu_{ik} = \left(-\frac{\lambda_k}{m \cdot d_{ik}^2} \right)^{1/(m-1)} \quad (3.68)$$

en le remplaçant dans l'équation (3.67) on arrive à :

$$\left(-\frac{\lambda_k}{m} \right)^{1/(m-1)} = \frac{1}{\sum_{j=1}^c \left(\frac{1}{d_{jk}^2} \right)^{1/(m-1)}} \quad (3.69)$$

et donc en combinant les expressions (3.68) et (3.69) on obtient :

$$\mu_{ik} = \frac{(1/d_{ik}^2)^{1/(m-1)}}{\sum_{j=1}^c (1/d_{jk}^2)^{1/(m-1)}} \quad (3.70)$$

Cette expression traduit ainsi la relation de mise à jour des coefficients d'appartenance optimales (μ_{ik}^*) de la matrice de partition floue, qui peut encore s'exprimer sous la forme :

$$\mu_{ik}^* = \frac{1}{\sum_{j=1}^c (d_{ik}^2 / d_{jk}^2)^{1/(m-1)}}, \quad 1 \leq i \leq c, \quad 1 \leq k \leq N. \quad (3.71)$$

Pour trouver l'ensemble optimal des paramètres $\mathbf{P} = [\mathbf{V}, \mathbf{A}, \boldsymbol{\theta}]$ qui caractérisent chacun des clusters, il faut également calculer et annuler le gradient de $\bar{J}_{FER}$ par rapport à $\mathbf{P}$:

$$\frac{\partial \bar{J}_{FER}}{\partial \mathbf{P}_i} = \sum_{k=1}^N (\mu_{ik}^*)^m \frac{\partial d_{ik}^2}{\partial \mathbf{P}_i} = 0, \quad \forall i \quad (3.72)$$

Pour les centres des clusters dans l'espace d'entrée $\mathfrak{R}^p$, l'application de la condition nécessaire représentée par (3.72) conduit à l'expression :

$$\frac{\partial \bar{J}_{FER}}{\partial \mathbf{V}_i} = -2 \sum_{k=1}^N (\mu_{ik}^*)^m \mathbf{A}_{id} (\mathbf{x}_k - \mathbf{v}_{ix}^*) = 0, \quad \forall i \quad (3.73)$$

et donc on obtient

$$\mathbf{v}_{ix}^* = \frac{\sum_{k=1}^N (\mu_{ik}^*)^m \cdot \mathbf{x}_k}{\sum_{k=1}^N (\mu_{ik}^*)^m}, \quad 1 \leq i \leq c. \quad (3.74)$$

Cette expression permet l'actualisation des centres des clusters hyperellipsoïdaux dans l'espace d'entrée de dimension $\mathfrak{R}^p$.

A partir de l'expression (3.72), en annulant le gradient de $\bar{J}_{FER}$ par rapport à la matrice $\mathbf{A}_{id}$, la minimisation du critère J_{FER} conduit à l'expression suivante :

$$\frac{\partial \bar{J}_{FER}}{\partial \mathbf{A}_{id}} = \text{diag} \left(\sum_{k=1}^N (\mu_{ij})^m ([\mathbf{x}_k, y_k]^T - [\mathbf{v}_{ix}, f_i(\mathbf{x}_k; \boldsymbol{\theta}_i)]) \cdot ([\mathbf{x}_k, y_k]^T - [\mathbf{v}_{ix}, f_i(\mathbf{x}_k; \boldsymbol{\theta}_i)])^T \right) + \beta_i |\mathbf{A}_{id}^*| (\mathbf{A}_{id}^*)^{-1} = 0 \quad (3.75)$$

dans laquelle nous avons utilisées les propriétés suivantes [GUS79] :

$$\frac{\partial (\mathbf{x}^T \mathbf{A} \mathbf{x})}{\partial \mathbf{A}} = \mathbf{x}^T \mathbf{A} \mathbf{x}, \quad \frac{\partial |\mathbf{A}|}{\partial \mathbf{A}} = |\mathbf{A}| \mathbf{A}^{-1} \quad (3.76)$$

valables pour des matrices $\mathbf{A}$ non singulières et tout vecteur compatible $\mathbf{x}$. Notons également que pour notre cas particulier, la matrice $\mathbf{A}_{id}$ est une matrice diagonale.

Ainsi, à partir de l'expression (3.75), on obtient alors :

$$(\mathbf{A}_{id}^*)^{-1} = \frac{1}{\beta_i |\mathbf{A}_{id}^*|} \text{diag} \left(\sum_{k=1}^N (\mu_{ik})^m ([\mathbf{x}_k, y_k]^T - [\mathbf{v}_{ix}, f_i(\mathbf{x}_k; \boldsymbol{\theta}_i)]) \cdot ([\mathbf{x}_k, y_k]^T - [\mathbf{v}_{ix}, f_i(\mathbf{x}_k; \boldsymbol{\theta}_i)])^T \right) \quad (3.77)$$

D'autre part, il est évident que pour chaque cluster,

$$\frac{\partial \bar{J}_{FER}}{\partial \beta_i} = |\mathbf{A}_{id}| - \psi_i = 0 \quad (3.78)$$

Si maintenant on définit une matrice diagonale de la forme

$$\mathbf{F}_{id} = \frac{\text{diag} \left(\sum_{k=1}^N (\mu_{ik})^m ([\mathbf{x}_k, y_k]^T - [\mathbf{v}_{ix}, f_i(\mathbf{x}_k; \boldsymbol{\theta}_i)]) \cdot ([\mathbf{x}_k, y_k]^T - [\mathbf{v}_{ix}, f_i(\mathbf{x}_k; \boldsymbol{\theta}_i)])^T \right)}{\sum_{k=1}^N (\mu_{ik})^m} \quad (3.79)$$

alors en remplaçant cette expression et la condition (3.78) dans l'équation (3.77), on arrive finalement à :

$$(\mathbf{A}_{id}^*)^{-1} = \left(\frac{1}{\psi_i |\mathbf{F}_{id}|} \right)^{1/n} \mathbf{F}_{id} \quad (3.80)$$

En décomposant la matrice décrite par l'équation (3.79) sous la forme diagonale donnée par l'expression (3.58) on obtient la matrice de covariance dans l'espace d'entrée :

$$\mathbf{F}_{ixd} = \text{diag} \left(\frac{\sum_{k=1}^N (\mu_{ik})^m (\mathbf{x}_k - \mathbf{v}_{ix})(\mathbf{x}_k - \mathbf{v}_{ix})^T}{\sum_{k=1}^N (\mu_{ik})^m} \right) \quad (3.81)$$

et la composante de sortie de la matrice de covariance :

$$f_{iy} = \frac{\sum_{k=1}^N (\mu_{ik})^m (y_k - f_i(\mathbf{x}_k; \boldsymbol{\theta}_i))^2}{\sum_{k=1}^N (\mu_{ik})^m} \quad (3.82)$$

En remplaçant les équations (3.81) et (3.82) dans l'expression (3.80) on obtient alors :

$$\mathbf{A}_{ixd}^* = \left(\psi_i \mathbf{f}_{iy} | \mathbf{F}_{ixd} | \right)^{1/n} (\mathbf{F}_{ixd})^{-1} \quad (3.83)$$

et

$$\sigma_i^* = \left(\psi_i \mathbf{f}_{iy} | \mathbf{F}_{ixd} | \right)^{1/n} (\mathbf{f}_{iy})^{-1} \quad (3.84)$$

Finalement, pour minimiser la fonction objectif J_{FER} par rapport aux paramètres $\boldsymbol{\theta}_i = [\mathbf{a}_i^T, d_i]^T$ correspondants aux sous modèles linéaires associés aux clusters, on arrive à une formulation de type moindres carrés pondérés (WLS) (cf. paragraphe 3.3.4) dont la solution est donnée par l'expression (3.57). Dans ce cas, la matrice diagonale de pondération est donnée par l'expression suivante :

$$\mathbf{W}_i = \begin{bmatrix} \mu_{i1} & 0 & \cdots & 0 \\ 0 & \mu_{i2} & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \mu_{iN} \end{bmatrix} \quad (3.85)$$

L'algorithme de clustering avec régression floue ellipsoïdale (FER) peut se formuler de la manière suivante :

Algorithme FER - Fuzzy Ellipsoidal Regression

Etant donné un ensemble de données $Z = \{(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_N, y_N)\}$:

Fixer le nombre de *clusters* $1 < c < N$

Fixer la structure des modèles de régression $f_i(\mathbf{x}_k; \boldsymbol{\theta}_i) = \mathbf{a}_i^T \mathbf{x}_k + d_i$

Fixer le degré de flou $m > 1$ et la tolérance de fin d'algorithme $\varepsilon > 0$

Initialiser la matrice $\mathbf{U} = [\mu_{ik}]$ de partition floue

Répéter

Pas 1. Calculer les centres des clusters $\mathbf{v}_{ix}$ dans l'espace d'entrée (Eq. (3.74))

Pas 2. Calculer les matrices de covariance de cluster $\mathbf{F}_{ixd}$ dans l'espace d'entrée en utilisant (Eq. (3.81))

Pas 3. Calculer les matrices de pondération $\mathbf{W}_i$ (Eq. (3.85)) et les paramètres des sous modèles linéaires $\boldsymbol{\theta}_i = [\mathbf{a}_i^T, d_i]^T$ pour chaque cluster (Eq. (3.57))

Pas 4. Calculer les facteurs $\mathbf{f}_{iy}$ (composants de la matrice $\mathbf{F}_{id}$) associées à chacune des sorties (Eq. (3.82))

Pas 5. Calculer les matrices diagonales de norme induite $\mathbf{A}_{id}$ par l'intermédiaire des expressions $\mathbf{A}_{ixd}$ (Eq. (3.83)) et σ_i (Eq. (3.84))

Pas 6. Calculer les distances d_{ik}^2 pour $1 \leq i \leq c$ et $1 \leq k \leq N$ (Eq. (3.64))

Pas 7. Mettre à jour la matrice $\mathbf{U}$ de partition floue (Eq. (3.71))

Jusqu'à obtenir la stabilité de la partition (Eq. (3.44))

3.3.6. Groupage robuste en présence de bruit (algorithme RoFER)

Jusqu'à présent, nous avons considéré des algorithmes de clustering flou qui d'une part, nécessitent l'établissement à l'avance du nombre de clusters et d'autre part, qui sont orientés pour le traitement des données non bruitées dans l'espace de représentation. Ces aspects constituent deux problèmes majeurs du "monde réel" dans le cadre des méthodes de coalescence floue par partition. Dans ce paragraphe nous présentons une version améliorée de l'algorithme FER, nommé algorithme robuste de groupage flou avec régression ellipsoïdale, en anglais *Robust Fuzzy Ellipsoidal Regression* (RoFER) [BAR03] [GRI04]. L'algorithme est basé sur l'approche « robuste d'agglomération compétitive » (RCA) utilisée en vision artificielle et reconnaissance de formes [FRI99] et combine la définition de la mesure de distance utilisée pour l'algorithme FER. L'algorithme RoFER permet non seulement d'établir "automatiquement" le nombre de clusters *linéaires* dans un but de modélisation floue, mais aussi d'avoir des caractéristiques de robustesse par rapport au bruit et aux points aberrants (*outliers*). Dans cette approche on surestime initialement le nombre de clusters puis l'algorithme tend vers la valeur "optimale" de ce nombre lors de son processus itératif.

L'algorithme de clustering *Robust Competitive Agglomeration* (RCA), proposé par Frigui et Krishnapuram [FRI96b] [FRI96a] [FRI99], minimise une fonction objectif basée sur des prototypes flous et incorpore le concept de fonctions de poids de la *statistique robuste* [HUB81] pour réduire les effets des points aberrants ainsi que le concept d'*agglomération compétitive* [FRI97] afin de déterminer un nombre approprié de clusters pour la partition. L'algorithme combine les avantages des techniques de coalescence floue hiérarchique et par partition. La fonction objectif inclut un composant d'agglomération afin d'incorporer les bénéfices du clustering hiérarchique, ainsi qu'une contrainte sur les degrés d'appartenance afin de favoriser la compétition parmi les clusters. L'algorithme génère ainsi une séquence de partitions avec un nombre de clusters qui diminue progressivement. Quand l'algorithme se termine on obtient une partition finale qui possède le nombre "optimal" de clusters. L'algorithme commence par un grand nombre de petits clusters, comparativement à d'autres algorithmes traditionnels de clustering par partition, il a été montré que le résultat final est moins sensible à l'initialisation et aux minimum locaux [FRI97].

D'autre part, étant donné que l'algorithme RCA est seulement un algorithme de clustering, il n'a pas la capacité de régression pour la détermination directe des paramètres des fonctions dans les termes conséquents du modèles flous TS. Dans le cadre de la modélisation floue des systèmes, différentes propositions ont été développées afin de pallier cet inconvénient [CHU01] [REN03]. Pour sa part, l'algorithme RoFER⁶⁶ permet, en incorporant la mesure de distance proposée dans l'algorithme FER, de déterminer simultanément avec une approche robuste d'agglomération compétitive, une partition floue pour l'espace des antécédents et en même temps, les paramètres des fonctions linéaires dans les termes conséquents des modèles flous affines du type Takagi-Sugeno.

Etant donné un ensemble de données $Z = \{(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_N, y_N)\}$, l'algorithme RoFER minimise la fonction objectif définie par l'expression :

⁶⁶ Antérieurement, l'algorithme a été initialement cité sous le nom RPCA (Robust Parallel Competitive Agglomerative) [GRI05]. Le nom RoFER (*Robust Fuzzy Ellipsoidal Regression*) a été introduit par Grisales dans le cadre de cette thèse et c'est en fait, le nom retenu au cours du développement de la boîte à outils graphique FMIDg sous Matlab pour la modélisation et l'identification floue TS de systèmes [GRI07].

$$J_{RoFER}(Z; \mathbf{U}, \mathbf{P}) = \sum_{i=1}^c \sum_{k=1}^N (\mu_{ik})^2 \rho_i(d_{ik}^2) - \alpha \sum_{i=1}^c \left(\sum_{k=1}^N w_{ik} \mu_{ik} \right)^2 \quad (3.86)$$

Soumise à la contrainte de normalisation d'inspiration probabiliste $\sum_{i=1}^c \mu_{ik} = 1$ donné par l'expression (3.37) et dans laquelle la mesure de distance d_{ik}^2 correspond à celle de l'algorithme FER (GKP+FCRM) établie par l'équation (3.64). Ici, $\mathbf{P} = [\mathbf{V}, \mathbf{A}, \mathbf{\theta}]$ représente l'ensemble des matrices associées aux paramètres qui caractérisent chacun des clusters (centres de clusters dans $\mathfrak{R}^p$, matrices de norme induite et jeu des paramètres des sous modèles linéaires associé dans $\mathfrak{R}^{p+1}$). Notons également que grâce au processus d'agglomération, le nombre de clusters c dans la fonction objectif (3.86) est mise à jour dynamiquement dans l'algorithme RoFER.

De la même façon que dans l'algorithme RCA⁶⁷, l'algorithme RoFER présente la particularité d'utiliser simultanément deux ensembles de poids ayant deux interprétations en termes d'appartenance aux classes. Le premier est constitué des coefficients d'appartenance classiques μ_{ik} (approche probabiliste) qui traduisent le degré d'appartenance des observations aux classes. Le second est formé des possibilités d'appartenance w_{ik} (approche possibiliste) [KRI93] [DAV91] qui traduisent la "typicité" des observations aux classes. Les coefficients μ_{ik} permettent de forcer la partition en un nombre c de classes dans l'espace de représentation, alors que les coefficients w_{ik} visent à améliorer la modélisation de celles-ci en conférant aux observations éloignées une pondération réduite qui dépend de la distance ($w_{ik} = 0$ pour les observations aberrantes).

L'idée fondamentale d'incorporer des fonctions de la statistique robuste⁶⁸ au niveau de la définition de la fonctionnelle au lieu de faire le calcul en utilisant directement la mesure de distance quadratique, a pour but de réduire les effets nocifs du bruit et des points aberrants en permettant une estimation *robuste* des paramètres. Pour cela, le critère d'optimisation J_{RoFER} est défini en utilisant une fonction de perte ρ_i propre à chaque cluster c , avec $1 \leq i \leq c$, ainsi qu'une fonction de poids associée, définie comme :

$$w_{ik} = w_i(d_{ik}^2) = \partial \rho_i(d_{ik}^2) / \partial d_{ik}^2 \quad (3.87)$$

où w_{ik} représente la typicité d'un point $\mathbf{z}_k = (\mathbf{x}_k, y_k)$, $1 \leq k \leq N$, par rapport au cluster i .

La fonction objectif (3.86) dépend de deux termes, le premier terme est associé à la fonction de perte ρ_i . Le minimum global de ce terme est atteint quand le nombre de clusters c est égal au nombre de "bons" points échantillon (i.e., chaque cluster contient un seul bon point de données). Les degrés d'appartenance μ_{ik} sont utilisés dans ce terme pour générer la partition floue des données. De plus, la contrainte de normalisation permet que les clusters partagent les points, en favorisant en même temps un environnement de *compétition* [FRI96b]. Le second terme du critère dans l'expression (3.86) est utilisé pour maximiser le nombre de "bons" points dans chaque cluster. Les poids w_{ik} (à déterminer) représentent les degrés d'affinité ou de typicité des points par rapport aux prototypes des clusters et ils sont

⁶⁷ Inspiré de l'algorithme *Robust c-Prototypes* (RCP) [FRI96a].

⁶⁸ Dans l'expression (3.86) le terme $\rho_i(\cdot)$ correspond à la fonction de perte utilisée dans les M-estimateurs de la statistique robuste et $w_i(\cdot)$ représente la fonction de poids (influence) correspondante [FRI99] [HUB81].

utilisés pour générer des estimations robustes des paramètres des clusters. La fonction de poids w_i réduit l'effet des points aberrants au moyen du calcul des cardinalités⁶⁹ robustes. Le minimum global de ce terme (en incluant le signe moins) est atteint quand tous les points sont concentrés dans un cluster et tous les autres clusters sont vides. Ce terme est associé alors à l'*agglomération* des clusters. Quand les deux termes sont combinés avec un choix approprié du paramètre d'agglomération α , la partition finale minimisera la somme des distances intra-cluster en obtenant une partition floue de l'ensemble des données avec un nombre réduit de clusters [FRI96b].

Choix des fonctions robustes ρ_i et w_i

Diverses fonctions robustes de perte (ρ_i) et de poids (w_i) ont été utilisées pour des algorithmes robustes de clustering dans la littérature [DAV97]. Pour l'approximation de courbes ou dans le cas de la régression linéaire, il est raisonnable de supposer que les résidus ont une distribution symétrique autour de zéro. Néanmoins, quand des distances positives sont utilisées (plutôt que des résidus), cette supposition n'est plus valable. Dans ce cas, des auteurs comme Frigui et Krishnapuram [FRI99] ont proposé l'utilisation d'une fonction de perte ρ_i et des pondérations w_{ik} qui traduisent pour chaque classe c une fonction de pondération $w_i(d^2) : \mathbb{R}^+ \rightarrow [0, 1]$ qui décroît avec la distance au prototype des classes jusqu'à atteindre la valeur zéro à partir d'un certain seuil. Cette fonction de poids est définie de façon à respecter les propriétés suivantes :

- $w_i(d^2)$ est monotone décroissante,
- $w_i(d^2) = 0$ pour $d^2 > T_i + \delta S_i$,
- $w_i(0) = 1$, $w_i(T_i) = 0,5$ et $w_i'(0) = 0$.

Dans les expressions précédentes δ est un facteur constant et les termes T_i et S_i correspondent respectivement à deux estimateurs couramment employés : la valeur médiane (Med) et l'écart moyen absolu (MAD⁷⁰) des distances entre les observations affectées à chaque cluster et son prototype respectif, selon les expressions :

$$T_i = \text{Med}_i(d_{ik}^2) \quad (3.88)$$

$$S_i = \text{MAD}_i(d_{ik}^2) = \text{Med}_i(|d_{ik}^2 - \text{Med}_i|) \quad (3.89)$$

A chaque itération du processus d'apprentissage, l'estimation des valeurs de T_i et S_i pour chaque cluster c conditionne l'allure des fonctions de pondération.

Un exemple de couple de fonctions robustes de perte $\rho_i(d_{ik}^2)$ et de poids $w_i(d_{ik}^2) = \partial \rho_i(d_{ik}^2) / \partial d_{ik}^2$ remplissant les différentes contraintes citées ci-dessus est défini au travers des expressions suivantes [FRI99] :

⁶⁹ Rappelons que la cardinalité d'un ensemble flou $A = \{\mu_A(x_j)/x_j \mid j=1, \dots, n\}$, notée par Σ_A est définie comme la somme de ses degrés d'appartenance respectifs : $\Sigma_A = \sum_{j=1}^n \mu_A(x_j)$.

⁷⁰ MAD : *Mean Absolute Deviation*. Moyenne arithmétique de la valeur absolue des écarts de chacune des données d'une série d'observations par rapport aux prévisions.

Pour la fonction de poids on a :

$$w_i(d^2) = \begin{cases} 1 - \frac{d^4}{2T_i^2} & \text{si } d^2 \in [0, T_i] \\ \frac{(d^2 - (T_i + \delta S_i))^2}{2\delta^2 S_i^2} & \text{si } d^2 \in (T_i, T_i + \delta S_i] \\ 0 & \text{si } d^2 > T_i + \delta S_i \end{cases} \quad (3.90)$$

Il est simple de démontrer que la fonction de perte associée est définie par :

$$\rho_i(d^2) = \begin{cases} d^2 - \frac{d^6}{6T_i^2} & \text{si } d^2 \in [0, T_i] \\ \frac{(d^2 - (T_i + \delta S_i))^3}{6\delta^2 S_i^2} + \frac{5T_i + \delta S_i}{6} & \text{si } d^2 \in (T_i, T_i + \delta S_i] \\ \frac{5T_i + \delta S_i}{6} + K_i & \text{si } d^2 > T_i + \delta S_i \end{cases} \quad (3.91)$$

Dans l'expression (3.91) le facteur K_i est une constante d'intégration définie pour chaque i -ème cluster et calculée selon l'expression :

$$K_i = \max_{1 \leq j \leq c} \left\{ \frac{5T_j + \delta S_j}{6} \right\} - \frac{5T_i + \delta S_i}{6}, \quad 1 \leq i \leq c \quad (3.92)$$

Cette constante permet d'éviter l'affectation d'observations bruitées à une classe particulière en forçant chaque fonction ρ_i à la même valeur maximale.

La Figure 3.5 montre de telles fonctions robustes de poids w_i et de perte ρ_i associées à chacun des clusters [FRI99] :

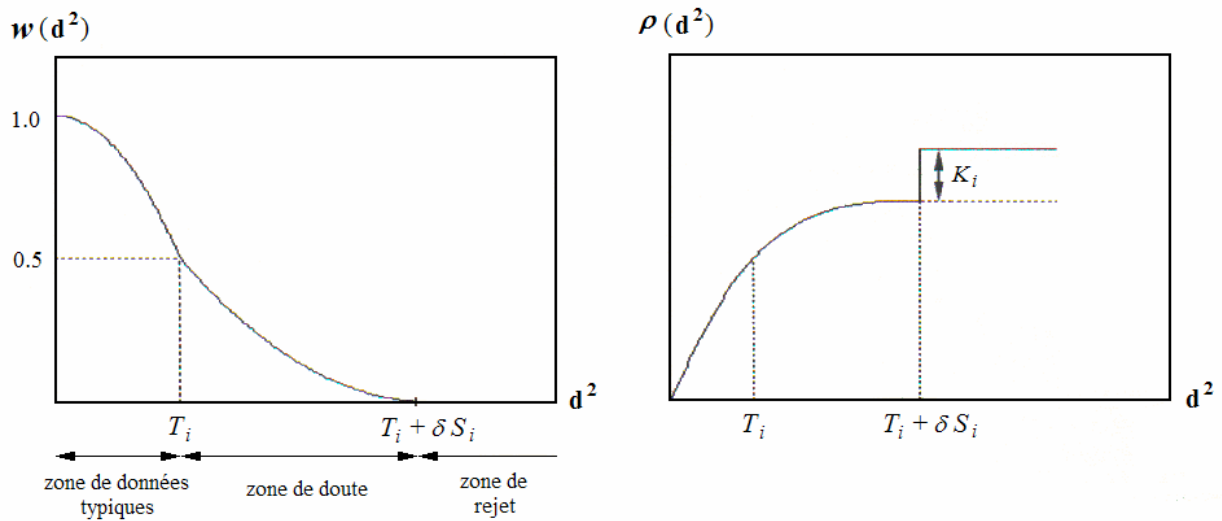


Figure 3.5. Illustration des fonctions robustes w_i et ρ_i associées à chacun des clusters.

Les trois zones repérées sur la Figure 3.5 correspondent respectivement à une zone de "données typiques" où les poids (possibilités) w_{ik} d'appartenance au i -ème cluster sont supérieures à 0,5, une zone intermédiaire de doute où les possibilités d'appartenance au cluster sont inférieures à 0,5 et une zone de rejet qui pénalise les observations trop éloignées du prototype du cluster. Il est clair que pour les points aberrants ce poids est $w_{ik} = 0$.

Par ailleurs, le coefficient δ dans les expressions (3.90) à (3.92) est un facteur de réglage qui permet d'agir sur la zone d'influence des fonctions de pondération. Celui-ci prend généralement une valeur entre 4 et 12. Si δ est trop grand, les fonctions de pondération w_i ont une région intermédiaire plus dilatée, c'est-à-dire que les observations éloignées ont des poids faibles mais non nuls ce qui perturbe l'estimation des paramètres des classes. D'un autre côté, si δ est trop petit, peu d'observations sont considérées au niveau du processus d'estimation des paramètres et l'algorithme risque de converger prématurément vers un minimum local. Pour pallier cette difficulté, [FRI99] considère qu'il faut établir un compromis, en initialisant δ avec une valeur importante puis faire décroître ce facteur lors des itérations suivantes. Ceci a pour effet de permettre une plus forte mobilité lors des premières itérations, puis au cours des dernières itérations de stabiliser les estimations des paramètres des clusters. La fonction de décroissance du coefficient δ entre deux itérations consécutives est alors établie de la façon suivante :

$$\delta^{(t)} = \max(\delta_{\min}, \delta^{(t-1)} - \Delta\delta) \quad (3.93)$$

avec $\{\delta^{(0)} = 12, \delta_{\min} = 4, \Delta\delta = 1\}$.

Mise à jour du paramètre d'agglomération α

Comme nous l'avons déjà indiqué, le paramètre α permet d'établir une balance entre les deux termes qui constituent la fonction objectif donnée par (3.86). D'après [FRI96b], le processus d'agglomération contrôlé par α devrait être lent au début pour encourager la formation de petits clusters. Ensuite, ce paramètre devrait être augmenté graduellement afin de favoriser l'agglomération. Après quelques itérations, quand le nombre de clusters devient proche d'une valeur appropriée (idéalement l'"optimale"), la valeur de α devrait de nouveau diminuer lentement pour permettre à l'algorithme d'atteindre la convergence. Donc, un choix appropriée de α à l'itération (t) est donnée par l'expression suivante [FRI99] :

$$\alpha^{(t)} = \eta^{(t)} \frac{\sum_{i=1}^c \sum_{k=1}^N (\mu_{ik}^{(t-1)})^2 \rho_i(d_{ik}^2)^{(t)}}{\sum_{i=1}^c \left(\sum_{k=1}^N w_{ik}^{(t)} \mu_{ik}^{(t-1)} \right)^2} \quad (3.94)$$

dans laquelle le facteur $\eta^{(t)}$ est donnée par :

$$\eta^{(t)} = \begin{cases} \eta_0 e^{-|t_0 - t|/\tau} & \text{si } t > 0 \\ 0 & \text{si } t = 0 \end{cases} \quad (3.95)$$

où η_0 est la valeur initiale, τ est une constante de temps et t_0 est le nombre d'itération à partir duquel le facteur η commence à diminuer. Typiquement les valeurs pour ces paramètres sont $\eta_0 = 1$, $\tau = 10$ et $t_0 = 5$. Avec une initialisation appropriée de l'algorithme, ces valeurs sont raisonnablement indépendantes de l'application [FRI99].

D'autre part, il faut remarquer que les valeurs utilisées dans l'expression (3.94) sont toutes connues à l'itération (t). En conséquence, α est une valeur fixe, qui est vue comme une constante lors de la procédure d'optimisation suivante.

Obtention des paramètres des prototypes des clusters

Pour minimiser la fonctionnelle J_{RoFER} donnée par l'expression (3.86) soumise à la contrainte de normalisation (3.37), nous utilisons le même type de procédure itérative que celle décrite pour l'algorithme FER à travers les conditions du premier ordre. En appliquant les multiplicateurs de Lagrange on obtient :

$$\bar{J}_{RoFER}(Z; \mathbf{U}, \mathbf{P}, \boldsymbol{\lambda}) = \sum_{i=1}^c \sum_{k=1}^N (\mu_{ik})^2 \rho_i(d_{ik}^2) - \alpha \sum_{i=1}^c \left(\sum_{k=1}^N w_{ik} \mu_{ik} \right)^2 - \sum_{k=1}^N \lambda_k \left[\sum_{i=1}^c \mu_{ik} - 1 \right] \quad (3.96)$$

où $\mathbf{P} = [\mathbf{V}, \mathbf{A}, \boldsymbol{\theta}]$ représente l'ensemble des matrices associées aux paramètres qui caractérisent les clusters et les termes λ_k correspondent aux multiplicateurs de Lagrange. Les conditions nécessaires pour la minimisation de J_{RoFER} sont alors :

$$\frac{\partial \bar{J}_{RoFER}}{\partial \mu_{ik}} = 2 \mu_{ik} \rho_i(d_{ik}^2) - 2\alpha \sum_{j=1}^N w_{ij} \mu_{ij} - \lambda_k = 0, \quad \forall i, k \quad (3.97)$$

$$\frac{\partial \bar{J}_{RoFER}}{\partial \mathbf{P}_i} = \sum_{k=1}^N (\mu_{ik})^2 w_{ik} \frac{\partial d_{ik}^2}{\partial \mathbf{P}_i} = 0, \quad \forall i \quad (3.98)$$

$$\frac{\partial \bar{J}_{RoFER}}{\partial \lambda_k} = \sum_{i=1}^c \mu_{ik} - 1 = 0, \quad \forall k \quad (3.99)$$

L'équation (3.97) et la contrainte de normalisation (3.37) représentent un ensemble de $N \times c + N$ équations linéaires avec $N \times c + N$ inconnues (μ_{ik} et λ_k). Une simplification du calcul est obtenue en considérant que les degrés d'appartenance μ_{ik} ne changent pas significativement entre deux itérations consécutives. Dès lors, le terme $\sum_{j=1}^N w_{ij} \mu_{ij}$ dans l'expression (3.97) est calculé en utilisant la valeur du degré d'appartenance de l'itération précédente. On obtient alors l'expression suivante :

$$\mu_{ik} = \frac{2\alpha \times \Sigma_i + \lambda_k}{2\rho_i(d_{ik}^2)} \quad (3.100)$$

où la *cardinalité robuste* du i -ème cluster, notée Σ_i , est donnée par l'expression :

$$\Sigma_i = \sum_{j=1}^N w_{ij} \mu_{ij} \quad (3.101)$$

Cette variable est d'une grande importance pour le processus d'agglomération, car elle va déterminer si le cluster considéré peut être fusionné avec son cluster adjacent ou si le cluster va disparaître. En effet, quand la cardinalité robuste Σ_i est plus petite qu'un certain seuil $\Sigma_{\mathcal{E}}$ préétabli dans l'algorithme, le cluster respectif est alors écarté et le nombre de clusters est mise à jour. La valeur $\Sigma_{\mathcal{E}}$ est une constante que nous appelons le *seuil de cardinalité robuste*. En utilisant l'expression (3.100) et la contrainte (3.37) on obtient :

$$\lambda_k = 2 \left(1 - \sum_{i=1}^c \alpha \Sigma_i / \rho_i(d_{ik}^2) \right) / \sum_{j=1}^c 1 / \rho_j(d_{jk}^2) \quad (3.102)$$

Et en remplaçant cette expression dans (3.100) on obtient finalement :

$$\mu_{ik} = \frac{1/\rho_i(d_{ik}^2)}{\sum_{j=1}^c 1/\rho_j(d_{jk}^2)} + \frac{\alpha}{\rho_i(d_{ik}^2)} (\Sigma_i - \bar{\Sigma}_k) \quad (3.103)$$

où le terme $\bar{\Sigma}_k = \sum_{j=1}^c \Sigma_j / \rho_j(d_{jk}^2) / \sum_{j=1}^c 1/\rho_j(d_{jk}^2)$ est la moyenne pondérée des cardinalités.

Le premier terme de l'expression (3.103) indique le degré avec lequel le cluster i partage la donnée $\mathbf{z}_k = (\mathbf{x}_k, y_k)$ (en utilisant des distances robustes) ; le deuxième terme dépend de la différence entre la cardinalité robuste du cluster d'intérêt et la moyenne pondérée des cardinalités. Le fonctionnement de l'algorithme est alors établi sur un principe d'agglomération compétitive entre clusters, basé sur des cardinalités robustes.

Par ailleurs, en ce qui concerne l'actualisation des paramètres des prototypes des clusters $\mathbf{P} = [\mathbf{V}, \mathbf{A}, \boldsymbol{\theta}]$, on retrouve les mêmes types d'expressions que celles de l'algorithme FER, sauf qu'elles sont affectées par le facteur robuste de poids w_{ik} correspondant. En considérant la contrainte d'hypervolume du cluster (3.62), à partir de l'expression (3.98) et de la définition de la mesure de distance (3.64), les équations de mise à jour des paramètres caractéristiques des clusters pour chaque itération (t) sont données par les expressions suivantes :

Centres des clusters hyperellipsoïdaux dans l'espace d'entrée de dimension $\mathfrak{R}^p$:

$$\mathbf{v}_{ix}^{(t)} = \frac{\sum_{k=1}^N w_{ik}^{(t)} (\mu_{ik}^{(t)})^m \cdot \mathbf{x}_k}{\sum_{k=1}^N w_{ik}^{(t)} (\mu_{ik}^{(t)})^m}, \quad 1 \leq i \leq c. \quad (3.104)$$

Matrice (robuste) de poids pour l'estimation des paramètres des sous modèles linéaires :

$$\mathbf{W}_i^{(t)} = \text{diag}[(\mu_{i1}^{(t)})^m w_{i1}^{(t)}, (\mu_{i2}^{(t)})^m w_{i2}^{(t)}, \dots, (\mu_{iN}^{(t)})^m w_{iN}^{(t)}] \quad (3.105)$$

Vecteur des paramètres du sous modèle linéaire associé à chaque cluster :

$$\boldsymbol{\theta}_i^{(t)} = (\mathbf{X}_e^T \mathbf{W}_i^{(t)} \mathbf{X}_e)^{-1} \mathbf{X}_e^T \mathbf{W}_i^{(t)} \mathbf{y} \quad (3.106)$$

Matrice de covariance (robuste) dans l'espace d'entrée :

$$\mathbf{F}_{ixd}^{(t)} = \text{diag} \left(\frac{\sum_{k=1}^N (\mu_{ik}^{(t)})^m w_{ik}^{(t)} (\mathbf{x}_k - \mathbf{v}_{ix}^{(t)}) (\mathbf{x}_k - \mathbf{v}_{ix}^{(t)})^T}{\sum_{k=1}^N (\mu_{ik}^{(t)})^m w_{ik}^{(t)}} \right) \quad (3.107)$$

Composante de sortie de la matrice de covariance :

$$f_{iy}^{(t)} = \frac{\sum_{k=1}^N (\mu_{ik}^{(t)})^m w_{ik}^{(t)} (y_k - f_i(\mathbf{x}_k; \boldsymbol{\theta}_i^{(t)}))^2}{\sum_{k=1}^N (\mu_{ik}^{(t)})^m w_{ik}^{(t)}} \quad (3.108)$$

Matrice de norme induite dans l'espace d'entrée :

$$\mathbf{A}_{ixd}^{(t)} = \left(\psi_i \mathbf{f}_{iy}^{(t)} \middle| \mathbf{F}_{ixd}^{(t)} \right)^{1/n} \left(\mathbf{F}_{ixd}^{(t)} \right)^{-1} \quad (3.109)$$

Composante de sortie de la matrice de norme induite :

$$\sigma_i^{(t)} = \left(\psi_i \mathbf{f}_{iy}^{(t)} \middle| \mathbf{F}_{ixd}^{(t)} \right)^{1/n} \left(\mathbf{f}_{iy}^{(t)} \right)^{-1} \quad (3.110)$$

L'algorithme de clustering robuste avec régression floue ellipsoïdale (RoFER) peut se formuler de la manière suivante :

Algorithme RoFER – Robust Fuzzy Ellipsoidal Regression

Etant donné un ensemble de données $Z = \{(\mathbf{x}_1, y_1), \dots, (\mathbf{x}_N, y_N)\}$:

Fixer le nombre *maximal* de clusters $c_{\max}$

Fixer le *seuil de cardinalité robuste* $\Sigma_{\mathcal{E}}$ et le coefficient $\delta^{(0)}$

Fixer le degré de flou $m > 1$ et la tolérance de fin d'algorithme $\varepsilon > 0$

Initialiser les prototypes des clusters $[\mathbf{V}, \mathbf{A}, \boldsymbol{\theta}]$ et la matrice $\mathbf{U} = [\mu_{ik}]$ de partition floue avec un algorithme approprié (e.g. avec l'algorithme FER)

Répéter

Pas 1. Calculer les distances d_{ik}^2 pour $1 \leq i \leq c$ et $1 \leq k \leq N$ (Eq. (3.64))

Pas 2. Calculer les paramètres T_i (Eq. (3.88)) et S_i (Eq. (3.89))

Pas 3. Mettre à jour les fonctions robustes w_i (Eq. (3.90)) et ρ_i (Eq. (3.91))

Pas 4. Mettre à jour le paramètre d'agglomération α (Eq. (3.94))

Pas 5. Calculer la cardinalité robuste Σ_i (Eq. (3.101))

Pas 6. Mettre à jour la matrice $\mathbf{U}$ de partition floue (Eq. (3.103))

Pas 7. Comparer la cardinalité robuste Σ_i de chaque cluster avec le seuil de cardinalité robuste $\Sigma_{\mathcal{E}}$ défini lors de l'initialisation de l'algorithme

Si $\Sigma_i < \Sigma_{\mathcal{E}}$ alors écarter l' i -ème cluster, $\forall i$, avec $1 \leq i \leq c$

Pas 8. Mettre à jour le nombre de clusters c , avec $c_{\text{optimal}} \leq c \leq c_{\max}$

Pas 9. Mettre à jour les paramètres des prototypes des clusters :

Les centres des clusters $\mathbf{v}_{ix}$ dans l'espace d'entrée (Eq. (3.104))

Les matrices de pondération $\mathbf{W}_i$ (Eq. (3.105)) pour l'estimation des paramètres des sous modèles linéaires $\boldsymbol{\theta}_i = [\mathbf{a}_i^T, d_i]^T$ pour chaque cluster (Eq. (3.106))

Les matrices de covariance de cluster $\mathbf{F}_{ixd}$ dans l'espace d'entrée (Eq. (3.107)), ainsi que la composante de sortie $\mathbf{f}_{iy}$ (Eq. (3.108))

Les matrices diagonales de norme induite $\mathbf{A}_{id}$ par l'intermédiaire des expressions $\mathbf{A}_{ixd}$ (Eq. (3.109)) et σ_i (Eq. (3.110))

Pas 10. Mettre à jour le coefficient de réglage δ (Eq. (3.93))

Jusqu'à obtenir la stabilité de la partition (Eq. (3.44))

Remarques

- R 3.7** Comparativement aux algorithmes classiques de clustering (e.g. FCM, GK, FCRM), l'algorithme RoFER peut impliquer un temps de calcul supérieur dû aux différentes équations à mettre à jour lors de chaque itération. Cependant, comme nous le verrons dans des exemples ultérieurs, l'algorithme considéré peut avoir une performance supérieure dans la mesure qu'il fournit directement une estimation des paramètres des sous modèles linéaires avec des caractéristiques de robustesse vis-à-vis du bruit et des points aberrants. De plus, avec une sélection appropriée du seuil de cardinalité robuste, il est possible d'obtenir "automatiquement" un nombre convenable de clusters pour la partition floue. Nous considérons que ces deux caractéristiques sont intéressantes dans le cadre de la modélisation floue de systèmes, compte tenu qu'un tel processus est réalisé hors ligne, où le temps de calcul n'est pas un facteur critique.
- R 3.8** Comme nous l'avons déjà fait remarquer, dans l'algorithme RoFER le seuil de cardinalité robuste Σ_{ε} , est un paramètre qui doit être fixé par l'utilisateur. Ce paramètre a une importance majeure dans le nombre final de clusters obtenus au niveau de la partition floue des données. En effet, ce seuil est une valeur constante qui conditionne la possibilité de survivre des clusters lors du processus d'agglomération compétitive. Pour des algorithmes similaires trouvées dans la littérature [CHU01], cette valeur est ajustée en fonction de l'application visée, comportant normalement des valeurs basses (e.g., $\Sigma_{\varepsilon} < 10$) pour des applications de petite taille (e.g., $N \leq 400$). Grisales au cours de ces travaux de thèse a développé une expression expérimentale lors de la modélisation floue de type Takagi-Sugeno pour diverses fonctions non linéaires. La valeur initiale proposée est alors donné par l'expression :

$$\Sigma_{\varepsilon} = \frac{N}{6\sqrt{c_{\max}}} \quad (3.111)$$

où la variable N est le nombre de données (paires entrée-sortie) disponibles pour la modélisation et $c_{\max}$ est la valeur d'initialisation (la valeur maximale) du nombre de clusters dans l'algorithme.

Nous abordons maintenant une description de la méthodologie générale pour la construction de modèles flous Takagi-Sugeno à partir des données entrée-sortie des systèmes.

3.4. Construction de modèles TS à partir de données : méthodologie générale

Nous l'avons vu, les techniques de clustering flou sont des outils puissants pour la reconnaissance des formes à partir des données. Dans le cadre de la modélisation floue des systèmes, notre intérêt porte sur l'obtention des modèles Takagi-Sugeno qui permettent une décomposition automatique d'un système non linéaire dans un ensemble de régions linéaires (clusters). Pour cela, nous appliquons des techniques de clustering flou dans l'espace produit d'entrée-sortie des données. Les algorithmes présentés précédemment ont tous des caractéristiques particulières qui peuvent être réunies dans une seule vue d'ensemble. Nous abordons ainsi une description d'une méthodologie générale pour la construction des modèles flous du type TS, en mettant l'accent sur les besoins communs qui sont : la validation du nombre de clusters, la génération des fonctions d'appartenance des antécédents, l'obtention des paramètres des conséquents et la validation numérique du modèle final. Ces aspects seront

présentés dans les paragraphes qui suivent, avant d'aborder quelques exemples illustratifs à la fin du chapitre.

La Figure 3.6 illustre la procédure de construction des modèles flous de type Takagi-Sugeno à partir de données, en utilisant des techniques de clustering. La procédure est itérative, dans la mesure où, dans une même session de modélisation, on peut répéter certains pas afin de tester différents choix pour plusieurs paramètres.

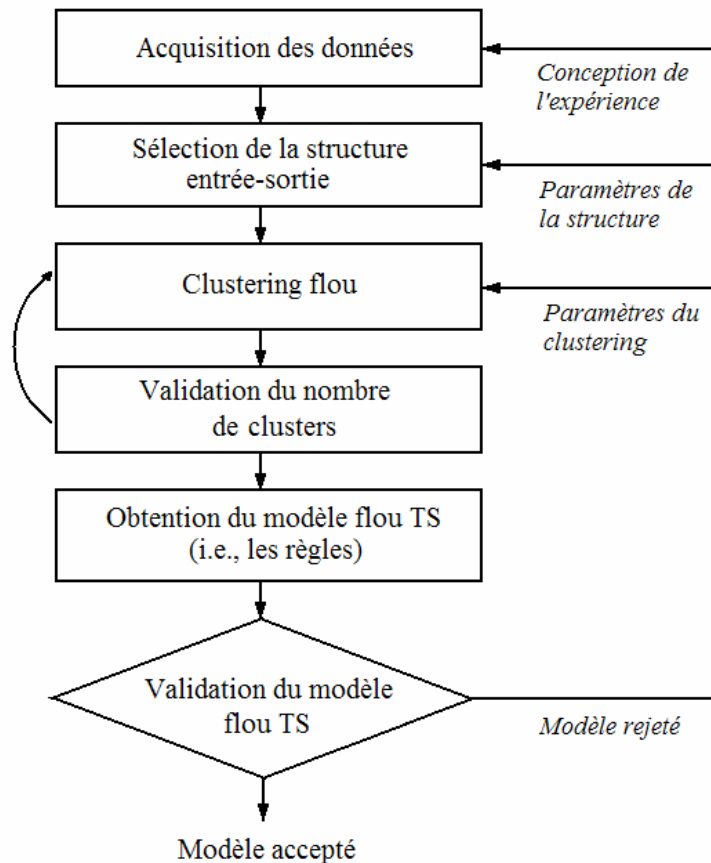


Figure 3.6. Vue générale de la méthodologie d'identification non linéaire basée sur l'approche de clustering flou

Nous allons décrire chaque étape de la méthodologie pour la construction de modèles flous du type Takagi-Sugeno.

Conception de l'expérience d'identification et acquisition des données

C'est l'étape initiale de n'importe quelle méthode d'identification, puisqu'elle détermine le contenu d'information des données d'identification. Contrairement aux techniques linéaires, les signaux du type séquence binaire pseudo-aléatoire (PRBS⁷¹) ne sont pas appropriés en général pour l'identification non linéaire et en particulier pour le clustering flou [BAB98a]. Bien que le choix du signal d'excitation puisse être dépendant du problème, le signal d'entrée devrait de préférence exciter le système dans toute la gamme de fonctionnement, tant en amplitude qu'en fréquence, pour les variables considérées. Le signal PRBS n'est pas approprié puisqu'il contient seulement deux valeurs d'amplitude. Des choix typiques sont des signaux multi-sinusoïdaux ou bien des signaux multi-échelon avec

⁷¹ PRBS : *Pseudo-Random Binary Sequence*.

amplitude et largeur aléatoires [GOD93]. Un bruit blanc de petite amplitude est souvent ajouté à ces signaux pour garantir l'excitation appropriée de la dynamique du procédé. Le choix d'une période d'échantillonnage ainsi que la durée des expériences sont aussi des aspects importants pour la conception de l'expérience.

Sélection de la structure entrée-sortie

Le propos de cette étape est la détermination des entrées et sorties prépondérantes par rapport au but final de la modélisation. Quand il s'agit de l'identification des systèmes dynamiques, il faut choisir la structure et l'ordre du modèle dynamique. Par exemple, pour le cas des modèles dynamiques non linéaires du type NARX (cf. expression (3.25)), cela se traduit par la définition des régresseurs et des valeurs pour les constantes n_y , n_u et n_d associées à l'ordre du système. Une fois les régresseurs établis, le choix de la structure permet de traduire l'identification d'un système dynamique en un problème de régression qui peut être résolu d'une façon *statique*⁷² [LEO85]. C'est pourquoi les modèles flous de type Takagi-Sugeno, caractérisés par leurs capacités d'approximation des fonctions non linéaires ainsi que de représentation en utilisant des règles, sont des structures candidates flexibles pour la modélisation.

A l'aide des données, la structure peut être choisie de façon automatique, en comparant différentes structures potentielles en termes d'une certaine mesure de performance [HE93] [BOM98] [RHO98] [HAD02] [FEI04] [SIN04] [UNR07]. Dans la plupart des cas, un choix raisonnable peut être fait par l'utilisateur, en se basant sur la connaissance préalable du procédé et/ou le but même de la modélisation⁷³. Dans la pratique, un compromis entre l'exactitude et la complexité est recherché. Ainsi par exemple, dans un but de prédiction, usuellement le modèle peut contenir plus de paramètres afin de chercher une excellente performance numérique. D'un autre côté, dans le cas de la commande floue basée sur un modèle, et en particulier dans l'approche du type PDC (cf. Chapitre 4), l'objectif de la modélisation consistera plutôt à trouver un compromis entre la complexité du modèle et sa performance numérique au niveau de la représentation du système dynamique. En effet, la complexité du contrôleur est directement associée à celle du modèle flou étudié.

Clustering flou

La sélection de la structure conduit à un problème de régression non linéaire statique, qui est alors approximé par une collection de sous modèles linéaires locaux. La localisation et les paramètres des sous modèles sont établis en partitionnant les données disponibles en clusters hyperplanaires (cas des algorithmes FCRM, FER et RoFER) ou hyperellipsoïdaux (cas de l'algorithme GK). Chacun des clusters définit une région floue (par l'intermédiaire des antécédents) dans laquelle le système peut être approximé localement par un sous modèle linéaire (au moyen des conséquents). Dans le cas des trois premiers algorithmes qui possèdent des clusters hyperplanaires dans l'espace de sortie, l'estimation des paramètres des sous modèles linéaires affines fait parti du processus de clustering. Par contre, dans le cas de l'algorithme de clustering GK, l'obtention des paramètres des conséquents du modèle TS correspondant est faite lors d'une étape postérieure au processus de clustering.

⁷² Cela signifie que la sortie du système dynamique non linéaire peut être vue comme une hypersurface statique dans l'espace produit d'entrée-sortie, qui est le même espace que celui des régresseurs-sortie. Du point de vue géométrique, le concept garde une certaine analogie avec celui de la "surface de commande" associé aux modèles flous destinés à la commande.

⁷³ Typiquement, le choix d'une structure pauvre (trop peu de régresseurs) entraîne une modélisation inexacte des non linéarités et de la dynamique du procédé. Par contre, le choix d'une structure plus riche que nécessaire (trop de régresseurs), amène à des problèmes d'estimation mal conditionnés et à une surestimation des données.

D'autre part, il faut remarquer aussi que pour les algorithmes GK, FCRM et FER, l'utilisateur doit définir à l'avance le nombre c de clusters (nombre souhaité des régions linéaires). Ce nombre reste constant pendant toute l'exécution de l'algorithme. Pour l'algorithme RoFER, le nombre initial de clusters doit être surdimensionné ; il correspond au nombre maximal de clusters, qui au cours des itérations va diminuer progressivement jusqu'à atteindre (idéalement) la valeur optimale.

Validation du nombre de clusters

Le nombre approprié de clusters peut être influencé par le but même de la modélisation (e.g., contrôle de la complexité,...). Pour l'algorithme RoFER, grâce au surdimensionnement du nombre de clusters et à l'approche d'agglomération compétitive, il est supposé que le nombre attendu des clusters corresponde à celui qui permet d'établir le meilleur compromis entre le nombre de classes et la partition floue obtenue pour les données. Pour les algorithmes GK, FCRM et FER, il est plutôt conseillé d'appliquer des mesures de validation du nombre de clusters (cf. paragraphe 3.4.1) [BAB98a] et/ou de fusion de clusters [KAY95] afin de déterminer le nombre de clusters approprié pour une application particulière. Typiquement cet étape du processus de modélisation implique plusieurs répétitions de l'étape précédente (clustering flou) afin d'évaluer les résultats pour un nombre différent de clusters ainsi que pour différentes initialisations des paramètres dans les algorithmes respectifs.

Obtention du modèle flou Takagi-Sugeno

Le clustering flou divise les données disponibles dans des groupes où des relations linéaires existent entre les entrées et la sortie. Afin d'obtenir un modèle approprié pour des propos de prédiction ou de synthèse d'un contrôleur, un modèle flou (de structure définie) basé sur des règles est extrait de la matrice de partition floue disponible ainsi que des prototypes des clusters. Les règles, les fonctions d'appartenance et d'autres paramètres qui constituent le modèle flou sont obtenus de façon automatique. Des mesures de similarité floue [SET98] peuvent être employées afin de réduire la base de règles initiale ou bien pour améliorer l'interprétation linguistique des fonctions d'appartenance. La procédure d'obtention du modèle flou TS correspondant passe typiquement par la génération des fonctions d'appartenance des antécédents (cf. paragraphe 3.4.2) et dans certains cas (e.g., pour l'algorithme GK), par l'obtention des paramètres des conséquents (cf. paragraphe 3.4.3).

Validation du modèle flou TS

Par intermédiaire de la validation, le modèle flou obtenu est considéré comme approprié pour le propos donné ou bien, il est rejeté. Dans ce dernier cas, certaines étapes de la méthodologie générale d'identification présenté sur la Figure 3.6 peuvent être répétés ; cette démarche est usuelle aussi dans d'autres approches d'identification linéaire et non linéaire des systèmes [LJU87] [MUR97]. En plus de la validation numérique habituelle au moyen de la simulation, l'interprétation des modèles flous joue un rôle important dans l'étape de validation. La couverture de l'espace d'entrée par les règles peut être analysée, et pour une base de règles incomplète, des règles supplémentaires peuvent être fournies en se basant sur la connaissance préalable (linéarisation locale ou par des modèles issus des principes ou lois physiques).

Pour terminer cette vue d'ensemble de la construction des modèles flous de type Takagi-Sugeno, nous abordons plus en détail les aspects concernant la validation du nombre de clusters, l'obtention de la base de règles (i.e., la génération des fonctions d'appartenance des antécédents et l'obtention des paramètres des conséquents) et la validation finale du modèle.

3.4.1. Validation du nombre de clusters

L'utilisation des mesures de validation des clusters est une approche standard pour la détermination d'un nombre approprié de clusters dans un ensemble de données. Cette approche s'avère utile spécialement pour les méthodes de clustering qui nécessitent la sélection *a priori* du nombre de classes. L'analyse de validité des clusters est réalisée en exécutant l'algorithme de clustering pour différents valeurs de c ⁷⁴, à plusieurs reprises pour chaque c avec une initialisation différente. La mesure de validité est calculée pour chaque exécution et le nombre de clusters qui optimise la mesure est alors choisi comme étant le nombre "convenable" de clusters pour l'ensemble des données. Les mesures de validité peuvent être aussi évaluées pour différentes structures du modèle (différent choix des entrées et de l'ordre du modèle). Il faut noter que l'utilisation des mesures de validité est lourde en temps de calcul car le processus de clustering doit être répété plusieurs fois.

La plupart des mesures de validité sont conçues pour quantifier la séparation et la compacité des clusters. Néanmoins, comme le souligne Bezdek [BEZ81], le concept de validité de clusters est ouvert à l'interprétation et peut être formulé de façons différentes. En conséquence, plusieurs mesures de validité ont été introduites dans la littérature (voir par exemple [BEZ81] [GAT89] [PAL95]).

Nous présentons par la suite quelques mesures de validité des clusters couramment utilisés. Chaque critère est supposé atteindre sa valeur optimale lorsque la valeur de c correspond au nombre de clusters réels. Cependant, la recherche d'une partition optimale n'est pas du tout une tâche simple parce que, bien que disposant du meilleur critère possible, il reste un bon nombre de partitions envisageables pour le même jeu de données. En plus, en général elles sont sensibles au bruit. En pratique, ces critères permettent à l'utilisateur d'avoir un guide pour le choix convenable du nombre de clusters pour son application particulière.

*Hypervolume flou (C_{HF})*⁷⁵

Il est associé aux matrices $\mathbf{F}_i$ de covariance floue des clusters. Une bonne partition est indiquée par des valeurs petites du critère, défini par l'expression (3.112) :

$$C_{HF}(c) = \sum_{i=1}^c [\det(\mathbf{F}_i)]^{1/2} \quad (3.112)$$

Coefficient de partition (C_{CP})

Proposé par Bezdek [BEZ81], il mesure la quantité de "chevauchement" (*overlapping*) entre clusters. Le principal inconvénient du critère C_{CP} est le manque de rapport direct avec certaine propriété des données. La valeur optimale du nombre de clusters correspond au maximum du critère ; il est défini comme :

$$C_{CP}(c) = \frac{1}{N} \sum_{i=1}^c \sum_{k=1}^N (\mu_{ik})^2 \quad (3.113)$$

Entropie de classification (C_{EC})

Elle est similaire au coefficient de partition précédent dans le sens que cette mesure est associée au degré de "flou" (*fuzzyness*) de la partition. Le critère est donné par l'expression suivante :

⁷⁴ Dans le cas de l'algorithme RoFER, pour différentes valeurs du seuil de cardinalité robuste Σ_c .

⁷⁵ Dans cette notation, la lettre C correspond au mot critère. Le petit c correspond (comme toujours dans notre notation) au nombre de clusters.

$$C_{EC}(c) = -\frac{1}{N} \sum_{i=1}^c \sum_{k=1}^N \mu_{ik} \log(\mu_{ik}) \quad (3.114)$$

Index de partition (C_{IP})

Ce critère tient compte du rapport de la compacité et la séparation des clusters. Il correspond à la somme des mesures de validité individuelles, normalisées par la division de la cardinalité floue Σ_i de chaque cluster, selon l'expression donné par [BENS96] :

$$C_{IP}(c) = \sum_{i=1}^c \frac{\sum_{k=1}^N (\mu_{ik})^m \|\mathbf{x}_k - \mathbf{v}_i\|^2}{\Sigma_i \cdot \sum_{j=1}^c \|\mathbf{v}_j - \mathbf{v}_i\|^2} \quad (3.115)$$

Ce critère est utile aussi pour la comparaison de différentes partitions avec le même nombre de clusters. Une valeur basse du critère C_{IP} indique une meilleure partition.

Index de séparation (C_{IS})

A l'inverse du critère précédent, cet index utilise une séparation de distance minimale pour la validation de la partition [BENS96]. Le critère C_{IS} est donné par :

$$C_{IS}(c) = \sum_{i=1}^c \frac{\sum_{k=1}^N (\mu_{ik})^m \|\mathbf{x}_k - \mathbf{v}_i\|^2}{\Sigma_i \cdot \min_{i,k} \|\mathbf{v}_k - \mathbf{v}_i\|^2} \quad (3.116)$$

Nous abordons par la suite l'obtention de la base de règles du modèle flou Takagi-Sugeno (TS). Cela implique la génération des fonctions d'appartenance des antécédents et l'obtention des paramètres des conséquents.

Obtention de la base de règles du modèle flou Takagi-Sugeno

En général, pour les algorithmes de coalescence floue présentés dans la section précédente (cf. section 3.3) excepté l'algorithme FCM⁷⁶, chaque cluster obtenu dans l'espace produit d'entrée-sortie⁷⁷ des données peut être vu comme une approximation locale linéaire de l'hypersurface de régression. Le modèle flou global peut être représenté convenablement comme une collection de règles du type Takagi-Sugeno affines (cf. paragraphe 3.1.2), de la forme donnée par l'expression suivante :

$$R_i : \quad \text{Si } \mathbf{x} \text{ est } A_i \quad \text{Alors } y_i = \mathbf{a}_i^T \mathbf{x} + d_i, \quad i = 1, \dots, r \quad (3.117)$$

où le nombre r de règles correspond au nombre c de classes (clusters) issu du processus de clustering, $\mathbf{x} \in \mathfrak{R}^p$ est le vecteur d'entrée correspondant à une observation et $y \in \mathfrak{R}$ est la variable de sortie. A_i est le sous-ensemble flou de l'antécédent de l' i -ème règle, définie, dans ce cas, par une fonction d'appartenance multidimensionnelle. Pour sa part, le conséquent de chacune des règles est formé par le vecteur des paramètres $\mathbf{a}_i \in \mathfrak{R}^p$ et le scalaire d_i . Les conclusions des règles dans ce modèle sont alors des hyperplans dans l'espace $\mathfrak{R}^{p+1}$. Les ensembles flous A_i peuvent être décrits dans leur forme naturelle multidimensionnelle, ou bien dans leur forme décomposée, en projetant la partition floue sur les variables de l'antécédent. Les paramètres $\mathbf{a}_i$ et d_i des conséquents sont estimés à partir des données en utilisant des méthodes des moindres carrés (cas des algorithmes FCRM, FER et RoFER), ou bien ils peuvent être extraits à partir de la structure propre des matrices de covariance des clusters (cas de l'algorithme GK), comme on le verra par la suite.

⁷⁶ Ce n'est pas exactement le cas avec l'algorithme FCM, car il conduit à des clusters hypersphériques.

⁷⁷ Espace d'entrée-sortie de dimension $\mathfrak{R}^{p+1}$.

3.4.2. Génération des fonctions d'appartenance des antécédents

Les fonctions d'appartenance des antécédents peuvent être obtenues en calculant les degrés d'appartenance directement dans l'espace produit des variables de l'antécédent. Elles peuvent aussi être extraites à partir de la matrice de partition floue $\mathbf{U}$ en appliquant le *mécanisme de projection* sur ces variables. Nous décrivons ces méthodes par la suite.

Fonctions d'appartenance multidimensionnelles des antécédents

Dans cette approche, les fonctions d'appartenance des antécédents sont représentées analytiquement en calculant l'inverse de la distance par rapport au prototype du cluster. Le degré d'appartenance est calculé directement pour la totalité du vecteur d'entrée $\mathbf{x} \in \mathbb{R}^p$, il n'y a pas de décomposition. Les antécédents des règles du modèle flou TS sont alors des propositions simples avec des ensembles flous multidimensionnels (cf. expression (3.117)), dans lesquels le degré d'accomplissement des règles est donné par $\beta_i(\mathbf{x}) = \mu_{A_i}(\mathbf{x})$. Quand les données ont une partition floue stricte, le degré d'accomplissement pour chaque règle est déterminé en utilisant l'expression [BAB98a] :

$$\beta_i(\mathbf{x}) = \frac{1}{\sum_{j=1}^c \left(D_{A_{ix}}(\mathbf{x}, \mathbf{v}_{ix}) / D_{A_{ix}}(\mathbf{x}, \mathbf{v}_{jx}) \right)^{2/(m-1)}} \quad (3.118)$$

Génération des fonctions d'appartenance des antécédents par projection

Le principe de cette méthode est de projeter pour chaque règle, les ensembles flous multidimensionnels⁷⁸ définis point par point dans la matrice de partition floue $\mathbf{U}$ sur les variables individuelles des antécédents. Dans le cas où ces variables sont les entrées utilisées pour faire le clustering, il s'agit d'une projection orthogonale des données. On projette donc la matrice de partition floue sur chacun des axes des variables (i.e., sur les régresseurs) de l'antécédent x_j , avec $1 \leq j \leq p$. Les règles correspondantes du modèle flou TS sont alors exprimées sous la forme conjonctive suivante :

$$R_i : \quad \text{Si } x_1 \text{ est } A_{i1} \text{ et } \dots \text{ et } x_p \text{ est } A_{ip} \quad \text{Alors } y_i = \mathbf{a}_i^T \mathbf{x} + d_i, \quad i = 1, \dots, r \quad (3.119)$$

Pour obtenir les fonctions d'appartenance des sous-ensembles flous A_{ij} de l'antécédent, l'ensemble flou multidimensionnel défini point par point par l' i -ème ligne de la matrice de partition est projeté sur les régresseurs x_j . Cette opération de *projection* est représentée par l'expression suivante puis illustré sur la Figure 3.7a :

$$\mu_{A_{ij}}(x_{jk}) = \text{proj}_j(\mu_{ik}) \quad (3.120)$$

Définition (projection) : Soit un sous-ensemble flou A défini sur un univers $X_1 \times X_2$, produit cartésien de deux univers X_1 et X_2 . On définit la *projection* de A sur X_1 comme le sous-ensemble flou de X_1 , noté $\text{proj}_{X_1}(A)$, de fonction d'appartenance :

$$\mu_{\text{proj}_{X_1}(A)}(x_1) = \sup_{x_2 \in X_2} \mu_A((x_1, x_2)), \quad \forall x_1 \in X_1 \quad (3.121)$$

On définit de manière analogue la projection de A sur X_2 . Pour des raisons de simplicité on a défini cette projection sur deux univers, mais de la même façon on peut le faire sur p univers.

⁷⁸ Un ensemble flou défini sur un univers de discours multidimensionnel est appelé un *ensemble flou multidimensionnel*.

Comme nous l'avons déjà dit, le degré d'accomplissement $\beta_i(\mathbf{x}) = \mu_{A_i}(\mathbf{x})$ pour l' i -ème règle doit être calculé comme une combinaison des degrés d'appartenance de chaque proposition individuelle en utilisant des opérateurs logiques flous. On reconstruit donc le cluster dans l'espace des entrées en appliquant l'opérateur intersection⁷⁹ dans l'espace Cartésien des variables de l'antécédent, de la façon suivante :

$$\beta_i = \mu_{A_{i1}}(x_1) \wedge \mu_{A_{i2}}(x_2) \wedge \dots \wedge \mu_{A_{ip}}(x_p) \quad (3.122)$$

Il faut noter que cette reconstruction n'est en général pas exacte. En effet, elle peut entraîner une *erreur de décomposition* dans le cas où les clusters sont obliques par rapport aux axes de projection (voir Figure 3.7b, adaptée d'après [GAW02]). L'erreur de décomposition est représentée dans cette figure par la différence des surfaces entre le rectangle et l'ellipse dans le support de l'ensemble flou dénoté A_1 . Cette erreur peut être partiellement compensé par une estimation des paramètres des conséquents en utilisant la méthode des moindres carrés globales [BAB98a]. Si l'erreur est trop importante, on peut également envisager un changement de variables et donc de repère en faisant une projection des p premiers vecteurs propres de la matrice de covariances des clusters. À partir de cette matrice on peut dans ce cas profiter de l'information sur l'orientation des clusters afin de réduire l'erreur de décomposition.

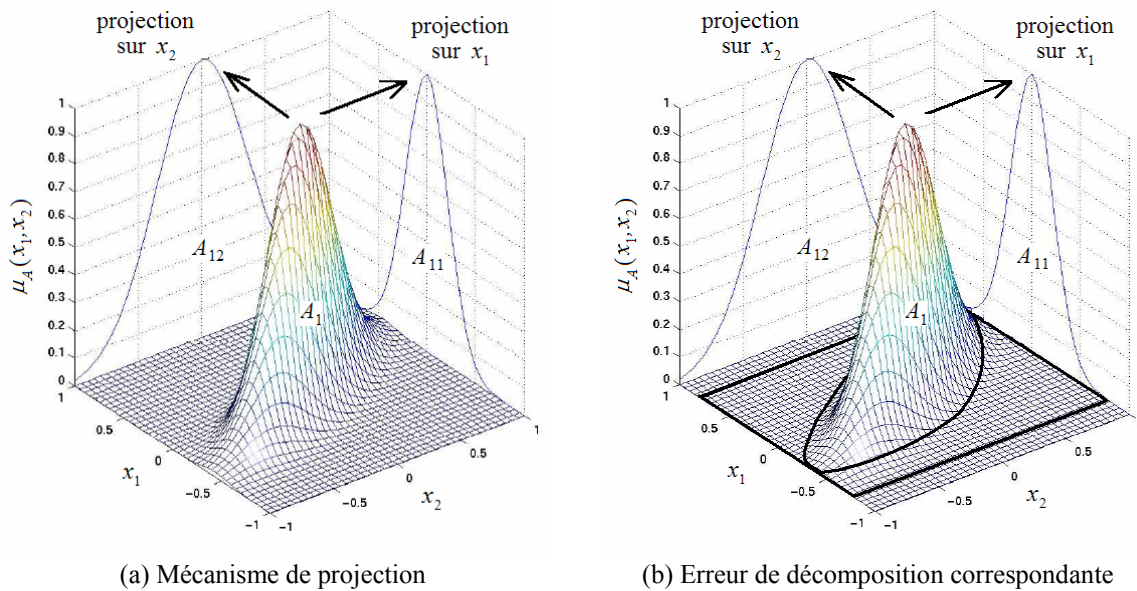


Figure 3.7. (a) Illustration du mécanisme de projection et (b) de l'erreur de décomposition au moment de la recomposition (intersection floue) dans l'espace de l'antécédent.

Paramétrisation des fonctions d'appartenance

Une définition point par point de l'ensemble flou A_{ij} est obtenue en projetant l' i -ème ligne μ_i de la matrice de partition floue $\mathbf{U}$ sur la variable x_j de l'antécédent. Afin d'obtenir un modèle dans un but de prédiction ou de commande, les fonctions d'appartenance de l'antécédent doivent être exprimées sous une forme qui permet le calcul des degrés d'appartenance, même pour des données d'entrée non contenues dans l'ensemble des données Z . Cela est réalisé en approximant la fonction d'appartenance définie point par point par une fonction paramétrique appropriée (e.g., les expressions (3.3) et (3.4)), comme illustré sur la Figure 3.8 (d'après [BAB98a]) :

⁷⁹ On peut utiliser d'autres t -normes, c'est-à-dire des opérateurs d'intersection, comme le produit au lieu de l'opérateur min ($\wedge$).

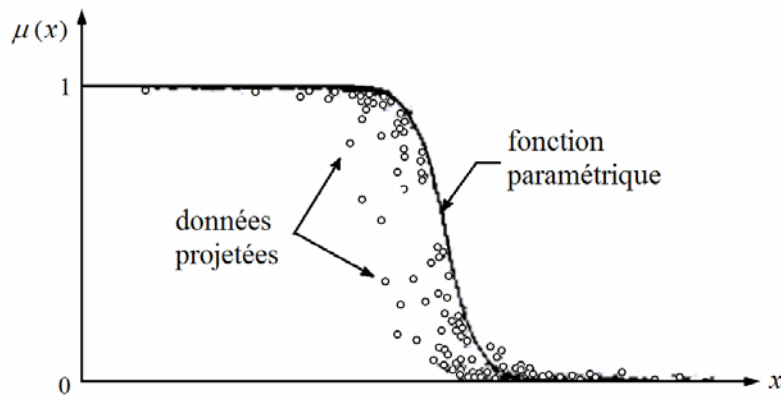


Figure 3.8. Approximation des données projetées par une fonction d'appartenance paramétrique.

La fonction d'appartenance est adaptée à l'enveloppe des données projetées en optimisant numériquement ses paramètres. Un avantage de cette méthode par rapport aux fonctions d'appartenance multidimensionnelles vient du fait que les fonctions d'appartenance projetées peuvent toujours être approximées par des ensembles flous convexes. De plus, des fonctions d'appartenance asymétriques peuvent être utilisées pour refléter la partition réelle du problème de régression non linéaire considéré. Nous trouverons ce type de fonctions d'appartenance dans le chapitre 5, lors de l'application des techniques floues pour l'identification et la commande du bioréacteur aérobique considéré dans le cadre de nos travaux.

Remarque

R 3.9 Même si par la projection orthogonale sur les axes on peut perdre une certaine information, cette méthode de génération des fonctions d'appartenance des antécédents s'avère très utile, car elle donne la possibilité d'interpréter le modèle flou en utilisant directement les régresseurs sous la forme décomposée de règles. Dans ce sens, il faut rappeler que les algorithmes FER et RoFER présentés à la fin de la section 3.3, partagent la caractéristique d'avoir une mesure de distance avec une matrice *diagonale* de norme induite, ce qui conduit à une réduction de l'erreur de décomposition car les clusters obtenus sont forcément parallèles aux axes des variables de l'antécédent. Cette caractéristique contribue à l'amélioration de l'interprétabilité du modèle flou de la part de l'utilisateur. De plus, ces deux algorithmes déterminent simultanément les paramètres des conséquents et la partition floue des données.

3.4.3. Obtention des paramètres des conséquents

Les paramètres $\mathbf{a}_i \in \mathbb{R}^p$ et d_i du modèle TS donné par l'expression (3.119) peuvent être déterminés à partir de la structure géométrique des clusters (cas de l'algorithme GK) ou bien en général, par des techniques d'estimation de type moindres carrés. Nous abordons ces approches dans les paragraphes suivants.

Obtention des paramètres des conséquents par interprétation géométrique

Dans le cas de l'algorithme GK, caractérisé par l'utilisation d'une matrice de covariance floue dans l'espace d'entrée-sortie (de dimension $\mathbb{R}^{p+1}$), on peut se baser sur l'interprétation

géométrique des clusters⁸⁰ afin d'obtenir les paramètres des conséquents dans le *modèle affine Takagi-Sugeno*. Pour illustrer la procédure, supposons qu'une collection de c clusters approxime la surface de régression qui établie la correspondance entre l'entrée (multivariable) et la sortie du système. Ces clusters peuvent être vus comme des sous-espaces linéaires p -dimensionnels de l'espace de régression. Le vecteur propre ϕ_{in} , correspondant à la plus petite valeur propre λ_{in} de la matrice de covariance floue, détermine le vecteur normal à l'hyperplan traversé par les vecteurs restants de ce i -ème cluster. Ceci est illustré par la Figure 3.9 :

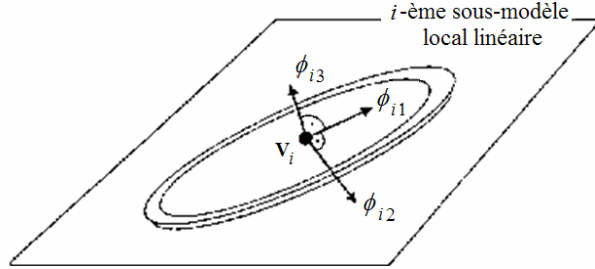


Figure 3.9. Interprétation géométrique de l'orientation du cluster à partir de la matrice de covariance floue (cas de l'algorithme GK).

Afin de simplifier la notation, le plus petit vecteur propre de l' i -ème cluster sera dénoté ϕ_i , en omettant l'indice n . Rappelons que $\mathbf{z}^T = [\mathbf{x}^T, y]$ est le vecteur des données et $\mathbf{v}_i$ est le prototype (centre) de l' i -ème cluster. La forme normale implicite de l'hyperplan du conséquent est donnée alors par l'expression :

$$\phi_i^T \cdot (\mathbf{z} - \mathbf{v}_i) = 0 \quad (3.123)$$

Cette expression établie que le produit intérieur du vecteur normal ϕ_i avec n'importe quel vecteur appartenant à l'hyperplan est nul. Pour continuer le développement, il convient de diviser le prototype $\mathbf{v}_i$ en un vecteur $\mathbf{v}_{ix}$ correspondant au régresseur $\mathbf{x}$, et un scalaire v_{iy} correspondant à la sortie y , c'est-à-dire, $\mathbf{v}_i^T = [(\mathbf{v}_{ix})^T, v_{iy}]$. Le plus petit vecteur propre est aussi divisé de la même façon, i.e., $\phi_i^T = [(\phi_{ix})^T, \phi_{iy}]$. L'équation (3.123) peut s'écrire maintenant de la façon suivante :

$$[(\phi_{ix})^T, \phi_{iy}] \cdot ([\mathbf{x}^T, y]^T - [(\mathbf{v}_{ix})^T, v_{iy}]^T) = 0 \quad (3.124)$$

En réalisant le produit intérieur, on obtient l'égalité suivante :

$$(\phi_{ix})^T (\mathbf{x} - \mathbf{v}_{ix}) + \phi_{iy} (y - v_{iy}) = 0 \quad (3.125)$$

Expression à partir de laquelle, par une manipulation algébrique simple, on obtient l'équation explicite de l'hyperplan [BAB98a] :

$$y = \underbrace{\frac{-1}{\phi_{iy}} (\phi_{ix})^T \mathbf{x}}_{\mathbf{a}_i^T} + \underbrace{\frac{1}{\phi_{iy}} \phi_i^T \mathbf{v}_i}_{d_i} \quad (3.126)$$

En comparant cette expression avec le conséquent de la i -ème règle du *modèle affine TS* représenté par l'équation (3.119), on obtient directement les expressions pour les paramètres $\mathbf{a}_i \in \mathbb{R}^p$ et le scalaire d_i de la façon suivante :

⁸⁰ Cependant, cette interprétation est très liée à l'estimation à l'aide des moindres carrés. En fait, la solution issue de l'interprétation géométrique correspond à la solution des moindres carrés totaux pondérés avec la linéarisation locale autour du centre du cluster, comme l'a fait remarquer Babuška [BAB98a].

$$\mathbf{a}_i = \frac{-1}{\phi_{iy}} \phi_{ix} = \frac{-1}{\phi_{ip+1}} [\phi_{i1}, \phi_{i2}, \dots, \phi_{ip}]^T \quad (3.127)$$

$$d_i = \frac{1}{\phi_{iy}} \phi_i^T \mathbf{v}_i \quad (3.128)$$

Obtention des paramètres des conséquents par la méthode des moindres carrés

Les paramètres des conséquents peuvent être établis par la technique des moindres carrés (en particulier pour les algorithmes FCRM, FER et RoFER), en utilisant comme facteurs de pondération des données, les degrés d'appartenance de la matrice de partition floue $\mathbf{U}$ issus du processus de clustering. Cette approche conduit à une formulation de c problèmes indépendants de type moindres carrés pondérés (WLS⁸¹) dans laquelle les degrés d'appartenance expriment l'importance de la paire de données $(\mathbf{x}_k, y_k)$ par rapport à chaque i -ème sous-modèle linéaire local, avec $1 \leq i \leq c$.

Les données d'identification entrée-sortie $\mathbf{z}_k = [\mathbf{x}_k^T, y_k]^T$, avec $1 \leq k \leq N$, et les degrés d'appartenance μ_{ik} de la matrice de partition floue sont regroupés dans les matrices suivantes :

$$\mathbf{X} = \begin{bmatrix} \mathbf{x}_1^T \\ \mathbf{x}_2^T \\ \vdots \\ \mathbf{x}_N^T \end{bmatrix}, \quad \mathbf{y} = \begin{bmatrix} y_1 \\ y_2 \\ \vdots \\ y_N \end{bmatrix}, \quad \mathbf{W}_i = \begin{bmatrix} \mu_{i1} & 0 & \cdots & 0 \\ 0 & \mu_{i2} & \cdots & 0 \\ \vdots & \vdots & \ddots & \vdots \\ 0 & 0 & \cdots & \mu_{iN} \end{bmatrix} \quad (3.129)$$

Les paramètres des conséquents $\mathbf{a}_i$ et d_i appartenant à la règle correspondant à l' i -ème cluster sont concaténés dans un seul vecteur de paramètres $\boldsymbol{\theta}_i$, donné par :

$$\boldsymbol{\theta}_i = [\mathbf{a}_i^T, d_i]^T \quad (3.130)$$

Afin de faciliter le calcul, la matrice de régression $\mathbf{X}$ est augmentée en ajoutant un vecteur-colonne unitaire, selon l'expression :

$$\mathbf{X}_e = [\mathbf{X}, \mathbf{1}] \quad (3.131)$$

Si les colonnes de $\mathbf{X}_e$ sont linéairement indépendantes et $\mu_{ik} > 0$ pour $1 \leq k \leq N$, alors la solution des moindres carrés de $\mathbf{y} = \mathbf{X}_e \boldsymbol{\theta} + \varepsilon$ où la k -ième paire de données $(\mathbf{x}_k, y_k)$ est pondérée par μ_{ik} , est donnée finalement par l'expression :

$$\boldsymbol{\theta}_i = (\mathbf{X}_e^T \mathbf{W}_i \mathbf{X}_e)^{-1} \mathbf{X}_e^T \mathbf{W}_i \mathbf{y} \quad (3.132)$$

Les paramètres $\mathbf{a}_i$ et d_i sont donnés respectivement par :

$$\mathbf{a}_i = [\theta_1, \theta_2, \dots, \theta_p], \quad d_i = \theta_{p+1} \quad (3.133)$$

3.4.4. Validation numérique du modèle flou

Pour évaluer la qualité de l'approximation (performance numérique) obtenue par les modèles flous Takagi-Sugeno, nous utilisons les critères suivants :

⁸¹ WLS : *Weighted Least Squares*.

RMSE, Erreur quadratique moyenne (en anglais *Root Mean Square Error*) : c'est une mesure globale sur le nombre total de points de l'écart par rapport à la valeur attendue. Sa valeur optimale est zéro; ce critère est défini par l'expression (3.134) :

$$RMSE = \sqrt{\frac{1}{N} \sum_{k=1}^N (y_k - \hat{y}_k)^2} \quad (3.134)$$

où $1 \leq k \leq N$ est le nombre de points considérés pour la modélisation, y est la sortie mesurée et $\hat{y}$ est la sortie du modèle.

VAF, Comptabilisation en pourcentage de la variance (en anglais *Variance Accounting For*) : introduit par Babuška *et al.* [BAB98b], ce critère permet d'évaluer en pourcentage la qualité d'un modèle en mesurant l'écart normalisé de la variance entre deux signaux. Sa valeur optimale est 100% quand les deux signaux sont égaux, s'ils sont différents, le VAF est inférieur. Le critère VAF est donné par l'expression (3.135) :

$$VAF = 100\% \left[1 - \frac{\text{var}(y - \hat{y})}{\text{var}(y)} \right] \quad (3.135)$$

Tcal, Temps de calcul : correspond au temps machine⁸² requis pour les calculs de construction du modèle flou à partir des données.

3.4.5. Outil logiciel pour la modélisation-identification floue de type TS

Dans le cadre du développement de nos travaux de thèse et grâce aux relations avec le laboratoire LAMIC⁸³ de l'Université Distrital en Colombie, nous avons développé le logiciel FMIDgraph [GRI07] [ORO07]. C'est une boîte à outils graphique sous MATLAB® pour l'identification floue de type Takagi-Sugeno à partir des données entré-sortie des systèmes. Elle est basée sur le logiciel FMID⁸⁴ développé à l'Université Technologique de Delft [BAB02] que nous avons amélioré par diverses fonctionnalités. Cet outil intègre dans un environnement graphique facile à utiliser, plusieurs outils de la méthodologie de modélisation abordée dans ce chapitre. Cinq catégories de fonctions peuvent être distinguées :

Algorithmes de clustering : FCM, GK, FER et RoFER, (cf. section 3.3).

Sélection de la structure (ordre) du modèle [GRI07] [UNR07], (cf. section 3.4).

Visualisation 2D et 3D du processus de modélisation (fonctions d'appartenance "en ligne" et approximation de systèmes non linéaires par des sous-modèles linéaires.

Validation de la partition floue et du modèle Takagi-Sugeno, (cf. sections 3.4.1 et 3.4.4).

Exemples de modélisation et identification floue de type TS de systèmes non linéaires statiques (cf. section 3.4.6) et dynamiques (cf. Chapitre 5).

Pour finaliser ce chapitre, nous présentons quelques exemples illustratifs que nous avons intégrés sous la boîte à outils FMIDgraph.

⁸² Disponible grâce à la commande `etime` sous MATLAB® v.6.5

⁸³ LAMIC : *Laboratorio de Automática, Microelectrónica e Inteligencia Computacional*, Universidad Distrital "Francisco José de Caldas", Facultad de Ingeniería, Bogotá, Colombie.

⁸⁴ FMID : *Fuzzy Modeling and Identification Toolbox* for use with Matlab.

3.4.6. Exemples

Afin d'illustrer la procédure de modélisation Takagi-Sugeno (TS) et les résultats obtenus pour les algorithmes GK et RoFER en présence de bruit et de points aberrants, nous considérons ici deux exemples qui sont des fonctions non linéaires statiques sur des espaces à deux et trois dimensions. La modélisation des systèmes dynamiques sera abordée dans le chapitre 5 en considérant l'application au bioprocédé de traitement des eaux usées.

Exemple 3.1. Modélisation TS de la fonction "sinus cardinal"

Le premier exemple considérée est une fonction non linéaire $y(x) = f(x)$ à une entrée et une sortie à laquelle a été ajouté du bruit et des points aberrants : $y(x) = f(x) + \xi$. Pour cet exemple nous avons choisi la fonction sinus cardinal, $f(x) = \text{sinusc}(x)$ représenté sur la Figure 3.10 et définie par (3.136) :

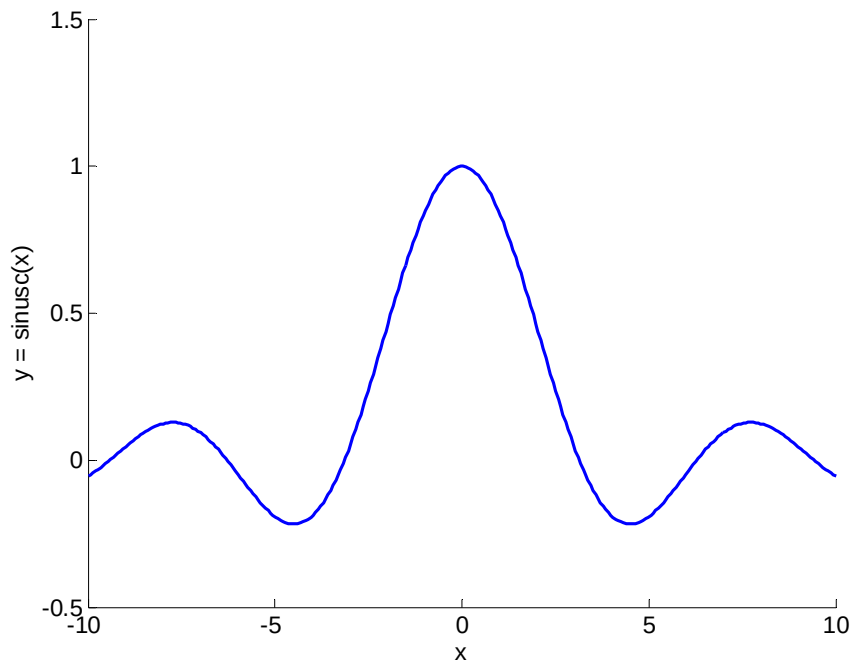


Figure 3.10. Fonction $y(x) = \text{sinusc}(x)$

$$f(x) = \text{sinusc}(x) = \begin{cases} \frac{\sin(x)}{x} & \text{si } x \neq 0 \\ 1 & \text{si } x = 0 \end{cases} \quad (3.136)$$

La présence de bruit et de points hors tendance est représentée par le modèle d'erreur grossière ξ , définie par l'expression suivante (3.137) :

$$\begin{aligned} \xi &= (1 - \varepsilon)G + \varepsilon H \\ G &\sim N(0, \sigma_1^2) \\ H &\sim N(0, \sigma_2^2) \end{aligned} \quad (3.137)$$

G et H sont des distributions de bruit blanc gaussien, ayant des écarts types $\sigma_2 > \sigma_1$. Le terme G permet de représenter la présence de bruit et le terme H la présence de points aberrants. Le paramètre ε permet d'ajuster la proportion de points aberrants dans l'ensemble

de la distribution. Pour cet exemple les valeurs des paramètres sont $\sigma_1 = 0,07$, $\sigma_2 = 5 \sigma_1$ et $\varepsilon = 0,15$.

En présence de bruit et de points hors tendance on considère alors l'expression $y(x) = \sinusc(x) + \xi$. Dans ce contexte, l'objectif de la modélisation floue Takagi-Sugeno est d'approximer la fonction originale $y(x) = \sinusc(x)$ en minimisant les effets perturbateurs introduits par le modèle d'erreur grossière ξ . Pour la modélisation une base de données de 300 points équidistants a été considérée avec $x \in [-10, 10]$, parmi lesquelles 30 points sont considérés comme des points aberrants. D'après la Figure 3.10 nous avons considéré une approximation de la fonction originale non linéaire par 6 régions linéaires, ce qui correspond au nombre de clusters pour l'algorithme GK. Pour sa part, l'algorithme RoFER a été initialisé avec 20 clusters. Avec un seuil de cardinalité de 15,5 nous avons obtenu un nombre final de clusters de 6, ce qui correspond au même nombre de clusters que celui fixé pour l'algorithme GK.

La Figure 3.11 illustre l'ensemble des données disponibles pour la modélisation Takagi-Sugeno. Les points considérés comme aberrants sont entourés par des cercles.

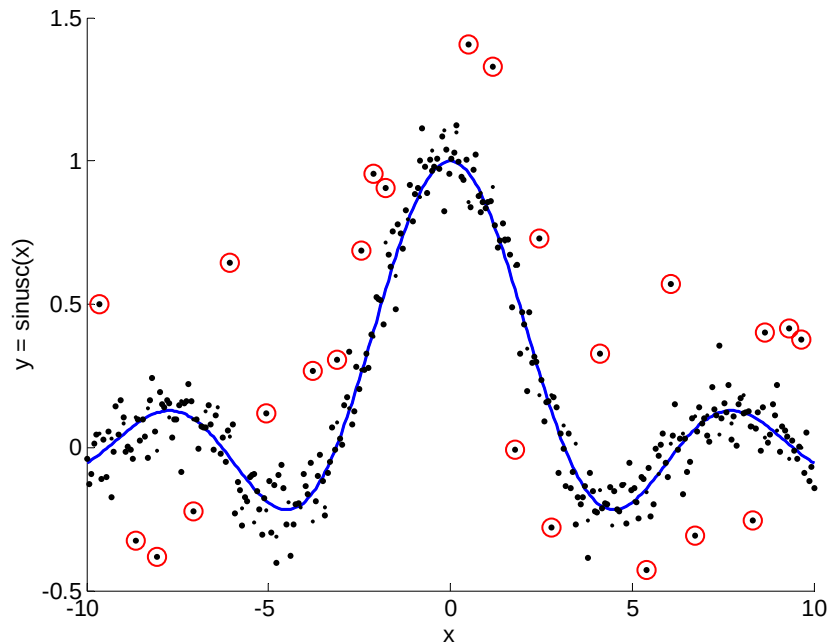


Figure 3.11. Fonction $y(x) = \sinusc(x)$ avec bruit et points aberrants (cercles)

Les résultats obtenus pour la modélisation floue Takagi-Sugeno en utilisant les algorithmes GK et RoFER sont résumés ci-dessous. Nous présentons l'approximation résultante pour la fonction originale (Figure 3.10) et celle avec bruit et points aberrants (Figure 3.11), les fonctions d'appartenance correspondantes, ainsi que les bases de règles obtenues et le positionnement des centres pour chacun des modèles (cf. Tableaux 3.2 à 3.5).

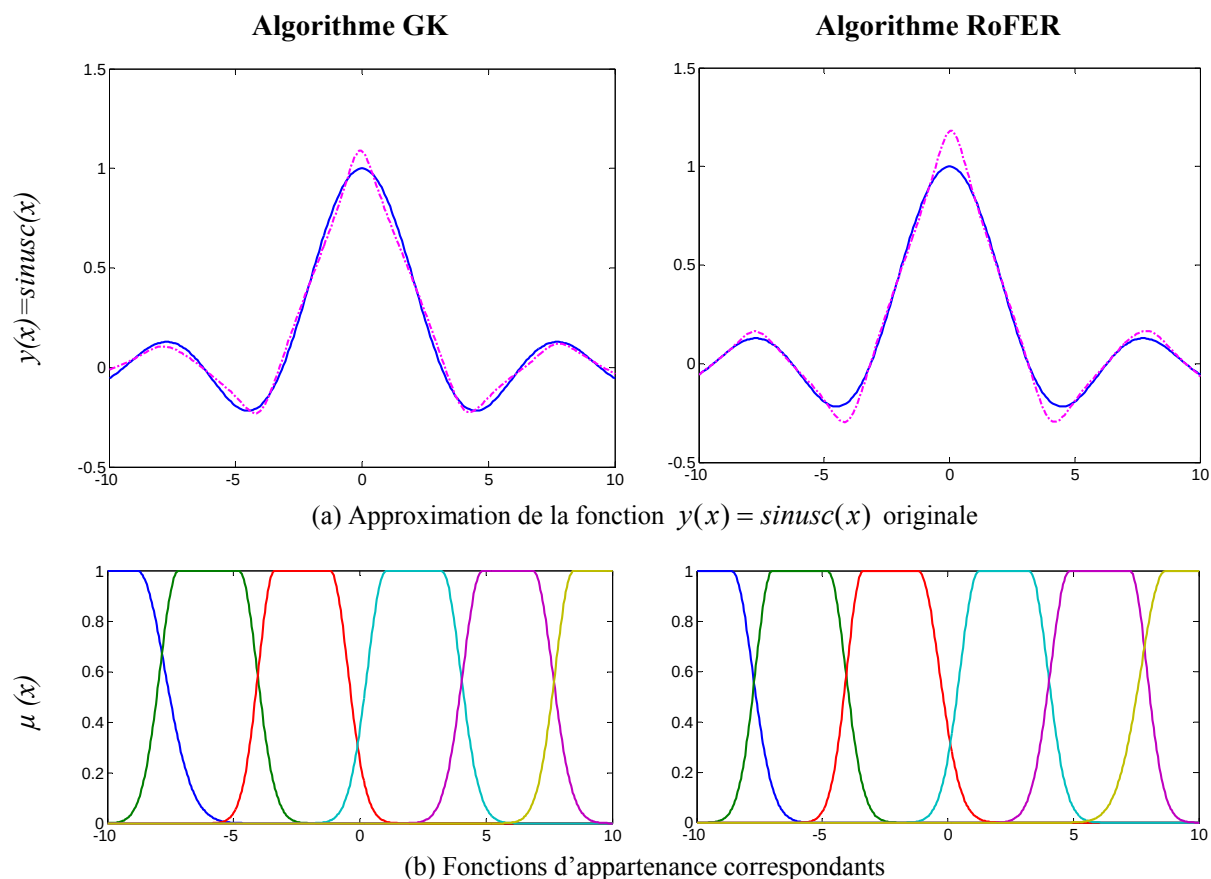


Figure 3.12. Approximation TS de la fonction $y(x) = \sinusc(x)$ - algorithmes GK et RoFER

Tableau 3.2. Règles du modèle flou TS pour la fonction $y(x) = \sinusc(x)$
algorithmes GK et RoFER

cluster	Antécédent		Conséquent GK	Conséquent RoFER
1	Si x est A_1	Alors	$y(x) = 4.75 \cdot 10^{-2}x + 4.64 \cdot 10^{-1}$	$y(x) = 9.60 \cdot 10^{-2}x + 9.04 \cdot 10^{-1}$
2	Si x est A_2	Alors	$y(x) = -8.77 \cdot 10^{-2}x - 5.69 \cdot 10^{-1}$	$y(x) = -1.21 \cdot 10^{-1}x - 7.69 \cdot 10^{-1}$
3	Si x est A_3	Alors	$y(x) = 3.28 \cdot 10^{-1}x + 1.10 \cdot 10^0$	$y(x) = 3.62 \cdot 10^{-1}x + 1.17 \cdot 10^0$
4	Si x est A_4	Alors	$y(x) = -3.18 \cdot 10^{-1}x + 1.08 \cdot 10^0$	$y(x) = -3.68 \cdot 10^{-1}x + 1.19 \cdot 10^0$
5	Si x est A_5	Alors	$y(x) = 8.88 \cdot 10^{-2}x - 5.73 \cdot 10^{-1}$	$y(x) = 1.16 \cdot 10^{-1}x - 7.42 \cdot 10^{-1}$
6	Si x est A_6	Alors	$y(x) = -6.84 \cdot 10^{-2}x + 6.53 \cdot 10^{-1}$	$y(x) = -1.06 \cdot 10^{-1}x + 9.97 \cdot 10^{-1}$

Tableau 3.3. Centres des clusters du modèle flou $y(x) = \sinusc(x)$ – algorithmes GK et RoFER

GK	RoFER
v_x	v_x
-7,99	-8,85
-6,00	-5,86
-2,30	-2,18
2,18	2,30
5,86	6,00
8,85	7,98

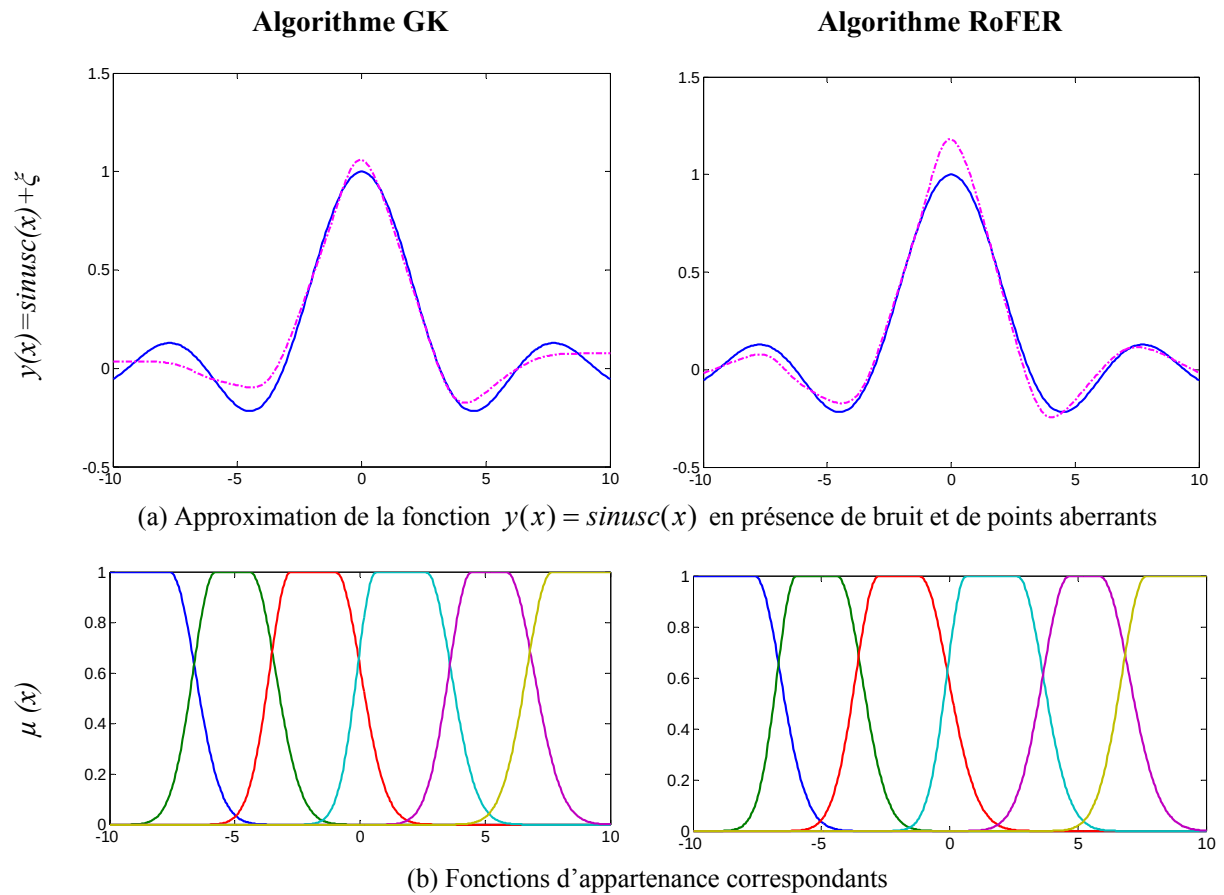


Figure 3.13. Approximation TS de la fonction $y(x) = \sinusc(x) + \xi$ - algorithmes GK et RoFER

Tableau 3.4. Règles du modèle flou TS pour la fonction $y(x) = \sinusc(x) + \xi$
algorithmes GK et RoFER

cluster	Antécédent	Conséquent GK	Conséquent RoFER
1	Si x est A_1	Alors $y(x) = 7.82 \cdot 10^{-4}x + 4.14 \cdot 10^{-2}$	$y(x) = 4.92 \cdot 10^{-2}x + 4.73 \cdot 10^{-1}$
2	Si x est A_2	Alors $y(x) = 2.85 \cdot 10^{-4}x - 8.05 \cdot 10^{-2}$	$y(x) = 1.03 \cdot 10^{-3}x - 1.45 \cdot 10^{-1}$
3	Si x est A_3	Alors $y(x) = 2.87 \cdot 10^{-1}x + 1.06 \cdot 10^0$	$y(x) = 3.31 \cdot 10^{-1}x + 1.15 \cdot 10^0$
4	Si x est A_4	Alors $y(x) = -3.14 \cdot 10^{-1}x + 1.06 \cdot 10^0$	$y(x) = -3.86 \cdot 10^{-1}x + 1.21 \cdot 10^0$
5	Si x est A_5	Alors $y(x) = 5.64 \cdot 10^{-2}x - 3.84 \cdot 10^{-1}$	$y(x) = 6.92 \cdot 10^{-2}x - 4.75 \cdot 10^{-1}$
6	Si x est A_6	Alors $y(x) = 2.41 \cdot 10^{-3}x + 5.29 \cdot 10^{-2}$	$y(x) = -6.27 \cdot 10^{-2}x + 6.11 \cdot 10^{-1}$

Tableau 3.5. Centres des clusters du modèle flou $y(x) = \sinusc(x) + \xi$ - algorithmes GK et RoFER

GK	RoFER
v_x	v_x
-8,24	-8,25
-5,04	-5,10
-1,55	-1,61
1,83	1,83
5,25	5,29
8,40	8,44

Pour évaluer la qualité de l'approximation obtenue par les modèles flous Takagi-Sugeno, nous utilisons les critères définis dans la section 3.4.4. Les résultats obtenus pour l'évaluation de l'approximation de la fonction $y(x) = \sinusc(x)$ sans et en présence de bruit et de points aberrants sont détaillés dans le Tableau 3.6.

Tableau 3.6. Evaluation de la qualité de l'approximation de la fonction $y(x) = \sinusc(x)$

Critère	$y(x) = \sinusc(x)$		$y(x) = \sinusc(x) + \xi$	
	GK	RoFER	GK	RoFER
RMSE ($\times 10^{-2}$)	3,47	4,57	6,01	5,49
VAF	99,03%	98,33%	97,17%	97,64%
Tcal ⁸⁵ [s]	0,621	10,785	0,621	10,785

En considérant les résultats précédents, on peut noter qu'en présence de bruit et de points aberrants, l'algorithme RoFER donne de meilleurs résultats au niveau numérique, mais aussi au niveau de l'interprétabilité locale, garantissant ainsi une meilleure qualité de l'approximation. En effet, comme on peut le constater aux extrémités droite et gauche de la Figure 3.13 (a) et (b)), les sous-modèles locaux obtenus en utilisant l'algorithme GK sont fortement biaisés par rapport à la fonction originale $y(x) = \sinusc(x)$. Par contre, les résultats obtenus en utilisant l'algorithme RoFER démontrent une sensibilité moindre. En contre partie, le temps de calcul requis pour la convergence de l'algorithme RoFER est considérablement plus grand que celui de l'algorithme GK (de l'ordre de 15 fois).

Remarques

- R 3.10** Dans la Figure 3.12 (a) on peut s'apercevoir que pour les deux algorithmes il y a un dépassement considérable de l'approximation sur les régions courbes. Ce dépassement est toutefois normal, il est dû au mécanisme standard d'inférence de la sortie globale dans le modèle Takagi-Sugeno. Si bien qu'une amélioration sur le mécanisme d'inférence/interpolation a été proposée [BAB96a] (en remplaçant le calcul de la moyenne pondérée par une fonction maximum adoucie), cette méthode demande l'ajustement des certains paramètres associés. Dans le but de simplifier les calculs et travailler dans des conditions standards, nous utilisons le mécanisme classique d'inférence pour le calcul de la sortie globale du modèle Takagi-Sugeno.
- R 3.11** Si la méthode d'inférence a une influence sur l'approximation obtenue, il faut noter aussi l'effet d'autres facteurs sur la qualité de l'approximation. Le mécanisme de projection, commun aux 2 méthodes de clustering, introduit un biais sur les sous-modèles linéaires. La paramétrisation des fonctions d'appartenance afin d'obtenir des ensembles flous unidimensionnels sur les variables d'entrée peut entraîner un changement de forme, généralement avec une couverture plus grande sur l'univers de discours que les ensembles multidimensionnels correspondants. L'effet final est une partition qui n'est pas forcément du type Ruspini (en effet, dans certains cas la somme des degrés d'appartenance peut être supérieure à 1). Cela peut se traduire par des biais sur les valeurs attendues des paramètres des conséquents. Voir par exemple la Figure 3.13 (b) pour le cas $x = -5$.

⁸⁵ Pour les résultats reportés dans ce manuscrit nous avons utilisé une machine Pentium 4 3,02 GHz avec 256 Mo de mémoire RAM et système d'exploitation Windows XP.

Exemple 3.2. Modélisation TS de la fonction "vallée de Rosenbrock"

Le deuxième exemple considérée est une fonction non linéaire dans un espace à trois dimensions $z(x, y) = f(x, y) + \xi$ à laquelle a été ajouté du bruit et des points hors tendance symbolisés par ξ . Pour cet exemple nous avons choisi la fonction vallée de Rosenbrock, $f(x, y) = \text{rosenbrock}(x, y)$ représenté sur la Figure 3.14 et définie par (3.138) :

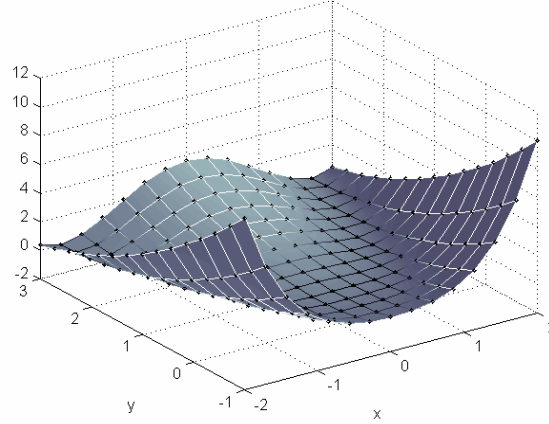


Figure 3.14. Fonction vallée de Rosenbrock

$$f(x, y) = \text{rosenbrock}(x, y) = \text{scale} * \left(\frac{100 \cdot (y - x^2)^2 + (1 - x)^2}{2509} \right) \quad (3.138)$$

avec $\text{scale} = 10$, $x \in [-2, 2]$ et $y \in [-1, 3]$. Le facteur 2509 permet la normalisation de l'expression (3.138) dans l'espace d'entrée-sortie $X \times Y$ choisi. La présence de bruit et de points aberrants est représentée par le modèle d'erreur grossière ξ , définie par l'expression (3.137). Pour cet exemple les valeurs des paramètres sont $\sigma_1 = 0,5$, $\sigma_2 = 5 \sigma_1$ et $\varepsilon = 0,2$. Pour la modélisation une base de données de 225 points équidistants sur (x, y) avec 25 points aberrants a été considérée. L'algorithme RoFER a été exécuté avec un nombre initial de clusters de 15. Avec un seuil de cardinalité de 4,9 nous avons obtenu un nombre final de clusters de 6, qui correspond au même nombre de clusters que dans l'algorithme GK. La Figure 3.15 illustre l'ensemble des données disponibles pour la modélisation Takagi-Sugeno. Les points considérés comme aberrants sont entourés par des cercles.

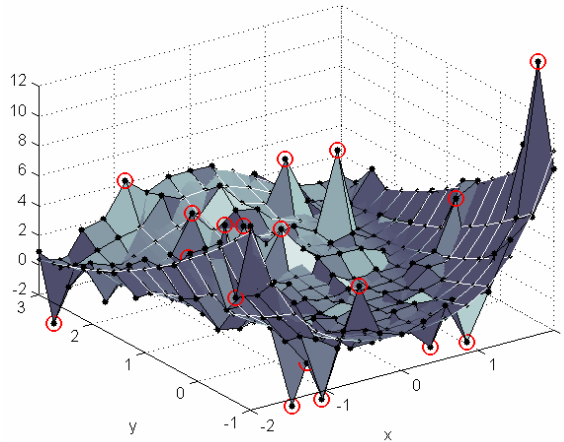


Figure 3.15. Fonction vallée de Rosenbrock avec bruit et points aberrants (cercles)

Les résultats obtenus pour la modélisation floue Takagi-Sugeno sans et avec bruit et points aberrants en utilisant les algorithmes GK et RoFER sont résumés ci-dessous.

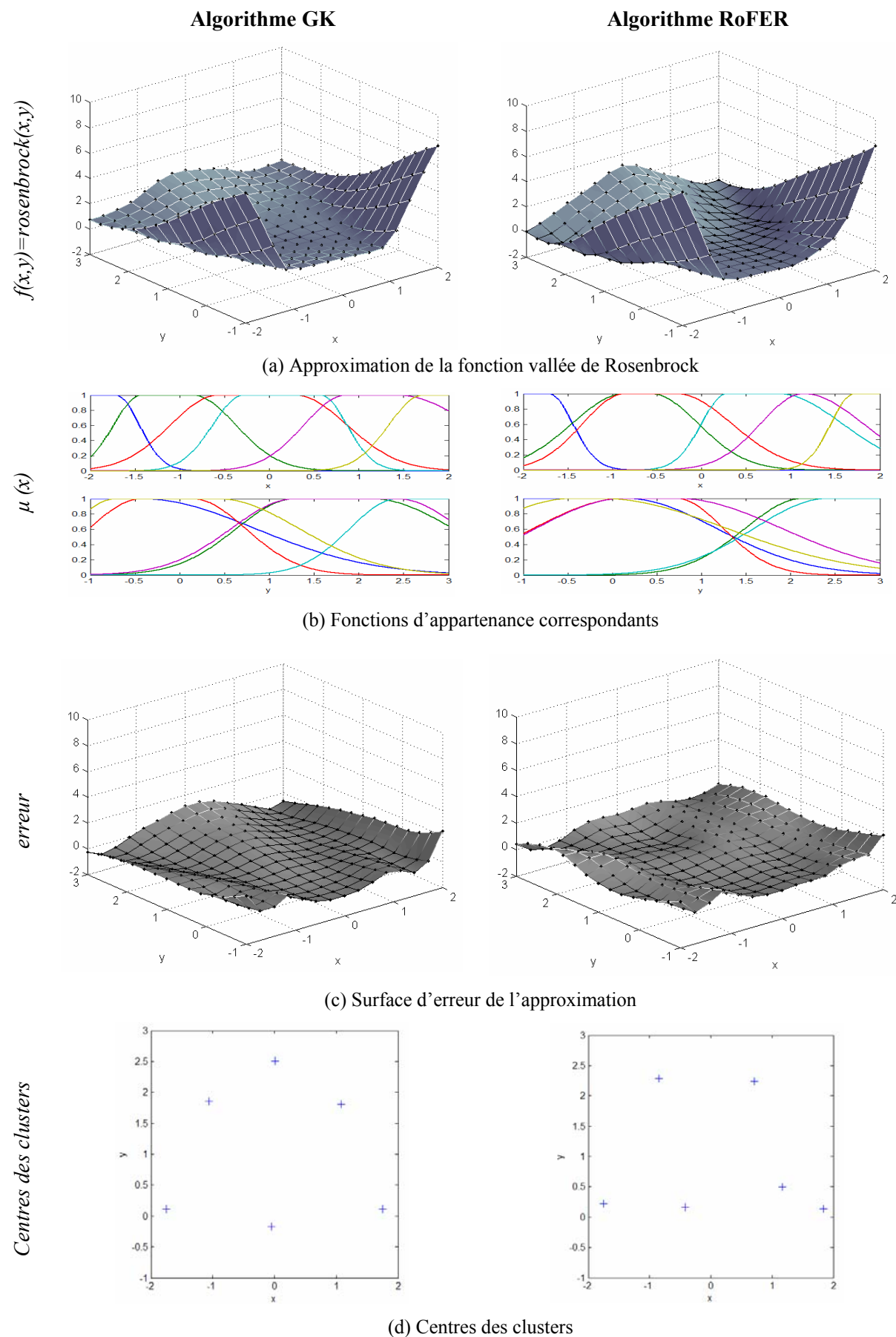


Figure 3.16. Approximation TS de la fonction vallée de Rosenbrock

Tableau 3.7. Règles du modèle flou TS pour la fonction vallée de Rosenbrock – algorithme GK

cluster	Antécédent			Conséquent GK
1	Si x est A_{11}	et y est A_{12}	Alors	$z(x, y) = -7.37 \cdot 10^0 x - 2.09 \cdot 10^0 y - 8.47 \cdot 10^0$
2	Si x est A_{21}	et y est A_{22}	Alors	$z(x, y) = 2.14 \cdot 10^{-1} x + 2.48 \cdot 10^{-1} y + 5.72 \cdot 10^{-1}$
3	Si x est A_{31}	et y est A_{32}	Alors	$z(x, y) = -1.83 \cdot 10^{-1} x - 3.95 \cdot 10^{-1} y + 4.62 \cdot 10^{-1}$
4	Si x est A_{41}	et y est A_{42}	Alors	$z(x, y) = -2.75 \cdot 10^{-1} x + 1.58 \cdot 10^0 y - 1.85 \cdot 10^0$
5	Si x est A_{51}	et y est A_{52}	Alors	$z(x, y) = 2.91 \cdot 10^{-1} x - 6.11 \cdot 10^{-2} y + 7.12 \cdot 10^{-1}$
6	Si x est A_{61}	et y est A_{62}	Alors	$z(x, y) = 5.77 \cdot 10^0 x - 1.88 \cdot 10^0 y - 5.73 \cdot 10^0$

Tableau 3.8. Règles du modèle flou TS pour la fonction vallée de Rosenbrock – algorithme RoFER

cluster	Antécédent			Conséquent ROFER
1	Si x est A_{11}	et y est A_{12}	Alors	$z(x, y) = -8.25 \cdot 10^0 x - 2.40 \cdot 10^0 y - 1.01 \cdot 10^1$
2	Si x est A_{21}	et y est A_{22}	Alors	$z(x, y) = 1.49 \cdot 10^0 x + 1.12 \cdot 10^0 y - 2.48 \cdot 10^{-1}$
3	Si x est A_{31}	et y est A_{32}	Alors	$z(x, y) = -1.79 \cdot 10^{-1} x - 1.47 \cdot 10^{-1} y + 1.64 \cdot 10^{-1}$
4	Si x est A_{41}	et y est A_{42}	Alors	$z(x, y) = -1.61 \cdot 10^0 x + 1.20 \cdot 10^0 y - 3.69 \cdot 10^{-1}$
5	Si x est A_{51}	et y est A_{52}	Alors	$z(x, y) = 1.67 \cdot 10^0 x - 7.10 \cdot 10^{-1} y - 1.02 \cdot 10^0$
6	Si x est A_{61}	et y est A_{62}	Alors	$z(x, y) = 9.71 \cdot 10^0 x - 2.63 \cdot 10^0 y - 1.29 \cdot 10^1$

Tableau 3.9. Centres des clusters du modèle flou pour la fonction vallée de Rosenbrock sans bruit

GK		RoFER	
v_x	v_y	v_x	v_y
$-1,76 \cdot 10^0$	$1,12 \cdot 10^{-1}$	$-1,76 \cdot 10^0$	$2,18 \cdot 10^{-1}$
$-1,06 \cdot 10^0$	$1,85 \cdot 10^0$	$-8,47 \cdot 10^{-1}$	$2,29 \cdot 10^0$
$-5,25 \cdot 10^{-2}$	$-1,70 \cdot 10^{-1}$	$-4,15 \cdot 10^{-1}$	$1,59 \cdot 10^{-1}$
$1,27 \cdot 10^{-2}$	$2,51 \cdot 10^0$	$7,12 \cdot 10^{-1}$	$2,23 \cdot 10^0$
$1,08 \cdot 10^0$	$1,81 \cdot 10^0$	$1,17 \cdot 10^0$	$4,96 \cdot 10^{-1}$
$1,75 \cdot 10^0$	$1,11 \cdot 10^{-1}$	$1,84 \cdot 10^0$	$1,36 \cdot 10^{-1}$

Les résultats obtenus pour l'évaluation de l'approximation de la fonction $z(x, y) = \text{rosenbrock}(x, y)$ sans bruit sont détaillés dans le Tableau 3.10 :

Tableau 3.10. Evaluation de la qualité de l'approximation de la fonction $z(x, y) = \text{rosenbrock}(x, y)$

Critère	$z(x, y) = \text{rosenbrock}(x, y)$	
	GK	RoFER
RMSE ($\times 10^{-1}$)	5,24	5,66
VAF	92,60%	95,48%
Tcal [s]	0,661	19,98

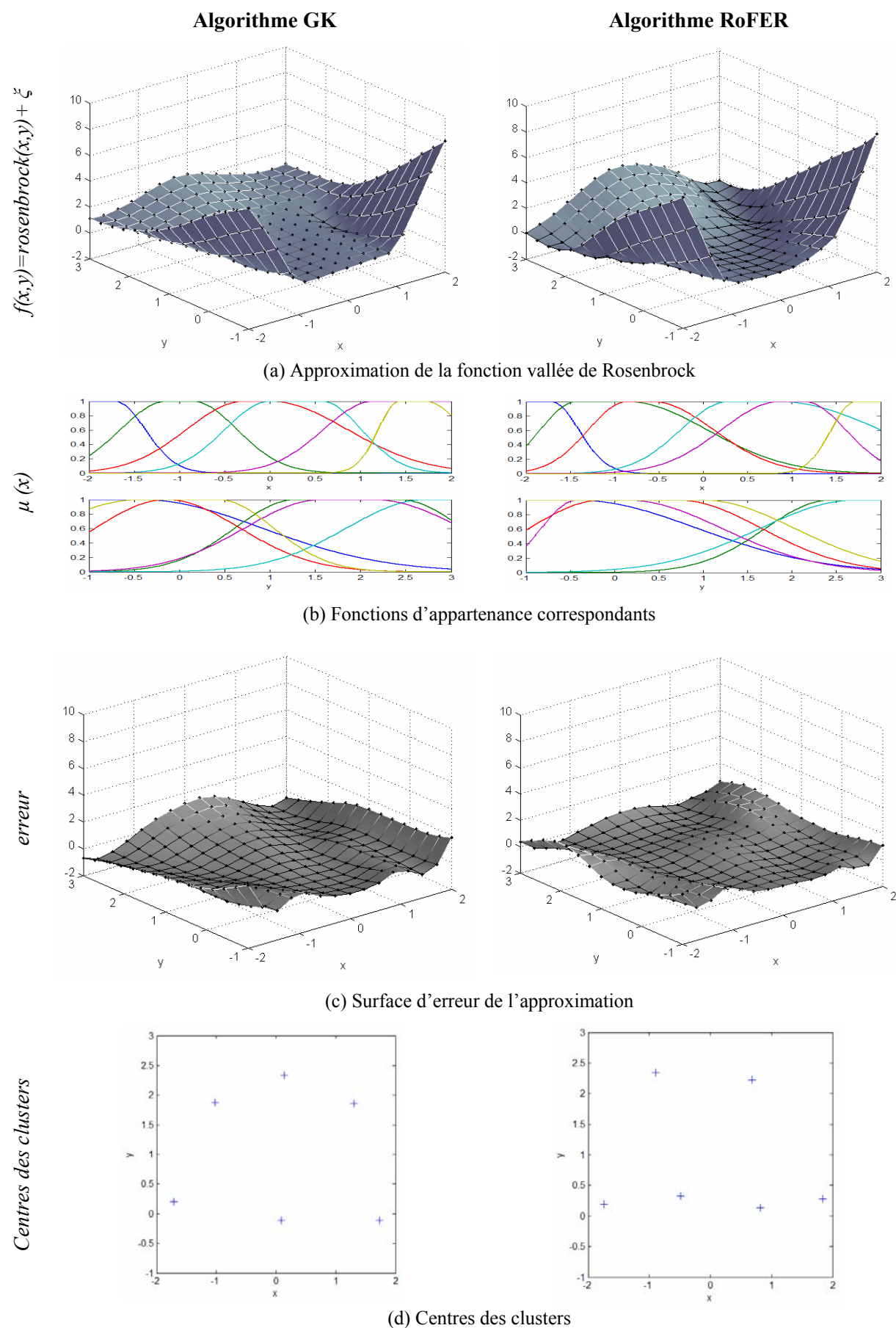


Figure 3.17. Approximation TS de la fonction vallée de Rosenbrock avec bruit et points aberrants

Tableau 3.11. Règles du modèle flou TS pour la fonction vallée de Rosenbrock avec bruit et points aberrants– algorithme GK

cluster	Antécédent		Conséquent GK	
1	Si x est A_{11}	et y est A_{12}	Alors	$z(x, y) = -5.81 \cdot 10^0 x - 1.69 \cdot 10^0 y - 5.91 \cdot 10^0$
2	Si x est A_{21}	et y est A_{22}	Alors	$z(x, y) = 1.86 \cdot 10^{-1} x + 2.54 \cdot 10^{-1} y + 8.16 \cdot 10^{-1}$
3	Si x est A_{31}	et y est A_{32}	Alors	$z(x, y) = 2.21 \cdot 10^{-1} x - 1.22 \cdot 10^{-1} y + 3.89 \cdot 10^{-1}$
4	Si x est A_{41}	et y est A_{42}	Alors	$z(x, y) = -7.20 \cdot 10^{-1} x + 1.18 \cdot 10^0 y - 8.31 \cdot 10^{-1}$
5	Si x est A_{51}	et y est A_{52}	Alors	$z(x, y) = 4.17 \cdot 10^{-1} x - 1.88 \cdot 10^{-1} y + 6.91 \cdot 10^{-1}$
6	Si x est A_{61}	et y est A_{62}	Alors	$z(x, y) = 7.17 \cdot 10^0 x - 2.36 \cdot 10^0 y - 8.14 \cdot 10^0$

Tableau 3.12. Règles du modèle flou TS pour la fonction vallée de Rosenbrock avec bruit et points aberrants– algorithme RoFER

cluster	Antécédent		Conséquent ROFER	
1	Si x est A_{11}	et y est A_{12}	Alors	$z(x, y) = -7.70 \cdot 10^0 x - 2.25 \cdot 10^0 y - 9.25 \cdot 10^0$
2	Si x est A_{21}	et y est A_{22}	Alors	$z(x, y) = 1.71 \cdot 10^0 x + 1.15 \cdot 10^0 y + 7.04 \cdot 10^{-2}$
3	Si x est A_{31}	et y est A_{32}	Alors	$z(x, y) = -1.15 \cdot 10^{-1} x - 2.71 \cdot 10^{-2} y + 2.75 \cdot 10^{-1}$
4	Si x est A_{41}	et y est A_{42}	Alors	$z(x, y) = -1.60 \cdot 10^0 x + 9.03 \cdot 10^{-1} y + 4.28 \cdot 10^{-1}$
5	Si x est A_{51}	et y est A_{52}	Alors	$z(x, y) = 1.66 \cdot 10^0 x - 7.86 \cdot 10^{-1} y - 9.71 \cdot 10^{-1}$
6	Si x est A_{61}	et y est A_{62}	Alors	$z(x, y) = 8.39 \cdot 10^0 x - 2.56 \cdot 10^0 y - 1.01 \cdot 10^1$

Tableau 3.13. Centres des clusters du modèle flou pour la fonction vallée de Rosenbrock avec bruit et points aberrants

GK		RoFER	
v_x	v_y	v_x	v_y
$-1,76 \cdot 10^0$	$1,12 \cdot 10^{-1}$	$-1,74 \cdot 10^0$	$1,88 \cdot 10^{-1}$
$-1,06 \cdot 10^0$	$1,85 \cdot 10^0$	$-8,91 \cdot 10^{-1}$	$2,34 \cdot 10^0$
$-5,25 \cdot 10^{-2}$	$-1,70 \cdot 10^{-1}$	$-4,85 \cdot 10^{-1}$	$3,32 \cdot 10^{-1}$
$1,27 \cdot 10^{-2}$	$2,51 \cdot 10^0$	$6,83 \cdot 10^{-1}$	$2,22 \cdot 10^0$
$1,08 \cdot 10^0$	$1,81 \cdot 10^0$	$8,17 \cdot 10^{-1}$	$1,30 \cdot 10^{-1}$
$1,75 \cdot 10^0$	$1,11 \cdot 10^{-1}$	$1,84 \cdot 10^0$	$2,79 \cdot 10^{-1}$

Les résultats obtenus pour l'évaluation de l'approximation de la fonction $z(x, y) = \text{rosenbrock}(x, y) + \xi$ avec bruit et points aberrants sont détaillés dans le Tableau 3.14 :

Tableau 3.14. Evaluation de la qualité de l'approximation de la fonction $z(x, y) = \text{rosenbrock}(x, y) + \xi$

Critère	$z(x, y) = \text{rosenbrock}(x, y) + \xi$	
	GK	RoFER
RMSE ($\times 10^{-1}$)	6,02	4,89
VAF	90,23%	95,64%
Tcal [s]	0,661	19,98

3.5. Conclusions

Ce chapitre nous a permis de faire le tour d'horizon des principes théoriques nécessaires à la compréhension de la modélisation et de l'identification floues de systèmes en utilisant l'approche par clustering à partir des données entrée-sortie.

Après d'avoir introduit des concepts de base concernant la structure générale et les différents types de modèles flous, nous avons étudié plus particulièrement le modèle de type Takagi-Sugeno (TS). En effet, au cours des dix dernières années la discipline a évolué d'une façon graduelle mais décidée, vers une utilisation pratiquement exclusif des systèmes flous dans lesquels le conséquent des règles utilise des variables numériques sous la forme des fonctions (modèle de type TS) plutôt que des variables linguistiques (modèle de type Mamdani) [GAL01] [SAL05]. Ces modèles conservent les caractéristiques de transparence et d'interprétabilité linguistique qui distinguent les modèles flous d'autres approches similaires de type boîte noire. Dans le cas de l'identification floue des systèmes, le formalisme Takagi-Sugeno est mieux adapté à une démarche plus systématique pour la construction de modèles non linéaires multivariés, grâce à leur bonne capacité d'interpolation numérique et d'"apprentissage" à partir de données.

Dans le cadre de nos travaux, les modèles flous de type TS sont considérés sous l'optique d'une approche de modélisation multi-linéaire, qui essaye de résoudre un problème complexe de modélisation en le décomposant en plusieurs sous-problèmes plus simples. La théorie des ensembles flous offre alors un excellent outil pour représenter l'incertitude associée à la tâche de décomposition, en fournissant des transitions douces entre les sous-modèles linéaires et afin d'intégrer divers types de connaissance dans un même cadre. A ce propos, nous avons considéré des méthodes de coalescence floue (clustering) basées sur la minimisation itérative d'une fonction objectif (critère de classification), afin de générer automatiquement la décomposition à partir des données entrée-sortie du système. Dans cette approche, le nombre de clusters correspond à celui des sous-modèles linéaires dans le modèle TS.

En ce qui concerne les méthodes de clustering, afin de trouver une partition floue appropriée, nous avons abordé plusieurs algorithmes avec *prototypes* :

- o *vectoriels* (centres) de la même dimension que les objets de données (algorithmes FCM (clusters hypersphériques) et GK (clusters hyperellipsoïdaux)),
- o des objets géométriques "de plus haut niveau", tels que des *sous-espaces linéaires* (algorithme FCRM),
- o ou bien des prototypes "*mixtes*", tels que l'algorithme de *régression floue ellipsoïdal* (FER) qui combine les mesures de distance des algorithmes GK des axes parallèles (GKP) et FCRM, ainsi que sa version *robuste* en présence de bruit et de points aberrants (algorithme RoFER) [BAR03] [GRI04].

Ce dernier est caractérisé par l'utilisation d'une méthode d'"agglomération compétitive" [FRI99] qui à partir d'un nombre surestimé de clusters à l'initialisation et la définition d'un certain seuil de cardinalité robuste, permet de trouver "automatiquement" un nombre convenable de clusters pour la partition. En même temps l'algorithme fournit les paramètres des conséquents des règles du modèle floue *affine* TS.

Ensuite, nous avons abordé la méthodologie générale d'identification de modèles TS à partir des données, en décrivant les différentes étapes depuis la conception de l'expérience et

l'acquisition des données jusqu'à la validation du modèle final. En particulier nous avons mis l'accent sur quelques aspects généraux et communs aux algorithmes traités, afin de passer de l'application des techniques de clustering à l'obtention du modèle flou. Ces aspects comprennent les pas suivants :

Validation du nombre de clusters, en utilisant les mesures de validité de la partition.

Génération des fonctions d'appartenance des antécédents, qui permet l'obtention des règles avec des ensembles flous multidimensionnels, ou bien des règles dans leur forme décomposée en utilisant le mécanisme de projection sur les axes d'entrée.

Obtention des paramètres des conséquents, à partir de la structure propre de la matrice de covariance floue (cas de l'algorithme GK) ou en général, en utilisant des techniques de type moindres carrés.

Validation numérique du modèle flou.

Dans le cadre de nos recherches et en se basant sur le logiciel FMID [BAB02], nous avons développé FMIDgraph [GRI07] [ORO07] [UNR07], une boîte à outils graphique sous MATLAB® pour l'identification floue de type TS, qui intègre dans un environnement graphique facile à utiliser, plusieurs éléments et fonctionnalités (e.g. visualisation 2D et 3D) concernant la méthodologie de modélisation abordée dans ce chapitre.

Finalement, nous avons présenté quelques exemples de simulation qui ont permis d'illustrer l'approche suivie pour la modélisation floue TS de systèmes non linéaires en utilisant des méthodes de clustering à partir des données entré-sortie des systèmes. Les résultats ont montré l'applicabilité des algorithmes de clustering comme GK et RoFER, avec une meilleure qualité de l'approximation *locale* du système pour ce dernier en présence de bruits et de points aberrants, malgré le temps de calcul requis pour la convergence (de l'ordre de 15 fois plus élevé). Nous considérons que l'utilisation de tels algorithmes pourrait être intéressante pour des applications de modélisation hors-ligne où le temps de calcul ne constitue pas un facteur critique.

Sans vouloir généraliser les résultats expérimentaux obtenus, il paraît raisonnable de considérer que la caractérisation d'un système non linéaire par une approche multimodèle est un chemin prometteur pour un rapprochement entre la théorie mathématique et la pratique industrielle de la commande automatique. En effet, en obtenant une représentation simple par zones d'opération définies par les fonctions d'appartenance dans les antécédents, le modèle flou TS permet la représentation des systèmes non linéaires statiques ou dynamiques par une concaténation de sous-modèles linéaires associés aux conséquents des règles. Cette approche s'avère intéressante du point de vue de la commande et sera exploitée à ce propos, comme nous le verrons dans le chapitre suivant.

Chapitre 4

Commande floue TS basée sur modèle

Dans ce chapitre, nous considérons différentes possibilités qu'offrent les modèles flous de type Takagi-Sugeno pour la commande des systèmes non linéaires en utilisant la philosophie de commande du type compensation parallèle distribuée (PDC). Nous commençons par un aperçu général sur la synthèse (stabilisation) des contrôleurs flous basée sur un modèle flou homogène TS du système. Ensuite nous développons spécifiquement une commande sous-optimale linéaire quadratique par retour d'état basée sur un modèle affine TS en temps discret. Pour cela, nous proposons trois lois de commande qui tiennent compte du terme indépendant caractéristique du modèle affine. Les lois de commande sont de type régulateur, régulateur avec compensation de gain statique et régulateur avec action intégrale, ce qui permet l'annulation des erreurs stationnaires de poursuite par rapport à des références de type échelon dans la commande.

Sommaire

4. Commande floue TS basée sur modèle.....	143
4.1. Stabilité et stabilisation à partir de modèles flous	146
4.1.1. Modeles flous dynamiques de type Takagi Sugeno (TS)	146
4.1.2. Stabilité des modèles flous.....	148
4.1.3. Stabilisation des modèles flous TS.....	150
4.2. Commande PDC sous optimale par retour d'état	153
4.2.1. Commande LQR pour des modèles affines à temps discret	154
4.2.2. Commande PDC avec compensation de gain statique	158
4.2.3. Commande PDC avec ajout d'intégrateurs.....	160
4.2.4. Choix des matrices de pondération Q et R.....	163
4.3. Conclusions	165

A partir de la disponibilité d'un modèle flou, celui-ci peut être utilisé pour la synthèse d'un contrôleur de deux manières [SAL05]. En premier lieu, n'importe quelle technique (non linéaire) *basée sur un modèle* peut être appliquée. Parmi les techniques les plus fréquemment utilisées dans cette approche, on trouve la linéarisation entrée-sortie [BOU03] et par rétroaction, la commande prédictive et les techniques basées sur l'inversion du modèle [BAB98a]. Deuxièmement, le contrôleur lui-même peut être un système flou dont la structure correspond à la structure du modèle flou du procédé. Cette idée, appelée dans le cas de systèmes flous de type TS "compensation parallèle distribuée" (PDC)⁸⁶ [WAN95], est très fructueuse. En fait, le plus grand nombre de résultats, concernant la stabilité des modèles Takagi-Sugeno tant en temps continu que dans le cas discret, sont basés sur des commandes de type PDC. Dans ce chapitre nous donnons dans la première partie, largement basée sur l'ouvrage [GUE03], un bref aperçu général sur la stabilisation des contrôleurs flous basée sur un modèle flou TS du système. Les développements mathématiques rapportés correspondent aux modèles flous *homogènes* TS, tant en temps continu qu'en temps discret. Dans la seconde partie, en utilisant l'approche PDC, nous élaborons une commande sous-optimale linéaire quadratique par retour d'état basée sur des modèles *affines* TS en temps discret. Nous proposons le développement de trois lois de commande qui prennent en compte explicitement le terme indépendant dans le conséquent du modèle. Pour une étude plus vaste concernant la commande floue de systèmes, nous invitons le lecteur à se référer sur les ouvrages suivantes : [DRI93] [WAN94] [BAB98a] [FOU03a].

La philosophie de la commande de type compensation parallèle distribuée (PDC), consiste à calculer une loi de commande linéaire par retour d'état, pour chaque sous modèle du modèle flou. La détermination d'une loi de commande revient à déterminer pour chaque modèle local des gains matriciels, par exemple en utilisant une synthèse quadratique ou des LMI. La Figure 4.1 illustre le concept de ce type de commande :

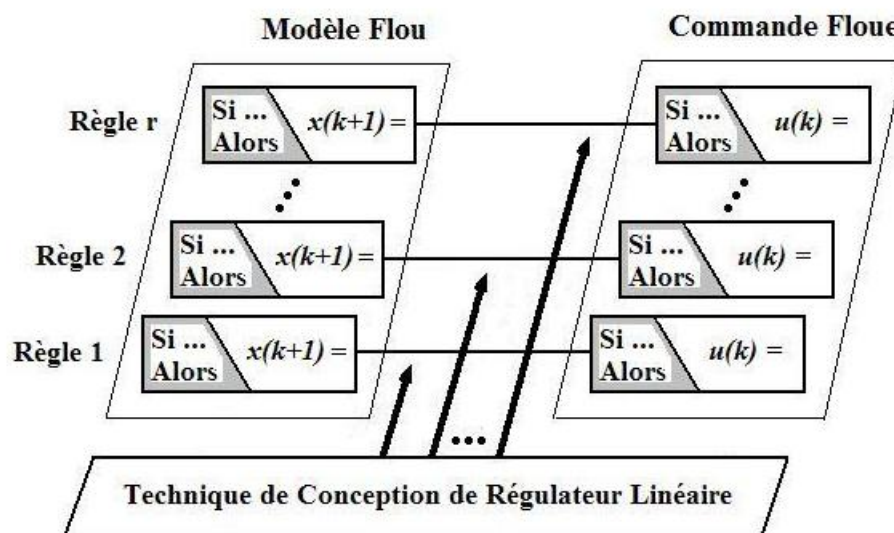


Figure 4.1. Représentation du concept de compensation parallèle distribuée (PDC)

La commande de type PDC a pour but d'intégrer dans une seule loi de commande globale les lois de commande individuelles issues de l'approche multi-modèle. La partie antécédent des règles reste la même que pour le modèle flou Takagi-Sugeno tandis que la partie conséquent est remplacée par une loi de commande par retour d'état. Dans ce contexte,

⁸⁶ PDC: *Parallel Distributed Compensation*.

chaque sous modèle du modèle flou de type Takagi-Sugeno est stabilisé localement par une loi linéaire. La loi de commande globale, qui en général est non linéaire, est une fusion floue des lois de commande linéaires. Pour l'application de la commande de type PDC il est nécessaire que tous les sous modèles linéaires soient stabilisables ; cette hypothèse est supposée vérifiée dans toute la suite de manuscrit.

4.1. Stabilité et stabilisation à partir de modèles flous

La stabilité de systèmes de type Takagi-Sugeno a été étudiée à partir de plusieurs points de vue : approche de type Lyapunov [TAN92], analogie avec des systèmes linéaires variant dans le temps [THA97], mise sous la forme d'un système linéaire avec des incertitudes de modélisation [KIR98], etc. Les conditions de stabilité obtenues dans ces travaux ne sont que suffisantes, puisqu'elles ne prennent pas en compte la partie antécédent des règles. Des travaux proposent de prendre en compte cette partie antécédent pour aboutir à des résultats moins conservatifs [KIM95] [MAR95]. Dans le cas de la stabilité quadratique, des conditions nécessaires et suffisantes ont été proposées [MAR95]. La plupart des travaux font appel, pour vérifier les conditions de stabilité, aux outils LMI (*Linear Matrix Inequalities*) [BOY94].

Dans tous les cas, le problème est de pouvoir déterminer une loi de commande stabilisante. On peut faire la synthèse d'une loi de commande linéaire [KIM95] [KIR98], mais, très souvent une approche de type PDC (*Parallel Distributed Compensation*) est utilisée pour calculer une loi de commande sur ces systèmes [MA98] [TAN95] [TAN98] [ZHA96]. Elle consiste à utiliser la même partie antécédent des règles que celle du système flou de type TS, et, dans la partie conclusion, des commandes locales par retour d'état. Dans ce contexte, chaque sous-système d'un système flou de type TS est stabilisé localement par une loi linéaire [TAN95]. Des approches utilisant la représentation d'état et des observateurs flous ont également été proposées [FEN97] [TAN94] et, dans le cas où les prémisses de règles utilisent des variables mesurables, un principe de séparation a également été démontré [MA98].

Nous allons faire un bref rappel sur les modèles flous continus et discrets afin de faciliter la formulation des théorèmes de stabilité.

4.1.1. Modèles flous dynamiques de type Takagi Sugeno (TS)

Avec les notations R^i : ($i=1,2,\dots,r$) la i -ème règle, r étant le nombre de règles « si...alors », F_j^i : ($j=1,2,\dots,r$) les sous-ensembles flous des prémisses, $x(t) \in R^n$ le vecteur d'état, $u(t) \in R^m$ le vecteur des entrées, $y(t) \in R^q$ le vecteur des sorties, $A_i \in R^{n \times n}$, $B_i \in R^{n \times m}$, $C_i \in R^{q \times n}$ et $z_1(t) \sim z_p(t)$ les variables des prémisses (variables dépendant de l'état et/ou des entrées), les modèles flous TS sont représentés sous la forme suivante, où MFC et MFD indiquent, respectivement, les modèles flous continus et les modèles flous discrets.

Modèles flous continus (MFC)

Règle R^i du modèle :

$$\text{Si } z_1(t) \text{ est } F_1^i \text{ et ...et } z_p(t) \text{ est } F_p^i \text{ alors } \begin{cases} \dot{x}(t) = A_i x(t) + B_i u(t) \\ y(t) = C_i x(t) \end{cases} \quad i = 1, 2, \dots, r \quad (4.1)$$

Chaque conséquent de règle représentée par des relations dans l'espace d'état $x(t) = A_i x(t) + B_i u(t)$ est appelée un « sous- modèle ». A chaque règle R^i est attribué un poids $w_i(z(t))$ qui dépend de la valeur de vérité (ou degré d'appartenance) des $z_j(t)$ aux sous-ensembles flous F_j^i , notée $F_j^i(z_j(t))$, et du choix de la modélisation du connecteur (opérateur) « et » reliant les prémisses. Le connecteur « et » est souvent choisi comme étant le produit, d'où :

$$w_i(z(t)) = \prod_{j=1}^p F_j^i(z_j(t)), i = 1, 2, \dots, r \text{ avec } w_i(z(t)) \geq 0, \text{ pour tout } t \quad (4.2)$$

puisque les fonctions d'appartenance prennent leur valeur dans l'intervalle $[0,1]$.

A partir des poids attribués à chaque règle, les sorties finales des modèles flous sont inférées de la manière suivante, qui correspond à une défuzzification barycentrique :

$$\dot{x}(t) = \frac{\sum_{i=1}^r w_i(z(t))(A_i x(t) + B_i u(t))}{\sum_{i=1}^r w_i(z(t))}, \quad y(t) = \frac{\sum_{i=1}^r w_i(z(t))C_i x(t)}{\sum_{i=1}^r w_i(z(t))} \quad (4.3)$$

qui peuvent être réécrites :

$$\dot{x}(t) = \sum_{i=1}^r h_i(z(t))(A_i x(t) + B_i u(t)), \quad y(t) = \sum_{i=1}^r h_i(z(t))C_i x(t) \quad (4.4)$$

avec : $\frac{w_i(z(t))}{\sum_{i=1}^r w_i(z(t))}$ vérifiant une propriété de somme convexe, c'est-à-dire :
 $\sum_{i=1}^r h_i(z(t)) = 1$ et $w_i(z(t)) \geq 0$, pour tout t .

De façon tout à fait analogue, les modèles flous TS discrets sont définis comme suit.

Modèles flous discrets (MFD)

Règle $R^i, i = 1, 2, \dots, r$, du modèle :

$$\text{Si } z_i(t) \text{ est } F_1^i \text{ et } \dots \text{ et } z_p(t) \text{ est } F_p^i \text{ alors } \begin{cases} x(t+1) = A_i x(t) + B_i u(t) \\ y(t) = C_i x(t) \end{cases} \quad (4.5)$$

ou, de façon plus compacte, avec la même définition des $h_i(z(t))$:

$$x(t+1) = \sum_{i=1}^r h_i(z(t))(A_i x(t) + B_i u(t)), \quad y(t) = \sum_{i=1}^r h_i(z(t))C_i x(t) \quad (4.6)$$

Une extension possible du modèle flou présenté ci-dessus permet de prendre en compte des incertitudes dans les paramètres des prémisses [TAN94], dans le cas discret. De plus, il est également utilisé pour les modèles non linéaires incertains, en introduisant les incertitudes du modèle dans les parties conclusion des règles floues.

Le potentiel offert par ce type de représentation permet l'utilisation de toutes les techniques associées. Elle permet, notamment, de traiter les problèmes multivariables et d'aborder le problème de la stabilité à l'aide de la méthode de Lyapunov.

4.1.2. Stabilité des modèles flous

On s'intéresse à la stabilité et à la stabilisation quadratique des modèles flous TS, c'est-à-dire en utilisant la fonction de Lyapunov $V(x(t)) = x^T(t)Px(t)$ avec $P > 0$ [BOY94]. Les théorèmes suivants basés sur la seconde méthode de Lyapunov donnent les conditions suffisantes permettant de garantir la stabilité de modèles flous continus et discrets décrits respectivement par les expressions (4.4) et (4.6).

THEOREME 4.1 (MFC) [TAN92].- *L'équilibre d'un modèle flou, continu, décrit par (4.4) (en boucle ouverte), est asymptotiquement stable s'il existe une matrice P commune définie positive telle que :*

$$A_i^T P + P A_i < 0, \quad i = 1, 2, \dots, r \quad (4.7)$$

THEOREME 4.2 (MFD) [TAN92].- *L'équilibre d'un modèle flou, discret, décrit par (4.6), est asymptotiquement stable s'il existe une matrice P commune définie positive telle que :*

$$A_i^T P A_i - P < 0, \quad i = 1, 2, \dots, r \quad (4.8)$$

Les conditions de stabilité obtenues sont évidemment conservatives puisque la partie la partie prémisses des règles n'est pas prise en compte. D'autres voies ont été explorées, certaines sont reprises dans la suite.

L'approche utilisée dans [KIM95] considère un modèle flou comme un système linéaire avec des incertitudes de modélisation. Toutes les paires (A_i, B_i) intervenant dans les règles peuvent être réécrites en fonction d'une paire (A_0, B_0) unique : $\forall_i, A_i = A_0 + \delta A_i, B_i = B_0 + \delta B_i$ d'où (4.1) peut se réécrire avec $\sum_{i=1}^r h_i(z(t)) = 1$:

$$x(t+1) = \left(A_0 + \sum_{i=1}^r h_i(z) \delta A_i \right) x(t) + \left(B_0 + \sum_{i=1}^r h_i(z) \delta B_i \right) u(t) \quad (4.9)$$

Le modèle (A_0, B_0) doit être choisi avec attention pour que les normes de δA_i et δB_i soient faibles. On peut envisager, par exemple, de prendre le modèle obtenu par moyennage des autres modèles.

Le modèle (A_0, B_0) est supposé stable ; il existe, donc deux matrices $P > 0$ et $Q > 0$ telles que : $A_0^T P A_0 - P = -Q$ Définissons les matrices symétriques :

$$D_i = (A_0 + \delta A_i)^T P (A_0 + \delta A_i) - P = \delta A_i^T P \delta A_i + \delta A_0^T P A_0 + A_0^T P \delta A_i - Q.$$

De plus, chaque partie antécédent d'une règle R^i peut s'écrire « Si $x(t)$ est F^i », où F^i est considéré comme un ensemble flou multidimensionnel pour le règle, c'est-à-dire, $x(t) \in \text{Supp}(F^i)$ signifie que $\forall j \in \{1, \dots, n\}, x_j(t) \in \text{Supp}(F_j^i)$.

THEOREME 4.3 (MFD) [KIM95].- *Le modèle TS (4.1) autonome ($u=0$) est stable si pour une matrice A_0 Hurwitz, il existe une matrice $Q > 0$ telle que, pour toutes les règles $i, x(t)^T D_i x(t) < 0$ pour tout $x(t) \neq 0, x(t) \in \{x(t) / x(t) \in \text{Supp}(F^i)\}$.*

Afin de déterminer une matrice $Q > 0$ vérifiant les conditions du théorème 4.3, les auteurs proposent un algorithme basé sur la descente du gradient. En ce qui concerne la partie

commande, les auteurs se sont limités à travailler avec un retour d'état linéaire $u(t) = -Fx(t)$. L'intérêt de cette approche est de réduire la conservativité des résultats qui ne sont basés que sur la partie conséquence des règles. Les auteurs donnent un exemple où il est possible de prouver la stabilité d'un modèle flou comprenant deux matrices A^i non Hurwitz.

Le fait d'utiliser un support sans propriétés particulières ne permet pas une formulation sous la forme de LMI du problème. Il est alors possible de remédier à ce problème en utilisant des propriétés sur ce support.

Dans le cas où le modèle contient des termes constants (4.10), il est possible d'utiliser les travaux de [MAR95]. Considérons un modèle flou, discret, autonome qui admet dans sa partie conclusion des termes constants. Règle R^i du modèle :

$$\text{Si } x(t) \text{ est } F^i \text{ alors } x(t+1) = A_i x(t) + b_i, \quad i = 1, 2, \dots, r \quad (4.10)$$

En utilisant la connaissance des supports des fonctions d'appartenance qui peuvent être décrits par des contraintes de type quadratique :

$$x \in \text{Supp}(F^i) \Leftrightarrow F_j^i(x) = x^T T_j^i x + 2x^T u_j^i + v_j^i \leq 0, \quad j = 1, \dots, n_i$$

on peut définir la P-Stabilité-Quadratique d'une règle de la façon suivante.

DÉFINITION 4.1 [MAR95].- *La règle R^i est P-Quadratiquement –Stable si, et seulement si, il existe une matrice $P > 0$, et un réel $\varepsilon_i > 0$ tel que : $\forall x(t) \in \text{Supp}(F^i), x(t) \neq 0$:*

$$\Delta V_i(x(t)) \leq -\varepsilon_i \|x(t)\|^2 \text{ et } \Delta V_i(x(t)) = 0 \text{ si } 0 \in \text{Supp}(F^i)$$

avec : $V_i(x) = x^T P x$ et $\Delta V_i(x) = x^T (A_i^T P A_i - P)x + 2x^T A_i^T P b_i + b_i^T P b_i$.

Et par l'utilisation de la S- procédure [BOY94], on peut déduire le théorème suivant.

THEOREME 4.4 [MAR95].- Si $\exists \tau_j^i \geq 0, j = 1, \dots, n_i, \varepsilon_i > 0$ et $P > 0$ tels que :

$$\begin{bmatrix} A_i^T P A_i - P + \varepsilon_i I & A_i^T P b_i \\ b_i^T A_i & b_i^T P b_i \end{bmatrix} - \sum_j \tau_j^i \begin{bmatrix} T_j^i & u_j^i \\ u_j^{iT} & v_j^i \end{bmatrix} \leq 0 \quad (4.11)$$

alors la règle R^i est P- Quadratiquement – Stable.

Enfin le théorème 4.5 pour le modèle flou complet.

THEOREME 4.5 (MFD) [MAR95].- *Le modèle flou décrit par l'équation (4.10) est quadratiquement stable s'il existe une matrice $P = P^T > 0$ telle que toutes les règles soient P- quadratiquement stables.*

Les conditions (4.11) du théorème 4.4 sont des LMI. Notons que si $b_i = 0$ et que le support contient 0, alors obligatoirement, les conditions deviennent celles de Tanaka.

Après avoir présenté différentes approches permettant d'étudier la stabilité de modèles flous, la partie suivante traite de la stabilisation de modèles flous.

4.1.3. Stabilisation des modèles flous TS

La réalisation d'un régulateur flou est basée sur le modèle flou TS représentant le système non linéaire. A partir du concept PDC (*Parallel Distributed Compensation*), la détermination d'une loi de commande revient à déterminer pour chaque modèle local une matrice de gains, par exemple en utilisant une synthèse quadratique ou des LMI. L'idée principale de ce concept est de calculer une loi de commande linéaire par retour d'état, pour chaque sous-modèle du modèle flou. Dans toute la suite, on suppose que les sous-modèles sont commandables et observables.

Construction d'un régulateur flou

Le régulateur flou PDC partage les mêmes ensembles flous que le modèle, donc, il garde les mêmes parties prémisses ainsi que les mêmes fonctions d'appartenance. Pour les modèles flous continus (4.1) et discrets (4.5), la réalisation du régulateur se fait de la façon suivante [TAN95]. Règle R^i du régulateur :

$$\text{Si } z_i(t) \text{ est } F_1^i \text{ et...et } z_p(t) \text{ est } F_p^i \text{ alors } u(t) = -F_i x(t), \quad i = 1, 2, \dots, r \quad (4.12)$$

La sortie finale du régulateur flou est inférée de la manière suivante :

$$u(t) = -\sum_{i=1}^r h_i(z(t)) F_i x(t) \text{ avec } h_i(z(t)) = \frac{w_i(z(t))}{\sum_{i=1}^r w_i(z(t))} \quad (4.13)$$

vérifiant toujours la même propriété de somme convexe.

La réalisation du régulateur flou consiste à déterminer les gains de contre-réaction F_i dans les parties conclusions.

Conditions de stabilité

Pour obtenir l'expression de la boucle fermée, il suffit de substituer (4.13) à (4.4) (cas continu) et (4.13) à (4.6) (cas discret). Ainsi les expressions sont les suivantes :

$$\begin{aligned} \text{(MFC)} \quad \dot{x}(t) &= \sum_{i=1}^r h_i(z(t)) \left(A_i x(t) - B_i \sum_{j=1}^r h_j(z(t)) F_j x(t) \right) \\ &= \sum_{i=1}^r \sum_{j=1}^r h_i(z(t)) h_j(z(t)) (A_i - B_i F_j) x(t) \\ &= \sum_{i=1}^r h_i^2(z(t)) G_{ii} x(t) + 2 \sum_{i < j}^r h_i(z(t)) h_j(z(t)) \frac{G_{ij} + G_{ji}}{2} x(t) \end{aligned} \quad (4.14)$$

avec $G_{ij} = A_i - B_i F_j$

$$\text{(MFD)} \quad x(t+1) = \sum_{i=1}^r h_i^2(z(t)) G_{ii} x(t) + 2 \sum_{i < j}^r h_i(z(t)) h_j(z(t)) \frac{G_{ij} + G_{ji}}{2} x(t) \quad (4.15)$$

En appliquant les théorèmes 4.1 et 4.2 respectivement à (4.14) et (4.15), il est possible de tirer des conditions de stabilité pour les MFC et MFD. La démonstration du théorème 4.6 (respectivement 4.7) provient directement du théorème 4.1 (respectivement 4.2).

THEOREME 4.6 (MFC).- *L'équilibre du modèle flou continu (4.14) est asymptotiquement stable s'il existe une matrice $P > 0$ telle que :*

$$G_{ii}^T P + P G_{ii} < 0, \\ \left(\frac{G_{ij} + G_{ji}}{2} \right)^T P + P \left(\frac{G_{ij} + G_{ji}}{2} \right) < 0, i < j \quad (4.16)$$

pour tout $i, j = 1, 2, \dots, r$, exceptées les paires (i, j) telles que $h_i(z(t))h_j(z(t)) = 0, \forall t$.

THEOREME 4.7. (MFD).- *L'équilibre du modèle flou discret (4.15) est asymptotiquement stable s'il existe une matrice $P > 0$ telle que*

$$G_{ii}^T P G_{ii} - P < 0, \\ \left(\frac{G_{ij} + G_{ji}}{2} \right)^T P \left(\frac{G_{ij} + G_{ji}}{2} \right) - P < 0, i < j \quad (4.17)$$

pour tout $i, j = 1, 2, \dots, r$, exceptées les paires (i, j) telles que $h_i(z(t))h_j(z(t)) = 0, \forall t$.

Le fait d'utiliser la condition (4.16) pour les MFC, (4.17) pour les MFD, avec $i < j$ permet de réduire un peu la conservativité des résultats puisqu'il n'est pas obligatoire d'avoir tous les sous- modèles croisés $G_{ij} = A_i - B_i F_j$ stables.

L'obtention du régulateur flou PDC consiste donc à déterminer les matrices de gains de retour d'état $F_j (j = 1, 2, \dots, r)$ satisfaisant les conditions du théorème 4.6 (cas MFC) ou théorème 4.7 (cas MFD) pour une matrice P définie positive.

Pour calculer les matrices de retour d'état, il est possible d'utiliser une synthèse quadratique et vérifier ensuite qu'il existe une matrice $P > 0$ commune. Il est également possible d'utiliser un placement de pôles pour assurer que les valeurs propres de $G_{ij} = A_i - B_i F_j$ avec $i \in \{1, \dots, r\}$ tombent dans un domaine préspecifié. On peut, alors, raisonnablement penser que si les pôles sont proches pour les modèles bouclés $G_{ij} = A_i - B_i F_j$, il y a de fortes chances que les équations (4.16) dans le cas MFC et (4.17) dans le cas MFD, soient également vérifiées. Néanmoins, des exemples où les paires (A_i, B_i) n'ont pas la même forme montrent que ce résultat n'est pas garanti.

Il est également possible d'utiliser les théorèmes de [TAN94] donnant des conditions moins conservatives en tenant compte d'une propriété particulière sur le poids des règles, ces théorèmes étant tirés des théorèmes 4.6 et 4.7.

Une autre façon de déterminer la matrice P et les gains de commande $F_j (j = 1, 2, \dots, r)$ simultanément est l'utilisation des outils issus de l'optimisation convexe, et plus particulièrement des LMI (*Linear Matrix Inequalities*) [ELG97]. Certains outils LMI sont utilisables à l'aide du logiciel MATLAB et de la boîte à outils LMI *Control Toolbox* [GAH95].

Par exemple, dans le cas de la synthèse d'une loi de commande de type PDC pour des modèles flous MFD, théorème 4.7, la première partie de l'équation (4.17), non linéaire en P et F_j , s'écrit, après le changement de variables $X = P^{-1}$ et $M_i = F_i P^{-1}$:

$-(XA_i^T - M_i^T B_i^T)X^{-1}(A_i X - B_i M_i) + X > 0$ et, en utilisant le lemme de Schur [ELG97], on arrive au problème suivant. Trouver $X > 0$ et M_i tels que :

$$\forall_i \in \{1, \dots, r\}, \begin{bmatrix} X & XA_i^T - M_i^T B_i^T \\ A_i X - B_i M_i & X \end{bmatrix} \geq 0 \quad (4.18)$$

De même, la seconde partie des équations (4.17) peut se réécrire. Trouver $X > 0$ et M_i, M_j tels que $\forall i, j \in \{1, \dots, r\}, i < j$:

$$\begin{bmatrix} X & \frac{1}{2}(XA_i^T + XA_j^T - M_j^T B_j^T - M_i^T B_j^T) \\ \frac{1}{2}(A_i X + A_j X - B_i M_j - B_j M_i) & X \end{bmatrix} \geq 0 \quad (4.19)$$

puis : $P = X^{-1}, F_i = M_i P$.

Plusieurs problèmes sous forme de LMI appliquées à des modèles flous de type TS peuvent être trouvés dans [TAN98] [ZHA95] [ZHA96]. On peut, de façon additionnelle, rajouter des contraintes de type degré de stabilité préspecifié, contraintes de saturation sur le commande, contraintes de saturation sur la sortie [ELG97] [TAN98].

Enfin, d'autres travaux traitant de la stabilisation de modèles flous de type TS existent. Parmi eux, nous allons en présenter un certain nombre.

[THA97] : leurs travaux utilisent la similitude des modèles flous MFD de type TS avec les modèles LTV (linéaires variant dans le temps). En termes de stabilité, ils ont montré l'équivalence entre les modèles flous de type TS et une classe particulière de modèles LTV. Une condition suffisante pour la stabilité asymptotique, globale, basée sur le système LTV équivalent, est donnée. Un cas particulier de cette condition suffisante, basé sur une normalisation simultanée des matrices, est également abordé, ainsi qu'une procédure de construction permettant, dans certains cas, d'obtenir cette normalisation. Il n'y a pas de discussions ayant trait à la commande à partir de ces modèles.

[KIR98] : ils se placent dans le même type d'approche que [KIM95], c'est-à-dire, le modèle flou MFD est l'analyse comme un modèle linéaire soumis à une classe de perturbations non linéaires (4.9); les termes constants du type (4.10) sont également pris en compte. Ils obtiennent un test de stabilité sous la forme de LMI qui réduit la complexité des LMI du type (4.18). Pour la partie commande, les auteurs se sont limités à travailler avec un retour d'état linéaire $u(t) = -Fx(t)$.

[FEN97] : ces travaux concernent la stabilité des modèles flous MFC de type (4.1) et une commande qui n'utilise que le sous-modèle dominant, c'est-à-dire dont la valeur d'appartenance $h_i(z(t)), i \in \{1, \dots, r\}$ est la plus élevée. A chaque instant : $u(t) = -F_k x(t)$ avec $k = \arg \max_{i \in \{1, \dots, r\}} \{h_i(z(t))\}$. Cette commande ne pouvant satisfaire tous les cas, une commande additionnelle est rajoutée qui s'apparente à une commande à grand

gain qui utilise une division par $\left\| \sum_{i=1}^r h_i(z) B_i^T P x \right\|$, sans garantir que cette loi est continue quand cette quantité est nulle.

[GUE99] : cette approche s'inspire de la stabilisation quadratique, simultanée, d'une famille de systèmes linéaires SIMO [PET87] et est appelée SSF (stabilisation simultanée pour modèles flous). Des conditions suffisantes d'existence d'une loi, pour des systèmes flous de type SIMO, sont données.

[GUE01a] : dans le cas particulier des modèles flous MIMO dont les vecteurs B_i sont linéairement dépendants $\forall i \in \{1, \dots, r\} \exists k_i > 0, B_i = k_i B_1$ il est possible de réduire la conservativité des résultats précédents. Pour ce faire, une loi de commande dénommée CDF (compensation et division pour systèmes flous) est utilisée. La propriété remarquable de cette loi de commande vient du fait qu'elle conserve un modèle en boucle fermée à r règles, c'est-à-dire les modèles croisés ne sont plus présents.

[BLA01] : Plusieurs approches nouvelles sont testées. Parmi celles-ci on peut citer : des nouvelles conditions permettant de réduire la conservativité dans le cas où les $h_i(z(t))$ n'atteignent pas leurs bornes; des résultats concernant également des observateurs dans le cas de variables de prémisse non mesurées, et une intéressante liaison avec les modes glissants.

[FEN01] [JOH99] : Ils utilisent une fonction de Lyapunov par morceaux, l'idée étant de stabiliser chaque sous-modèle et de construire un régulateur global de manière à garantir la stabilité globale du système en boucle fermée. L'utilisation d'une telle fonction de Lyapunov impose le rajout de conditions restrictives aux frontières.

[GUE01b] [GUE01c] : dans le cas discret, une approche de type fonction de Lyapunov non quadratique permet, par l'intermédiaire de lois de commande différentes et de nouvelles transformations, de réduire la conservativité des résultats du cas quadratique.

Après cette étude bibliographique sur la stabilité des modèles flous, nous allons présenter une loi de commande de type PDC par minimisation d'un critère quadratique.

4.2. Commande PDC sous optimale par retour d'état

La philosophie de la commande de type compensation parallèle distribuée (PDC), consiste à calculer une loi de commande linéaire par retour d'état, pour chaque sous modèle du modèle flou. La détermination d'une loi de commande revient à déterminer pour chaque modèle local des gains matriciels, par exemple en utilisant une synthèse sous la forme de LMI ou par minimisation d'un critère quadratique. Nous nous sommes intéressés ici à cette dernière approche. La minimisation d'un critère quadratique constitue l'un des moyens de parvenir à la détermination d'une structure de commande par retour d'état pour les systèmes linéaires multidimensionnels. En utilisant cette approche, il semble intéressant d'élargir les résultats classiques de la commande optimale pour les sous modèles qui constituent le modèle *affine* Takagi-Sugeno, comme on le développera par la suite. On parle dans cette section de commande sous optimale par retour d'état car la loi de commande n'est optimale que pour chacun des sous modèles linéaires.

4.2.1. Commande LQR pour des modèles affines à temps discret

On considère ici un modèle sous forme d'état qui représente un système non linéaire en temps discret⁸⁷ décrit par l'équation (4.20) :

$$x_{k+1} = f_k(x_k, u_k) \quad (4.20)$$

où l'indexe $k = 0, 1, \dots, N$ indique l'instant d'échantillonnage, x_k est le vecteur d'état à l'instant k ($x_k \in \mathbb{R}^n$), u_k est le vecteur de commande à l'instant k ($u_k \in \mathbb{R}^m$) et x_0 est la condition initiale. Le système représenté par (4.20) est approximé par un modèle flou *affine* Takagi-Sugeno composé par la concaténation de sous modèles linéaires discrets de la forme :

$$\begin{aligned} x_{k+1} &= A_k x_k + B_k u_k + d_k \\ y_k &= C_k x_k \end{aligned} \quad (4.21)$$

où $y_k \in \mathbb{R}^p$ représente le vecteur de sortie, $A_k \in \mathbb{R}^{n \times n}$, $B_k \in \mathbb{R}^{n \times m}$, $C_k \in \mathbb{R}^{p \times n}$ et $d_k \in \mathbb{R}^n$ est un vecteur constant, propre au modèle affine.

La détermination de la commande optimale des sous modèles linéaires par minimisation d'un critère quadratique peut être traitée en utilisant plusieurs théories, comme la programmation dynamique [BEL57] et le Principe du Maximum de Pontryagin basé sur des méthodes variationnelles [PON62]. Dans ce document nous nous sommes focalisés sur cette deuxième approche. Afin d'avoir plus de détail sur la première approche, le lecteur peut se reporter à la référence [FOU79] qui considère divers cas avec systèmes linéaires affines.

Dans ce contexte, la position du problème du régulateur linéaire quadratique (LQR)⁸⁸, consiste à trouver la loi de commande permettant de faire évoluer le sous-système linéaire affine décrit par (4.21) à partir de l'état initial x_0 en minimisant le critère quadratique suivant [NAI03]:

$$J = \frac{1}{2} x_N^T S_N x_N + \frac{1}{2} \sum_{k=0}^{N-1} (x_k^T Q_k x_k + u_k^T R_k u_k) \quad (4.22)$$

où $S_k \in \mathbb{R}^{n \times n}$, $Q_k \in \mathbb{R}^{n \times n}$ sont des matrices symétriques semi définies positives et $R_k \in \mathbb{R}^{m \times m}$ est symétrique définie positive. La matrice S_k pondère l'état final, Q_k pondère l'évolution de l'état et R_k pondère l'énergie de commande. Afin d'optimiser le critère, J doit converger. Pour cela, le sous-système (4.21) doit être stable. Le but de l'optimisation est d'obtenir la séquence de commande $u_0, u_1, \dots, u_{N-1}$ qui minimise (4.22). Dans ce cas, le Hamiltonien est donné par :

$$\mathcal{H}(x_k, u_k, \lambda_{k+1}) = \frac{1}{2} (x_k^T Q_k x_k + u_k^T R_k u_k) + \lambda_{k+1}^T (A_k x_k + B_k u_k + d_k) \quad (4.23)$$

En suivant la théorie de la commande optimale basée sur le calcul des variations, on considère les expressions suivantes :

⁸⁷ Dans cette section, la notation a été changée par simplicité : $x(k)$ a été remplacée par x_k .

⁸⁸ LQR: *Linear Quadratic Regulator*.

Equation d'état

$$x_{k+1} = \frac{\partial \mathcal{H}}{\partial \lambda_{k+1}} = A_k x_k + B_k u_k + d_k \quad (4.24)$$

Equation de co-état

$$\lambda_k^* = \frac{\partial \mathcal{H}}{\partial x_k} = Q_k x_k^* + A_k^T \lambda_{k+1}^* \quad (4.25)$$

Equation de stationnarité

$$0 = \frac{\partial \mathcal{H}}{\partial u_k} = R_k u_k^* + B_k^T \lambda_{k+1}^* \quad (4.26)$$

d'après (4.26), la loi de commande optimale en boucle ouverte est donnée par :

$$u_k^* = -R_k^{-1} B_k^T \lambda_{k+1}^* \quad (4.27)$$

Comme R_k est définie positive, ceci assure son inversibilité. En remplaçant u_k^* dans l'équation d'état (4.24) on obtient :

$$x_{k+1}^* = A_k x_k^* - B_k R_k^{-1} B_k^T \lambda_{k+1}^* + d_k \quad (4.28)$$

Condition de contour

On considère ici le cas où l'état final x_N est libre. Par conséquent, la condition correspondante de contour est donnée par :

$$\lambda_N = \frac{\partial}{\partial x_N} \left(\frac{1}{2} x_N^T S_N x_N \right) = S_N x_N \quad (4.29)$$

Afin d'obtenir une loi de commande optimale en boucle fermée, il est nécessaire d'exprimer la fonction de co-état λ_{k+1}^* dans la commande optimale (4.27) en termes de la fonction d'état x_k^* . La condition finale ainsi que la présence du terme affine dans l'équation d'état (4.24), nous amène à proposer pour le terme λ_k^* l'expression suivante :

$$\lambda_k^* = P_k x_k^* + q_k \quad (4.30)$$

où la matrice $P_k \in \mathbb{R}^{n \times n}$ et le vecteur $q_k \in \mathbb{R}^n$ doivent être déterminés. En utilisant l'expression (4.30) dans les équations d'état (4.28) et de co-état (4.25), on obtient :

$$P_k x_k^* + q_k = Q_k x_k^* + A_k^T (P_{k+1} x_{k+1}^* + q_{k+1}) \quad (4.31)$$

et

$$x_{k+1}^* = A_k x_k^* - B_k R_k^{-1} B_k^T (P_{k+1} x_{k+1}^* + q_{k+1}) + d_k \quad (4.32)$$

en résolvant x_{k+1}^* d'après (4.32),

$$x_{k+1}^* = \left(I + B_k R_k^{-1} B_k^T P_{k+1} \right)^{-1} \left(A_k x_k^* - B_k R_k^{-1} B_k^T q_{k+1} + d_k \right) \quad (4.33)$$

où $I \in \mathbb{R}^{n \times n}$ est la matrice identité. En remplaçant (4.33) dans (4.31) on obtient :

$$P_k x_k^* + q_k = Q_k x_k^* + A_k^T P_{k+1} \left(\left(I + B_k R_k^{-1} B_k^T P_{k+1} \right)^{-1} (A_k x_k^* - B_k R_k^{-1} B_k^T q_{k+1} + d_k) \right) + A_k^T q_{k+1} \quad (4.34)$$

la vérification de cette expression pour tout x_k impose les deux relations suivantes :

$$P_k = Q_k + A_k^T P_{k+1} (I + B_k R_k^{-1} B_k^T P_{k+1})^{-1} A_k \quad (4.35)$$

et

$$q_k = A_k^T P_{k+1} (I + B_k R_k^{-1} B_k^T P_{k+1})^{-1} (d_k - B_k R_k^{-1} B_k^T q_{k+1}) + A_k^T q_{k+1} \quad (4.36)$$

L'expression (4.35) correspond à une équation matricielle aux différences de type Riccati (DRE)⁸⁹. La condition terminale pour résoudre l'équation matricielle DRE (4.35) est obtenue à partir de (4.29) et (4.30) comme suit :

$$P_N = S_N, \quad q_N = 0 \quad (4.37)$$

Dans l'équation matricielle DRE (4.35), le terme P_k se trouve du côté gauche et P_{k+1} est du côté droit, ce qui nécessite une résolution itérative en remontant le temps à partir de la condition finale (4.37). Une procédure similaire s'applique pour le calcul de q_k . Puisque les relations de récurrence sont définies en temps inverse, une réalisation pratique de la commande telle qu'elle a été précédemment définie par (4.27) et (4.30), nécessite de précalculer la matrice P_k et le vecteur q_k pour tous les instants de l'horizon considéré. Cela supposerait donc de reprendre les calculs et d'en mémoriser tous les résultats avant chaque changement de consigne. Il est évident que cela peut constituer une contrainte importante pour l'implémentation dans un ordinateur couplé au procédé.

Il est donc particulièrement intéressant, pour la majorité des applications pratiques en milieu industriel, de considérer l'horizon d'optimisation comme infini, ou tout au moins comme très grand par rapport à l'échelle de temps des phénomènes physiques considérés. Le critère à minimiser est alors de la forme :

$$J = \frac{1}{2} \sum_{k=0}^{\infty} (x_k^T Q_k x_k + u_k^T R_k u_k) \quad (4.38)$$

Dans le cas particulier d'un système stationnaire ($A_k = A$, $B_k = B$, $d_k = d$, $C_k = C$) et avec les conditions imposées aux matrices Q et R , lorsque N tend vers l'infini, il est montré que les équations récurrentes (4.35) et (4.36) convergent vers des solutions limites uniques P et q , respectivement, lorsque le sous modèle est commandable. Dans ce cas, il vient alors:

$$P = Q + A^T P (I + B R^{-1} B^T P)^{-1} A \quad (4.39)$$

en utilisant le lemme d'inversion matricielle⁹⁰ :

$$(M_1 + M_2 M_3 M_4)^{-1} = M_1^{-1} - M_1^{-1} M_2 (M_3^{-1} + M_4 M_1^{-1} M_2)^{-1} M_4 M_1^{-1} \quad (4.40)$$

avec $M_1 = I$, $M_2 = B$, $M_3 = R^{-1}$ et $M_4 = B^T P$ on obtient l'équation algébrique discrète de Riccati (DARE)⁹¹:

$$A^T P A - P - (A^T P B)(R + B^T P B)^{-1} (B^T P A) + Q = 0 \quad (4.41)$$

Pour le cas du vecteur q , d'après (4.36) et en utilisant la solution P de l'équation (4.41), il vient :

$$q = \left(I + A^T P (I + B R^{-1} B^T P)^{-1} B R^{-1} B^T - A^T \right)^{-1} (A^T P (I + B R^{-1} B^T P)^{-1} d) \quad (4.42)$$

⁸⁹ DRE: *Difference Riccati Equation*.

⁹⁰ Si la décomposition est judicieuse, elle permet de calculer l'inverse d'une matrice à l'aide d'une matrice plus facilement calculable (M_1^{-1}) et de l'inverse d'une matrice éventuellement plus petite ($M_3^{-1} + M_4 M_1^{-1} M_2$).

⁹¹ DARE: *Discrete Algebraic Riccati Equation*.

cette expression est de la forme

$$q = (G + FE)^{-1} F d, \quad \text{avec} \quad \begin{cases} E = BR^{-1}B^T \\ F = A^T P(I + EP)^{-1} = A^T (P^{-1} + E)^{-1} \\ G = I - A^T \end{cases} \quad (4.43)$$

d'après la commande optimale (4.27) et l'expression (4.30), il vient :

$$u_k^* = -R^{-1}B^T (P x_{k+1}^* + q) = -R^{-1}B^T (P(Ax_k^* + Bu_k^* + d) + (G + FE)^{-1} F d) \quad (4.44)$$

ce qui conduit à la commande optimale en boucle fermée suivante :

$$u_k^* = -(K x_k + N d), \quad \text{avec} \quad \begin{cases} K = (R + B^T P B)^{-1} B^T P A \\ N = (R + B^T P B)^{-1} B^T (P + (G + FE)^{-1} F) \end{cases} \quad (4.45)$$

où la matrice $N \in \mathbb{R}^{m \times n}$, $K \in \mathbb{R}^{m \times n}$ et le vecteur constant $d \in \mathbb{R}^n$.

Compte tenu que le régulateur flou PDC (Figure 4.1) partage les mêmes ensembles flous que le modèle Takagi-Sugeno, il garde donc les mêmes parties prémisses ainsi que les mêmes fonctions d'appartenance. Pour un modèle flou *affine* Takagi-Sugeno discret, la réalisation du régulateur flou PDC consiste à déterminer pour les $i = 1, 2, \dots, r$ règles du modèle, les gains de contre-réaction K_i et de biais N_i dans les parties conclusions selon l'expression (4.45), en utilisant des règles R^i sous la forme :

$$\text{Si } z_1(k) \text{ est } M_1^i \text{ et } \dots \text{ et } z_p(k) \text{ est } M_p^i \text{ alors } u(k) = -(K_i x(k) + N_i d), \quad i = 1, 2, \dots, r \quad (4.46)$$

où r est le nombre de règles. La sortie du régulateur flou est inférée de la façon suivante :

$$u(k) = -\sum_{i=1}^r h_i(z(k)) (K_i x(k) + N_i d), \quad \text{avec} \quad h_i(z(k)) = \frac{w_i(z(k))}{\sum_{i=1}^r w_i(z(k))} \quad (4.47)$$

vérifiant pour tout k , les propriétés de somme convexe et d'activation des fonctions d'appartenance :

$$\sum_{i=1}^r h_i(z(k)) = 1 \quad \text{et} \quad w_i(z(k)) \geq 0 \quad (4.48)$$

Le schéma général de la commande de type régulateur PDC, pour les sous modèles linéaires affines, ainsi obtenu est représenté par la Figure 4.2.

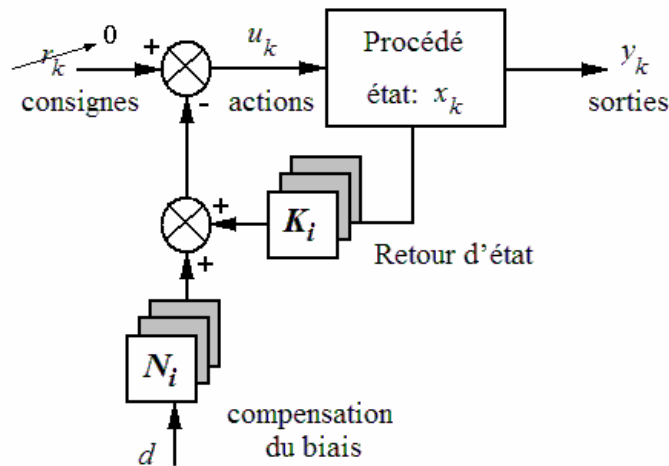


Figure 4.2. Schéma général de la commande du type régulateur PDC pour des sous modèles affines

Remarques

- R 4.1** Compte tenu des propriétés des matrices S , Q et R , la solution P_k de l'équation de Riccati est symétrique.
- R 4.2** Il faut noter que l'expression (4.45) constitue une généralisation de la solution de la commande optimale LQR pour des systèmes linéaires affines sous la forme (4.21). En effet, dans le cas où le terme de biais $d=0$, on retrouve naturellement l'équation algébrique de Riccati standard à temps discret.
- R 4.3** Les résultats précédents permettent de calculer la loi de commande optimale pour chacun des sous modèles locaux du modèle affine Takagi-Sugeno. La loi globale de commande, sous la forme du type commande PDC, est obtenue à partir de la défuzzification du modèle flou. Dans ce cas, bien que la procédure pour calculer les lois locales de commande soit systématique, il faut souligner qu'on ne peut pas garantir l'optimalité des résultats pour la commande globale du modèle. Cet aspect est susceptible d'être développé dans le futur pour le cas de modèles flous affines Takagi-Sugeno.

Jusqu'ici, on a considéré le cas général d'une loi de commande du type régulateur, pour les sous modèles linéaires affines. L'inconvénient majeur de cette structure de commande est de ne pas assurer des erreurs stationnaires égales à zéro lorsque les consignes imposées au procédé ne sont pas nulles. Cette caractéristique se révèle assez gênante pour une application industrielle, pour laquelle on désire généralement que les consignes fixées soient effectivement atteintes. Cet important problème d'utilisation pratique peut cependant être résolu de manière satisfaisante par différentes techniques. Nous allons aborder en premier lieu le cas de pré compensation du gain statique et ensuite, de l'ajout au système d'un certain nombre d'intégrateurs numériques.

4.2.2. Commande PDC avec compensation de gain statique

Une première technique envisagée, pour annuler les erreurs en régime permanent quand les consignes appliquées au procédé sont constantes, consiste à utiliser un pré compensateur qui impose un gain statique pour les sous modèles.

A partir des résultats obtenus pour la loi de commande optimale (4.45) et en ajoutant un terme pour l'entrée de consigne, il vient alors

$$u_k = L r_k - (K x_k + N d) \quad (4.49)$$

où $r_k \in \mathbb{R}^m$ représente les entrées ou consignes, $L \in \mathbb{R}^{m \times m}$ est une matrice de gain dite de pré compensation et $N \in \mathbb{R}^{m \times n}$ est la matrice associée au vecteur de biais d .

En remplaçant (4.49) dans l'expression de la dynamique du sous modèles (4.21) et en calculant la transformée en z à partir des conditions initiales nulles, on obtient l'expression suivante pour la sortie du sous modèle :

$$Y(z) = C(zI - A + BK)^{-1} B L r(z) + C z^{-1} d \quad (4.50)$$

A partir de cette expression il est clair que pour obtenir un erreur nulle pour une consigne constante il faut compenser le terme due au biais d . Pour cela, on propose d'ajouter un terme matriciel V dans l'expression de la commande (4.49), il vient alors :

$$u_k = Lr_k - (Kx_k + (N + V)d) \quad (4.51)$$

Dans ce cas, l'application de la transformée en z à partir des conditions initiales nulles, conduit à l'expression suivante pour la sortie du sous modèle :

$$Y(z) = C(zI - A + BK)^{-1}(BLr(z) - (BN + BV - I)d) \quad (4.52)$$

Pour une consigne constante $r_k = r$, le but est d'obtenir une erreur stationnaire nulle :

$$\lim_{t \rightarrow \infty} y_k = \lim_{z \rightarrow 1} (z - 1)Y(z) = r \quad (4.53)$$

d'après (4.52), l'évaluation de la limite conduit à l'expression

$$\lim_{t \rightarrow \infty} y_k = r = \lim_{z \rightarrow 1} (z - 1)C(zI - A + BK)^{-1} \left(BL \frac{r z}{z - 1} - (BN + BV - I) \frac{d z}{z - 1} \right) \quad (4.54)$$

qui vaut

$$r = C(I - A + BK)^{-1}(BLr - (BN + BV - I)d) \quad (4.55)$$

afin d'obtenir une erreur stationnaire nulle, l'évaluation de l'expression (4.55) impose alors les deux conditions suivantes

$$\begin{cases} L = (MB)^{-1} \\ M(BN + BV - I)d = 0 \end{cases}, \text{ avec } M = C(I - A + BK)^{-1} \quad (4.56)$$

où la matrice $M \in \mathfrak{R}^{p \times n}$. Compte tenu des dimensions des matrices $A \in \mathfrak{R}^{n \times n}$, $B \in \mathfrak{R}^{n \times m}$, $C \in \mathfrak{R}^{p \times n}$ et $K \in \mathfrak{R}^{m \times n}$, la matrice MB est inversible à condition que les sous modèles aient les mêmes dimensions d'entrées que de sorties ($p = m$) et que la matrice résultante ne soit pas singulière. Puisque la structure de commande considérée assure la stabilité du sous modèle bouclée, la validité de cette condition est donc toujours implicitement supposée ce qui permet de faire les calculs à partir des expressions matricielles (4.56).

Concernant la deuxième condition de (4.56), elle conduit à l'expression matricielle

$$V = LM - N \quad (4.57)$$

en remplaçant (4.57) dans (4.51) l'équation de la loi de commande devient alors

$$u_k = Lr_k - (Kx_k + Td), \text{ avec } T = LM = (MB)^{-1}M \quad (4.58)$$

où la matrice $K \in \mathfrak{R}^{m \times n}$, $T \in \mathfrak{R}^{m \times n}$ et le vecteur constant $d \in \mathfrak{R}^n$.

Pour un modèle flou Takagi-Sugeno *affine* discret, la réalisation du régulateur flou PDC avec pré compensateur du gain statique, consiste à déterminer pour les $i = 1, 2, \dots, r$ règles du modèle, les gains de compensation L_i , de contre-réaction K_i et de biais T_i dans les parties conclusions selon l'expression (4.58), en utilisant des règles R^i sous la forme :

$$\text{Si } z_1(k) \text{ est } M_1^i \text{ et } \dots \text{ et } z_p(k) \text{ est } M_p^i \text{ alors } u(k) = L_i r(k) - (K_i x(k) + T_i d), \quad i = 1, 2, \dots, r \quad (4.59)$$

où r est le nombre des règles. La sortie finale du régulateur flou est inférée comme suit :

$$u(k) = \sum_{i=1}^r h_i(z(k)) (L_i r(k) - (K_i x(k) + T_i d)), \text{ avec } h_i(z(k)) = \frac{w_i(z(k))}{\sum_{i=1}^r w_i(z(k))} \quad (4.60)$$

vérifiant pour tout k , les propriétés de somme convexe et d'activation des fonctions d'appartenance (cf. expression (4.48)).

En tenant compte de l'expression (4.58), le schéma général de la structure de commande PDC avec pré compensation de gain statique, pour les sous modèles linéaires affines, ainsi obtenu peut se représenter par la Figure 4.3.

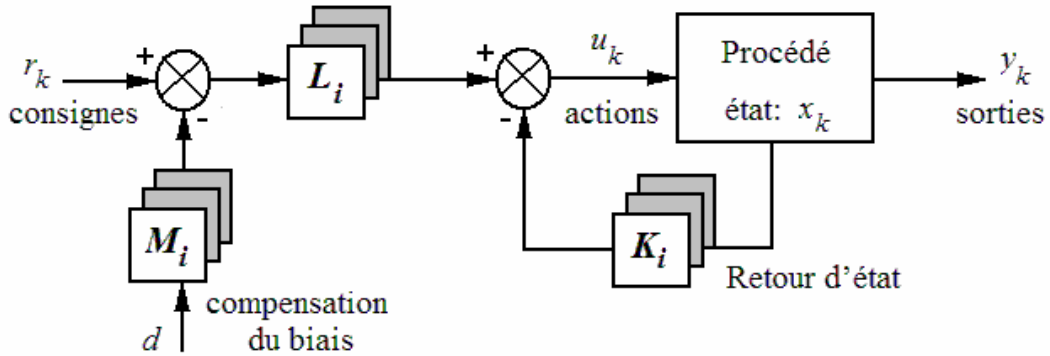


Figure 4.3. Schéma général de la commande du type régulateur PDC avec pré compensateur

L'expression (4.58) permet de calculer une loi de commande optimale pour annuler les erreurs en régime permanent quand les consignes appliquées au procédé sont constantes. Cependant, la principale limitation de la technique de pré compensation par gain statique est le manque de robustesse vis-à-vis des variations paramétriques dans le modèle, ce qui peut empêcher l'annulation totale de l'erreur en régime permanent. Une solution mieux adaptée est développée par la suite.

4.2.3. Commande PDC avec ajout d'intégrateurs

Une deuxième approche pour traiter le problème de poursuite des références constantes, consiste à ajouter au système un certain nombre d'intégrateurs numériques et puis, calculer les gains associés pour chacun des sous modèles affines. La dynamique peut alors s'exprimer à partir d'une représentation augmentée, dans laquelle les intégrateurs ont pour entrées respectives les écarts entre consignes et mesures correspondants.

En utilisant l'approximation d'une intégrale par la méthode des rectangles, les intégrateurs numériques sont de la forme :

$$v_{k+1} = v_k + T_e \cdot e_k = v_k + T_e \cdot (r_k - y_k) \quad (4.61)$$

où le vecteur $v_k \in \mathbb{R}^p$, T_e est la période d'échantillonnage et $e_k \in \mathbb{R}^p$ est le vecteur d'écart entre le vecteur de consignes $r_k \in \mathbb{R}^m$ et de mesures $y_k \in \mathbb{R}^p$. D'après l'expression (4.61), on peut noter que si le système à commander possède un nombre d'entrées m et de sorties p , alors la poursuite des consignes constantes ne se fera pour les p sorties que dans le cas $p \leq m$. Dans ces conditions, les développements mathématiques présentés ici supposent

que chaque sortie est associée à un intégrateur⁹² et que le vecteur de consignes a une dimension correspondante à celui de mesures, c'est-à-dire, on considérera ici $r_k \in \mathbb{R}^p$.

En prenant en compte la définition du vecteur v_k , alors le système augmenté associé à chacun des sous modèles a ainsi pour équations d'état :

$$\begin{bmatrix} x_{k+1} \\ v_{k+1} \end{bmatrix} = \begin{bmatrix} A_{[n \times n]} & 0_{[n \times p]} \\ \bar{C}_{[p \times n]} & I_{[p \times p]} \end{bmatrix} \begin{bmatrix} x_k \\ v_k \end{bmatrix} + \begin{bmatrix} B_{[n \times m]} \\ 0_{[p \times m]} \end{bmatrix} u_k + \begin{bmatrix} d \\ \bar{r}_k \end{bmatrix}, \text{ avec } \begin{cases} \bar{C} = -T_e \cdot C \\ \bar{r}_k = T_e \cdot r_k \end{cases} \quad (4.62)$$

$$[y_k] = [C_{[p \times n]} \quad 0_{[p \times p]}] \begin{bmatrix} x_k \\ v_k \end{bmatrix}$$

où les indices des matrices indiquent leur dimensions correspondantes, les vecteurs d'état, de sortie, de termes indépendants et d'entrées ont pour dimensions respectivement $x_k \in \mathbb{R}^n$, $v_k \in \mathbb{R}^p$, $y_k \in \mathbb{R}^p$, $d \in \mathbb{R}^n$, $u_k \in \mathbb{R}^m$ et $\bar{r}_k \in \mathbb{R}^p$.

En considérant l'expression (4.62), cette représentation peut être mise sous la forme condensée suivante :

$$\begin{cases} \tilde{x}_{k+1} = \tilde{A} \tilde{x}_k + \tilde{B} u_k + \tilde{d} \\ \tilde{y}_k = \tilde{C} \tilde{x}_k \end{cases}, \text{ avec } \begin{cases} \tilde{A} = \begin{bmatrix} A & 0 \\ \bar{C} & I \end{bmatrix}, \quad \tilde{B} = \begin{bmatrix} B \\ 0 \end{bmatrix}, \quad \tilde{d} = \begin{bmatrix} d \\ \bar{r}_k \end{bmatrix}, \quad \tilde{C} = [C \quad 0] \\ \tilde{x}_k = \begin{bmatrix} x_k \\ v_k \end{bmatrix}, \quad \tilde{y}_k = [y_k] \end{cases} \quad (4.63)$$

où les éléments matriciels ont respectivement pour dimensions $\tilde{A} \in \mathbb{R}^{(n+p) \times (n+p)}$, $\tilde{B} \in \mathbb{R}^{(n+p) \times m}$, $\tilde{d} \in \mathbb{R}^{(n+p)}$ et $\tilde{C} \in \mathbb{R}^{p \times (n+p)}$. De leur côté, les nouveaux vecteurs d'état, d'entrée, de termes constantes et de sortie pour le système augmenté ont pour dimensions respectivement $\tilde{x}_k \in \mathbb{R}^{n+p}$, $u_k \in \mathbb{R}^m$, $\tilde{d} \in \mathbb{R}^{n+p}$ et $\tilde{y}_k \in \mathbb{R}^p$.

Si on considère un horizon infini, le critère quadratique à minimiser a alors pour expression :

$$\tilde{J} = \frac{1}{2} \sum_{k=0}^{\infty} (\tilde{x}_k^T \tilde{Q} \tilde{x}_k + u_k^T R u_k) \quad (4.64)$$

où $\tilde{Q} \in \mathbb{R}^{(n+p) \times (n+p)}$ est une matrice symétrique semidéfinie positive et $R \in \mathbb{R}^{m \times m}$ est symétrique définie positive. Afin d'optimiser le critère, $\tilde{J}$ doit converger, pour cela, le sous-système (4.63) doit être stable. La matrice $\tilde{Q}$ est définie de la façon suivante :

⁹² S'il n'en est pas ainsi, la méthode proposée demeure bien sûr valable, en prenant soin simplement de constituer correctement les matrices d'état du sous système global sur lequel s'effectue la détermination des divers éléments de la structure de commande.

$$\tilde{Q} = \begin{bmatrix} Q_{[n \times n]} & 0_{[n \times p]} \\ 0_{[p \times n]} & Q_v_{[p \times p]} \end{bmatrix} \quad (4.65)$$

Le critère $\tilde{J}$ peut donc encore s'écrire en utilisant la définition précédente de la matrice $\tilde{Q}$

$$\tilde{J} = \frac{1}{2} \sum_{k=0}^{\infty} (x_k^T Q x_k + v_k^T Q_v v_k + u_k^T R u_k) \quad (4.66)$$

Il est alors possible de calculer pour le système augmenté (4.63) la commande optimale u_k^* minimisant le critère quadratique $\tilde{J}$ en suivant une procédure similaire à celle développée précédemment pour l'obtention de l'expression (4.45). La commande optimale u_k^* en boucle fermée ainsi déterminée pour chacun des sous modèles est de la forme :

$$u_k^* = -(\tilde{K} \tilde{x}_k + \tilde{N} \tilde{d}) = \tilde{L} \tilde{r}_k - (\tilde{K} x_k + \tilde{K}_v v_k + \tilde{N} d) \quad (4.67)$$

Il convient de remarquer que dans l'expression (4.67) les matrices $\tilde{K}$ et $\tilde{N}$ ne sont pas les mêmes que les matrices K et N obtenues dans l'expression (4.45) lors de la synthèse de la commande du système ne comportant pas les intégrateurs additionnels. C'est pourquoi la notation a été changée en conséquence. D'autre part, les matrices $\tilde{K} \in \mathfrak{R}^{m \times (n+p)}$ et $\tilde{N} \in \mathfrak{R}^{m \times (n+p)}$ ont été décomposées en $\tilde{K} = [\tilde{K} \quad \tilde{K}_v]$ où $\tilde{K} \in \mathfrak{R}^{m \times n}$ et $\tilde{K}_v \in \mathfrak{R}^{m \times p}$, et en $\tilde{N} = [\tilde{N} \quad \tilde{L}]$ où $\tilde{N} \in \mathfrak{R}^{m \times n}$ et $\tilde{L} \in \mathfrak{R}^{m \times p}$. Le vecteur d'état augmenté est $\tilde{x}_k = [x_k \quad v_k]^T \in \mathfrak{R}^{n+p}$ et le vecteur lié aux entrées et aux termes de biais est $\tilde{d} = [d \quad \tilde{r}_k]^T \in \mathfrak{R}^{n+p}$.

Pour calculer la commande optimale u_k^* qui annule l'erreur stationnaire (entre consigne et mesure) pour chacun des sous modèles à partir de la représentation d'état augmenté (4.63), la méthodologie développée préalablement consiste à suivre la procédure suivante :

1. Trouver la matrice $\tilde{P} \in \mathfrak{R}^{(n+p) \times (n+p)}$ solution de l'équation DARE de Riccati :

$$\tilde{A}^T \tilde{P} \tilde{A} - \tilde{P} - (\tilde{A}^T \tilde{P} \tilde{B})(R + \tilde{B}^T \tilde{P} \tilde{B})^{-1} (\tilde{B}^T \tilde{P} \tilde{A}) + \tilde{Q} = 0 \quad (4.68)$$

2. Calculer les matrices $\tilde{E}$, $\tilde{F}$ et $\tilde{G} \in \mathfrak{R}^{(n+p) \times (n+p)}$ suivantes :

$$\begin{cases} \tilde{E} = \tilde{B} R^{-1} \tilde{B}^T \\ \tilde{F} = \tilde{A}^T \tilde{P} (I + \tilde{E} \tilde{P})^{-1} = \tilde{A}^T (\tilde{P}^{-1} + \tilde{E})^{-1} \\ \tilde{G} = I - \tilde{A}^T \end{cases} \quad (4.69)$$

où la matrice identité $I \in \mathfrak{R}^{(n+p) \times (n+p)}$.

3. D'après les résultats (4.68) et (4.69), calculer la loi de commande en déterminant les matrices $\tilde{K} \in \mathfrak{R}^{m \times (n+p)}$ et $\tilde{N} \in \mathfrak{R}^{m \times (n+p)}$ selon les expressions :

$$u_k^* = -(\tilde{K} \tilde{x}_k + \tilde{N} \tilde{d}), \text{ avec } \begin{cases} \tilde{K} = (R + \tilde{B}^T \tilde{P} \tilde{B})^{-1} \tilde{B}^T \tilde{P} \tilde{A} \\ \tilde{N} = (R + \tilde{B}^T \tilde{P} \tilde{B})^{-1} \tilde{B}^T (\tilde{P} + (\tilde{G} + \tilde{F} \tilde{E})^{-1} \tilde{F}) \end{cases} \quad (4.70)$$

En reprenant l'expression (4.67), on peut noter que la loi de commande optimale obtenue pour les sous modèles linéaires affines consiste en un retour d'état $-\tilde{K} x_k$, une compensation

du terme de biais $-\tilde{N}d$, un terme d'anticipation lié à la consigne $\tilde{L}r_k$ et une action intégrale $-\tilde{K}_v v_k$ agissant sur l'erreur entre consignes et mesures. Par analogie avec les régulateurs classiques comportant une action « proportionnelle » et une action « intégrale », le résultat obtenu peut être vue comme une loi de commande multidimensionnelle du type P.I. avec compensation du terme de biais.

Pour un modèle flou *affine* Takagi-Sugeno discret, la réalisation du régulateur flou PDC avec action intégrale, consiste à déterminer pour les $i = 1, 2, \dots, r$ règles du modèle, les gains $\tilde{K}_i$ et $\tilde{N}_i$ du système augmenté (4.67), dans les parties conclusions selon l'expression (4.70), en utilisant des règles R^i sous la forme :

$$\text{Si } z_1(k) \text{ est } M_1^i \text{ et } \dots \text{ et } z_p(k) \text{ est } M_p^i \text{ alors } u(k) = -(\tilde{K}_i \tilde{x}(k) + \tilde{N}_i \tilde{d}), \quad i = 1, 2, \dots, r \quad (4.71)$$

où r est le nombre des règles. La sortie finale du régulateur flou est inférée comme suit :

$$u(k) = -\sum_{i=1}^r h_i(z(k)) (\tilde{K}_i \tilde{x}(k) + \tilde{N}_i \tilde{d}), \quad \text{avec } h_i(z(k)) = \frac{w_i(z(k))}{\sum_{i=1}^r w_i(z(k))} \quad (4.72)$$

vérifiant pour tout k , les propriétés de somme convexe et d'activation des fonctions d'appartenance (cf. expression (4.48)).

En prenant en compte l'expression (4.67), le schéma général de la structure de commande PDC avec action intégrale, pour les sous modèles linéaires affines, ainsi obtenu est représenté par la Figure 4.4.

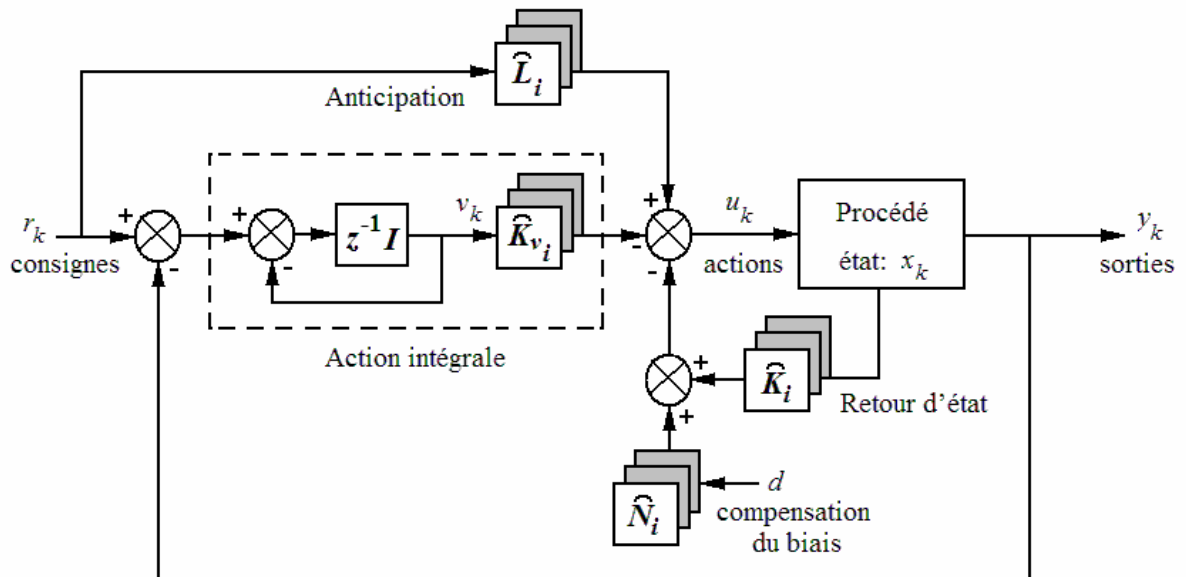


Figure 4.4. Schéma général de la commande PDC avec action intégrale pour des sous modèles affines

4.2.4. Choix des matrices de pondération Q et R

Il est clair que les valeurs des matrices $\tilde{K}$ et $\tilde{N}$ dans la loi de commande (4.70) dépendent du choix des matrices $\tilde{Q}$ et R qui doit être fait *a priori*. Afin de pondérer l'état élargie $\tilde{x}_k$ et l'entrée u_k , ces deux matrices doivent être définies positives.

Pour simplifier le test de positivité, on pourra éventuellement se limiter au choix de deux matrices diagonales, avec tous les éléments de la diagonale principale positifs, de la forme :

$$Q = \text{diag}(q_1, q_2, \dots, q_n), \quad R = \text{diag}(r_1, r_2, \dots, r_m) \quad (4.73)$$

Dans ces conditions :

- Plus q_i est grand, plus le retour vers zéro de la variable $\tilde{x}_i$ correspondante sera rapide.
- Plus r_j est grand, plus la commande u_j correspondante sera de faible amplitude.

Si le i -ème élément de $\tilde{x}_k$ et le j -ème élément de u_k ont un maximum permet, respectivement $\max(\tilde{x}_i)$ et $\max(u_j)$, alors on peut adopter le choix initial suivant [BRY69]:

$$q_i = (1/\max(\tilde{x}_i))^2, \quad r_j = (1/\max(u_j))^2 \quad (4.74)$$

Il faut remarquer que la méthode proposée pour calculer la structure de commande suppose en effet que les changements de consignes sont effectués sous forme d'échelons. Pour des applications industrielles, il s'avère qu'un réglage fin des matrices de pondération est souhaitable pour conférer au système, lors des changements de consignes, une évolution dynamique correcte en même temps que des amplitudes d'actions compatibles avec une utilisation pratique. Dans ce cas, les études en simulation se montrent un outil précieux pour évaluer et améliorer le comportement attendu du système.

Lorsque la loi de commande doit être utilisée pour une fonction d'asservissement (avec des modifications imposées au vecteur de consignes), il semble probable que la maîtrise du système dans ces conditions soit beaucoup plus difficile. Un ajustement des matrices $\tilde{Q}$ et R intervenant dans le critère quadratique pourrait s'appliquer dans certains cas, mais il faut noter que le réglage de ces matrices doit satisfaire un compromis entre les amplitudes supportées par les actionneurs et la performance globale de la boucle. Ainsi par exemple, une augmentation des valeurs des coefficients de la matrice R par rapport à ceux de la matrice $\tilde{Q}$, pénaliserait les entrées, mais en même temps cela pourrait donner une faible pondération des écarts entre consignes et mesures. A ce moment, la compensation des perturbations pourrait devenir très lente par rapport aux performances potentielles de la structure de commande choisie.

Pour conserver un fonctionnement correct dans les différents cas envisagés, une solution intéressante consiste à introduire dans la structure de commande un modèle de référence ayant pour entrées le vecteur de consignes r_k et pour sorties la forme d'évolution réellement souhaitée pour le vecteur des sorties y_k du procédé [FOU79]. L'inclusion d'un modèle de référence dans la structure de commande permet une bonne séparation entre le problème du choix des boucles de réaction et le problème des performances en asservissement.

4.3. Conclusions

Le but de ce chapitre n'était pas d'être exhaustif, mais d'exposer quelques principes de base concernant les différentes possibilités qu'offrent les modèles flous de type Takagi-Sugeno pour la commande des systèmes non linéaires multivariables en utilisant la philosophie de commande du type compensation parallèle distribuée (PDC). Cette approche consiste à utiliser la même partie antécédent des règles que celle du modèle flou de type Takagi-Sugeno, et, dans la partie conclusion, des commandes locales, par retour d'état. Dans ce contexte, chaque sous modèle du modèle flou de type Takagi-Sugeno est stabilisé localement par une loi linéaire.

Dans la première partie du chapitre, nous avons donné un aperçu général en considérant les problèmes de la stabilité (analyse) et de la stabilisation quadratique (synthèse d'un régulateur flou) des modèles flous *homogènes* de type TS. La plupart des travaux font appel, pour vérifier les conditions de stabilité, à la méthode de Lyapunov. Au niveau de la synthèse, la détermination de la loi de commande revient à déterminer pour chaque sous modèle local une matrice de gains de retour d'état. L'un des aspects intéressants de ce type de problèmes consiste en ce que sa formulation peut se faire sous une forme de LMI (*Linear Matrix Inequalities*), qui peuvent être résolues d'une façon efficace par des algorithmes⁹³ d'optimisation convexe sous contraintes.

Dans la deuxième partie du chapitre, nous avons abordé le problème de synthèse quadratique de la commande. En utilisant l'approche PDC, nous avons développé une commande sous-optimale linéaire quadratique par retour d'état basée sur un modèle *affine* TS en temps discret. Pour cela, nous avons proposé trois lois de commande qui tiennent compte du terme indépendant caractéristique (biais) dans la partie conséquent du modèle affine. Les lois de commande sont de type régulateur, régulateur avec compensation de gain statique et régulateur avec action intégrale, ce qui permet l'annulation des erreurs stationnaires de poursuite par rapport à des références de type échelon dans la commande. L'approche retenue permet de calculer d'une façon analytique, la loi de commande multivariable pour des modèles flous affines TS issus de l'application des méthodes d'identification par clustering flou que nous avons abordé dans le chapitre précédent.

La suite de ce rapport est dédiée à l'application de nos travaux de recherche concernant la modélisation et la commande floue de type Takagi-Sugeno, sur le bioprocédé de traitement des eaux usées considéré.

⁹³ Par exemple, la boîte à outils *LMI Toolbox* de Matlab ou des interfaces LMI telles que SeDuMi (développée au LAAS) ou YALMIP. Pour plus d'information sur des outils logiciels pour l'optimisation en commande, voir par exemple la page Web du projet OLOCEP du LAAS : <http://www.laas.fr/OLOCEP/logiciels.htm>.

Chapitre 5

Identification et Commande floues TS du bioréacteur

Dans ce chapitre, nous développons l'application des méthodes floues d'identification et de commande au bioprocédé considéré : un bioréacteur aérobique multivariable en mode continu pour le traitement des eaux usées issues d'une industrie papetière. Après avoir reprécisé le modèle obtenu à partir des bilans de matière, nous développons un modèle flou Takagi-Sugeno pour le bioréacteur basé sur des données expérimentales entrée-sortie en appliquant les techniques abordées dans le chapitre 3. Le système dynamique multi-variable est ainsi représentée comme un modèle à base de règles qui approxime la dynamique non linéaire comme une concaténation d'un ensemble de sous modèles localement linéaires sous la forme d'auto-régression non-linéaire (NARX) dans l'espace d'état. En appliquant la philosophie de commande du type PDC basée sur un modèle, présentée dans le chapitre 4, nous faisons la synthèse de la loi de commande la mieux adaptée au bioprocédé. Nous nous sommes intéressés, en particulier, à la commande sous-optimale linéaire quadratique par retour d'état pour les modèles TS affines en temps discret, afin de réguler les concentrations de biomasse et de substrat xénobiotique du bioprocédé.

Sommaire

5. Identification et Commande floues TS du bioréacteur.....	167
5.1. Description entrée-sortie du bioprocédé de traitement des eaux usées.....	169
5.1.1. Le modèle de bilan matière.....	169
5.1.2. Les paramètres du modèle	170
5.1.3. Les objectifs de conduite du procédé	171
5.1.4. Les contraintes sur les actionneurs.....	172
5.2. Modélisation et identification floue du bioprocédé à partir des données.....	172
5.2.1. Elaboration du jeu des données entrée-sortie.....	173
5.2.2. Utilisation du clustering flou Gustafson-Kessel et RoFER	176
5.3. Commande floue de type TS du bioprocédé basée sur le modèle flou	188
5.4. Conclusions	195

5.1. Description entrée-sortie du bioprocédé de traitement des eaux usées

Dans cette section nous reprenons le bioprocédé de dépollution des eaux résiduaires par lagunage aérée développé dans la section 2.5, en nous focalisant sur une représentation du comportement dynamique non linéaire entrée-sortie du procédé.

5.1.1. Le modèle de bilan matière

Comme nous l'avons déjà indiqué, le modèle considéré dans le cadre de notre étude est un bioréacteur aérobie multivariable en mode continu pour le traitement des eaux usées issues d'une industrie papetière. Ce procédé est alimenté par un mélange bi-polluant composé par des substrats énergétique et xénobiotique, dans lequel la biomasse, constituée par une population mixte de bactéries qui dégrade les polluants, exhibe des phénomènes tels que la mortalité endogène et des interactions microbiennes (l'activation/inhibition compétitive dans le mélange bi-substrat (cf. paragraphe 2.5.2)). Le substrat énergétique composé principalement de cellulose est caractérisé par une forte biodégradabilité alors que le substrat xénobiotique constitué d'un mélange de dérivés ligneux et de lignosulfonates, est très peu biodégradable. Ce substrat est considéré comme le principal polluant. Le schéma de principe du bioprocédé est présenté sur la Figure 5.1.

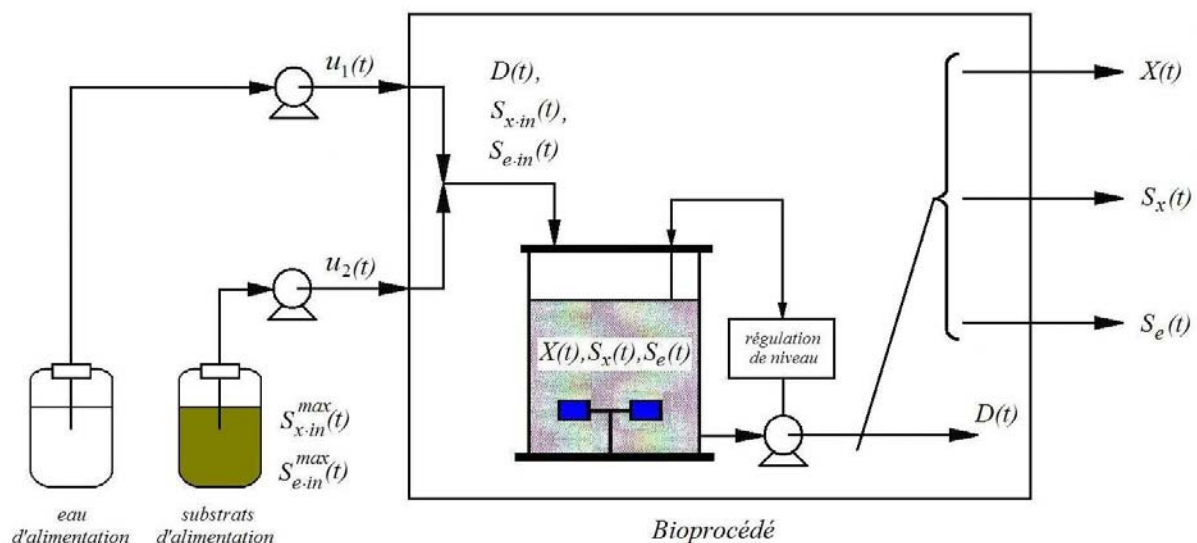


Figure 5.1. Schéma de principe entrée-sortie du bioprocédé de traitement des eaux usées.

Le système comporte deux entrées (variables de commande) : $u_1(t)$, le taux de dilution en eau et $u_2(t)$, le taux de dilution en substrats et trois variables de sorties qui correspondent aux concentrations de biomasse $X(t)$ et de substrats xénobiotique $S_x(t)$ et énergétique $S_e(t)$. Nous rappelons que le modèle mathématique que nous avons retenu, développé sous Matlab®, est une version améliorée par rapport celui du projet européen FAMIMO, pour lequel nous avons proposé l'ajout du terme de mortalité endogène au niveau de la dynamique bactérienne. En établissant le bilan de matière pour les principaux composants (cf. section 2.5.4), nous obtenons l'ensemble des équations différentielles (5.1) à (5.3), qui décrivent le comportement dynamique du bioprocédé :

$$\frac{dX(t)}{dt} = \left\{ \mu_{S_x}^{\max} \frac{S_x(t)}{K_{S_x} + S_x(t) + a_{S_e} S_e(t)} + \mu_{S_e}^{\max} \frac{S_e(t)}{K_{S_e} + S_e(t) + a_{S_x} S_x(t)} \right\} X(t) - (u_1(t) + u_2(t))X(t) - bX(t) \quad (5.1)$$

$$\frac{dS_x(t)}{dt} = -\frac{\mu_{S_x}^{\max}}{Y_{X/S_x}} \frac{S_x(t)}{K_{S_x} + S_x(t) + a_{S_e} S_e(t)} X(t) - u_1(t)S_x(t) + (S_{x-in}^{\max} - S_{x-in}(t))u_2(t) \quad (5.2)$$

$$\frac{dS_e(t)}{dt} = -\frac{\mu_{S_e}^{\max}}{Y_{X/S_e}} \frac{S_e(t)}{K_{S_e} + S_e(t) + a_{S_x} S_x(t)} X(t) - u_1(t)S_e(t) + (S_{e-in}^{\max} - S_{e-in}(t))u_2(t) \quad (5.3)$$

Comme nous l'avons indiqué précédemment, ce modèle de connaissance nous servira de base pour comparer dans la section 5.2, les performances d'un modèle flou représentant le comportement dynamique non-linéaire entrée-sortie du système. En utilisant la technique de compensation parallèle distribuée, nous aborderons également dans la section 5.3 la synthèse d'un contrôleur flou de type Takagi-Sugeno basé sur ce modèle.

5.1.2. Les paramètres du modèle

Les différents paramètres et variables qui interviennent dans les équations différentielles (5.1) à (5.3) et qui décrivent la dynamique du bioprocédé ont été explicités dans le chapitre 2. Ces paramètres et variables sont:

$X(t)$	concentration en (g/L) de biomasse
$S_x(t)$	concentration en (g/L) de substrat xénobiotique
$S_e(t)$	concentration en (g/L) de substrat énergétique
$\mu_{S_x}^{\max}$	taux spécifique de croissance maximal du au substrat xénobiotique
$\mu_{S_e}^{\max}$	taux spécifique de croissance maximal du au substrat énergétique
Y_{X/S_x}	rendement de conversion du substrat xénobiotique en biomasse
Y_{X/S_e}	rendement de conversion du substrat énergétique en biomasse
K_{S_x}, K_{S_e}	constantes de Michaëlis-Menten pour les substrats xénobiotique, énergétique ⁹⁴
a_{S_x}	constante d'inhibition du substrat énergétique sur le substrat xénobiotique
a_{S_e}	constante d'inhibition du substrat xénobiotique sur le substrat énergétique
b	terme de mortalité bactérienne due à la respiration endogène
$u_1(t)$	taux de dilution d'eau d'alimentation (eau propre)
$u_2(t)$	taux de dilution de substrats d'alimentation (eau résiduaire) ⁹⁵
$S_{x-in}(t)$	concentration en substrat xénobiotique du milieu d'alimentation
$S_{e-in}(t)$	concentration en substrat énergétique du milieu d'alimentation
$S_{x-in}^{\max}(t), S_{e-in}^{\max}(t)$	concentrations maximums en substrats dans le milieu d'alimentation

Les valeurs des différents paramètres du bioprocédé sont données dans le Tableau 5.1.

⁹⁴ Ces constantes correspondent à l'affinité de la population bactérienne mixte pour chaque substrat polluant.

⁹⁵ Le taux de dilution $D(t)=u_1(t)+u_2(t)$ correspondant au rapport entre le débit d'alimentation et le volume (constant) du bioréacteur qui n'apparaît pas explicitement dans le modèle. A sa place interviennent $u_1(t)$ et $u_2(t)$.

Description	Paramètre	Valeur	Unités
Rendements de conversion des substrats en biomasse	$Y_{X/Sx}$	0,7	g/g
	$Y_{X/Se}$	0,2	g/g
Taux spécifiques maximums de croissance	$\mu_{Sx}^{\max}$	0,1	h ⁻¹
	$\mu_{Se}^{\max}$	0,2	h ⁻¹
Mortalité bactérienne endogène	b	0,005	h ⁻¹
Constantes d'affinité de Michaëlis-Menten	K_{Sx}	1,5	g/L
	K_{Se}	1	g/L
Constantes d'inhibition/activation	a_{Sx}	0,1	g/g
	a_{Se}	10	g/g

Tableau 5.1. Paramètres stœchiométriques et cinétiques du bioprocédé

5.1.3. Les objectifs de conduite du procédé

Traditionnellement l'objectif principal en traitement d'eaux usées par voie biologique est la dégradation d'un maximum de polluants. Ainsi, l'augmentation de l'activité de biodégradation en assurant des conditions pour avoir une bonne croissance de la biomasse est une solution attrayante et naturelle. Néanmoins, le besoin de contrôler la prolifération de micro-organismes est devenu une nécessité pour préserver l'écosystème des nuisances bactériologiques. C'est en particulier le cas des industries papetières qui, malheureusement, affectent sérieusement l'équilibre naturel dans l'environnement proche du lieu du traitement. Ainsi, la synthèse du contrôleur pour le bioprocédé est multivariable : il doit réguler la concentration de biomasse (afin d'éviter la prolifération de micro-organismes) et celle du substrat xénobiotique (afin de dégrader le substrat polluant) sur des consignes pre-spécifiées. Les valeurs des concentrations de référence sont fixées en tenant compte des effets toxiques des bactéries et du substrat polluant sur le milieu récepteur naturel.

Ainsi, nous avons retenu le scénario de commande établi dans le benchmark de Ben Youssef [BEN97], utilisé lors du projet FAMIMO. Le scénario consiste en les consignes suivantes sur le substrat xénobiotique $S_x^*(t)$ et la biomasse $X^*(t)$ suivantes :

$$\begin{cases} S_x^*(t) = 0,3 \text{ g/L} & \forall t, \quad 0 \leq t < 180 \text{ h} \\ S_x^*(t) = 0,2 \text{ g/L} & \forall t, \quad 180 \leq t \leq 300 \text{ h} \end{cases} \quad (5.4)$$

$$\begin{cases} X^*(t) = 0,6 \text{ g/L} & \forall t, \quad 0 \leq t < 80 \text{ h} \\ X^*(t) = 0,7 \text{ g/L} & \forall t, \quad 80 \leq t < 220 \text{ h} \\ X^*(t) = 0,6 \text{ g/L} & \forall t, \quad 220 \leq t \leq 300 \text{ h} \end{cases} \quad (5.5)$$

Remarque

R 5.1 Si le bioréacteur évolue dans des conditions normales, telles que celles décrites dans le scénario ci-dessus, la concentration du substrat énergétique est résiduelle du fait de sa haute biodégradabilité. Ceci explique pourquoi on ne tient pas compte de la régulation du substrat énergétique au niveau du contrôleur.

5.1.4. Les contraintes sur les actionneurs

Quand le taux de dilution $D(t)$ est nul ou très faible et en absence du terme b de mortalité endogène, la biomasse suit une trajectoire de croissance exponentielle, qui est caractéristique d'un comportement instable. Pour les procédés biologiques fonctionnant en mode continu, l'objectif de commande est basé sur la régulation du procédé, à la différence des bioprocédés fed batch dans lesquels le but est généralement l'accumulation d'un produit de la réaction. En conséquence, l'imposition de valeurs minimales de taux de dilution et de concentrations dans les substrats d'alimentation donne une marge appréciable de sûreté pour éviter un comportement potentiellement instable du bioprocédé [BEN97]. De plus, une période relativement longue sans alimentation des micro-organismes pourrait aussi entraîner la mortalité microbienne car leur fragile coexistence dépend d'interactions complexes pour la consommation des substrats.

D'autre part, on a un phénomène de lavage de la biomasse du bioréacteur si le taux de dilution $D(t)$ est beaucoup plus grand que le taux spécifique de croissance des micro-organismes $\mu_x(t)$. Ainsi, les deux entrées (variables de commande) $u_1(t)$ et $u_2(t)$ sont soumises aux limitations supérieures d'entrée dues aux limitations mécaniques des pompes.

Les valeurs numériques des variables qui interviennent dans le système d'alimentation du modèle du bioprocédé sont regroupées dans le Tableau 5.2.

Description	Paramètre	Valeur	Unités
Limite supérieur d'entrée du taux de dilution	$u^{\max}$	0,04	h^{-1}
Limite inférieur d'entrée du taux de dilution	$u_{\min}$	0,0001	h^{-1}
Concentrations maximums de substrats d'alimentation	$S_{x-in}^{\max}(t)$	6	g/L
	$S_{e-in}^{\max}(t)$	6	g/L

Tableau 5.2. Bornes des variables dans le système d'alimentation du bioprocédé

5.2. Modélisation et identification floue du bioprocédé à partir des données

Comme nous l'avons déjà indiqué, le modèle complet du bioprocédé représenté par des équations différentielles (5.1) à (5.3) nous servira de référence pour la modélisation floue de type Takagi-Sugeno. Le modèle flou est construit à partir des données entrée-sortie du système obtenues en simulation en utilisant le logiciel Matlab®. Le système dynamique multi-variable est ainsi représentée comme un modèle à base de règles qui approxime le comportement dynamique non linéaire comme une concaténation d'un ensemble de sous modèles localement linéaires sous la forme d'auto-régression non-linéaire (NARX).

Nous considérons par la suite l'élaboration du jeu des données ainsi que l'utilisation du clustering flou Gustafson-Kessel et RoFER pour la construction du modèle Takagi-Sugeno du système. Nous abordons également l'obtention d'un modèle flou final avec une base commune de règles (partition unique) pour les trois sorties, afin d'utiliser le modèle pour la détermination de la commande dans l'espace d'état.

5.2.1. Elaboration du jeu des données entrée-sortie

Le premier pas vers la construction du modèle flou de type Takagi-Sugeno en utilisant les techniques de clustering flou consiste en la conception expérimentale de l'essai pour produire et acquérir les données entrée-sortie du procédé. Pour générer le jeu de données, nous supposons que toutes les variables sont accessibles à la mesure (mesure de l'état complet). En réalité dans le benchmark initial proposé par [BEN97], le substrat xénobiotique n'est pas accessible directement. Un observateur asymptotique réalisé par Ben Youssef a été utilisé lors du projet FAMIMO. Compte tenu de la présence du terme de mortalité endogène ajouté à la dynamique bactérienne, cet observateur n'est plus valable avec la même performance d'estimation. Bien que nous considérons les trois variables de sortie disponibles pour la mesure, il est tout à fait envisageable de développer un observateur flou basé sur ce modèle, en profitant de la structure de type Takagi-Sugeno obtenue.

Afin de disposer d'un jeu de données approprié pour l'identification non linéaire, il est préférable d'exciter le système dans toute la gamme des variables considérés tant en amplitude qu'en fréquence. En même temps, afin de rester dans l'espace défini par les références souhaitées, on est obligé de restreindre la commande $u_2(t)$. Normalement du fait de la modélisation des actionneurs, on est obligé de restreindre les variations des variables de commande entre 0,0001 et 0,04 h⁻¹. Un choix convenable de la gamme de $u_2(t)$ est de la restreindre de 0,0001 à 0,01 h⁻¹. On applique alors au processus des signaux multisinusoïdal ayant ces bornes. Un bruit blanc de petite taille est ajouté afin de garantir une excitation correcte de la dynamique du procédé pour les différents points de fonctionnement. On choisira une période d'échantillonnage $T_e = 2$ heures, avec un horizon de temps de 2000 h. On obtient le jeu de données suivant

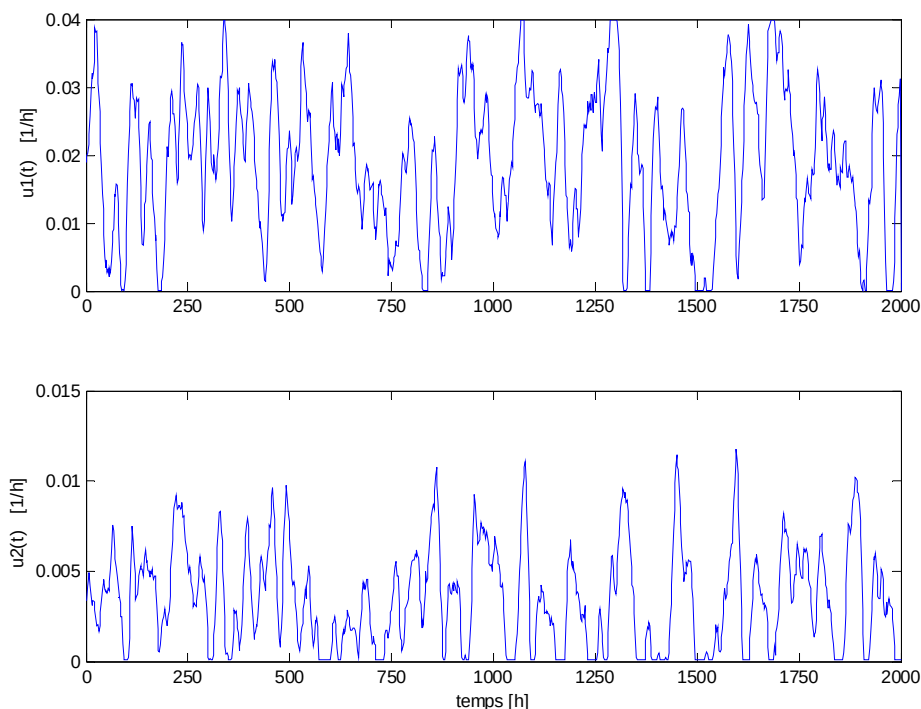


Figure 5.2. Signaux d'entrée $u_1(t)$ et $u_2(t)$ pour l'identification du bioprocédé

La répartition du jeu de données peut aussi être visualisé dans l'espace des entrées : les taux de dilution d'eau et de substrats d'alimentation - $u_1 u_2$ – (voir Figure 5.3).

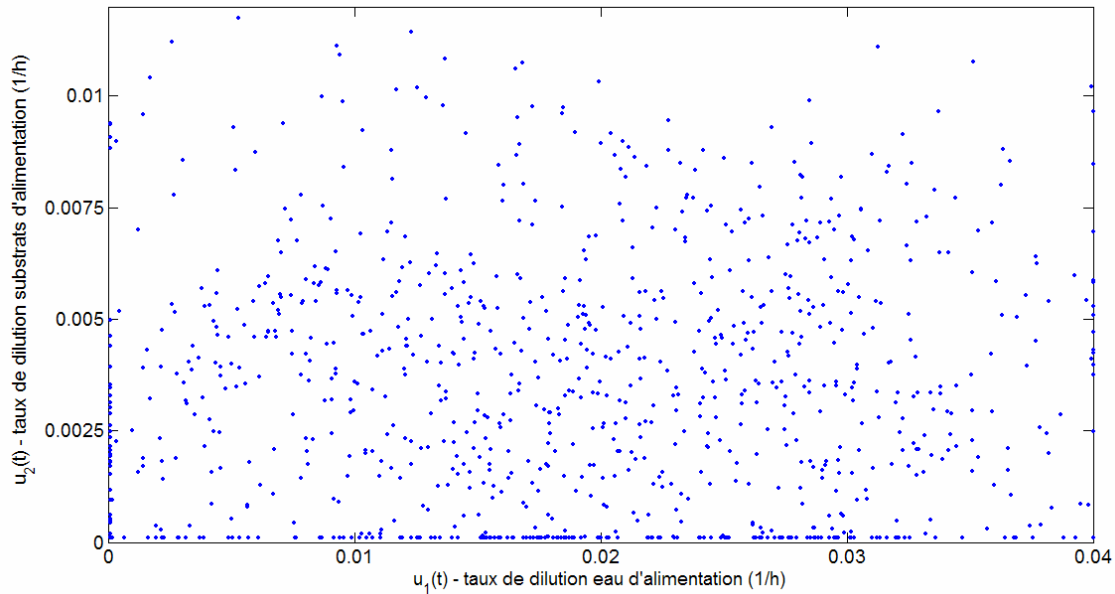


Figure 5.3. Répartition des données dans l'espace des entrées du procédé.

On peut noter que les entrées sont relativement bien distribuées dans la plage des variations possibles. Toutefois, il faut rappeler que l'obtention d'un jeu de données assez général pour l'identification n'est pas évidente. En effet, la conception de l'expérience d'identification doit tenir compte du niveau de l'excitation et des variations fréquentielles dans la gamme de réponse du procédé. De plus il faut être certain, avec ce jeu de valeurs, d'atteindre les différents points de fonctionnement. Le temps de réponse typique des bioprocédés ainsi que le modèle de connaissance nous ont aidés à dimensionner les signaux appropriés à notre cas d'étude.

Il est intéressant de souligner ici l'importance de la compréhension des phénomènes représentés par le modèle et sa relation avec les conditions de simulation. En effet, bien que la concentration de biomasse puisse atteindre une valeur nulle, il est évident que cette condition doit être impérativement évitée car sans biomasse, aucune dégradation ne peut continuer à se produire. Le respect de cette condition a été alors pris en compte au niveau de notre expérience d'identification. Le jeu des données d'identification a été obtenu avec $[X(0) \ S_x(0) \ S_e(0)]^T = [0,01 \ 0,3 \ 0,3]^T$ afin de tenir compte des situations réelles du comportement du bioréacteur, par exemple, à partir d'une phase de démarrage. Les données ont été prises en mode de fonctionnement continu, elles couvrent toute la région d'intérêt de notre application.

La Figure 5.4 montre la répartition correspondante des données dans l'espace des sorties des variables à commander, les concentrations de biomasse et du substrat xénobiotique :

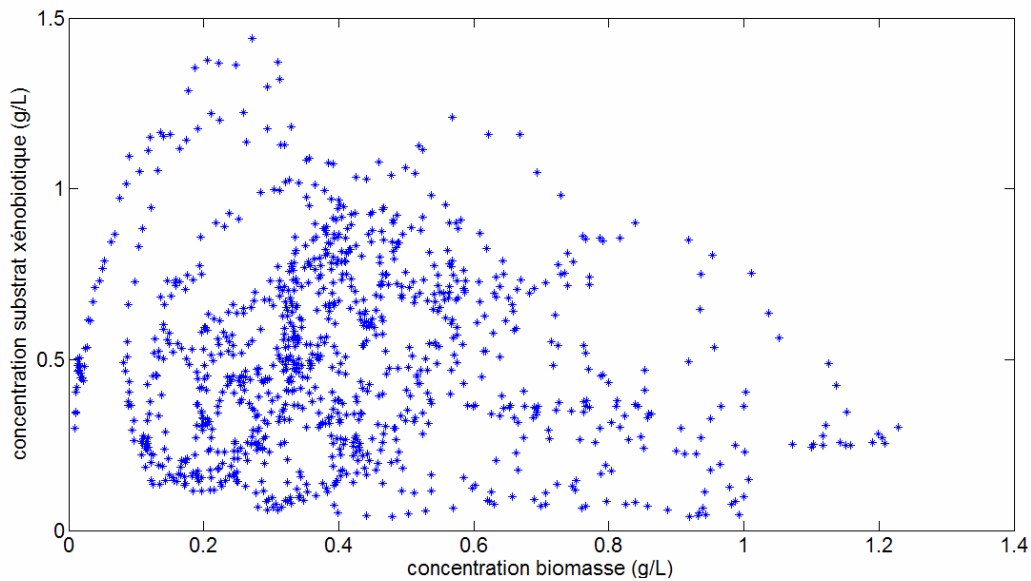


Figure 5.4. Répartition des données dans l'espace des sorties du procédé

Il est important de noter qu'un nombre important de points relativement bien distribués se trouve autour de la zone de contrôle d'intérêt. Le comportement des variables de sortie du bioprocédé pour l'expérience d'identification est visualisée sur la Figure 5.5.

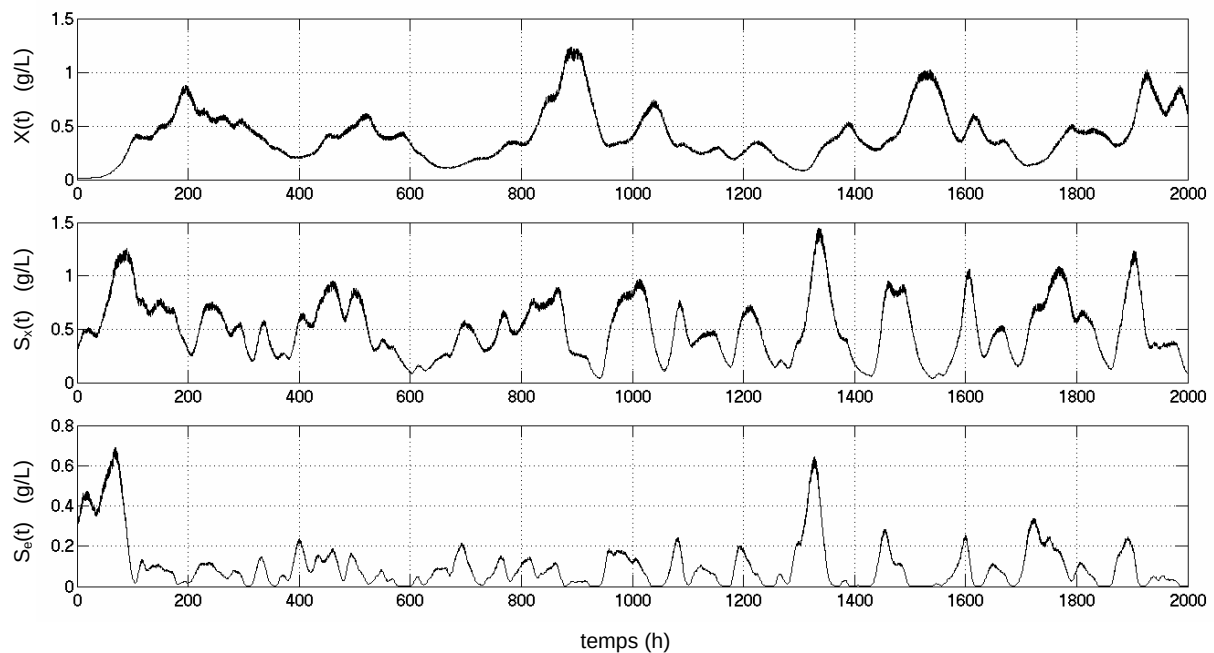


Figure 5.5. Evolution des sorties du procédé – données d'identification

L'augmentation de la concentration des substrats xénobiotique et énergétique dans les premières heures de l'expérience s'explique par la faible concentration de biomasse lors du démarrage. Ensuite, la concentration de biomasse est suffisante pour parvenir au processus de biodégradation. Les variations des taux de dilution d'entrée influencent le processus en faisant varier les variables de sortie du bioprocédé dans toute la gamme d'intérêt. La haute biodégradabilité du substrat énergétique par rapport au substrat xénobiotique est évidente en rapport aux faibles concentrations que ce premier atteint. Une valeur minimale de concentration de la biomasse est aussi respectée tout au long de l'expérience.

5.2.2. Utilisation du clustering flou Gustafson-Kessel et RoFER

Avec les données entrée-sortie du bioprocédé et afin de modéliser le comportement dynamique du système en appliquant les techniques de clustering flou [GRI05], on décompose le système MIMO à deux entrées ($u_1(k)$ et $u_2(k)$) et trois sorties ($X(k)$, $S_x(k)$ et $S_e(k)$) en trois modèles flous MISO de type Takagi-Sugeno qui conservent le même type de structure au niveau de leurs règles. À partir de la connaissance préalable du bioprocédé, nous avons choisi une dynamique pour les modèles flous TS telle que chaque sortie à l'instant $k+1$ dépende des trois sorties et des deux entrées à l'instant k . Ce choix est dû d'un côté, à la structure des variables couplées dans le modèle de bilan de matières donné par les expressions (5.1)-(5.3) (on peut supposer que le système est couplé) et d'autre part, afin de maintenir un compromis entre une complexité réduite du modèle et une bonne capacité d'approximation.

Dans ces conditions, l'ordre du système, lié aux constantes (cf. expression (3.25) et paragraphe subséquent) n_y (ordre de la sortie), n_u (ordre de l'entrée) et n_d (retard pur), est exprimé alors par les structures suivantes :

$$n_y = \begin{bmatrix} 1 & 1 & 1 \\ 1 & 1 & 1 \\ 1 & 1 & 1 \end{bmatrix}, \quad n_u = \begin{bmatrix} 1 & 1 \\ 1 & 1 \\ 1 & 1 \end{bmatrix}, \quad n_d = \begin{bmatrix} 1 & 1 \\ 1 & 1 \\ 1 & 1 \end{bmatrix} \quad (5.6)$$

où les facteurs '1' indiquent la présence d'une dépendance fonctionnelle entre les variables. Ainsi, les équations dynamiques du type NARX pour chaque règle R_i , avec $1 \leq i \leq c$, associée à chaque sortie $y_j(k)$ à l'instant k , avec $j = \{1,2,3\}$, - (notées ici $R_i^{y_j(k)}$) dans les trois modèles flous prennent alors la forme générale suivante :

$$\begin{aligned} & R_i^{y_j(k)} : \\ & \text{Si } y_1(k) \text{ est } A_{i1} \text{ et } y_2(k) \text{ est } A_{i2} \text{ et } y_3(k) \text{ est } A_{i3} \text{ et } u_1(k) \text{ est } A_{i4} \text{ et } u_2(k) \text{ est } A_{i5} \end{aligned} \quad (5.7)$$

Alors

$$y(k+1) = a_{i1}^{y_j(k)} y_1(k) + a_{i2}^{y_j(k)} y_2(k) + a_{i3}^{y_j(k)} y_3(k) + a_{i4}^{y_j(k)} u_1(k) + a_{i5}^{y_j(k)} u_2(k) + d_i^{y_j(k)}$$

Où $y_j(k)$ représente n'importe laquelle des sorties d'intérêt : $X(k)$, $S_x(k)$ ou $S_e(k)$.

Au niveau du processus de clustering, parmi les algorithmes présentés au Chapitre 3, on a décidé d'utiliser les algorithmes qui nous ont semblé les plus représentatifs et les mieux adaptés pour la tâche de modélisation : les algorithmes GK et RoFER. Cela nous permettra de faire des comparaisons pour l'obtention des modèles en utilisant d'un côté des clusters hyperellipsoïdaux avec obtention *a posteriori* des paramètres des conséquents (algorithme GK) et d'autre part, des clusters hyperplanaires avec estimation directe (simultanée avec la partition) des paramètres des conséquents (algorithme RoFER). Cela nous permettra aussi d'élargir le champ d'application montré dans les exemples à la fin du Chapitre 3, sur le bioprocédé actuellement considéré.

Rappelons que le nombre de règles dans le modèle flou de type Takagi-Sugeno correspond au nombre de clusters défini (cas de l'algorithme GK) ou obtenu (cas de l'algorithme RoFER). En utilisant les critères de validité du nombre de clusters rapportés dans

la section 3.4.1, on divise l'espace de représentation⁹⁶ en quatre clusters ($c = 4$) pour chaque sortie et le paramètre flou de la partition m est pris égal à 2. Pour l'algorithme RoFER le nombre initial (nombre maximal) de clusters a été fixé en 10. Avec un seuil de cardinalité robuste $\Sigma_{\mathcal{E}}$ proche de 21 pour chacune des trois sorties, nous avons obtenu le même nombre de clusters qu'avec l'algorithme GK. En utilisant le mécanisme de projection et la paramétrisation des fonctions d'appartenance exponentielles par morceaux, on extrait aussi les fonctions d'appartenance des ensembles flous multidimensionnels (clusters), de façon à avoir directement une interprétation du modèle par rapport aux variables des régresseurs, et donc celles des prémisses. On obtient alors des règles sur la forme décomposée.

Nous avons généré deux jeux de $N = 500$ données de validation, afin de tester la performance de prédiction du modèle flou TS dans des diverses scenarios [GRI05]. Un bruit blanc d'amplitude 5% a été ajouté à chaque sortie. Le premier jeu de données, illustré sur la Figure 5.6, correspond à un mode d'opération que nous avons appelé "mixte", dans la mesure qu'il possède des données d'entrée pour les modes d'opération *batch* (e.g. entre 0-180 h et 900-1000 h), mode *continu* (e.g., entre 180-400 h et 500-900 h), ainsi qu'une courte période de prédominance du taux spécifique de dilution de l'eau propre $u_1(t)$ (e.g., 400-450 h). Pour le mode batch, nous avons fait cette considération en prenant en compte des cas où les deux taux de dilution $u_1(t)$ et $u_2(t)$ sont nuls, par exemple, lors d'une action de commande imposée ou lors d'une procédure de démarrage du bioréacteur. Les conditions initiales pour ce test sont $[X(0) S_x(0) S_e(0)]^T = [0,5 \ 0,45 \ 0,2]^T$ (g/L). Cette expérience de validation est orienté afin de tester les capacités de généralisation du modèle.

Les jeux des données ainsi que les réponses obtenues sur les sorties du bioprocédé sont illustrées dans les figures suivantes :

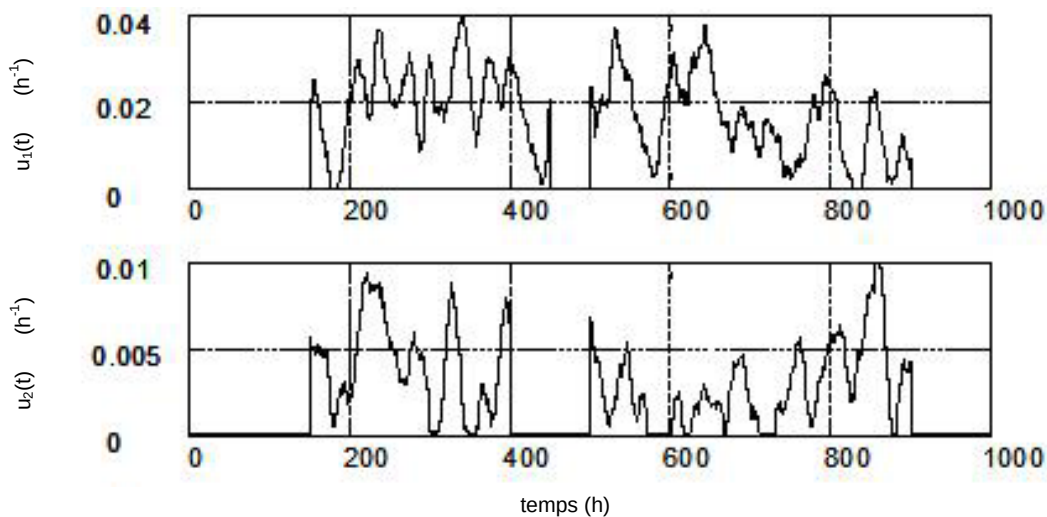


Figure 5.6. Signaux d'entrée $u_1(t)$ et $u_2(t)$ pour le premier test de validation du bioprocédé

Les résultats du premier test de validation (mode "mixte") du modèle flou TS pour chacune des sorties sont montrés sur la figure suivante :

⁹⁶ Espace de représentation de dimension 5, correspondant au nombre choisi de régresseurs dans le modèle dynamique : $y_1(k), y_2(k), y_3(k), u_1(k)$ et $u_2(k)$.

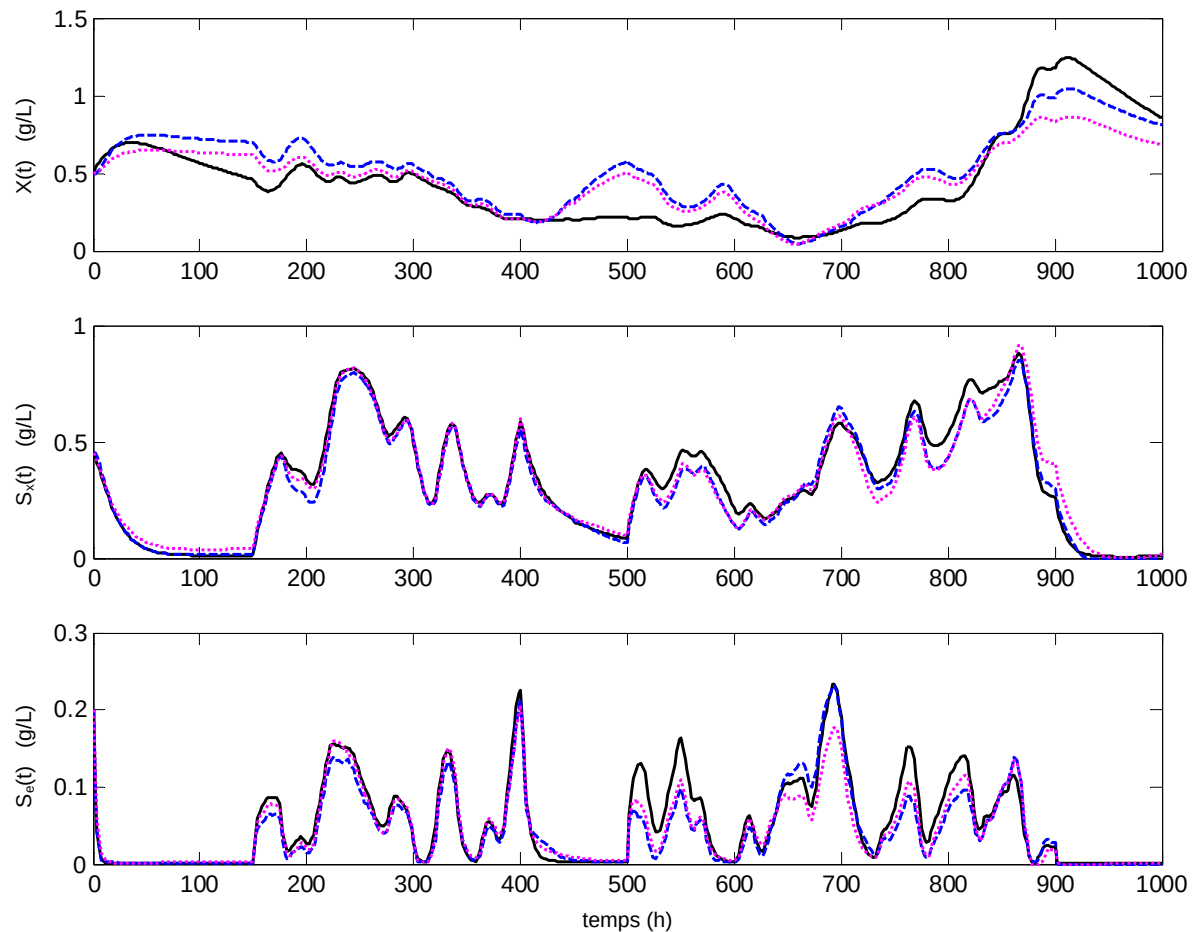


Figure 5.7. Premier test de validation du modèle flou TS du bioprocédé avec données d'opération en mode "mixte" (batch/continu). Sortie du procédé (ligne en trait plein), sortie du modèle avec clustering GK (ligne pleine) et sortie du modèle avec clustering RoFER (ligne pointillée).

Comme on peut le voir sur la Figure 5.7, pour les données du mode "mixte" de validation (batch/continu) en général la dynamique pour chacune des trois sorties du bioprocédé est raisonnablement bien approximé. Cependant, les capacités de généralisation du modèle hors des conditions établis lors du processus d'identification (réalisé en mode continu) sont assez limitées, en particulier pour reproduire la dynamique propre du système dans le mode batch pour de longues périodes. Dans notre cas, cela ne constitue en général pas un problème, car l'opération du bioréacteur pour le traitement des eaux usées se déroule en mode continu. Néanmoins, si le mode batch était envisagée, il vaudrait mieux proposer un modèle flou TS plus approprié que celui proposé actuellement, c'est-à-dire, en supprimant directement les régresseurs correspondants aux entrées $u_1(t)$ et $u_2(t)$ pour éviter des divergences numériques quand le mécanisme de projection est utilisé pour l'obtention des fonctions d'appartenance des règles. Si les deux modes d'opération étaient envisagés pour une application particulière, nous considérons qu'une solution à ce problème de modélisation pourrait être la conception d'un modèle hybride qui utilise des modèles flous dynamiques Takagi-Sugeno pour décrire le fonctionnement pour chacun des modes d'opération du bioprocédé. Cette perspective est à explorer dans une cadre plus général d'application de la modélisation flou de type Takagi-Sugeno aux procédés biotechnologiques.

Remarque

R 5.2 Pour le mode d'opération batch les taux de dilution $u_1(t)$ et $u_2(t)$ doivent être égaux à zéro, cependant en pratique, il est nécessaire d'imposer une valeur très faible au niveau de la simulation pour éviter des divergences numériques. Pour le cas du benchmark du projet FAMIMO, la valeur établie expérimentalement à partir de plusieurs simulations est de $0,0001 \text{ h}^{-1}$, ce qui correspond aussi à la limite inférieure d'entrée du taux de dilution dans le système d'alimentation du bioréacteur.

Un deuxième jeu de données est utilisé pour un deuxième test de validation du modèle. Dans ce cas, le jeu des données correspond au mode d'opération continu du bioprocédé, pour lequel nous avons utilisés les mêmes conditions initiales que dans le cas précédent. Cette expérience de validation est faite pour tester les capacités de prédiction du modèle dans le mode continu, qui est le mode d'opération du bioprocédé de dépollution.

Les résultats du second test de validation (mode continu) du modèle flou TS pour chacune des sorties sont montrés sur la Figure 5.8 :

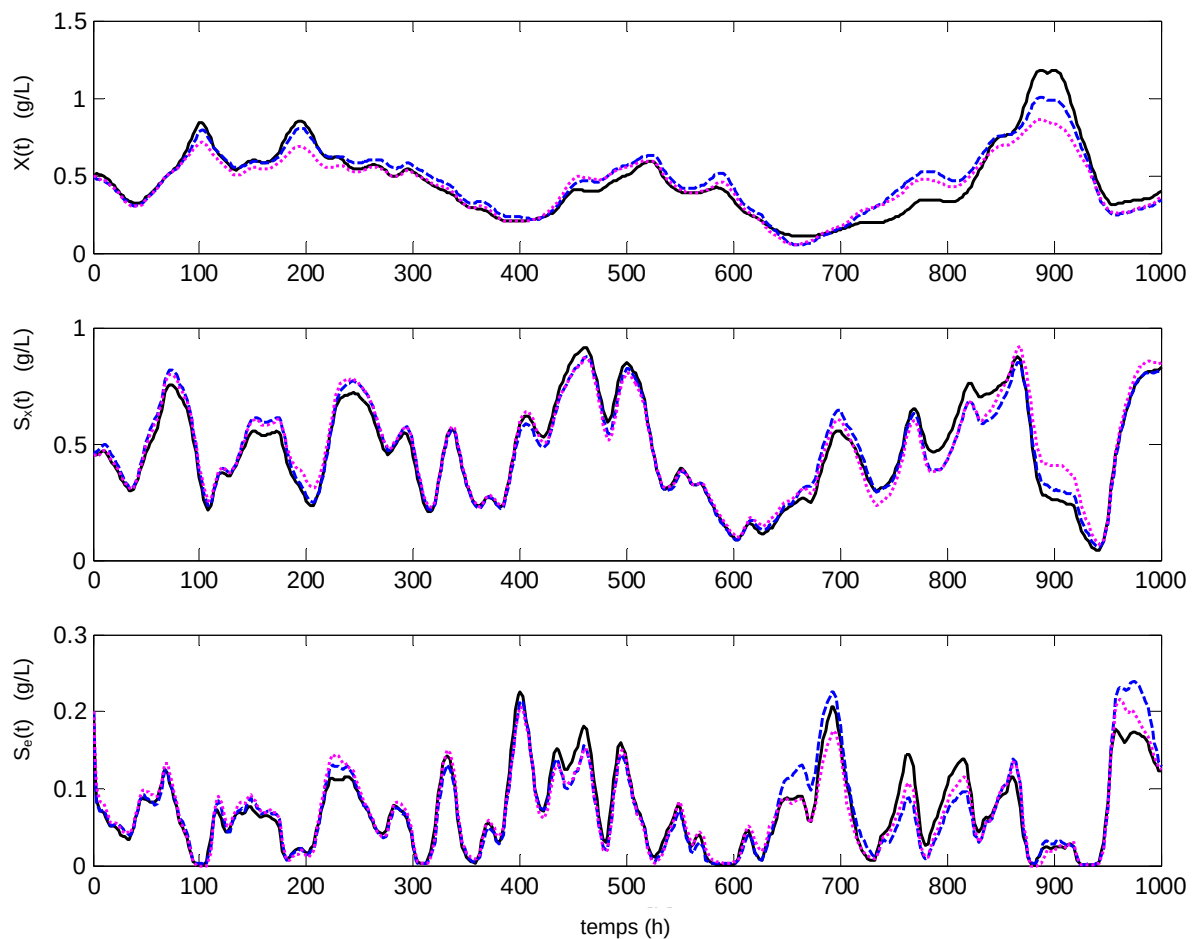


Figure 5.8. Deuxième test de validation du modèle flou TS du bioprocédé avec données d'opération en mode continue. Sortie du procédé (ligne en trait plein), sortie du modèle avec groupement GK (ligne pleine) et sortie du modèle avec groupement RoFER (ligne pointillée).

La qualité numérique de l'approximation est mesuré en utilisant les critères RMSE et VAF décrits dans la section 3.4.4. Les valeurs idéales pour ces deux critères sont

respectivement 0% et 100%. La performance numérique comparative des deux approximations est regroupée dans les Tableaux suivants :

Variable	Test en mode "mixte"		Test en mode continu	
	RMSE	VAF	RMSE	VAF
$X(k)$	0,4902	77,35%	0,4902	91,31%
$S_x(k)$	0,4888	95,55%	0,4760	96,11%
$S_e(k)$	0,5000	87,29%	0,4695	84,58%

Tableau 5.3. Performance numérique de la validation du modèle flou TS pour le bioréacteur – résultats avec l'algorithme GK

Variable	Test en mode "mixte"		Test en mode continu	
	RMSE	VAF	RMSE	VAF
$X(k)$	0,4902	79,29%	0,4902	91,87%
$S_x(k)$	0,4902	95,52%	0,4778	95,12%
$S_e(k)$	0,4987	92,71%	0,4687	95,26%

Tableau 5.4. Performance numérique de la validation du modèle flou TS pour le bioréacteur – résultats avec l'algorithme RoFER

En présence du bruit ajouté l'algorithme RoFER est légèrement meilleur que l'algorithme GK. Néanmoins, il faut rappeler que dans le premier algorithme l'estimation des paramètres se fait simultanément avec la partition des données, ce qui conduit à une meilleure approximation locale, utile pour notre propos de commande basée sur le modèle. En contrepartie, l'algorithme GK est beaucoup plus performant au niveau temps de calcul.

Remarque

R 5.3 Les données expérimentales (dans de milieu industriel) sont entachées de bruit et éventuellement de points hors tendance. Dans ce contexte l'algorithme RoFER aurait montré un des avantages qui est la robustesse au bruit et points aberrants.

Nous donnons à présent les fonctions d'appartenance obtenues avec chacun des deux algorithmes pour le cas du test de validation du modèle flou TS du bioprocédé avec données d'opération en mode continu. Les Figures 5.9 à 5.11 correspondent aux résultats avec l'algorithme GK et les Figures 5.12 à 5.14 aux résultats avec l'algorithme RoFER.

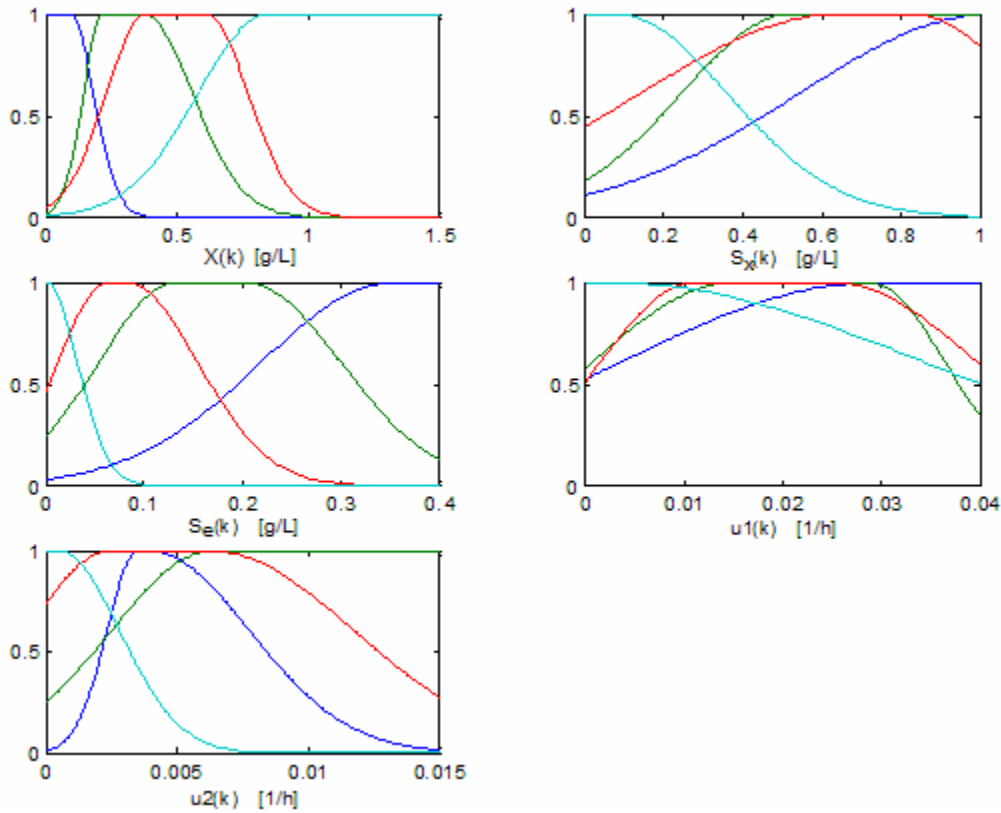


Figure 5.9. Partition $X(k)$ avec GK données d'opération en mode continue

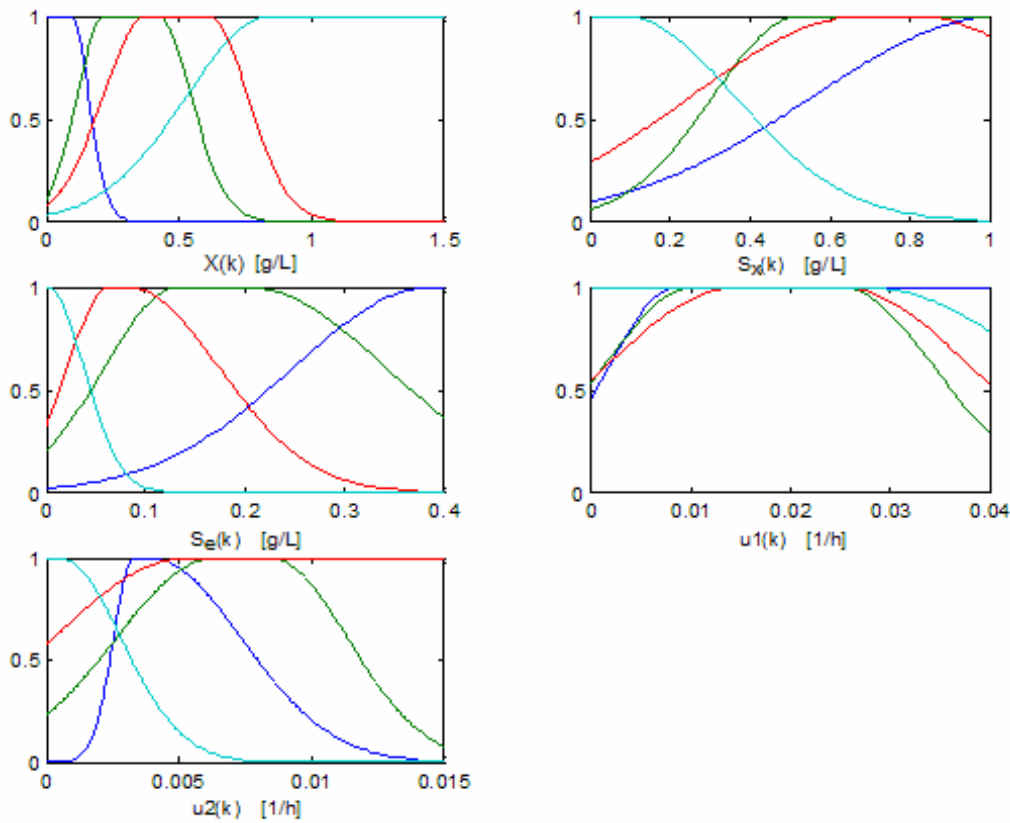


Figure 5.10. Partition $S_x(k)$ avec GK données d'opération en mode continue

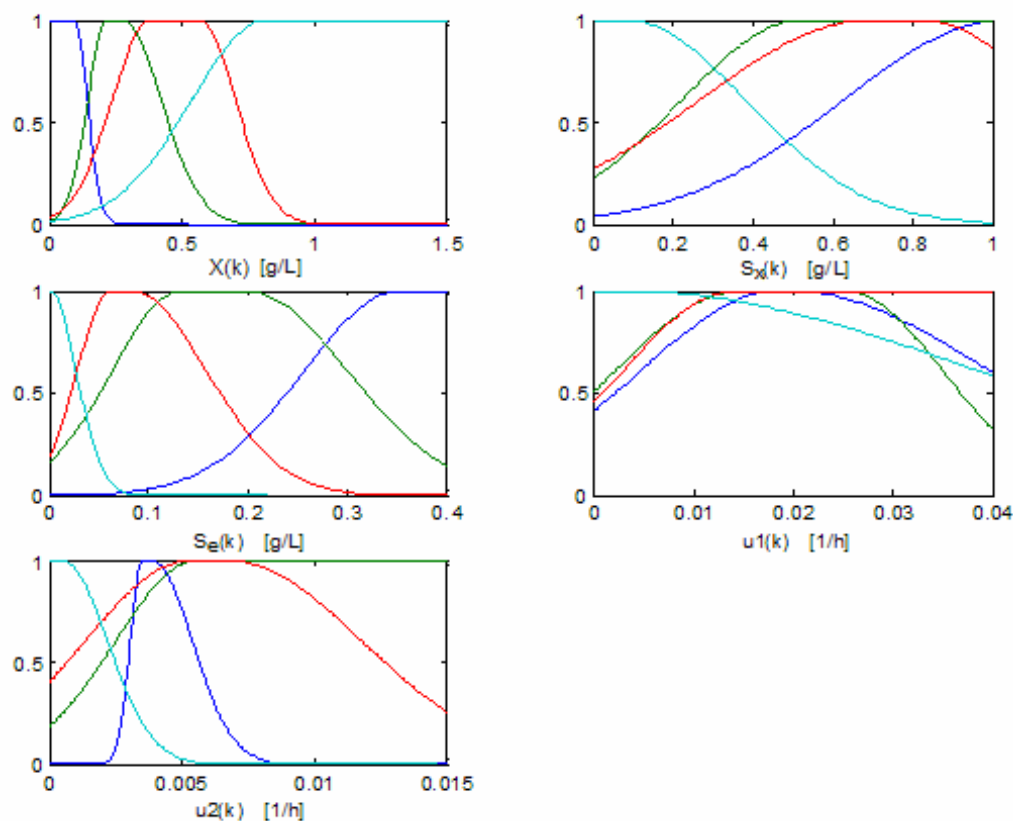


Figure 5.11. Partition $S_e(k)$ avec GK données d'opération en mode continue

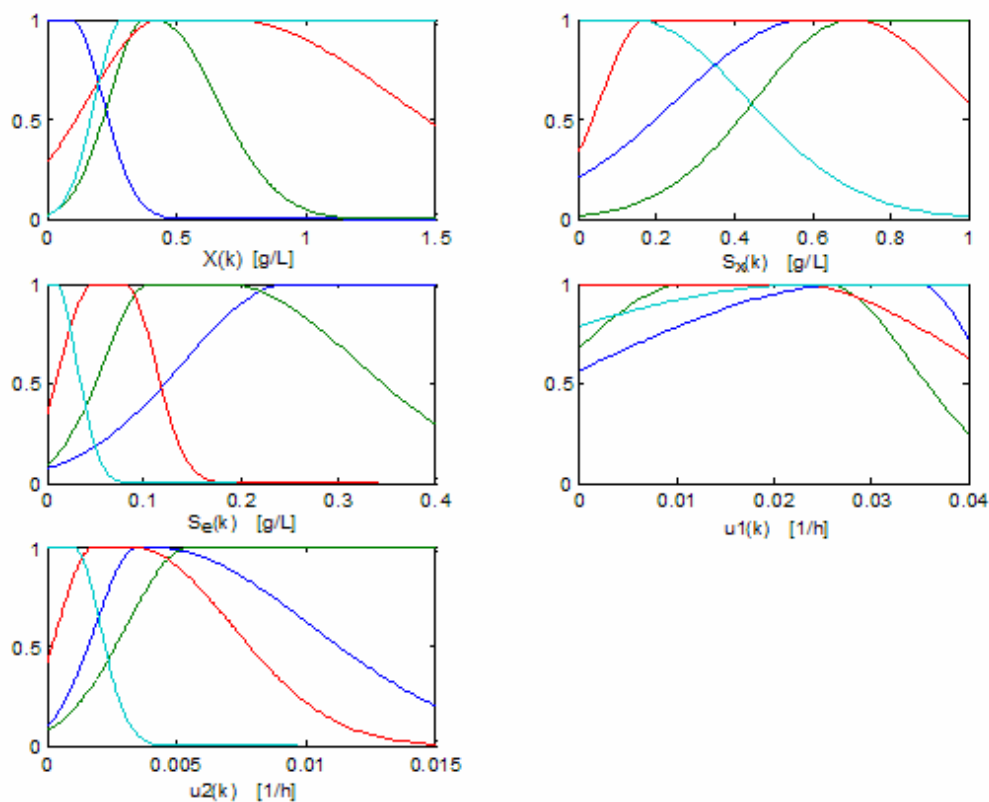


Figure 5.12. Partition $X(k)$ avec RoFER données d'opération en mode continue

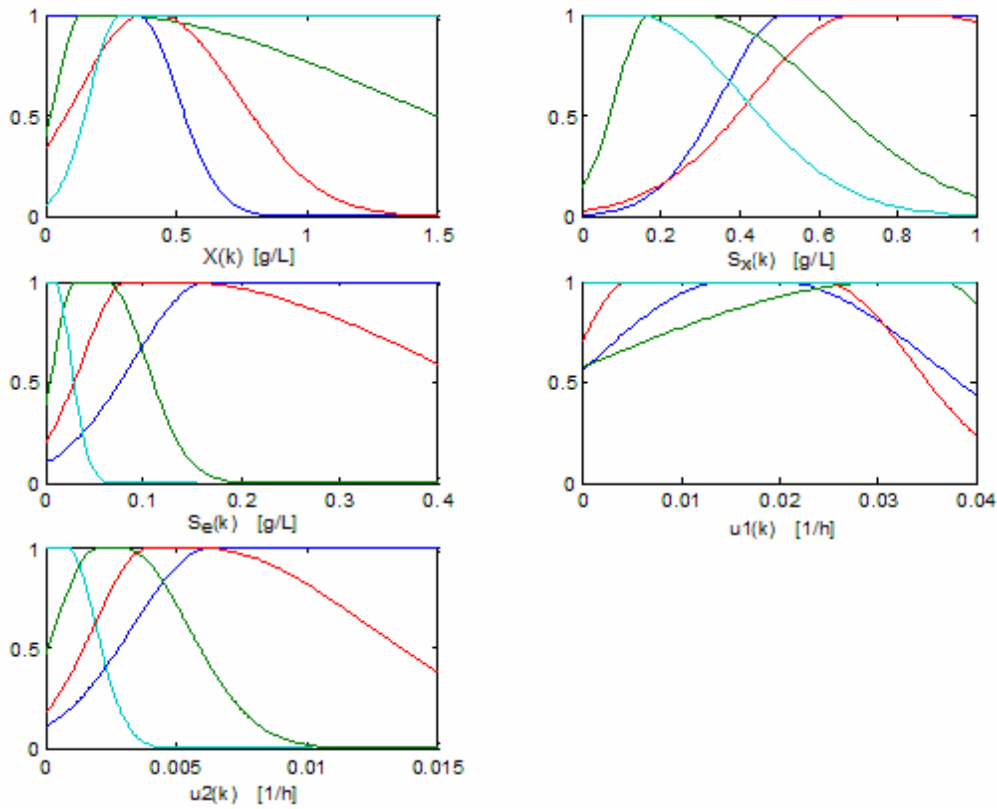


Figure 5.13. Partition $S_x(k)$ avec RoFER données d'opération en mode continue

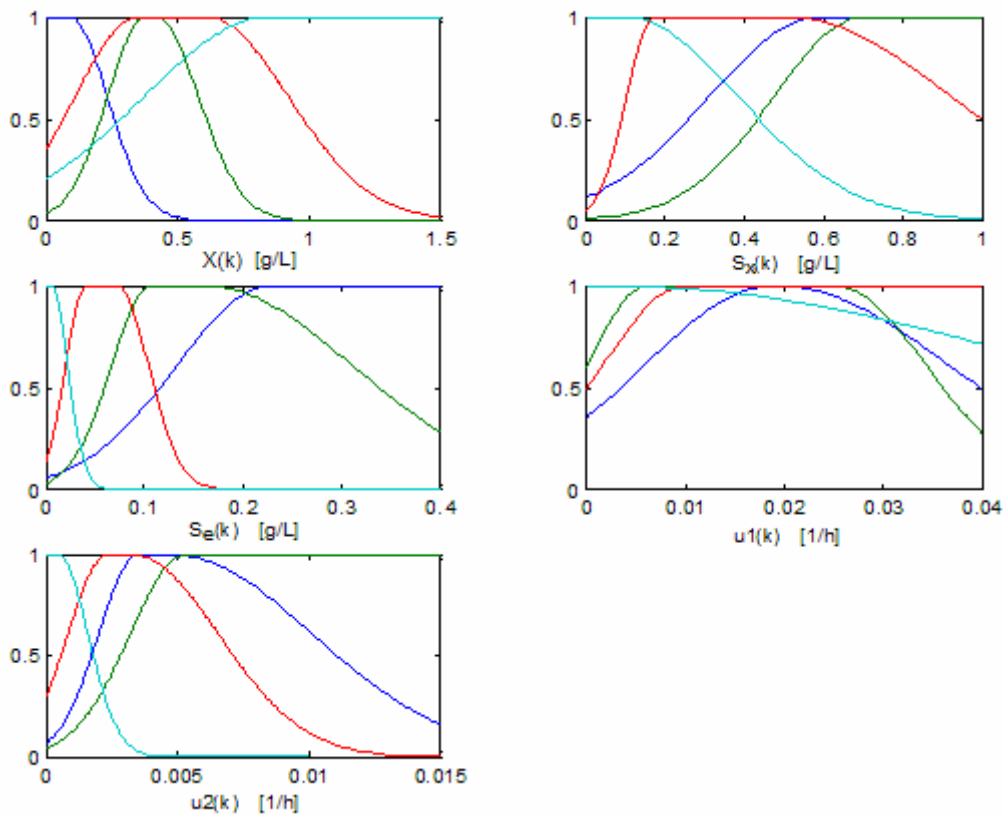


Figure 5.14. Partition $S_e(k)$ avec RoFER données d'opération en mode continue

Maintenant nous allons présenter les règles obtenues pour chacune des trois sorties des modèles flous Takagi-Sugeno.

La forme des règles pour la sortie concentration de biomasse $X(k)$ tant en utilisant l'algorithme GK qu'en utilisant l'algorithme RoFER, est donnée par les expressions suivantes :

$R_1^{X(k)}$:

Si $X(k)$ est A_{11} **et** $S_x(k)$ est A_{12} **et** $S_e(k)$ est A_{13} **et** $u_1(k)$ est A_{14} **et** $u_2(k)$ est A_{15}
Alors $X(k+1) = a_{11}^{X(k)} X(k) + a_{12}^{X(k)} S_x(k) + a_{13}^{X(k)} S_e(k) + a_{14}^{X(k)} u_1(k) + a_{15}^{X(k)} u_2(k) + d_1^{X(k)}$

$R_2^{X(k)}$:

Si $X(k)$ est A_{21} **et** $S_x(k)$ est A_{22} **et** $S_e(k)$ est A_{23} **et** $u_1(k)$ est A_{24} **et** $u_2(k)$ est A_{25}
Alors $X(k+1) = a_{21}^{X(k)} X(k) + a_{22}^{X(k)} S_x(k) + a_{23}^{X(k)} S_e(k) + a_{24}^{X(k)} u_1(k) + a_{25}^{X(k)} u_2(k) + d_2^{X(k)}$

$R_3^{X(k)}$:

Si $X(k)$ est A_{31} **et** $S_x(k)$ est A_{32} **et** $S_e(k)$ est A_{33} **et** $u_1(k)$ est A_{34} **et** $u_2(k)$ est A_{35}
Alors $X(k+1) = a_{31}^{X(k)} X(k) + a_{32}^{X(k)} S_x(k) + a_{33}^{X(k)} S_e(k) + a_{34}^{X(k)} u_1(k) + a_{35}^{X(k)} u_2(k) + d_3^{X(k)}$

$R_4^{X(k)}$:

Si $X(k)$ est A_{41} **et** $S_x(k)$ est A_{42} **et** $S_e(k)$ est A_{43} **et** $u_1(k)$ est A_{44} **et** $u_2(k)$ est A_{45}
Alors $X(k+1) = a_{41}^{X(k)} X(k) + a_{42}^{X(k)} S_x(k) + a_{43}^{X(k)} S_e(k) + a_{44}^{X(k)} u_1(k) + a_{45}^{X(k)} u_2(k) + d_4^{X(k)}$

Les coefficients respectifs a_{il} du modèle de régression, avec $l = \{1, 2, \dots, 5\}$, en utilisant l'algorithme GK sont regroupés dans le Tableau 5.5 :

Tableau 5.5. Paramètres des conséquents du modèle flou TS pour $X(k+1)$ - algorithme GK

règle	$X(k)$	$S_x(k)$	$S_e(k)$	$u_1(k)$	$u_2(k)$	biais
$R_1^{X(k)}$	$1,06 \cdot 10^0$	$5,22 \cdot 10^{-3}$	$-2,63 \cdot 10^{-3}$	$-1,61 \cdot 10^{-1}$	$-2,78 \cdot 10^{-1}$	$3,50 \cdot 10^{-3}$
$R_2^{X(k)}$	$9,58 \cdot 10^{-1}$	$2,53 \cdot 10^{-2}$	$-1,58 \cdot 10^{-2}$	$-5,33 \cdot 10^{-1}$	$2,10 \cdot 10^0$	$-1,58 \cdot 10^{-3}$
$R_3^{X(k)}$	$9,70 \cdot 10^{-1}$	$3,84 \cdot 10^{-2}$	$-7,39 \cdot 10^{-2}$	$-9,54 \cdot 10^{-1}$	$9,96 \cdot 10^{-1}$	$1,89 \cdot 10^{-2}$
$R_4^{X(k)}$	$9,82 \cdot 10^{-1}$	$7,69 \cdot 10^{-2}$	$3,11 \cdot 10^{-1}$	$-1,70 \cdot 10^0$	$-3,59 \cdot 10^0$	$1,11 \cdot 10^{-2}$

Les centres des clusters dans l'espace des régresseurs en utilisant l'algorithme GK sont indiqués dans le Tableau 5.6 :

Tableau 5.6. Centres des clusters du modèle flou TS pour $X(k+1)$ - algorithme GK

règle	$X(k)$	$S_x(k)$	$S_e(k)$	$u_1(k)$	$u_2(k)$
$R_1^{X(k)}$	$1,31 \cdot 10^{-1}$	$6,86 \cdot 10^{-1}$	$3,50 \cdot 10^{-1}$	$1,90 \cdot 10^{-2}$	$4,72 \cdot 10^{-3}$
$R_2^{X(k)}$	$3,44 \cdot 10^{-1}$	$5,81 \cdot 10^{-1}$	$1,16 \cdot 10^{-1}$	$1,88 \cdot 10^{-2}$	$4,77 \cdot 10^{-3}$
$R_3^{X(k)}$	$4,28 \cdot 10^{-1}$	$5,41 \cdot 10^{-1}$	$7,16 \cdot 10^{-2}$	$1,91 \cdot 10^{-2}$	$3,74 \cdot 10^{-3}$
$R_4^{X(k)}$	$7,29 \cdot 10^{-1}$	$2,94 \cdot 10^{-1}$	$1,58 \cdot 10^{-2}$	$1,65 \cdot 10^{-2}$	$1,62 \cdot 10^{-3}$

Les coefficients respectifs a_{il} du modèle de régression, avec $l = \{1, 2, \dots, 5\}$, en utilisant l'algorithme RoFER sont regroupés dans le Tableau 5.7 :

Tableau 5.7. Paramètres des conséquents du modèle flou TS pour $X(k+1)$ - algorithme RoFER

règle	$X(k)$	$S_x(k)$	$S_e(k)$	$u_1(k)$	$u_2(k)$	biais
$R_1^{X(k)}$	$9,61 \cdot 10^{-1}$	$2,51 \cdot 10^{-2}$	$-2,44 \cdot 10^{-2}$	$-1,71 \cdot 10^{-1}$	$2,78 \cdot 10^{-1}$	$2,98 \cdot 10^{-3}$
$R_2^{X(k)}$	$9,72 \cdot 10^{-1}$	$2,52 \cdot 10^{-2}$	$-4,84 \cdot 10^{-2}$	$-8,88 \cdot 10^{-1}$	$1,15 \cdot 10^0$	$1,56 \cdot 10^{-2}$
$R_3^{X(k)}$	$9,68 \cdot 10^{-1}$	$3,43 \cdot 10^{-2}$	$-3,05 \cdot 10^{-2}$	$-7,84 \cdot 10^{-1}$	$1,18 \cdot 10^0$	$1,02 \cdot 10^{-2}$
$R_4^{X(k)}$	$9,78 \cdot 10^{-1}$	$5,20 \cdot 10^{-2}$	$-6,23 \cdot 10^{-2}$	$-1,15 \cdot 10^0$	$5,80 \cdot 10^{-1}$	$1,26 \cdot 10^{-2}$

Les centres des clusters dans l'espace des régresseurs en utilisant l'algorithme RoFER sont indiqués dans le Tableau 5.8 :

Tableau 5.8. Centres des clusters du modèle flou TS pour $X(k+1)$ - algorithme RoFER

règle	$X(k)$	$S_x(k)$	$S_e(k)$	$u_1(k)$	$u_2(k)$
$R_1^{X(k)}$	$1,44 \cdot 10^{-1}$	$5,49 \cdot 10^{-1}$	$2,71 \cdot 10^{-1}$	$2,22 \cdot 10^{-2}$	$4,68 \cdot 10^{-3}$
$R_2^{X(k)}$	$3,92 \cdot 10^{-1}$	$7,00 \cdot 10^{-1}$	$1,29 \cdot 10^{-1}$	$1,85 \cdot 10^{-2}$	$5,81 \cdot 10^{-3}$
$R_3^{X(k)}$	$3,98 \cdot 10^{-1}$	$4,69 \cdot 10^{-1}$	$5,97 \cdot 10^{-2}$	$2,02 \cdot 10^{-2}$	$3,10 \cdot 10^{-3}$
$R_4^{X(k)}$	$4,76 \cdot 10^{-1}$	$2,64 \cdot 10^{-1}$	$9,74 \cdot 10^{-3}$	$1,98 \cdot 10^{-2}$	$6,03 \cdot 10^{-4}$

La forme des règles pour la sortie concentration de substrat xénobiotique $S_x(k)$ tant en utilisant l'algorithme GK qu'en utilisant l'algorithme RoFER, est donnée par les expressions suivantes :

$R_1^{S_x(k)} :$

Si $X(k)$ est A_{11} **et** $S_x(k)$ est A_{12} **et** $S_e(k)$ est A_{13} **et** $u_1(k)$ est A_{14} **et** $u_2(k)$ est A_{15}

Alors $X(k+1) = a_{11}^{S_x(k)} X(k) + a_{12}^{S_x(k)} S_x(k) + a_{13}^{S_x(k)} S_e(k) + a_{14}^{S_x(k)} u_1(k) + a_{15}^{S_x(k)} u_2(k) + d_1^{S_x(k)}$

$R_2^{S_x(k)} :$

Si $X(k)$ est A_{21} **et** $S_x(k)$ est A_{22} **et** $S_e(k)$ est A_{23} **et** $u_1(k)$ est A_{24} **et** $u_2(k)$ est A_{25}

Alors $X(k+1) = a_{21}^{S_x(k)} X(k) + a_{22}^{S_x(k)} S_x(k) + a_{23}^{S_x(k)} S_e(k) + a_{24}^{S_x(k)} u_1(k) + a_{25}^{S_x(k)} u_2(k) + d_2^{S_x(k)}$

$R_3^{S_x(k)} :$

Si $X(k)$ est A_{31} **et** $S_x(k)$ est A_{32} **et** $S_e(k)$ est A_{33} **et** $u_1(k)$ est A_{34} **et** $u_2(k)$ est A_{35}

Alors $X(k+1) = a_{31}^{S_x(k)} X(k) + a_{32}^{S_x(k)} S_x(k) + a_{33}^{S_x(k)} S_e(k) + a_{34}^{S_x(k)} u_1(k) + a_{35}^{S_x(k)} u_2(k) + d_3^{S_x(k)}$

$R_4^{S_x(k)} :$

Si $X(k)$ est A_{41} **et** $S_x(k)$ est A_{42} **et** $S_e(k)$ est A_{43} **et** $u_1(k)$ est A_{44} **et** $u_2(k)$ est A_{45}

Alors $X(k+1) = a_{41}^{S_x(k)} X(k) + a_{42}^{S_x(k)} S_x(k) + a_{43}^{S_x(k)} S_e(k) + a_{44}^{S_x(k)} u_1(k) + a_{45}^{S_x(k)} u_2(k) + d_4^{S_x(k)}$

Les coefficients respectifs a_{il} du modèle de régression, avec $l = \{1, 2, \dots, 5\}$, en utilisant l'algorithme GK sont regroupés dans le Tableau 5.9 :

Tableau 5.9. Paramètres des conséquents du modèle flou TS pour $S_x(k+1)$ - algorithme GK

règle	$X(k)$	$S_x(k)$	$S_e(k)$	$u_1(k)$	$u_2(k)$	biais
$R_1^{S_x(k)}$	$1,25 \cdot 10^0$	$6,45 \cdot 10^{-1}$	$3,92 \cdot 10^{-1}$	$-1,59 \cdot 10^0$	$4,61 \cdot 10^0$	$2,00 \cdot 10^{-2}$
$R_2^{S_x(k)}$	$2,29 \cdot 10^{-2}$	$9,18 \cdot 10^{-1}$	$5,44 \cdot 10^{-2}$	$-2,05 \cdot 10^0$	$1,04 \cdot 10^1$	$3,23 \cdot 10^{-2}$
$R_3^{S_x(k)}$	$-5,87 \cdot 10^{-2}$	$8,95 \cdot 10^{-1}$	$7,10 \cdot 10^{-2}$	$-8,20 \cdot 10^{-1}$	$1,08 \cdot 10^1$	$4,99 \cdot 10^{-2}$
$R_4^{S_x(k)}$	$-2,72 \cdot 10^{-2}$	$8,82 \cdot 10^{-1}$	$-5,72 \cdot 10^{-1}$	$-2,15 \cdot 10^{-1}$	$1,47 \cdot 10^1$	$2,00 \cdot 10^{-2}$

Les centres des clusters dans l'espace des régresseurs en utilisant l'algorithme GK sont indiqués dans le Tableau 5.10 :

Tableau 5.10. Centres des clusters du modèle flou TS pour $S_x(k+1)$ - algorithme GK

règle	$X(k)$	$S_x(k)$	$S_e(k)$	$u_1(k)$	$u_2(k)$
$R_1^{S_x(k)}$	$1,20 \cdot 10^{-1}$	$7,15 \cdot 10^{-1}$	$3,84 \cdot 10^{-1}$	$1,79 \cdot 10^{-2}$	$4,77 \cdot 10^{-3}$
$R_2^{S_x(k)}$	$3,35 \cdot 10^{-1}$	$6,09 \cdot 10^{-1}$	$1,26 \cdot 10^{-1}$	$1,85 \cdot 10^{-2}$	$4,84 \cdot 10^{-3}$
$R_3^{S_x(k)}$	$4,25 \cdot 10^{-1}$	$5,69 \cdot 10^{-1}$	$8,10 \cdot 10^{-2}$	$1,92 \cdot 10^{-2}$	$4,10 \cdot 10^{-3}$
$R_4^{S_x(k)}$	$6,62 \cdot 10^{-1}$	$2,79 \cdot 10^{-1}$	$1,67 \cdot 10^{-2}$	$1,78 \cdot 10^{-2}$	$1,59 \cdot 10^{-3}$

Les coefficients respectifs a_{il} du modèle de régression, avec $l = \{1, 2, \dots, 5\}$, en utilisant l'algorithme RoFER sont regroupés dans le Tableau 5.13 :

Tableau 5.11. Paramètres des conséquents du modèle flou TS pour $S_x(k+1)$ - algorithme RoFER

règle	$X(k)$	$S_x(k)$	$S_e(k)$	$u_1(k)$	$u_2(k)$	biais
$R_1^{S_x(k)}$	$-4,05 \cdot 10^0$	$9,15 \cdot 10^{-1}$	$1,28 \cdot 10^{-2}$	$-1,29 \cdot 10^0$	$1,13 \cdot 10^1$	$4,54 \cdot 10^{-2}$
$R_2^{S_x(k)}$	$-4,38 \cdot 10^{-2}$	$9,01 \cdot 10^{-1}$	$-8,13 \cdot 10^{-2}$	$-5,79 \cdot 10^{-1}$	$1,24 \cdot 10^1$	$3,31 \cdot 10^{-2}$
$R_3^{S_x(k)}$	$-7,06 \cdot 10^{-2}$	$9,10 \cdot 10^{-1}$	$2,56 \cdot 10^{-2}$	$-1,55 \cdot 10^0$	$1,16 \cdot 10^1$	$6,14 \cdot 10^{-2}$
$R_4^{S_x(k)}$	$-3,04 \cdot 10^{-2}$	$8,92 \cdot 10^{-1}$	$8,88 \cdot 10^{-2}$	$-3,56 \cdot 10^{-1}$	$1,17 \cdot 10^1$	$2,18 \cdot 10^{-2}$

Les centres des clusters dans l'espace des régresseurs en utilisant l'algorithme RoFER sont indiqués dans le Tableau 5.12 :

Tableau 5.12. Centres des clusters du modèle flou TS pour $S_x(k+1)$ - algorithme RoFER

règle	$X(k)$	$S_x(k)$	$S_e(k)$	$u_1(k)$	$u_2(k)$
$R_1^{S_x(k)}$	$3,11 \cdot 10^{-1}$	$6,20 \cdot 10^{-1}$	$1,54 \cdot 10^{-1}$	$1,97 \cdot 10^{-2}$	$5,76 \cdot 10^{-3}$
$R_2^{S_x(k)}$	$4,12 \cdot 10^{-1}$	$3,23 \cdot 10^{-1}$	$4,50 \cdot 10^{-2}$	$2,23 \cdot 10^{-2}$	$2,70 \cdot 10^{-3}$
$R_3^{S_x(k)}$	$4,33 \cdot 10^{-1}$	$6,78 \cdot 10^{-1}$	$9,97 \cdot 10^{-2}$	$1,78 \cdot 10^{-2}$	$4,88 \cdot 10^{-3}$
$R_4^{S_x(k)}$	$4,93 \cdot 10^{-1}$	$2,47 \cdot 10^{-1}$	$8,01 \cdot 10^{-3}$	$1,94 \cdot 10^{-2}$	$4,90 \cdot 10^{-4}$

La forme des règles pour la sortie concentration de substrat énergétique $S_e(k)$ tant en utilisant l'algorithme GK qu'en utilisant l'algorithme RoFER, est donnée par les expressions suivantes :

$R_1^{S_e(k)}$:

Si $X(k)$ est A_{11} **et** $S_x(k)$ est A_{12} **et** $S_e(k)$ est A_{13} **et** $u_1(k)$ est A_{14} **et** $u_2(k)$ est A_{15}
Alors $X(k+1) = a_{11}^{S_e(k)} X(k) + a_{12}^{S_e(k)} S_x(k) + a_{13}^{S_e(k)} S_e(k) + a_{14}^{S_e(k)} u_1(k) + a_{15}^{S_e(k)} u_2(k) + d_1^{S_e(k)}$

$R_2^{S_e(k)}$:

Si $X(k)$ est A_{21} **et** $S_x(k)$ est A_{22} **et** $S_e(k)$ est A_{23} **et** $u_1(k)$ est A_{24} **et** $u_2(k)$ est A_{25}
Alors $X(k+1) = a_{21}^{S_e(k)} X(k) + a_{22}^{S_e(k)} S_x(k) + a_{23}^{S_e(k)} S_e(k) + a_{24}^{S_e(k)} u_1(k) + a_{25}^{S_e(k)} u_2(k) + d_2^{S_e(k)}$

$R_3^{S_e(k)}$:

Si $X(k)$ est A_{31} **et** $S_x(k)$ est A_{32} **et** $S_e(k)$ est A_{33} **et** $u_1(k)$ est A_{34} **et** $u_2(k)$ est A_{35}
Alors $X(k+1) = a_{31}^{S_e(k)} X(k) + a_{32}^{S_e(k)} S_x(k) + a_{33}^{S_e(k)} S_e(k) + a_{34}^{S_e(k)} u_1(k) + a_{35}^{S_e(k)} u_2(k) + d_3^{S_e(k)}$

$R_4^{S_e(k)}$:

Si $X(k)$ est A_{41} **et** $S_x(k)$ est A_{42} **et** $S_e(k)$ est A_{43} **et** $u_1(k)$ est A_{44} **et** $u_2(k)$ est A_{45}
Alors $X(k+1) = a_{41}^{S_e(k)} X(k) + a_{42}^{S_e(k)} S_x(k) + a_{43}^{S_e(k)} S_e(k) + a_{44}^{S_e(k)} u_1(k) + a_{45}^{S_e(k)} u_2(k) + d_4^{S_e(k)}$

Les coefficients respectifs a_{il} du modèle de régression, avec $l = \{1, 2, \dots, 5\}$, en utilisant l'algorithme GK sont regroupés dans le Tableau 5.13 :

Tableau 5.13. Paramètres des conséquents du modèle flou TS pour $S_e(k+1)$ - algorithme GK

règle	$X(k)$	$S_x(k)$	$S_e(k)$	$u_1(k)$	$u_2(k)$	biais
$R_1^{S_e(k)}$	$-7,18 \cdot 10^{-1}$	$6,70 \cdot 10^{-2}$	$7,69 \cdot 10^{-1}$	$-1,09 \cdot 10^0$	$1,10 \cdot 10^1$	$7,81 \cdot 10^{-2}$
$R_2^{S_e(k)}$	$-2,16 \cdot 10^{-1}$	$-4,19 \cdot 10^{-4}$	$-8,27 \cdot 10^{-1}$	$-3,20 \cdot 10^{-1}$	$1,01 \cdot 10^1$	$3,70 \cdot 10^{-2}$
$R_3^{S_e(k)}$	$-9,11 \cdot 10^{-2}$	$6,38 \cdot 10^{-3}$	$4,45 \cdot 10^{-1}$	$-1,34 \cdot 10^{-1}$	$9,38 \cdot 10^0$	$4,86 \cdot 10^{-2}$
$R_4^{S_e(k)}$	$-8,10 \cdot 10^{-3}$	$7,20 \cdot 10^{-3}$	$6,97 \cdot 10^{-3}$	$-1,77 \cdot 10^{-2}$	$7,83 \cdot 10^0$	$5,54 \cdot 10^{-3}$

Les centres des clusters dans l'espace des régresseurs en utilisant l'algorithme GK sont indiqués dans le Tableau 5.14 :

Tableau 5.14. Centres des clusters du modèle flou TS pour $S_e(k+1)$ - algorithme GK

règle	$X(k)$	$S_x(k)$	$S_e(k)$	$u_1(k)$	$u_2(k)$
$R_1^{S_e(k)}$	$1,01 \cdot 10^{-1}$	$7,13 \cdot 10^{-1}$	$4,00 \cdot 10^{-1}$	$1,89 \cdot 10^{-2}$	$5,02 \cdot 10^{-3}$
$R_2^{S_e(k)}$	$3,13 \cdot 10^{-1}$	$5,69 \cdot 10^{-1}$	$1,23 \cdot 10^{-1}$	$1,85 \cdot 10^{-2}$	$4,89 \cdot 10^{-3}$
$R_3^{S_e(k)}$	$4,33 \cdot 10^{-1}$	$5,85 \cdot 10^{-1}$	$8,19 \cdot 10^{-2}$	$1,96 \cdot 10^{-2}$	$4,33 \cdot 10^{-3}$
$R_4^{S_e(k)}$	$6,80 \cdot 10^{-1}$	$3,05 \cdot 10^{-1}$	$1,41 \cdot 10^{-2}$	$1,69 \cdot 10^{-2}$	$1,35 \cdot 10^{-3}$

Les coefficients respectifs a_{il} du modèle de régression, avec $l = \{1, 2, \dots, 5\}$, en utilisant l'algorithme RoFER sont regroupés dans le Tableau 5.15 :

Tableau 5.15. Paramètres des conséquents du modèle flou TS pour $S_e(k+1)$ - algorithme RoFER

règle	$X(k)$	$S_x(k)$	$S_e(k)$	$u_1(k)$	$u_2(k)$	biais
$R_1^{S_e(k)}$	$-9,77 \cdot 10^{-2}$	$-5,72 \cdot 10^{-2}$	$9,70 \cdot 10^{-1}$	$-5,54 \cdot 10^{-1}$	$6,74 \cdot 10^0$	$3,87 \cdot 10^{-2}$
$R_2^{S_e(k)}$	$-1,09 \cdot 10^{-1}$	$-7,47 \cdot 10^{-4}$	$6,35 \cdot 10^{-1}$	$-3,14 \cdot 10^{-1}$	$8,82 \cdot 10^0$	$4,26 \cdot 10^{-2}$
$R_3^{S_e(k)}$	$-6,21 \cdot 10^{-2}$	$-1,31 \cdot 10^{-3}$	$5,28 \cdot 10^{-1}$	$-1,22 \cdot 10^{-1}$	$9,13 \cdot 10^0$	$2,74 \cdot 10^{-2}$
$R_4^{S_e(k)}$	$-6,30 \cdot 10^{-3}$	$-1,88 \cdot 10^{-3}$	$5,68 \cdot 10^{-1}$	$-8,78 \cdot 10^{-3}$	$5,10 \cdot 10^0$	$4,60 \cdot 10^{-3}$

Les centres des clusters dans l'espace des régresseurs en utilisant l'algorithme RoFER sont indiqués dans le Tableau 5.16 :

Tableau 5.16. Centres des clusters du modèle flou TS pour $S_e(k+1)$ - algorithme RoFER

règle	$X(k)$	$S_x(k)$	$S_e(k)$	$u_1(k)$	$u_2(k)$
$R_1^{S_e(k)}$	$1,62 \cdot 10^{-1}$	$5,39 \cdot 10^{-1}$	$2,37 \cdot 10^{-1}$	$2,26 \cdot 10^{-2}$	$4,90 \cdot 10^{-3}$
$R_2^{S_e(k)}$	$3,97 \cdot 10^{-1}$	$7,08 \cdot 10^{-1}$	$1,27 \cdot 10^{-1}$	$1,83 \cdot 10^{-2}$	$5,84 \cdot 10^{-3}$
$R_3^{S_e(k)}$	$3,98 \cdot 10^{-1}$	$4,55 \cdot 10^{-1}$	$5,63 \cdot 10^{-2}$	$2,06 \cdot 10^{-2}$	$3,03 \cdot 10^{-3}$
$R_4^{S_e(k)}$	$5,64 \cdot 10^{-1}$	$2,78 \cdot 10^{-1}$	$8,20 \cdot 10^{-3}$	$1,79 \cdot 10^{-2}$	$5,64 \cdot 10^{-4}$

Remarque

R 5.4 Bien que les résultats obtenus lors de la modélisation floue du bioprocédé sont encourageants, il faut dire que la génération des données d'identification non linéaire en parcourant toute la gamme des variables d'entrée reste une contrainte pour l'applicabilité de la technique sur des procédés biologiques industriels. En pratique on doit se contenter d'utiliser les données disponibles. Si elles ont été obtenues correctement, on trouvera un modèle plus performante.

5.3. Commande floue de type TS du bioprocédé basée sur le modèle flou

Dans le cas de notre bioréacteur pour le traitement des eaux résiduaires, lorsque l'objectif fixé est uniquement la régulation des variables de sortie, la synthèse d'un régulateur de type "boite noire" utilisant une caractérisation du comportement entrée/sortie du système suffit. La complexité des dynamiques internes de ce type de procédé ne permettent pas en général d'atteindre des objectifs de commande satisfaisants en considérant une relation linéaire entrée/sortie. L'expérience a montré que l'application de régulateurs classiques linéaires de type PID à ces systèmes complexes est limitée [DAH92a] [DAH92b]. Il est par conséquent nécessaire de s'intéresser à des techniques de commande plus sophistiquées tenant compte par exemple des caractéristiques non-linéaires et non-stationnaires de systèmes.

Dans le cadre de notre étude nous nous sommes intéressés à la modélisation floue TS du comportement non linéaire du bioprocédé. Ce type d'approche "boîte grise" permet la représentation du système dynamique comme un modèle à base de règles qui approxime la dynamique non linéaire comme une concaténation des sous modèles localement linéaires sous la forme d'auto-régression non-linéaire (NARX). Afin de tester la validité de la loi de commande, nous considérons dans notre étude en simulation que toutes les variables de sortie sont mesurables.

A partir du modèle flou obtenu dans la section précédente, nous allons réaliser la synthèse d'un contrôleur flou MIMO dans l'espace d'état du procédé, en utilisant la philosophie de commande du type compensation parallèle distribuée (PDC). Nous nous sommes intéressés en particulier, à la commande sous-optimale linéaire quadratique par retour d'état avec annulation des erreurs stationnaires (en utilisant des intégrateurs numériques), basée sur des modèles TS affines en temps discret (cf. section 4.2.1). Pour cela, nous ajoutons le même nombre d'intégrateurs numériques dans la boucle de commande que de variables de sortie à commander (dans notre cas 2). Nous sommes intéressés par la régulation des concentrations du substrat xénobiotique, $S_x(k)$, et de la biomasse, $X(k)$, par rapport à des consignes préspecifiées données respectivement par les expressions (5.4) et (5.5). Le nombre de clusters issus de la procédure de modélisation (dans ce cas $c = 4$) correspond au nombre de règles de la loi de commande.

Comme on l'a indiqué précédemment, au niveau de la modélisation on a décomposé le système MIMO en trois sous systèmes MISO, chacun correspondant à la modélisation dynamique pour chacune des sorties du bioprocédé. Cependant, pour réaliser la synthèse de la loi de commande, il est nécessaire d'avoir une base de règles "commune" entre les trois partitions obtenues dans l'espace des régresseurs. Dans un but de commande et en sachant que le contrôleur fonctionnera en boucle fermée, on peut faire une approximation en considérant que les clusters obtenus pour chaque sortie dans l'espace des régresseurs peuvent être fusionnés en quatre clusters. En effet, les projections sur les axes des variables de la partie antécédent (régresseurs) de ces différents clusters sont assez proches (voir par exemple les Figures 5.9, 5.10 et 5.11 ou bien les Figures 5.12, 5.13 et 5.14). La fusion des clusters est assez simple, dans la mesure où les fonctions d'appartenance (exponentielles par morceaux) correspondantes sont des fonctions paramétrées (cf. expression (3.4)). Le critère que nous avons utilisé pour faire la fusion des clusters consiste à calculer la moyenne pondérée pour chacun des quatre paramètres $\{w_l, w_r, c_l \text{ et } c_r\}$ qui définissent chacune des fonctions d'appartenance, en utilisant comme facteurs de poids les critères VAF obtenus respectivement pour les partitions des variables $X(k)$ et $S_x(k)$ ⁹⁷ (cf. Tableau 5.4). Afin de simplifier la notation, nous noterons ici d'une façon générale ces paramètres comme '*par*'. On obtient alors l'expression suivante qui définit les paramètres des fonctions d'appartenance dans la base de règles *commune*⁹⁸ (notés '*par_{com}*') à des propos de la commande :

$$par_{com} = \frac{VAF_{X(k)} \cdot par_{X(k)} + VAF_{S_x(k)} \cdot par_{S_x(k)}}{VAF_{X(k)} + VAF_{S_x(k)}} \quad (5.8)$$

⁹⁷ Puisque ce sont les variables d'intérêt au niveau de la régulation du bioprocédé.

⁹⁸ Partition "commune" dans l'espace des antécédents (régresseurs) pour les trois sorties du bioprocédé.

Compte tenu de l'unification de la base de règles, les ensembles flous correspondants au niveau de l'écriture des règles seront notés $\bar{A}_{il}$ au lieu de A_{il} , avec $1 \leq i \leq 4$ et $1 \leq l \leq 5$ (rappelons que le nombre de clusters est $c = 4$ et le nombre de régresseurs est 5 : $\{X(k), S_x(k), S_e(k), u_1(k), u_2(k)\}$). Pour faire la commande du bioprocédé, nous avons retenu la base de règles commune issue des partitions obtenues avec l'algorithme RoFER. La procédure est similaire pour l'algorithme GK.

Une fois obtenue cette base de règles commune au niveau des partitions correspondantes aux trois sorties du bioprocédé, il est alors nécessaire de porter les expressions des règles sous la forme d'état, afin d'appliquer les formulations de la commande pour des systèmes MIMO développées dans le chapitre 4. Précédemment, on a vu que pour chaque sortie du système MISO il y a un ensemble de quatre règles qui décrivent leur dynamique associée. A partir de telles règles, il est simple de montrer qu'on peut passer à la représentation d'état désirée en regroupant les règles de la façon appropriée. Il suffit de conformer chaque i -ème règle de la représentation d'état à partir des i -èmes règles de chacune des sorties. Ainsi par exemple, la première règle de la représentation d'état MIMO est composé des premières règles de chacune des sorties du sous systèmes MISO, que nous reprenons ici à titre de rappel, de la façon suivante :

$R_1^{X(k)} :$

Si $X(k)$ est A_{11} **et** $S_x(k)$ est A_{12} **et** $S_e(k)$ est A_{13} **et** $u_1(k)$ est A_{14} **et** $u_2(k)$ est A_{15}
Alors $X(k+1) = a_{11}^{X(k)} X(k) + a_{12}^{X(k)} S_x(k) + a_{13}^{X(k)} S_e(k) + a_{14}^{X(k)} u_1(k) + a_{15}^{X(k)} u_2(k) + d_1^{X(k)}$

$R_1^{S_x(k)} :$

Si $X(k)$ est A_{11} **et** $S_x(k)$ est A_{12} **et** $S_e(k)$ est A_{13} **et** $u_1(k)$ est A_{14} **et** $u_2(k)$ est A_{15}
Alors $X(k+1) = a_{11}^{S_x(k)} X(k) + a_{12}^{S_x(k)} S_x(k) + a_{13}^{S_x(k)} S_e(k) + a_{14}^{S_x(k)} u_1(k) + a_{15}^{S_x(k)} u_2(k) + d_1^{S_x(k)}$

$R_1^{S_e(k)} :$

Si $X(k)$ est A_{11} **et** $S_x(k)$ est A_{12} **et** $S_e(k)$ est A_{13} **et** $u_1(k)$ est A_{14} **et** $u_2(k)$ est A_{15}
Alors $X(k+1) = a_{11}^{S_e(k)} X(k) + a_{12}^{S_e(k)} S_x(k) + a_{13}^{S_e(k)} S_e(k) + a_{14}^{S_e(k)} u_1(k) + a_{15}^{S_e(k)} u_2(k) + d_1^{S_e(k)}$

La première des quatre règles de la représentation d'état MIMO est alors exprimée sous la forme matricielle donnée par l'expression suivante :

$R_1^{MIMO} :$

Si $X(k)$ est $\bar{A}_{11}$ **et** $S_x(k)$ est $\bar{A}_{12}$ **et** $S_e(k)$ est $\bar{A}_{13}$ **et** $u_1(k)$ est $\bar{A}_{14}$ **et** $u_2(k)$ est $\bar{A}_{15}$
Alors
$$\begin{bmatrix} X(k+1) \\ S_x(k+1) \\ S_e(k+1) \end{bmatrix} = \begin{bmatrix} a_{11}^{X(k)} & a_{12}^{X(k)} & a_{13}^{X(k)} \\ a_{11}^{S_x(k)} & a_{12}^{S_x(k)} & a_{13}^{S_x(k)} \\ a_{11}^{S_e(k)} & a_{12}^{S_e(k)} & a_{13}^{S_e(k)} \end{bmatrix} \cdot \begin{bmatrix} X(k) \\ S_x(k) \\ S_e(k) \end{bmatrix} + \begin{bmatrix} a_{14}^{X(k)} & a_{15}^{X(k)} \\ a_{14}^{S_x(k)} & a_{15}^{S_x(k)} \\ a_{14}^{S_e(k)} & a_{15}^{S_e(k)} \end{bmatrix} \cdot \begin{bmatrix} u_1(k) \\ u_2(k) \end{bmatrix} + \begin{bmatrix} d_1^{X(k)} \\ d_1^{S_x(k)} \\ d_1^{S_e(k)} \end{bmatrix}$$

De manière similaire, on peut obtenir les trois autres règles de la représentation d'état MIMO. Elles sont données par les expressions suivantes :

R_2^{MIMO} :

Si $X(k)$ est $\bar{A}_{21}$ **et** $S_x(k)$ est $\bar{A}_{22}$ **et** $S_e(k)$ est $\bar{A}_{23}$ **et** $u_1(k)$ est $\bar{A}_{24}$ **et** $u_2(k)$ est $\bar{A}_{25}$

$$\textbf{Alors} \begin{bmatrix} X(k+1) \\ S_x(k+1) \\ S_e(k+1) \end{bmatrix} = \begin{bmatrix} a_{21}^{X(k)} & a_{22}^{X(k)} & a_{23}^{X(k)} \\ a_{21}^{S_x(k)} & a_{22}^{S_x(k)} & a_{23}^{S_x(k)} \\ a_{21}^{S_e(k)} & a_{22}^{S_e(k)} & a_{23}^{S_e(k)} \end{bmatrix} \cdot \begin{bmatrix} X(k) \\ S_x(k) \\ S_e(k) \end{bmatrix} + \begin{bmatrix} a_{24}^{X(k)} & a_{25}^{X(k)} \\ a_{24}^{S_x(k)} & a_{25}^{S_x(k)} \\ a_{24}^{S_e(k)} & a_{25}^{S_e(k)} \end{bmatrix} \cdot \begin{bmatrix} u_1(k) \\ u_2(k) \end{bmatrix} + \begin{bmatrix} d_2^{X(k)} \\ d_2^{S_x(k)} \\ d_2^{S_e(k)} \end{bmatrix}$$

R_3^{MIMO} :

Si $X(k)$ est $\bar{A}_{31}$ **et** $S_x(k)$ est $\bar{A}_{32}$ **et** $S_e(k)$ est $\bar{A}_{33}$ **et** $u_1(k)$ est $\bar{A}_{34}$ **et** $u_2(k)$ est $\bar{A}_{35}$

$$\textbf{Alors} \begin{bmatrix} X(k+1) \\ S_x(k+1) \\ S_e(k+1) \end{bmatrix} = \begin{bmatrix} a_{31}^{X(k)} & a_{32}^{X(k)} & a_{33}^{X(k)} \\ a_{31}^{S_x(k)} & a_{32}^{S_x(k)} & a_{33}^{S_x(k)} \\ a_{31}^{S_e(k)} & a_{32}^{S_e(k)} & a_{33}^{S_e(k)} \end{bmatrix} \cdot \begin{bmatrix} X(k) \\ S_x(k) \\ S_e(k) \end{bmatrix} + \begin{bmatrix} a_{34}^{X(k)} & a_{35}^{X(k)} \\ a_{34}^{S_x(k)} & a_{35}^{S_x(k)} \\ a_{34}^{S_e(k)} & a_{35}^{S_e(k)} \end{bmatrix} \cdot \begin{bmatrix} u_1(k) \\ u_2(k) \end{bmatrix} + \begin{bmatrix} d_3^{X(k)} \\ d_3^{S_x(k)} \\ d_3^{S_e(k)} \end{bmatrix}$$

R_4^{MIMO} :

Si $X(k)$ est $\bar{A}_{41}$ **et** $S_x(k)$ est $\bar{A}_{42}$ **et** $S_e(k)$ est $\bar{A}_{43}$ **et** $u_1(k)$ est $\bar{A}_{44}$ **et** $u_2(k)$ est $\bar{A}_{45}$

$$\textbf{Alors} \begin{bmatrix} X(k+1) \\ S_x(k+1) \\ S_e(k+1) \end{bmatrix} = \begin{bmatrix} a_{41}^{X(k)} & a_{42}^{X(k)} & a_{43}^{X(k)} \\ a_{41}^{S_x(k)} & a_{42}^{S_x(k)} & a_{43}^{S_x(k)} \\ a_{41}^{S_e(k)} & a_{42}^{S_e(k)} & a_{43}^{S_e(k)} \end{bmatrix} \cdot \begin{bmatrix} X(k) \\ S_x(k) \\ S_e(k) \end{bmatrix} + \begin{bmatrix} a_{44}^{X(k)} & a_{45}^{X(k)} \\ a_{44}^{S_x(k)} & a_{45}^{S_x(k)} \\ a_{44}^{S_e(k)} & a_{45}^{S_e(k)} \end{bmatrix} \cdot \begin{bmatrix} u_1(k) \\ u_2(k) \end{bmatrix} + \begin{bmatrix} d_4^{X(k)} \\ d_4^{S_x(k)} \\ d_4^{S_e(k)} \end{bmatrix}$$

A partir du modèle flou MIMO dans l'espace d'état donné par les règles R_i^{MIMO} , avec $1 \leq i \leq c$, et en utilisant la technique de la commande PDC, on obtient donc une structure de commande qui consiste à utiliser la même partie antécédent des règles que celle du système flou de type affine TS, et, dans la partie conclusion, des commandes locales prenant la forme donnée par l'expression (4.67)⁹⁹, que nous reprenons par la suite :

$$u_k^* = -(\tilde{K}_i \tilde{x}_k + \tilde{N}_i \tilde{d}) = \tilde{L}_i \bar{r}_k - (\bar{K}_i x_k + \bar{K}_{v_i} v_k + \bar{N}_i d_i) \quad (5.9)$$

La loi de commande optimale obtenue pour les sous modèles linéaires affines [GRI06] consiste en un retour d'état $-\bar{K}_i x_k$, une compensation du terme de biais $-\bar{N}_i d_i$, un terme d'anticipation lié à la consigne $\tilde{L}_i \bar{r}_k$ et une action intégrale $-\bar{K}_{v_i} v_k$ agissant sur l'erreur entre consignes et mesures. Par analogie avec les régulateurs classiques comportant une action « proportionnelle » et une action « intégrale », le résultat obtenu peut être vue comme une loi globale de commande non linéaire multidimensionnelle du type P.I. avec compensation du terme de biais d_i . Les matrices $\tilde{K}_i \in \mathbb{R}^{2 \times 5}$ et $\tilde{N}_i \in \mathbb{R}^{2 \times 5}$ ont été décomposées en $\tilde{K}_i = [\bar{K}_i \ \bar{K}_{v_i}]$ où $\bar{K}_i \in \mathbb{R}^{2 \times 3}$ et $\bar{K}_{v_i} \in \mathbb{R}^{2 \times 2}$, et en $\tilde{N}_i = [\bar{N}_i \ -\tilde{L}_i]$ où $\bar{N}_i \in \mathbb{R}^{3 \times 2}$ et $\tilde{L}_i \in \mathbb{R}^{2 \times 2}$.

Dans ce contexte, chaque sous modèle du système flou de type TS est stabilisé localement par une loi linéaire. La loi globale de commande est non linéaire. La forme de chacune des quatre règles ($c = 4$) du système de commande est donnée par l'expression suivante :

⁹⁹ Pour le cas du bioréacteur, le vecteur d'état augmenté, représenté par $\tilde{x}_k \in \mathbb{R}^5$, est composé par le vecteur d'état des variables de sortie du bioprocédé $x_k = [X(k) \ S_x(k) \ S_e(k)]^T$, et par les sorties des 2 intégrateurs.

$R_i^{commande}$:

Si $X(k)$ est $\bar{A}_{i1}$ et $S_x(k)$ est $\bar{A}_{i2}$ et $S_e(k)$ est $\bar{A}_{i3}$ et $u_1(k)$ est $\bar{A}_{i4}$ et $u_2(k)$ est $\bar{A}_{i5}$

Alors

$$\begin{bmatrix} u_1^*(k) \\ u_2^*(k) \end{bmatrix} = \underbrace{\begin{bmatrix} l_{11} & l_{12} \\ l_{21} & l_{22} \end{bmatrix}}_{\tilde{L}_i} \cdot \begin{bmatrix} \bar{r}_1(k) \\ \bar{r}_2(k) \end{bmatrix} - \left(\underbrace{\begin{bmatrix} k_{11} & k_{12} & k_{13} \\ k_{21} & k_{22} & k_{23} \end{bmatrix}}_{\tilde{K}_i} \cdot \begin{bmatrix} X(k) \\ S_x(k) \\ S_e(k) \end{bmatrix} + \underbrace{\begin{bmatrix} k_{v11} & k_{v12} \\ k_{v21} & k_{v22} \end{bmatrix}}_{\tilde{K}_{v_i}} \cdot \begin{bmatrix} v_1(k) \\ v_2(k) \end{bmatrix} + \underbrace{\begin{bmatrix} n_{11} & n_{12} & n_{13} \\ n_{21} & n_{22} & n_{23} \end{bmatrix}}_{\tilde{N}_i} \cdot \begin{bmatrix} d_{i1} \\ d_{i2} \\ d_{i3} \end{bmatrix} \right)$$

La réalisation du régulateur flou PDC avec action intégrale pour le bioprocédé, consiste à déterminer pour les règles $i=1, \dots, 4$ du modèle, les gains $\tilde{K}_i$ et $\tilde{N}_i$ dans les parties conclusions selon l'expression (5.9).

La sortie finale du régulateur flou de type PDC, qui en fait est une loi de commande non linéaire, est inférée comme suit :

$$u^*(k) = -\sum_{i=1}^4 h_i(z(k)) \left(\tilde{K}_i \tilde{x}(k) + \tilde{N}_i \tilde{d}_i \right), \text{ avec } h_i(z(k)) = \frac{w_i(z(k))}{\sum_{i=1}^4 w_i(z(k))} \quad (5.10)$$

avec $z(k) = [X(k), S_x(k), S_e(k), u_1(k), u_2(k)]$, vérifiant pour tout k , les propriétés de somme convexe et d'activation des fonctions d'appartenance (cf. expression (4.48)).

A partir du modèle flou MIMO *affine* Takagi-Sugeno discret obtenu, en choisissant les matrices de pondération Q et R selon les expressions (4.73) et (4.74) et en appliquant la procédure de synthèse décrite dans la section 4.2.3 (expressions (4.68) à (4.70)), on obtient alors la solution suivante pour les matrices $\tilde{K}_i$ et $\tilde{N}_i$ de la loi de commande [GRI06] :

Règle 1 :

$$\tilde{K}_1 = \begin{bmatrix} -0,4414 & -0,1328 & 0,1173 & 0,0120 & 0,0075 \\ -0,0464 & 0,0397 & 0,0134 & 0,0010 & -0,0070 \end{bmatrix} \quad (5.11)$$

$$\tilde{N}_1 = \begin{bmatrix} -5,9468 & -0,3248 & 0,7883 & 0,3057 & -0,0168 \\ -0,6544 & 0,0420 & 0,1050 & 0,0301 & -0,0329 \end{bmatrix} \quad (5.12)$$

Règle 2 :

$$\tilde{K}_2 = \begin{bmatrix} -0,2319 & -0,0200 & 0,0241 & 0,0120 & 0,0020 \\ -0,0114 & 0,0490 & -0,0040 & 0,0003 & -0,0074 \end{bmatrix} \quad (5.13)$$

$$\tilde{N}_2 = \begin{bmatrix} -1,1890 & 0,0533 & 0,0798 & 0,1268 & -0,0077 \\ -0,0558 & 0,0861 & -0,0009 & 0,0046 & -0,0295 \end{bmatrix} \quad (5.14)$$

Règle 3 :

$$\tilde{K}_3 = \begin{bmatrix} -0,2358 & -0,0560 & 0,0126 & 0,0114 & 0,0080 \\ -0,0320 & 0,0460 & 0,0033 & 0,0011 & -0,0069 \end{bmatrix} \quad (5.15)$$

$$\tilde{N}_3 = \begin{bmatrix} -1,4694 & 0,0918 & 0,0733 & 0,1361 & -0,0014 \\ -0,1949 & 0,0976 & 0,0109 & 0,0154 & -0,0307 \end{bmatrix} \quad (5.16)$$

Règle 4 :

$$\tilde{K}_4 = \begin{bmatrix} -0,1979 & -0,0188 & 0,0210 & 0,0119 & -0,0004 \\ -0,0056 & 0,0484 & 0,0059 & 0,0001 & -0,0073 \end{bmatrix} \quad (5.17)$$

$$\tilde{N}_4 = \begin{bmatrix} -0,8747 & 0,0089 & 0,0791 & 0,1083 & -0,0139 \\ -0,0261 & 0,0838 & 0,0072 & 0,0018 & -0,0294 \end{bmatrix} \quad (5.18)$$

Les résultats obtenus en appliquant la commande PDC sous-optimale linéaire quadratique par retour d'état pour le modèle affine MIMO TS en temps discret, afin de réguler les concentrations de biomasse et de substrat xénobiotique du bioprocédé en suivant les consignes établies dans le scénario défini par les expressions (5.4) et (5.5), sont illustrés sur la Figure 5.15 :

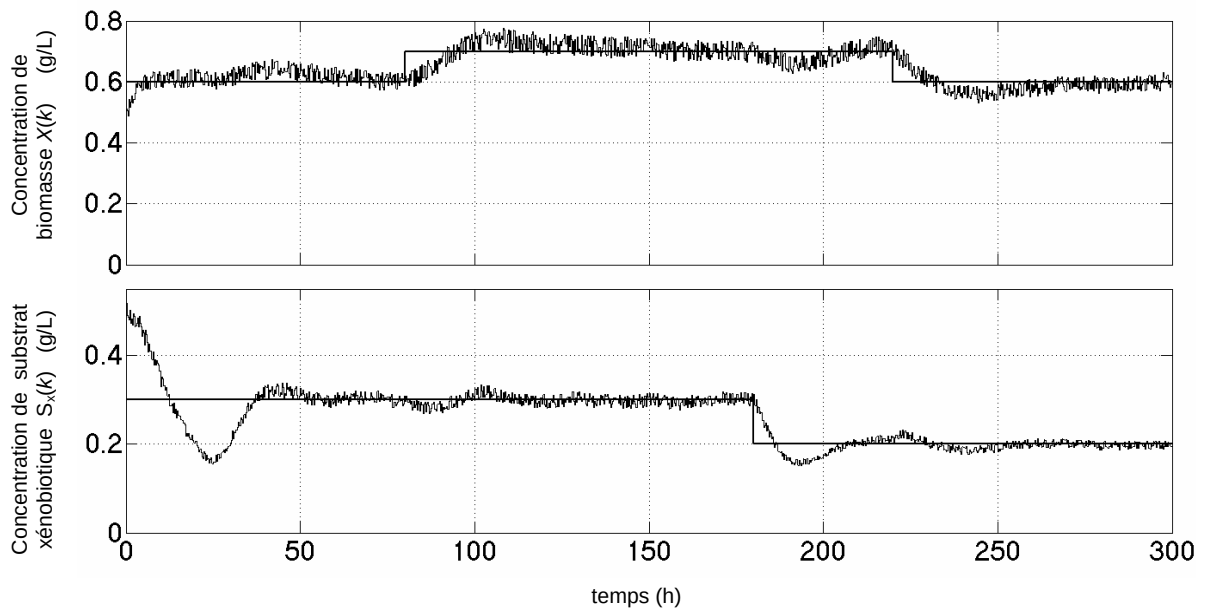


Figure 5.15. Evolution des concentrations $X(k)$ et $S_x(k)$ suivant le scénario de contrôle désiré.

La Figure 5.16 montre l'évolution correspondante aux variables de commande : les taux de dilution $u_1(k)$ (eau d'alimentation) et $u_2(k)$ (substrats d'alimentation).

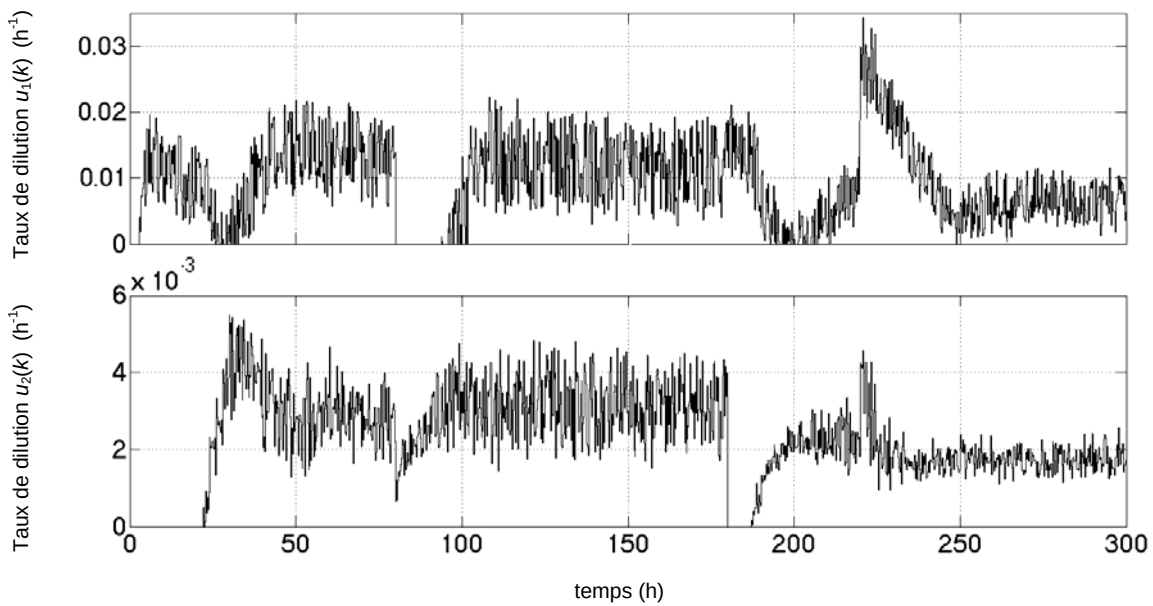


Figure 5.16. Evolution des variables de commande $u_1(k)$ et $u_2(k)$ du bioprocédé.

Un critère de performance est utilisé pour évaluer la qualité des performances du système en boucle fermée. Il correspond à l'erreur de contrôle et à la variabilité des entrées. Soit les erreurs de contrôle relatives pour les concentrations en biomasse et en substrat xénobiotique donnés par les expressions :

$$e_X(k) = \left(\frac{X(k) - X^*(k)}{X^*(k)} \right) \quad \text{et} \quad e_{S_x}(k) = \left(\frac{S_x(k) - S_x^*(k)}{S_x^*(k)} \right) \quad (5.19)$$

avec $X(k)$ et $S_x(k)$ les signaux pris en ligne sans le bruit du aux capteurs. Leur variabilité est définie par :

$$V_X = \frac{1}{N} \sum_{k=1}^N e_X(k)^2, \quad V_{S_x} = \frac{1}{N} \sum_{k=1}^N e_{S_x}(k)^2, \quad V_{u_1} = \frac{1}{N} \sum_{k=1}^N u_1(k)^2, \quad V_{u_2} = \frac{1}{N} \sum_{k=1}^N u_2(k)^2 \quad (5.20)$$

où N est le nombre de données recueillies. Les résultats de performance obtenus avec le contrôleur flou PDC sous optimal pour le bioprocédé sont regroupés dans le Tableau 5.17 :

Paramètre	Description	Valeur
V_X	Variabilité de la commande de la concentration de biomasse $X(k)$	$2,10 \times 10^{-3}$
V_{S_x}	Variabilité de la commande de la concentration du substrat xénobiotique $S_x(k)$	$2,04 \times 10^{-2}$
V_{u_1}	Variabilité du signal de commande $u_1(k)$: taux de dilution de l'eau d'alimentation	$1,34 \times 10^{-4}$
V_{u_2}	Variabilité du signal de commande $u_2(k)$: taux de dilution des substrats d'alimentation	$6,81 \times 10^{-6}$

Tableau 5.17. Critères des performances en boucle fermée du système de commande flou de type PDC basé sur le modèle dynamique flou affine de type Takagi-Sugeno.

Le contrôleur proposé présente de bons résultats pour la régulation autour des consignes imposées. Cependant, un réglage fin des matrices de pondération Q et R du critère quadratique pourrait améliorer les performances. Les résultats obtenus sont similaires à ceux reportés dans le cadre du projet européen FAMIMO¹⁰⁰ [ROU99] pour d'autres techniques de commande basées sur le modèle flou de type TS, telles que la commande prédictive ou adaptative. Nous croyons que ces résultats sont intéressants et encourageants pour la suite, compte tenu que pour des résultats similaires, la synthèse proposée du contrôleur flou basé sur des sous modèles linéaires affines locaux est méthodique et simple, ainsi que l'obtention du modèle flou Takagi-Sugeno à partir des données entrée-sortie du bioprocédé de traitement des eaux usées considéré.

Comparativement à d'autres approches de modélisations utilisées dans le domaine des bioprocédés, comme par exemple celle basée sur des bilans de matière, la modélisation floue de type Takagi-Sugeno ne faisant pas intervenir une connaissance préalable des équations de la dynamique du système (seulement des données), peut dans certains cas être un avantage. L'approche de modélisation floue de type Takagi-Sugeno, orientée vers la représentation du comportement entrée-sortie du procédé, peut constituer alors une alternative de modélisation performante quand on dispose de telles données. Néanmoins, du à la forme des règles du modèle flou TS qui décrivent la dynamique du procédé sous la forme de sous modèles autorégressifs de type ARX, les paramètres de tels sous modèles sont des coefficients qui n'ont pas la même signification physique que les paramètres des modèles de bilan de matières. Nous considérons que des applications focalisées sur le comportement entrée-sortie des bioprocédés, comme dans le cas de la commande par ordinateur, ne constitue pas un inconvénient. Nous pensons que la décomposition de la dynamique non linéaire du système par des sous modèles linéaires locaux peut être une alternative pour aborder des problèmes de modélisation et commande des systèmes non linéaires de façon relativement simple.

5.4. Conclusions

Ce chapitre a été consacré à l'application (en simulation) des méthodes d'identification et de commande floues de type Takagi-Sugeno sur un bioréacteur aérobique multivariable en mode continu pour le traitement des eaux usées issues d'une industrie papetière.

Après avoir reprécisé le modèle obtenu à partir des bilans de matière, nous avons développé un modèle flou dynamique de type Takagi-Sugeno pour le bioréacteur basé sur des données expérimentales entrée-sortie en appliquant les techniques de clustering flou abordées dans le Chapitre 3. En particulier, nous avons utilisé les algorithmes GK et RoFER qui nous ont semblé les plus représentatives et les mieux adaptés pour la tâche de modélisation. Nous avons obtenu de bons résultats avec ces deux algorithmes (cf. Tableaux 5.3 et 5.4), avec une légère supériorité de performance numérique globale de l'approximation (selon le critère VAF de l'ordre de 95%) en utilisant l'algorithme RoFER, en présence d'un bruit blanc de faible amplitude ajouté aux données d'identification. Dans le cadre de la modélisation à des fins de commande, les principaux avantages de l'utilisation de ce dernier algorithme pourraient être, malgré l'effort du calcul associé, la détermination "automatique" du nombre des clusters (régions linéaires), basée sur la définition du seuil de cardinalité robuste, ainsi que l'obtention des clusters parallèles aux axes des régresseurs (meilleure interprétabilité pour l'obtention des règles du modèle flou sous la forme décomposée).

¹⁰⁰ Waste-water Benchmark, in Fuzzy Algorithms for the Control of Multiple-Input, Multiple Output Processes (FAMIMO) project, funded by the European Commission (Esprit LTR 21911).

Le système dynamique MIMO avec deux entrées (les taux de dilution de l'eau d'alimentation $u_1(k)$ et des substrats d'alimentation $u_2(k)$) et trois sorties (les concentrations de biomasse $X(k)$ et des substrats xénobiotique $S_x(k)$ et énergétique $S_e(k)$) a été ainsi représenté comme trois sous systèmes de type MISO, représentés chacun par un modèle flou affine TS. Ces modèles flous de type TS, pour lesquels le nombre de règles correspond au nombre de clusters (dans notre cas quatre), approximent la dynamique non linéaire comme une concaténation des sous modèles localement linéaires sous la forme d'auto-régression non linéaire (NARX). Nous avons testé la validité des modèles obtenus dans l'espace des données de l'identification en obtenant des résultats remarquables, en particulier pour le mode d'opération continu du bioréacteur. A titre illustratif, nous avons aussi conduit des expériences de modélisation en incluant des données en mode batch, qui ont montré la limitation du modèle pour représenter ce mode de culture. Cela est dû à la forme conjonctive des règles dans lesquelles interviennent comme régresseurs $u_1(k)$ et $u_2(k)$, qui suggère l'utilisation d'un modèle spécifique pour chacun des modes d'opération, ou bien, l'utilisation de plusieurs modèles dynamiques (modèle hybride) pour des applications en biotechnologie qui demandent l'opération du bioprocédé en différents modes.

Nous avons également proposé une base de règles commune (partition commune dans l'espace des antécédents (régresseurs) pour les trois sorties du bioprocédé), conduisant à un modèle MIMO autorégressif dans l'espace d'état, afin de pouvoir réaliser la synthèse de la loi de commande proposée dans le Chapitre 4. En sachant que le contrôleur fonctionnera en boucle fermée, cette approximation dans un but de commande peut être valide, comme effectivement nous l'avons constaté lors de nos expériences de simulation. En appliquant la philosophie de commande de type PDC basée sur le modèle flou dynamique TS, nous avons fait la synthèse de la loi de commande la mieux adaptée au modèle du bioprocédé : une commande sous-optimale linéaire quadratique par retour d'état avec annulation des erreurs statiques (en utilisant des intégrateurs numériques), afin de réguler les concentrations de biomasse et de substrat xénobiotique du bioprocédé.

Le contrôleur proposé, qui peut être vu comme une sorte de contrôleur de type PI multidimensionnel avec retour d'état, présente de bons résultats pour la régulation autour des consignes imposées (cf. Figure 5.15 et Tableau 5.17). Cependant, un réglage fin des matrices de pondération Q et R du critère quadratique pourrait améliorer les performances. Nous pensons que les résultats obtenus sont intéressants, en effet la synthèse proposée du contrôleur flou basé sur des sous modèles linéaires *affines* locaux est méthodique et simple, ainsi que l'obtention du modèle flou Takagi-Sugeno à partir des données entrée-sortie du bioprocédé de traitement des eaux usées considéré.

Conclusions et perspectives

Les travaux présentés dans ce mémoire s'articulent autour des thèmes principaux de la modélisation et de la commande floue de type Takagi-Sugeno (TS) pour une application à une station d'épuration des effluents issus d'une industrie papetière. Ces travaux sont divisés en deux parties, une première concernant le traitement des eaux résiduaires et la seconde, sur la modélisation floue à partir des données et la commande floue de type PDC (compensation parallèle distribuée) avec application au bioprocédé de traitement des eaux usées.

Dans la première partie nous montrons d'abord comment les procédés biologiques sont utilisés pour restaurer la qualité de l'eau en dégradant les matières organiques. Ce sont des procédés complexes, car ils font intervenir des organismes vivants. Du point de vue de l'Automatique, de manière générale, ces procédés sont caractérisés par des systèmes multivariables, non linéaires et non stationnaires. Nous avons été confrontés aux deux premiers aspects pour le cas du procédé envisagé. Ainsi, dans la première partie de nos travaux, nous avons développé des aspects liés à la modélisation dynamique des bioprocédés en nous focalisant sur l'approche des bilans de matière des composants principaux pour caractériser la dynamique de tels procédés. Nous avons également positionné la problématique générale de l'Automatique des bioprocédés, en remarquant deux défis majeurs, qui sont le manque d'une instrumentation de mesure en ligne permettant de connaître leur fonctionnement interne et la difficulté liée à la modélisation de la dynamique (dans une perspective de commande).

Le procédé considéré qui est un bioréacteur aérobie multivariable en mode continu pour le traitement des eaux usées issues d'une industrie papetière a été la base d'une première étude de modélisation en utilisant l'approche de bilans de matières. Ce procédé est alimenté par un mélange bi-polluant (substrats énergétique et xénobiotique), dans lequel les micro-organismes exhibent des interactions telles que la compétition et l'activation/inhibition compétitive. Le modèle que nous avons retenu, développé sous Matlab®, est une version améliorée de celui utilisé lors du projet européen FAMIMO, pour lequel nous avons ajouté un terme de mortalité endogène bactérienne. Ce modèle de connaissance nous a servi de base pour comparer en simulation la performance d'un modèle flou de type Takagi-Sugeno pour représenter le comportement dynamique non-linéaire entrée-sortie du système.

Dans la deuxième partie, nous avons présenté d'abord une approche structurale de la modélisation et de l'identification floues en nous focalisant particulièrement sur le modèle de type TS à partir des données entrée-sortie. En effet, au cours des dix dernières années la discipline a évolué d'une façon graduelle, vers une utilisation pratiquement exclusive des systèmes flous dans lesquels le conséquent des règles utilise des variables numériques sous la forme de fonctions (modèle de type TS) plutôt que des variables linguistiques (modèle de type Mamdani). Ces modèles conservent les caractéristiques de transparence et d'interprétabilité linguistique qui distinguent les modèles flous d'autres approches similaires de type boîte noire. Dans le cas de l'identification floue des systèmes, le formalisme Takagi-Sugeno est mieux adapté à une démarche plus systématique pour la construction de modèles non linéaires

multivariables, grâce à leur bonne capacité d'interpolation numérique et d'"apprentissage" à partir de données. En utilisant ce formalisme, nous représentons le comportement non linéaire d'un système par une composition de règles du type « *Si-Alors* », concaténant un ensemble de sous-modèles localement linéaires sous la forme d'auto-régression non-linéaire (NARX). Les différentes zones d'opération sont définies par les fonctions d'appartenance dans les antécédents et les sous-modèles linéaires sont associés aux conséquents des règles.

Pour la construction de tels modèles nous abordons d'une part, différents algorithmes de clustering tels que les algorithmes FCM (clusters hypersphériques), GK (clusters hyperellipsoïdaux) et FCRM (clusters hyperplanaires), ainsi que les algorithmes baptisés FER¹⁰¹ et RoFER¹⁰², caractérisés par des prototypes mixtes. L'algorithme de *régression floue ellipsoïdal* (FER) combine les mesures de distance des algorithmes GK des axes parallèles (GKP) et FCRM. L'algorithme RoFER correspond à sa version *robuste* en présence de bruit et de points aberrants. Ce dernier est caractérisé par l'utilisation d'une méthode d'"agglomération compétitive" qui à partir d'un nombre surestimé de clusters à l'initialisation et la définition d'un certain seuil de cardinalité robuste, permet de trouver "automatiquement" un nombre convenable de clusters pour la partition. L'algorithme fournit aussi simultanément les paramètres des conséquents des règles du modèle floue *affine* TS. Nous avons présenté quelques exemples illustratifs d'application qui ont montré l'applicabilité des algorithmes de clustering comme GK et RoFER, avec une meilleure qualité de l'approximation *locale* du système pour ce dernier en présence de bruits et de points aberrants, malgré un temps de calcul plus important pour la convergence (de l'ordre de 15 fois plus élevé). De plus, dans le cadre de nos recherches nous avons développé FMIDgraph, une boîte à outils graphique sous MATLAB® pour l'identification floue de type TS des systèmes non linéaires, qui intègre dans un environnement graphique facile à utiliser, plusieurs éléments et fonctionnalités (e.g. visualisation 2D et 3D) concernant la méthodologie de modélisation abordée.

Comparativement à d'autres approches de modélisation utilisées dans le domaine des bioprocédés, la modélisation floue de type Takagi-Sugeno ne fait pas intervenir une connaissance préalable des équations de la dynamique du système (seulement des données). Nous pensons que la décomposition de la dynamique non linéaire du système par des sous modèles localement linéaires peut être une alternative pour aborder des problèmes de modélisation et de commande des systèmes non linéaires de façon relativement simple.

Après la modélisation, nous avons abordé la synthèse de la commande floue TS basée sur un modèle, en utilisant la philosophie de commande du type compensation parallèle distribuée (PDC), sous la forme de LMI et de synthèse quadratique. Nous nous sommes intéressés, en particulier, à la commande sous-optimale linéaire quadratique par retour d'état basée sur un modèle *affine* TS en temps discret. Pour cela, nous proposons trois lois de commande qui tiennent compte du terme indépendant caractéristique du modèle affine. Les lois de commande sont de type régulateur, régulateur avec compensation de gain statique et régulateur avec action intégrale, ce qui permet l'annulation des erreurs statiques.

Enfin, nous avons considéré l'application de ces techniques d'identification et de commande floues TS sur le modèle du bioréacteur, en particulier pour le mode d'opération continu. Nous avons présenté la démarche utilisée pour construire le modèle du procédé et la

¹⁰¹ FER : *Fuzzy Ellipsoidal Regression*.

¹⁰² RoFER : *Robust Fuzzy Ellipsoidal Regression*.

loi de commande la mieux adaptée au bioprocédé : une commande sous-optimale linéaire quadratique par retour d'état avec annulation des erreurs statiques (en utilisant des intégrateurs numériques), afin de réguler les concentrations de biomasse et de substrat xénobiotique du bioprocédé. Afin de pouvoir réaliser la synthèse de la loi de commande développée, nous avons également proposé la conformation d'une base de règles commune (partition commune dans l'espace des antécédents (régresseurs) pour les trois sorties du bioprocédé), conduisant à un modèle MIMO autorégressif dans l'espace d'état.

Le contrôleur proposé, qui peut être vu comme une sorte de contrôleur de type PI multidimensionnel avec retour d'état, présente de bons résultats pour la régulation autour des consignes imposées. Un réglage fin des matrices de pondération Q et R du critère quadratique est en général nécessaire afin d'améliorer les performances. Nous croyons que les résultats obtenus sont intéressants, compte tenu du fait que la synthèse proposée du contrôleur flou basé sur des sous modèles linéaires *affines* locaux est méthodique et simple. De plus, l'obtention du modèle flou Takagi-Sugeno se fait à partir des seules données entrée-sortie du bioprocédé de traitement des eaux usées considéré.

Nous considérons à présent, par la suite quelques-unes des perspectives qui nous semble les plus intéressantes à étudier :

- L'extension de l'analyse de la stabilité pour le système constitué du modèle flou-commande floue de type *affine* Takagi-Sugeno en temps discret. Bien que certains travaux ont été réalisés dans cette direction [MAR95] [KIM01] [FEN04], les hypothèses ou les développements considérés sont assez restrictives, tant en termes de conservativité de la loi de commande que sur l'hypothèse de considérer un unique point d'équilibre à l'origine du système.
- L'étude approfondi des équivalences, au niveau du comportement dynamique, entre le modèle flou de type Takagi-Sugeno et le système non linéaire approximé. A notre connaissance, très peu des publications ont été consacrées à ce sujet [JOH00]. Cette étude s'avérerait intéressante afin d'établir des conditions d'équivalence qui permettrait d'assurer le même type de comportements dynamiques entre la représentation Takagi-Sugeno et le système.
- Le développement d'un observateur flou basé sur le modèle affine Takagi-Sugeno, afin d'estimer les variables non mesurables du système, comme c'est typiquement le cas dans le domaine des bioprocédés.
- L'intégration dans le processus de clustering, probablement au niveau de la fonction objectif, de l'unification de la base de règles pour obtenir une partition commune dans le cas des systèmes MIMO à des fins de commande.

Références bibliographiques

- [ABO03] ABONYI, J., *Fuzzy Modeling Identification for Control*. Birkhäuser, Boston, 2003.
- [AND68] ANDREWS, J. F., A mathematical model for the continuous culture of microorganisms utilizing inhibitory substrate, *Biotechnology and Bioengineering*, vol.10, pp.707-723, 1968.
- [ASE96] ASENJO, J. A., SUN, W. H., SPENCER, J. L., Optimal control of batch processes involving simultaneous enzymatic and microbial reactions, *Bioprocess and Biosystems Engineering*, vol.14, n6, pp.323-329, 1996.
- [ATI05] ATINE, J.C., *Méthodes d'apprentissage flou :application à la segmentation d'images biologiques*. Thèse de Doctorat de l'Institut National des Sciences Appliquées, Rapport LAAS No. 05600, Toulouse, France, 2005.
- [BAB02] BABUŠKA, R., *Fuzzy Modeling and Identification Toolbox*. Control Engineering Laboratory, Faculty of Information Technology and Systems, Delft University of Technology, Delft, The Netherlands, version 3.3, ed. 2002. (Available at <http://www.dcsc.tudelft.nl/~babuska/>).
- [BAB95] BABUŠKA, R., VERBRUGGEN, H.B., Identification of composite linear models via fuzzy clustering, *Proceedings European Control Conference*, Rome, Italy, pp.1207-1212, 1995.
- [BAB96a] BABUŠKA, R. VAN CAN, H., VERBRUGGEN H., Fuzzy modelling of enzymatic Penicilin-G conversion, In *Preprints 13th IFAC World Congress*, San Francisco, USA, vol.N, pp.479-484, 1996.
- [BAB96a] BABUŠKA, R. FANTUZZI, C. KAYMAK, U., VERBRUGGEN H. B., Improved inference for Takagi-Sugeno models, In: *Proceedings of the IEEE International Conference on Fuzzy Systems*, New Orleans, USA, vol.1, pp.701-706, 1996.
- [BAB98a] BABUŠKA, R., *Fuzzy Modeling for Control*. Kluwer Academic Publishers, Mass., USA, 1998.
- [BAB98b] BABUŠKA, R., ROUBOS J. A. VERBRUGGEN, H. B., Identification of MIMO systems by input-output TS models, *IEEE International Conference on Fuzzy Systems*, Anchorage, USA, vol.1, pp.657-662, 1998.
- [BAI86] BAILEY, J. E., OLLIS, D. F., *Biochemical Engineering Fundamentals*. McGraw Hill, New York, 1986.
- [BAR03] BARATO, S., GONZALEZ, D.M., Diseño de un algoritmo de clustering con fines de modelamiento difuso para sistemas no lineales. Memoria de fin de estudios Ingeniería Electrónica, Laboratorio de Automática, Microelectrónica e Inteligencia Computacional-LAMIC, Universidad Distrital FJDC, Bogotá, Colombia, 2003.
- [BAS90] BASTIN, G. and DOCHAIN, D., *On-line Estimation and Adaptive Control of Bioreactors*. Elsevier, Amsterdam, 1990.
- [BAS95] BASTIN, G., VAN IMPE J. F., Nonlinear and adaptive control in biotechnology : A tutorial, *European Journal of Control*, vol.1, n1, pp.1-37, 1995.

- [BEL57] BELLMAN, R. E. *Dynamic Programming*. Princeton University Press, Princeton, NJ, 1957.
- [BEL99] BELANCHE, L. A., VALDES, J. J., COMAS, J., RODA, I. R., POCH, M., Towards a model of input-output behaviour of wastewater treatment plants using soft computing techniques, *Environmental Modelling & Software*, vol.14, pp.409-419, 1999.
- [BEN96] BEN YOUSSEF, C. *Filtrage, Estimation et Commande d'un Procédé de Traitements des Eaux Usées*. Thèse de Doctorat de l'Institut National Polytechnique de Toulouse, Rapport LAAS No. 96116, Toulouse, France, 1996.
- [BEN96] BEN YOUSSEF, C., DAHHOU, B., Multivariable adaptive predictive control of an aerated lagoon for a wastewater treatment, *Journal of Process Control*, vol.6, n5, pp.265-276, 1996.
- [BEN97] BEN YOUSSEF, C., *Benchmark of a waste water treatment process*, Rapport LAAS No. 97458, Toulouse, France, 1997.
- [BENS96] BENSaid, A.M., HALL, L.O., BEZDEK, J., CLARKE, L.P., SILBINGER, M.L., ARRINGTON, J.A., MURTAGH, R.F., Validity-guided (re)clustering with applications to image segmentation, *IEEE Transactions on Fuzzy Systems*, vol.4, pp.112-123, 1996.
- [BER01] BERNARD, O., QUEINNEC, I., Modèles dynamiques de procédés biochimiques. Propriétés des modèles. Dans : DOCHAIN, D., (Ed.), *Automatique des bioprocédés*. Hermès Science Publications, Paris, 2001.
- [BEZ80] BEZDEK, J., A convergence theorem for the fuzzy isodata clustering algorithms, *IEEE Transactions on Pattern Analysis and Machine Intelligence-PAMI*, vol.2, n1, pp.1-8, 1980.
- [BEZ81] BEZDEK, J., *Pattern Recognition with Fuzzy Objective Function*. Plenum Press, New York, 1981.
- [BEZ87] BEZDEK, J., HATHAWAY, R.J., SABIN, M.J., TUCKER, W.T., Convergence theory for fuzzy c-means: counterexamples and repairs, *IEEE Transactions on Systems, Man and Cybernetics*, vol.17, n5, pp.873-877, 1987.
- [BIE99] BIELA, P., *Classification automatique d'observations multidimensionnelles par réseaux de neurones compétitifs*. Thèse de Doctorat de l'Université des Sciences et Technologies de Lille, France, 1999.
- [BLA01] BLANCO, Y., *Stabilisation des modèles Takagi-Sugeno et leur usage pour la commande des systèmes non linéaires*, Thèse de doctorat, Ecole centrale de Lille, UST Lille 2001.
- [BOM98] BOMBERGER, J.D., SEBORG, D.E., Determination of model order for NARX models directly from input-output data, *Journal of Process Control*, vol.8, n5, pp.459-468, 1998.
- [BOR99] BORTOLET, P., PASSAQUAY, D., BOVERIE, S., TITLI, A., Iterative fuzzy modeling and control of nonlinear system using multidimensional fuzzy sets, *Proc. IFAC'99 World Congress*, 9 p., 1999.
- [BOR03] BORTOLET, P., TITLI, A., Modélisation floue : approche structurale. Dans : FOULLOY, L., GALICHET, S., TITLI, A., (Dir.), *Commande floue 2 – de l'approximation à l'apprentissage*. Série Systèmes Automatisés - Traité IC2, Hermès Science Publications - Lavoisier, Paris, pp.37-91, 2003.

- [BOU03] BOUKEZZOULA, R., GALICHET, S., FOULLOY, L., Linéarisation entrée-sortie floue. Dans : FOULLOY, L., GALICHET, S., TITLI, A., (Dir.), *Commande floue 1 – de la stabilisation à la supervision*. Série Systèmes Automatisés - Traité IC2, Hermès Science Publications - Lavoisier, Paris, 2003.
- [BOU93] BOUCHON-MEUNIER, B., *La logique floue*, Collection Que sais-je, no 2702, Presses Universitaires de France, 1993.
- [BOU95] BOUCHON-MEUNIER, B., *La logique floue et ses applications*, Addison-Wesley, Paris, 1995.
- [BOY94] BOYD, S., GHAOUI, L., FERON, E., *Linear Matrix Inequalities in Systems and Control Theory*. SIAM, Philadelphia, 1994.
- [BRY69] BRYSON, A. E. Jr., HO, Y-C. *Applied optimal control*. Blaisdell, Massachusetts, 1969.
- [CER03] CERRADA, M., *Sur les modèles flous adaptatifs dynamiques*. Thèse de Doctorat de l'Institut National des Sciences Appliquées, Rapport LAAS No. 03499, Toulouse, France, 2003.
- [CHA97] CHAUME, F., BÉTEAU, J. F., Model based selection of an appropriate control strategy – application to an anaerobic digester, *Proc. of the 2nd. MATH-MOD*, pp.853-858, Vienne, Autriche, 5-7 février 1997.
- [CHU01] CHUANG, C.C., SU, S.F., CHEN, S.S., Robust TSK fuzzy modeling for function approximation with outliers, *IEEE Transactions on Fuzzy Systems*, vol.9, n6, pp.810-821, 2001.
- [COL00] COLOMER, J., MELENDEZ, J., AYZA, J., *Sistemas de supervisión: introducción a la monitorización y supervisión experta de procesos: métodos y herramientas*. Cetisa Boixerau, Barcelona, 2000.
- [CON59] CONTOIS, D. E., Relationship between population density and specific growth rate of continuous culture, *Journal de Génie Microbiologique*, vol.21, pp.40-50, 1959.
- [CZO89] CZOGALA, E., RAWLIK, T., Modelling of a fuzzy controller with application to the control of biological processes, *Fuzzy Sets and Systems*, vol.31, pp.13-22, 1989.
- [DAH92a] DAHHOU, B., ROUX, G., CHAMILOTHORIS, G., Modelling and adaptive predictive control of a continuous fermentation process, *Applied Mathematical Modelling*, vol.16, pp.545-552, 1992.
- [DAH92b] DAHHOU, B., LAKRORI, M., QUEINNEC, I., FERRET, E., CHÉRUY, A., Control of a continuous fermentation process, *Journal of Process Control*, vol.2, n2, pp.103-111, 1992.
- [DAH96] DAHHOU, B., ROUX, G. and VASILACHE, A. Neural networks for biomass growth rate modelling and model changement detection in a yeast fermentation process, *4th IFAC Conference "System Structure and Control"*, Bucharest, Roumanie, pp. 469-474, 1997.
- [DAH06] DAHHOU, B., ROUX, G. and NAKKABI, Y. *Modelling, software sensors, control and supervision of fermentation process*. Rapport LAAS No. 06344, Toulouse, France, 2006.
- [DAV91] DAVÉ, R.N., Characterization and detection of noise in clustering, *Pattern Recognition Letters*, Vol.12, n11, pp.657-664, 1991.
- [DAV97] DAVÉ, R.N., KRISHNAPURAM, R., Robust clustering methods: a unified view, *IEEE Transactions on Fuzzy Systems*, vol.5, pp.270-293, 1997.

- [DEM86] DEMAİN, A. L., SOLOMON, N. A., *Manual of industrial microbiology and biotechnology*, American Society for Microbiology, Washington D.C., USA, 1986.
- [DEN88] DENAC, M., MIGUEL, A., DUNN, I. J., Modeling dynamic experiments on the anaerobic degradation of molasses wastewater, *Biotechnology and Bioengineering*, vol.31, pp.1-10, 1988.
- [DOC01] DOCHAIN, D. *Automatique des bioprocédés*. Hermès Science Publications, Paris, 2001.
- [DOC86] DOCHAIN, D., *On-line parameter estimation, adaptive state estimation and control of fermentation processes*. Thèse de Doctorat, Université catholique de Louvain, Belgique, 1986
- [DOC91] DOCHAIN, D., Design of adaptive controllers for non-linear stirred tank reactors : Extension to the MIMO situation, *Journal of Process Control*, vol.1, n1, p.41-48, 1991.
- [DOC94] DOCHAIN, D. *Contribution to the Analysis and Control of Distributed Parameter Systems with Application to (Bio)chemical Processes and Robotics*. Thèse d'agrégation de l'enseignement supérieur, Université catholique de Louvain, Belgique, 1994.
- [DOH85] DOHNAL, M., Fuzzy bioengineering models, *Biotechnology and Bioengineering*, vol.27, pp.1146-1151, 1985.
- [DON01] DONDO, R. G., MARQUES, D., Optimal control of a batch bioreactor: a study on the use of an imperfect model, *Process Biochemistry*, vol.37, n4, pp.379-385, 2001.
- [DRI93] DRIANKOV, D., HELLENDORF, H., REINFRANK, M., *An Introduction to Fuzzy Control*, Springer, Berlin, 1993.
- [DUB80] DUBOIS, D., PRADE, H., *Fuzzy Sets and Systems – Theory and Applications*. Academic Press, New York, 1980.
- [DUB87] DUBOIS, D., PRADE, H., *Théorie des possibilités*, Masson, 2ième édition, 1987.
- [DUN74] DUNN, J., A fuzzy relative to ISODATA process and its use in detecting compact well separated clusters, *Journal of Cybernetics*, vol.3, n3, pp.32-57, 1974.
- [DUQ03] DUQUE, M., GAUTHIER, A., VILLA, J.L. and RAKOTO, N., Modeling and Control of a Water Treatment Plant, *9th IEEE International Conference on Systems, Man and Cybernetics*, Washington, DC, pp. 171-176, 2003.
- [DUV88] DUVIVIER, E., *Optimisation des procédés semi-continus de fermentation par programmation dynamique*. Thèse de Doctorat de l'Université Paul Sabatier, Toulouse, France, 1988.
- [EDE97] EDELINE, F., *L'Épuration Biologique des Eaux : Théorie & Technologie des Réacteurs*. 4^{ème} édition, Editions Cebedoc, Liège, France, 1997.
- [ELG97] EL GHAOU, L., Approche LMI pour la commande: une introduction, *Ecole d'été d'Automatique de Grenoble, identification et commande robustes : approche LMI*, 1997.
- [EMM96] EMMANOUILIDES, C., PETROU, L., Identification and control of anaerobic digesters using adaptive, on-line trained neural networks, *Computers & Chemical Engineering*, vol.21, n1, pp.113-143, 1996.

- [EST95] ESTABEN, M., POLIT, M., LABAT, P., STEYER, J. P., Modélisation floue d'un procédé de digestion anaérobie, *Récents Progrès en Génie des Procédés*, vol.9, pp.137-142, 1995.
- [EST97] ESTABEN, M., *Analyse et Commande par Logique Floue d'un Procédé Biologique de Dépollution*. Thèse de Doctorat de l'Université de Perpignan, France, 1997.
- [FAN96] FANTUZZI, C., ROVATTI, R., On the approximation capabilities of the homogeneous Takagi-Sugeno model, in *Proc. Fifth IEEE International Conference on Fuzzy Systems*, New Orleans, USA, pp.1067-1072.
- [FEI04] FEIL, B., ABONYI, J., SZEIFERT, F., Model order selection of nonlinear input-output models – a clustering based approach, *Journal of Process Control*, vol.14, pp.593-602, 2004.
- [FEN97] FENG, G., CAO, S.G., REES, N.W., CHAK, C.K., Design of fuzzy control systems with guaranteed stability, *Fuzzy Sets and Systems*, vol. 2, pp. 1-10, 1997.
- [FEN01] FENG, G., WANG, L., Controller synthesis of fuzzy dynamic systems based on piecewise Lyapunov functions, *Fuzzy IEEE'2001*, Melbourne, Australie, 2001.
- [FEN04] FENG, G., Stability analysis of discrete-time fuzzy dynamic systems based on piecewise Lyapunov functions, *IEEE Transactions on Fuzzy Systems*, vol.12, n1, pp. 22-28, 2004.
- [FIL85] FILEV, P. D., KISHIMOTO, M., SENGUPTA, S., YOSHIDA, T., TAGUCHI, H., Application of the fuzzy set theory to simulation of batch fermentation, *Journal of Fermentation Technology*, vol.63, pp.545-553, 1985.
- [FIL96] FILEV, D., Model based fuzzy control, In: *Proc. Fourth European Congress on Intelligent Techniques and Soft Computing EUFIT'96*, Aachen, Germany, 1996.
- [FLA80] FLANAGAN, M. J., On the application of approximate reasoning to control of the activated sludge process, In : *Joint Automatic Control Conference*, San Francisco, USA, 1980.
- [FOU03a] FOULLOY, L., GALICHET, S., TITLI, A., (Dir.), *Commande floue 1 – de la stabilisation à la supervision*. Série Systèmes Automatisés - Traité IC2, Hermès Science Publications - Lavoisier, Paris, 2003.
- [FOU03b] FOULLOY, L., GALICHET, S., TITLI, A., (Dir.), *Commande floue 2 – de l'approximation à l'apprentissage*. Série Systèmes Automatisés - Traité IC2, Hermès Science Publications - Lavoisier, Paris, 2003.
- [FOU79] FOULARD, C., GENTIL, S., SANDRAZ, J.P., *Commande et régulation numérique par calculateur*. Editions Eyrolles, Deuxième édition, Paris, 1979.
- [FOU95] FOULLOY, L., GALICHET, S., Typology of fuzzy controllers. In: NGUYEN, H.T, SUGENO, M., TONG, R., YAGER, R., (Eds.), *Theoretical aspects of fuzzy control*, John Wiley & Sons, pp.65-90, 1995.
- [FRI96] FRIGUI, H., KRISHNAPURAM, R., A robust clustering algorithm based on competitive agglomeration and soft rejection of outliers, *Proc. IEEE Computer Society Conference on Computer Vision and Pattern Recognition-CVPR'96*, pp.550-555, 1996.
- [FRI96a] FRIGUI, H., KRISHNAPURAM, R., A robust algorithm for automatic extraction of an unknown number of clusters from noisy data, *Pattern Recognition Letters*, Vol.17, pp.1223-1232, 1996.

- [FRI96b] FRIGUI, H., KRISHNAPURAM, R., A robust clustering algorithm based on competitive agglomeration and soft rejection of outliers, *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, San Francisco, CA, USA, pp.550-555, 1996.
- [FRI97] FRIGUI, H., KRISHNAPURAM, R., Clustering by competitive agglomeration, *Pattern Recognition*, Vol.30, n7, pp.1109-1119, 1997.
- [FRI99] FRIGUI, H., KRISHNAPURAM, R., A robust competitive clustering algorithm with applications in computer vision, *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol.21, n5, pp.450-465, 1999.
- [FU98] FU C. S., POCH, M., Fuzzy model and decision of COD control for an activated sludge process, *Fuzzy Sets and Systems*, vol.93, n3, pp.281-292, 1998.
- [GAH95] GAHINET, P., NEMIROVSKI, A., LAUB, A., CHILALI, M., *LMI Control Toolbox*, The Math Works Inc., Natick Massachusetts, 1995.
- [GAL01] GALICHET, S., *Contrôle flou : de l'interpolation numérique au codage de l'expertise*. Habilitation à Diriger des Recherches, Laboratoire LAMII, Université de Savoie, Annecy, France, 2001.
- [GAT89] GATH, I., GEVA, A., Unsupervised optimal fuzzy clustering, *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol.7, pp.773-781, 1989.
- [GAU97] GAUTHIER, A., DUQUE, M. and GIRALDO, E., Multivariable identification and control for a high rate anaerobic reactor, *XIII International Congress on CAD/Cam & Robotics*, Pereira, Colombia, pp.1-8, 1997.
- [GAW02] GAWEDA, A.E., *Optimal data-driven rule extraction using adaptive fuzzy-neural models*. Ph.D Thesis, Department of Computer Engineering and Computer Science, University of Louisville, Kentucky, USA, 2002.
- [GEO96] GEORGIEVA, O., PATARINSKA, T., Modeling and control of batch fermentation processes under conditions of uncertainty, *Bioprocess and Biosystems Engineering*, vol.14, n6, pp.299-306, 1996.
- [GEO01] GEORGIEVA, O., WAGENKNECHT, M., HAMPEL, R., Takagi-Sugeno fuzzy model development of batch biotechnological processes, *International Journal of Approximate Reasoning*, vol.26, n3, pp.233-250, 2001.
- [GER00] GÉRIN, F., Modélisation mathématique pour la valorisation des déchets organiques solides. Publication technique, Conservatoire National des Arts et Métiers, Mémoire de fin d'études Ingénieur, Toulouse, France, 2000.
- [GIR86] GIRALDO, E., OROZCO, A., Tratamientos anaerobios de las aguas residuales, Tesis Magister Ingeniería Civil, Universidad de los Andes, Colombia, 1986.
- [GIR98] GIRALDO, E., *Reglamento de Agua Potable y Saneamiento Básico*. Bogotá, Colombia, 1998.
- [GLO03] GLORENNEC, P.Y., Apprentissage par renforcement – application aux systèmes d'inférence floue. Dans : FOULLOY, L., GALICHET, S., TITLI, A., (Dir.), *Commande floue I – de la stabilisation à la supervision*. Série Systèmes Automatisés - Traité IC2, Hermès Science Publications - Lavoisier, Paris, pp.149-190, 2003.
- [GOD93] GODFREY, K., (Ed.), *Perturbation Signals for System Identification*, Prentice Hall, New York, 1993.

- [GOM02] GOMEZ-QUINTERO, C., *Modélisation et estimation robuste pour un procédé boues activées en alternance de phases*, Thèse de Doctorat de l'Université Paul Sabatier, Toulouse, France, 2002.
- [GOT82] GOTTSCHAL, J.C., FREDE THINGSTAD, T., Mathematical description of competition between two and three bacterial species under dual substrate limitation in the chemostat; a comparison with experimental data, *Biotechnology and Bioengineering*, vol.24, pp.1403-1419, 1982.
- [GRI04] GRISALES, V.H., SORIANO, J.J., BARATO, S., GONZALEZ, D.M., Robust agglomerative clustering algorithm for fuzzy modeling purposes, *Proceedings of the American Control Conference-ACC'04*, Boston, MA, USA, pp.1782-1787, 2004.
- [GRI05] GRISALES, V. H., GAUTHIER, A., ROUX, G., Fuzzy model identification of a biological process based on input-output data clustering, *Proc. of the 2005 IEEE International Conference on Fuzzy Systems - FuzzIEEE*, Reno, USA, pp.927-932, 2005.
- [GRI06] GRISALES, V.H., GAUTHIER, A., ROUX, G., Fuzzy optimal control design for discrete affine Takagi-Sugeno fuzzy models: application to a biotechnological process, *Proceedings of the 2006 IEEE World Congress on Computational Intelligence – WCCI'06*, Vancouver, BC, Canada, July 16-21, 2006, pp. 10830-10837.
- [GRI07] GRISALES, V.H., OROZCO, P.R., UNRIZA, J.C., *FMIDg - Fuzzy modeling and identification graphical toolbox for use with Matlab - User's Manual v1.0*, Reporte técnico, Laboratorio de Automática, Microelectrónica e Inteligencia Computacional - LAMIC, Universidad Distrital FJDC, Bogotá, Colombia, 2007.
- [GUE99] GUERRA, T.M., VERMEIREN, L., DELMOTTE, F., BORNE, P., Lois de commande pour systèmes flous continus, *Journal Européen des Systèmes Automatisés*, vol. 33, n 4, pp. 489-527, 1999.
- [GUE01a] GUERRA, T.M., VERMEIREN, L., Control laws for continuous fuzzy systems, *Fuzzy Sets and Systems*, vol. 120, p. 95-108, 2001.
- [GUE01b] GUERRA, T.M., VERMEIREN, L., Conditions for non quadratic stabilization of discrete fuzzy models, *Advanced Fuzzy /Neural Control, IFAC AFNC'01*, Valence, Espagne pp. 15-20, 2001.
- [GUE01c] GUERRA, T.M., PERRUQUETTI, W., Non quadratic stabilization of discrete fuzzy models, *Fuzzy IEEE'2001*, Melbourne, Australie, 2001.
- [GUE03] GUERRA, T.M., VERMEIREN, L., Stabilité et stabilisation à partir de modèles flous. Dans : FOULLOY, L., GALICHET, S., TITLI, A., (Dir.), *Commande floue I – de la stabilisation à la supervision*. Série Systèmes Automatisés - Traité IC2, Hermès Science Publications - Lavoisier, Paris, 2003.
- [GUI01] GUILLAUME, S., Designing fuzzy inference systems from data: an interpretability-oriented review, *IEEE Transactions on Fuzzy Systems*, vol.9, n3, pp.426-443, 2001.
- [GUIA02] Ministerio del Medio Ambiente (MA), Gestión para el manejo, tratamiento y disposición final de las aguas residuales municipales, Colombia, Guía 2002.
- [GUS79] GUSTAFSON, D.E., KESSEL, W.C., Fuzzy clustering with a fuzzy covariance matrix, *Proceedings of the IEEE CDC*, San Diego, CA, USA, pp.761-766, 1979.

- [HAD02] HADJILI, M.L., WERTZ, V., Takagi–Sugeno fuzzy modeling incorporating input variables selection, *IEEE Transactions on Fuzzy Systems*, vol.10, n2, pp.728-742, 2002.
- [HAD99] HADJ-SADOK, Z., K., *Modélisation et estimation dans les bioréacteurs ; prise en compte des incertitudes : application au traitement de l'eau*. Thèse de Doctorat de l'Université de Nice – Sophia Antipolis, Nice, France, 1999.
- [HAT91] HATHAWAY, R., BEZDEK, J., Grouped coordinate minimization using Newton's method for inexact minimization in one vector coordinate, *Journal of Optimization Theory and Applications*, vol.71, n3, pp.503-516, 1991.
- [HAT93] HATHAWAY, R., BEZDEK, J., HATHAWAY, R., BEZDEK, J., Switching regression models and fuzzy clustering, *IEEE Transactions on Fuzzy Systems*, vol.1, n3, pp.195-204, 1993.
- [HE93] HE, X., ASADA, H., A new method for identifying orders of input-output models for nonlinear dynamic systems, *Proc. of the American Control Conference*, San Francisco, CA, pp.2520-2523, 1993.
- [HEL97] HELLENDORF, H., DRIANKOV, D., *Fuzzy model identification – selected approaches*, Springer-Verlag, Berlin, 1997.
- [HEN87] HENZE, M., LESLIE GRADY Jr, C.P., GUJER, W., MARAIS, G.v.R., MATSUO, T., *The activated sludge model no.1*, IAWQ Scientific and Technical Report, Londres, UK, 1987.
- [HEN94] HENZE, M., GUJER, W., MINO, T., MATSUO, WENTZEL, M. T., MARAIS, G.v.R., *The activated sludge model no.2 – biological phosphorus removal*, IAWQ Scientific and Technical Report, Londres, UK, 1994.
- [HEN02] HENZE, M., HARREMOËS, P., LA COUR JANSEN, J. and ARVIN, E., *Wastewater Treatment – Biological and Chemical Process*, 3rd Edition, Environmental Engineering Series, Springer-Verlag, Berlin, 2002.
- [HER06] HERNANDEZ DE LEON, H., *Supervision et diagnostic des procédés de production d'eau potable*. Thèse de Doctorat de l'Institut National des Sciences Appliquées, Rapport LAAS No. 06584, Toulouse, France, 2006.
- [HON00] HONDA, H., KOBAYASHI, T., Fuzzy control of bioprocess, *Journal of Bioscience and Bioengineering*, vol.89, n5, pp.401-408, 2000.
- [HOR95] HORIUCHI, J., KISHIMOTO, M., MOMOSE, H., Hybrid simulation of microbial behavior combining a statistical procedure and fuzzy identification of culture phases, *Journal of Fermentation Bioengineering*, vol.79, pp.297-299, 1995.
- [HOR02] HORIUCHI, J. I., Fuzzy modeling and control of biological processes, *Journal of Bioscience and Bioengineering*, vol.94, n6, pp.574-578, 2002.
- [HUB81] HUBER, P.J., *Robust Statistics*, John Wiley & Sons, New York, 1981.
- [ISE97] ISERMAN, R., Supervision, fault detection and fault diagnosis methods – an introduction, *Control Engineering Practice*, vol.5, n5, pp.639-652, 1997.
- [JAI88] JAIN, A.K., DUBES, R.C., *Algorithms for Clustering Data*. Prentice Hall, Englewood Cliffs, NJ, 1988.
- [JAN93] JANG, J.-S.R. ANFIS: Adaptive-network-based fuzzy inference systems, *IEEE Transactions on Systems, Man and Cybernetics*, vol.23, n3, pp.665-685, 1993.
- [JIN00] JIN, Y., Fuzzy modeling of high-dimensional systems: complexity reduction and interpretability improvement, *IEEE Transactions on Fuzzy Systems*, vol.8, pp.212-222, 2000.

- [JOH00] JOHANSEN, T. A., SHORTEN, R., MURRAY-SMITH, R., On the interpretation and identification of dynamic Takagi-Sugeno fuzzy models, *IEEE Transactions on Fuzzy Systems*, vol.8, n3, pp. 297-313, 2000.
- [JOH99] JOHANSSON, M., RANTZER, A., ARZEN, K., Piecewise quadratic stability of fuzzy systems, *IEEE Transactions on Fuzzy Systems*, vol. 7, pp. 713-722, 1999.
- [JUL97] JULIEN, S., *Modélisation et estimation pour le contrôle d'un procédé boues activées éliminant l'azote des eaux résiduaires urbaines*. Thèse de Doctorat de l'Université Paul Sabatier, Toulouse, France, 1997.
- [KAY95] KAYMAK, U., BABUŠKA, R., Compatible cluster merging for fuzzy modelling, In: *Proceedings of the FUZZ-IEEE/IFES'95*, Yokohama, Japan, pp.897-904, 1995.
- [KEM04] KEMPOWSKY, T. *Surveillance de procédés à base de méthodes de classification : conception d'un outil d'aide pour la détection et le diagnostic de défaillances*. Thèse de Doctorat de l'Institut National des Sciences Appliquées, Rapport LAAS No. 04748, Toulouse, France, 2004.
- [KHO97] KHODJA, L., *Contribution à la classification floue non supervisée*. Thèse de Doctorat de l'Université de Savoie, France, 1997.
- [KIE97] KIELY, G., TAYFUR, G., DOLAN, C., TANJI, K. K., Physical and mathematical modeling of anaerobic digestion of organic wastes, *Water Research*, vol.31, n3, pp.534-540, 1997.
- [KIM01] KIM, E., KIM, D., Stability analysis and synthesis for an affine fuzzy system via LMI and ILMI : discrete case, *IEEE Transactions on Systems, Man and Cybernetics-Part B : Cybernetics*, vol.31, n1, pp. 132-140, 2001.
- [KIM95] KIM, W.C., HAN, S.C., KWON, W.H., Stability analysis and stabilization of fuzzy state space models, *Fuzzy Sets and Systems*, vol. 71, pp. 131-142, 1995.
- [KIR98] KIRIAKIDIS, K., GRIVAS, A., TZES, A., Quadratic stability analysis of the Takagi-Sugeno fuzzy model, *Fuzzy Sets and Systems*, vol. 98, n 1, pp. 1-14, 1998.
- [KIS91] KISHIMOTO, M., KITTA, Y., TAKEUCHI, S., NAKAJIMA, M., YOSHIDA, T., Computer control of glutamic acid production based on fuzzy clusterization of culture phases, *Journal of Fermentation Bioengineering*, vol.72, pp.110-114, 1991.
- [KIT94] KITSUTA, Y., KISHIMOTO, M., Fuzzy supervisory control of glutamic acid production, *Biotechnology and Bioengineering*, vol.44, pp.87-94, 1994.
- [KLI95] KLIR, G., YUAN, B., *Fuzzy sets and fuzzy logic; theory and applications*. Prentice Hall, 1995.
- [KON89] KONSTANTINOV, K., YOSHIDA, T., Physiological state control of fermentation processes, *Biotechnology and Bioengineering*, vol.33, pp.1145-1156, 1989.
- [KON92] KONSTANTINOV, K., YOSHIDA, T., Knowledge-based control of fermentation processes, *Biotechnology and Bioengineering*, vol.39, pp.479-486, 1992.
- [KOS91] KOSKO, B., *Neural Networks and Fuzzy Systems*, Prentice Hall, Englewood Cliffs, NJ, 1991.
- [KOS94] KOSKO, B., Fuzzy systems as universal approximators , *IEEE Trans.Computers*, vol.43, pp.1329-1333, 1994.
- [KRI93] KRISHNAPURAM, R., KELLER, J., A possibilistic approach to clustering, *IEEE Transactions on Fuzzy Systems*, vol.1, n2, pp.98-110, 1993.

- [LAC03] LACROSE, V., TITLI, A., Contrôleurs flous à complexité réduite : fusion de variables et structuration de la commande. Dans : FOULLOY, L., GALICHET, S., TITLI, A., (Dir.), *Commande floue I – de la stabilisation à la supervision*. Série Systèmes Automatisés - Traité IC2, Hermès Science Publications - Lavoisier, Paris, 2003.
- [LEE90a] LEE, C., Fuzzy logic in control systems: fuzzy logic controller – part I, *IEEE Trans. on Systems, Man and Cybernetics*, vol.20, n2, pp.404-418, 1990.
- [LEE90b] LEE, C., Fuzzy logic in control systems: fuzzy logic controller – part II, *IEEE Trans. on Systems, Man and Cybernetics*, vol.20, n2, pp.419-435, 1990.
- [LEO85] LEONARITIS, I., BILLINGS, S., Input-output parametric models for nonlinear systems, *International Journal of Control*, vol.41, pp.303-344, 1985.
- [LEO99] LÉONARD, D., BEN YOUSSEF, C., DESTRUHAUT, C., LINDLEY, N. D., QUEINNEC, I., Phenol degradation by *Ralstonia eutropha* : colorimetric determination of 2-hydroximuconate semialdehyde accumulation to control feed strategy in fed-batch fermentations, *Biotechnology and Bioengineering*, vol.65, n4, pp.407-415, 1999.
- [LIN82] LINTON, J.D., DROZD, J.W., *Microbial interactions and communities in biotechnology*, In : Microbial interactions and communities, BULL, A.T., SLATER, J.H., (Eds.), Academic Press, London, 1982, pp.357-406.
- [LIN94] LIN, C., *Neural Fuzzy Control Systems with Structure and Parameter Learning*. World Scientific, Singapore, 1994.
- [LJU87] LJUNG, L., *System identification, Theory for the user*. Prentice Hall, New Jersey, 1987.
- [MA98] MA, X.J., SUN, Z.Q., HE, Y.Y., Analysis and design of fuzzy controller and fuzzy observer, *IEEE Transactions on Fuzzy Systems*, vol. 6, n 1, pp. 41-50, 1998.
- [MAM77] MAMDANI, E., Application of fuzzy logic to approximate reasoning using linguistic systems, *Fuzzy Sets and Systems*, vol.26, pp.1182-1191, 1977.
- [MAN98] MANESIS, S. A., SAPIDIS, D. J., KING, R. E., Intelligent control of wastewater treatment plants, *Artificial Intelligence in Engineering*, vol.12, pp.275-281, 1998.
- [MAR95] MARIN, J.P., Conditions nécessaires et suffisantes de stabilité quadratique d'une classe de systèmes flous, *Rencontres LFA'95*, pp. 240-247, Paris, novembre, 1995.
- [MEG72] MEGEE, R.D., DRAKE, J.F., FREDRICKSON, A.G., TSUCHIYA, H.M., Studies in intermicrobial symbiosis, *Saccharomyces cerevisiae* and *Lactobacillus casei*, *Canadian Journal of Microbiology*, vol.18, pp.1733-1742, 1972.
- [MET03] METCALF & EDDY Inc, TCHOBANOPOULOS, G., BURTON, F. and STENCEL, H., *Wastewater Engineering : Treatment and Reuse*, 4th Edition, McGraw-Hill, 2003.
- [MOJ05] MÓJICA, E., GAUTHIER, A., Desarrollo y simulación de un modelo para un proceso de cultivo celular en bioreactor, *II Congreso Colombiano de Bioingeniería e Ingeniería Biomédica*, Bogotá, Colombia, pp.1-6, 2005.
- [MOL86] MOLETTA, R., VERRIER, D., ALBAGNAC, G., Dynamic modelling of anaerobic digestion, *Water Research*, vol.20, n4, pp.427-434, 1986.
- [MON94] MONTAGUE, G., MORRIS, J., Neural network contribution in biotechnology, *TIBTECH*, vol.12, pp.312-323, 1994.

- [MOR92] MORENO, R., DE PRADA, C., LAFUENTE, J. M., POCH, M. MONTAGUE, G., Non-linear predictive control of dissolved oxygen in the activated sludge process. In : Proc. of the ICCAFT 5/IFAC-BIO 2 Conference, Keystone, USA, pp.289-293, 1992.
- [MUL97] MÜLLER, A., MARSILI-LIBELI, S., AIVASIDIS, A., LLOYD, T., KRONER, S., WANDREY, C., Fuzzy control of disturbances in a wastewater treatment process, *Water Research*, vol.31, n12, pp.3157-3167, 1997.
- [MUR95] MURALIKRISHNAN, G., CHIDAMBARAM, M., Control of bioreactors using a neural network model, *Bioprocess Engineering*, vol.12, pp.35-39, 1995.
- [MUR97] MURRAY-SMITH, R., JOHANSEN, T.A., (Eds.), *Multiple Model Approaches to Modelling and Control*, Taylor & Francis, London, 1997.
- [NAI03] NAIDU, D. S. *Optimal Control Systems*. CRC Press, USA, 2003.
- [NAK85] NAKAMURA, T., KURATANI, T., MORITA, Y., Fuzzy control: application to glutamic acid fermentation, *Proc. of IFAC Modeling and Control of Biotechnology Process*, pp.211-215, 1985.
- [NAK97a] NAKOULA, Y., *Apprentissage des modèles linguistiques flous, par jeu de règles pondérées*. Thèse de Doctorat, Université de Savoie, France, 1997.
- [NAK97b] NAKOULA, Y., GALICHET, S., FOULLOY, L., Identification of linguistic fuzzy models based on learning, In: HELLENDORF, H., DRIANKOV, D., (Eds.), *Fuzzy model identification – selected approaches*, Springer-Verlag, Berlin, pp.281-319, 1997.
- [NAR90] NARENDRA, K. S., PARTHASARATHY, K., Identification and control of dynamical systems using neural networks, *IEEE Trans. on Neural Networks*, vol.1, n1, pp.4-27, 1990.
- [NGU98] NGUYEN, H., SUGENO, M., (Dir.), *Fuzzy Systems: Modeling and Control*, The Handbooks of Fuzzy Sets Series. Kluwer Academic Publishers, 1998.
- [OBE99] OBENAU, F., ROSENWINKEL, K. H., ALEX, J., TSCHPETZKI, R., JUMAR, U., Components of a model-based operation system for wastewater treatment plants, *Water Science and Technology*, vol.39, n4, pp.103-111, 1999.
- [OHN76] OHNO, H., NAKANISHI, E., TAKAMATSU, T., Optimal control of a semibatch fermentation, *Biotechnology & Bioengineering*, vol.18, pp. 847-864, 1976.
- [ORO03] OROZCO, A., *Bioingeniería de Aguas Residuales*. Bogotá, Colombia, 2003.
- [ORO07] OROZCO, P.R., UNRIZA, J.C., GRISALES, V.H., FMIDGraph, Herramienta computacional para modelamiento e identificación difusa de sistemas, sometido al VII Congreso de la Asociación Colombiana de Automática-ACA, 6p., Cali, Marzo, 2007.
- [PAI97] PAILLES, D., Modélisation et commande robuste d'un procédé de traitement des eaux usées de type biofiltre. Mémoire de fin d'études Ingénieur, Publication technique, Conservatoire National des Arts et Métiers, Toulouse, France, 1997.
- [PAL95] PAL, N., BEZDEK, J., On cluster validity for the fuzzy c-means model, *IEEE Transactions on Fuzzy Systems*, vol.3, n3, pp.370-379, 1995.
- [PAR93] PARK, Y. S., SHI, Z. P., SHIBA, S., CAYUELA, C., IJIMA, S., KOBAYASHI, T., Application of fuzzy reasoning on control of glucose and ethanol concentrations in baker's yeast culture, *Applications of Microbiology and Biotechnology*, vol.38, pp.649-655, 1993.

- [PAT97] PATNAIK, P., A recurrent neural network for a fed-batch fermentation with recombinant *Escherichia coli* subject to inflow disturbances, *Process Biochemistry*, vol.32, n5, pp.391-400, 1997.
- [PAV94] PAVÉ, A., *Modélisation en biologie et en écologie*. Aléas, Lyon, 1994.
- [PED84] PEDRYCZ, W., An identification algorithm in fuzzy relational systems, *Fuzzy Sets and Systems*, vol.13, pp.153-167, 1984.
- [PED95] PEDRYCZ, W., *Fuzzy Sets Engineering*. CRC Press, Boca Raton, FL, 1995.
- [PET87] PETERSEN, I.R., A procedure for simultaneously stabilizing a collection of single input linear systems using non linear state feedback control, *Automatica*, vol.23, n1, pp.33-40, 1987.
- [POL01] POLIT, M., GENOVESI, A., CLAUDET, B., Fuzzy logic observers for a biological wastewater treatment process, *Applied Numerical Mathematics*, vol.39, pp.173-180, 2001.
- [POL02] POLIT, M., ESTABEN, M., LABAT, M., A fuzzy model for an anaerobic digester, comparison with experimental results, *Engineering Applications of Artificial Intelligence*, vol.15, pp.385-390, 2002.
- [PON62] PONTRYAGIN, L. S., BOLTYANSKII, V. G., GAMKRELIDZE, R. V., MISHCHENKO, E. F., *The Mathematical Theory of Optimal Processes*. Wiley-Interscience, New York, NY, 1962. (Traduit du Russe).
- [PRO90] PROSSER, J. I., Mathematical modeling of nitrification processes, *Advances in Microbial Ecology*, vol.11, pp.263-304, 1990.
- [PSI92] PSICHOGIOS, D., UNGAR, L., A hybrid neural-network – first principles approach to process modeling, *AIChE Journal*, vol.38, pp.1499-1511, 1992.
- [QUE00] QUEINNEC, I. *Contribution à la Commande de Procédés Biotechnologiques : Application au Traitement Biologique de la Pollution*. Habilitation à Diriger des Recherches, Laboratoire d'Analyse et d'Architecture des Systèmes, Rapport LAAS No. 00455, Toulouse, France, 2000.
- [QUE00a] QUEINNEC, I., LÉONARD, D., DENOU, G., State observation of phenol degradation by *Ralstonia eutropha* based on dissolved oxygen measurement, *Proc. of the IFAC Symposium Adchem*, Pise, Italie, 14-16 juin, 2000.
- [QUE00b] QUEINNEC, I., DOCHAIN, D., Modelling and simulation of steady-state secondary settlers in wastewater treatment plants, *Proc. of the 5th International Symposium on Systems Analysis and Computing in Water Quality Management (WATERMATEX)*, Gent, Belgique, 18-20 septembre, 2000.
- [QUE97] QUEINNEC, I., TROFINO-NETO, A., H_{∞} attenuation of a bioprocess with parameter uncertainty, In: *Proc. of the American Control Conference*, Albuquerque, USA, pp. 1066-1067, 1997.
- [QUE98] QUEINNEC, I., HARMAND, J., STEYER, J. P., BERNET, N., Mathematical analysis of a nitrification process, *Recent Research Developments in Fermentation and Bioengineering*, vol.1, pp.153-163, 1998.
- [RAG01] RAGOT, J., GRAPIN, G., CHATELLIER, P., COLIN, F., Modelling of a water treatment plant. A multi-model representation, *Environmetrics*, vol.12, pp.599-611, 2001.
- [RAM05] RAMASWAMY, S., CUTRIGHT, T. J., QAMMAR, H. K., Control of a continuous bioreactor using model predictive control, *Process Biochemistry*, vol.40, n8, pp.2763-2770, 2005.

- [RAP02] RAPAPORT, A., HARMAND, J., Robust regulation of a class of partially observed nonlinear continuous bioreactors, *Journal of Process Control*, vol.12, n2, pp.291-302, 2002.
- [REN03] REN, L., IRWIN, G.W., Robust fuzzy Gustafson-Kessel clustering for nonlinear system identification, *International Journal of Systems Science*, vol.34, n14-15, pp.787-803, 2003.
- [RHO98] RHODES, C., MORARI, M., Determining the model order for nonlinear input/output systems, *AIChE Journal*, vol.44, n1, pp.151-163, 1998.
- [ROD02] RODRIGUEZ, M., PAZ, A., HERNANDEZ, C., BELALCAZAR, L. and GIRALDO, E., Retention of granular sludge at high hydraulic loading rates in an anaerobic membrane bioreactor, *Water Science and Technology*, vol.45, n.10, pp.169-174, 2002.
- [ROD99] RODRIGO, M. A., SECO, A., FERRER, J., PENYA-ROJA, J. M., VALVERDE, J. L., Nonlinear control of an activated sludge aeration process : use of fuzzy techniques for tuning PID controllers, *ISA Transactions*, vol.38, n3, pp.231-241, 1999.
- [ROL95] ROLS, J.L., GOMA, G., FONADE, C., *Biotechnology and paper industry. Aerated lagoon for wastewater treatment*, Dans : Environmental Biotechnology : Principles and Practice, Kluwer Academic Publishers, 1995.
- [ROU92] ROUX, G. *Contribution à l'élaboration d'algorithmes de commande adaptative pour la conduite de procédés fermentaires*. Thèse de Doctorat de l'Université Paul Sabatier, Rapport LAAS No. 92513, Toulouse, France, 1992.
- [ROU94] ROUX, G., DAHHOU, B., QUEINNEC, I., Adaptive non-linear control of a continuous stirred tank bioreactor, *Journal of Process Control*, vol.4, n3, pp.121-126, 1994.
- [ROU99] ROUX, G., TITLI, A. (Eds.), *FAMIMO Waste-water Benchmark*, FAMIMO deliverable, Rapport LAAS-CNRS, Février, 1999, 58p.
- [ROU01] ROUX, G. *Contribution à l'estimation, l'observation et à la commande de procédés*. Habilitation à Diriger des Recherches, Laboratoire d'Analyse et d'Architecture des Systèmes, Rapport LAAS No. 01650, Toulouse, France, 2001.
- [ROV96] ROVATTI, R., Takagi-Sugeno models as approximators in Sobolev norms: the SISO case, in *Proc. Fifth IEEE International Conference on Fuzzy Systems*, New Orleans, USA, pp.1060-1066.
- [RUS69] RUSPINI, E., A new approach to clustering, *Information and Control*, vol.15, pp.22-32, 1969.
- [RYD72] RYDER, D.N., SINCLAIR, C.G., Model for the growth of aerobic microorganisms under oxygen limiting conditions, *Biotechnology and Bioengineering*, vol.14, pp.787-798, 1972.
- [SAL05] SALA, A., BABUŠKA, R., GUERRA, T.M., Perspectives of fuzzy systems and control, *Fuzzy Sets and Systems*, vol.156, pp.432-444, 2005.
- [SET98] SETNES, M., BABUŠKA, R., VERBRUGGEN, H.B., Rule-based modelling: precision and transparency, *IEEE Transactions on Systems, Man and Cybernetics - Part C: Applications and Reviews*, vol.28, n1, pp.165-169, 1998.
- [SHI94] SHIBA, S., NISHIDA, Y., PARK, Y. S., IJIMA, S., KOBAYASHI, T., Improvement of cloned α -amylase gene expression in fed-batch culture of recombinant *Saccharomyces cerevisiae* by regulating both glucose and ethanol

- concentrations using a fuzzy controller, *Biotechnology and Bioengineering*, vol.44, pp.1055-1063, 1994.
- [SII95] SIIMES, T., LINKO, P., VON NUMERS, C., NAKAJIMA, M., ENDO, I., Real-time fuzzy-knowledge-based control of baker's yeast production, *Biotechnology and Bioengineering*, vol.45, pp.135-143, 1995.
- [SIN04] ŠINDELÁŘ, R., BABUŠKA, R., Input selection for nonlinear regression models, *IEEE Transactions on Fuzzy Systems*, vol.12, n5, pp.688-696, 2004.
- [SPE98] SPÉRANDIO, M., *Développement d'une procédure de compartimentation d'une eau résiduaire urbaine et application à la modélisation dynamique de procédés à boues activées*. Thèse de Doctorat de l'Institut National des Sciences Appliquées, Toulouse, France, 1998.
- [STE01] STEYER, J. P., GÉNOVÉSI, A., HARMAND, J., Outils d'aide au diagnostic et détection de pannes. Dans : DOCHAIN, D., (Ed.), *Automatique des bioprocédés*. Hermès Science Publications, Paris, 2001.
- [STE98] STEYER, J.P. *Modélisation, Commande et Diagnostic des Procédés Biologiques de Dépollution*. Habilitation à Diriger des Recherches, Laboratoire de Biotechnologie – INRA, Narbonne, France, 1998.
- [SUG88] SUGENO, M., KANG, G.T., Structure identification of fuzzy model, *Fuzzy Sets and systems*, vol.28, pp.15-33, 1988.
- [SYK73] SYKES, R.M., Identification of the limiting nutrient and specific growth rate, *Journal of Water Pollution of Control Federation*, vol.45, pp.888-890, 1973.
- [TAK85] TAKAGI, T., SUGENO, M., Fuzzy identification of systems and its application to modeling and control, *IEEE Trans. on Systems, Man and Cybernetics*, vol.15, n1, pp.116-132, 1985.
- [TAK85] TAKAGI, T., SUGENO, M., Fuzzy identification of systems and its application to modeling and control, *IEEE Trans. on Systems, Man and Cybernetics*, vol.15, n1, pp.116-132, 1985.
- [TAN92] TANAKA, K., SUGENO, M., Stability analysis and design of fuzzy control systems, *Fuzzy Sets and Systems*, vol.45, n2, pp.135-156, 1992.
- [TAN94] TANAKA, K., SANO, M., A., Robust stabilization problem of fuzzy control systems and its application to backing up control of a truck-trailer, *IEEE Transactions on Fuzzy Systems*, vol. 2, n 2, pp. 119-134, 1994.
- [TAN95] TANAKA, K., SANO, M., Trajectory stabilization of model car via fuzzy control, *Fuzzy Sets and Systems*, vol. 70, pp. 155-170, 1995.
- [TAN96] TANAKA, K., IKEDA, T., WANG, H., Robust stabilization of a class of uncertain nonlinear systems via fuzzy control: Quadratic stability, H_∞ control theory and linear matrix inequalities, *IEEE Transactions on Fuzzy Systems*, vol.4, n1, pp.1-13, 1996.
- [TAN98] TANAKA, K., IKEDA, T., WANG, H.O., Fuzzy regulators and fuzzy observers: relaxed stability conditions and LMI-based designs, *IEEE Transactions on Fuzzy Systems*, vol. 6, n 2, pp. 1-16, 1998.
- [TAY00] TAY, J. H., ZHANG, X. A fast predicting neural fuzzy model for high-rate anaerobic wastewater treatment systems, *Water Research*, vol.34, n11, pp.2849-2860, 2000.
- [TEB98] TE BRAAKE, H., BABUŠKA, R., VAN CAN, E., HELLINGA, C., Predictive Control in Biotechnology using Fuzzy and Neural models. In: VAN IMPE, J.F.M., VANROLLEGHEM, P.A., ISERENTANT, D.M., (Eds.), *Advanced*

- Instrumentation, Data Interpretation, and Control of Biotechnological Processes*. Kluwer Academic Publishers, The Netherlands, 1998.
- [THA97] THATHACHAR M.A.L., VISWANATH P.; On the stability of fuzzy systems, *IEEE Transactions on Fuzzy Systems*, vol. 5, n 1, pp. 145-151, 1997.
- [TOM99] TOMIDA, S., HANAI, T., UEDA, N., HONDA, H., KOBAYASHI, T., Construction of COD simulation model for activated sludge process by fuzzy neural network, *Journal of Bioscience and Bioengineering*, vol.88, pp.215-220, 1999.
- [TON80] TONG, R. M., BECK, M. B., LATTEN, A., Fuzzy control of the activated sludge wastewater treatment process, *Automatica*, vol.16, pp.659-701, 1980.
- [TRA97] TRAVÉ-MASSUYÈS, L., DAGUE, P., GUERRIN, F., Le raisonnement qualitatif. Hermès, France, 1997.
- [TRA03] TRAORE, A., *Logique floue et Contrôle Supervisé des Procédés Biologiques de Dépollution*. Thèse de Doctorat de l'Université Montpellier II, France, 2003.
- [TSA96] TSAI, Y. P., OUYANG, C. F., WU, M. Y., CHIANG, W. L., Effluent suspended solids control of activated sludge process by fuzzy control approach, *Water Environment Research*, vol.68, n6, pp.1045-1053, 1996.
- [TSO95] TSONEVA, R. G., PATARINSKA, T. D., Optimal control of continuous fermentation processes, *Bioprocess and Biosystems Engineering*, vol.13, n4, pp.189-196, 1995.
- [TUC87] TUCKER, W.T., BEZDECK, J., Counterexamples to the convergence theorem for fuzzy ISODATA clustering algorithms, In the *Analysis of Fuzzy Information*, CRC Press, Boca Raton, FL, vol.3, ch.7, 1987.
- [UNR07] UNRIZA, J.C., OROZCO, P.R., GRISALES, V.H., Selección del orden de modelos dinámicos no lineales MIMO a partir de datos de entrada-salida, sometido al VII Congreso de la Asociación Colombiana de Automática-ACA, 6p., Cali, Marzo, 2007.
- [VAN98] VAN IMPE, J.F.M., VANROLLEGHEM, P.A., ISERENTANT, D.M., (Eds.), *Advanced Instrumentation, Data Interpretation, and Control of Biotechnological Processes*. Kluwer Academic Publishers, The Netherlands, 1998.
- [VAN98a] VAN IMPE, J.F.M., BASTIN, G., Optimal Adaptive Control of Fed Batch Fermentation Process. In: VAN IMPE, J.F.M., VANROLLEGHEM, P.A., ISERENTANT, D.M., (Eds.), *Advanced Instrumentation, Data Interpretation, and Control of Biotechnological Processes*. Kluwer Academic Publishers, The Netherlands, 1998.
- [VANL02] VAN LITH, P. F., BETLEM, B. H. L., ROFFEL, B., A structured modeling approach for dynamic hybrid fuzzy-first principles models, *Journal of Process Control*, vol.12, pp. 605-615, 2002.
- [VEN99] VENKATESWARLU, Ch., NAIDU, K. V. S., Dynamic fuzzy model based predictive controller for a biochemical reactor, *Bioprocess Engineering*, vol.23, pp.113-120, 2000.
- [VON94] VON NUMERS, C., NAKAJIMA, M., ASAMA, H., LINKO, P., ENDO, I., A knowledge based system using fuzzy inference for supervisory control of bioprocesses, *Journal of Biotechnology*, vol.34, pp.109-118, 1994.
- [WAN94] WANG, L.X., *Adaptive Fuzzy Systems and Control, Design and Stability Analysis*. Prentice Hall, New Jersey, 1994.

- [WAN95] WANG, H. O., TANAKA, K. and GRIFFIN, M., Parallel Distributed Compensation of Nonlinear Systems by Takagi-Sugeno Fuzzy Model, *Proc. IEEE International Conference on Fuzzy Systems - FUZZ-IEEE'95*, Yokohama, Japan, pp.531-538, 1995.
- [WAN96] WANG, K. W., BALTZIS, B. C., LEWANDOWSKI, G. A., Kinetics of phenol biodegradation in the presence of glucose, *Biotechnology and Bioengineering*, vol.51, pp.87-94, 1996.
- [WEN98] WEN, C. H., VASSILIADIS, C. A., Applying hybrid artificial intelligence techniques in wastewater treatment, *Engineering Applications of Artificial Intelligence*, vol.11, n6, pp.685-705, 1998.
- [YAG94] YAGER, R.R., FILEV, D.P., *Essentials of fuzzy modeling and control*, John Wiley & Sons, New York, 1994.
- [YEN99] YEN, J., WANG, L., Simplifying fuzzy rule-base models using orthogonal transformation methods, *IEEE Transactions on Systems, Man and Cybernetics - Part B: Cybernetics*, vol.29, n1, pp.13-24, 1999.
- [YI93] YI, S., CHUNG, M., Identification of fuzzy relational model and its application to control, *Fuzzy Sets and Systems*, vol.59, pp.25-33, 1993.
- [YOO77] YOON, H., KLINZING, G., BLANCH, H.W., Competition for mixed substrates by microbial populations, *Biotechnology and Bioengineering*, vol.19, pp.1193-1210, 1977.
- [YU04] YU, J., CHENG, Q., HUANG, H., Analysis of the weighting exponent in the FCM, *IEEE Transactions on Systems, Man and Cybernetics-Part B: Cybernetics*, vol.34, n1, pp.634-639, 2004.
- [YU98] YU, R. F., LIAW, S. L., CHANG, C. N., CHENG, W. Y., Applying real-time control to enhance the performance of nitrogen removal in the continuous-flow SBR system, *Water Science and Technology*, vol.38, n3, pp.271-280, 1998.
- [ZAD65] ZADEH, L., Fuzzy sets, *Inf. Control*, vol.8, pp.338-353, 1965.
- [ZAD73] ZADEH, L., Outline of a new approach to the analysis of complex systems and decision processes, *IEEE Trans. on Systems, Man, and Cybernetics*, vol.1, pp.28-44, 1973.
- [ZEN95] ZENG, X., SINGH, M., Approximation theory of fuzzy systems – MIMO case, *IEEE Trans. on Fuzzy Systems*, vol.3, n2, pp.219-235, 1995.
- [ZHA95] ZHAO, J., *System modelling, identification and control using fuzzy logic*, Thèse de doctorat, Université Catholique de Louvain, Belgique, 1995.
- [ZHA96] ZHAO, J., GOREZ, R., WERTZ, V., Fuzzy controllers with guaranteed robustness and performance, *EUFIT'96*, Aachen, Allemagne, pp. 1886-1890, 1996.
- [ZHA97] ZHAO, J., GOREZ, R., WERTZ, V., Synthesis of Fuzzy Control Systems Based on Linear Takagi-Sugeno Fuzzy Models. In : MURRAY-SMITH, R., JOHANSEN, T.A., (Eds.), *Multiple Model Approaches to Modelling and Control*, Taylor & Francis, London, pp.307-336, 1997.
- [ZHU00] ZHU, G. Y., ZAMAMIRI, A., HENSON, M. A., HJORTSØ, M. A., Model predictive control of continuous yeast bioreactors using cell population balance models, *Chemical Engineering Science*, vol.55, n24, pp.6155-6167, 2000.
- [ZIM87] ZIMMERMANN, H.J., *Fuzzy Sets, Decision Making and Expert Systems*. Kluwer Academic Publishers, Boston, USA, 1987.

Thèse Université Paul Sabatier-Toulouse III, France - Université de los Andes, Colombie
École Doctorale Systèmes, Spécialité Systèmes Automatiques

Monsieur Victor-Hugo GRISALES PALACIO, 22 février 2007

MODÉLISATION ET COMMANDE FLOUES DE TYPE TAKAGI-SUGENO APPLIQUÉES À UN BIOPROCÉDÉ DE TRAITEMENT DES EAUX USÉES

Résumé : Ce travail de thèse s'inscrit au carrefour de l'Automatique, de l'Intelligence Artificielle et des Biotechnologies. Il cherche à développer une méthodologie de modélisation et de commande qui repose sur une approche par logique floue. La première partie du travail présente une introduction aux principes et techniques mises en œuvre dans les stations d'épuration actuelles, et met en évidence la difficulté de modélisation des différents phénomènes mis en jeu. À partir de ce constat, dans une deuxième partie, focalisée sur la modélisation et la commande floues, nous développons d'abord l'identification de modèles flous affines de type Takagi-Sugeno (TS) à partir de données entrées-sorties. Nous considérons différentes méthodes de coalescence floue et une méthode d'agglomération compétitive, robuste en présence de bruit. Ce type d'approche « boîte grise » permet une représentation à base de règles qui approxime la dynamique non linéaire comme une concaténation de sous-modèles localement linéaires sous la forme d'auto-régression non-linéaire (NARX). De plus, nous avons développé une version graphique de la boîte à outils pour la modélisation floue des systèmes (FMIDg). Ensuite, nous proposons une commande floue TS sous-optimale linéaire quadratique adaptée à la structure du modèle flou identifié, en utilisant la philosophie de commande du type compensation parallèle distribuée (PDC). La méthodologie globale est finalement testée et validée en simulation sur un bioprocédé aérobie de dépollution des eaux usées.

Mots-clés: Modèle flou Takagi-Sugeno, identification des systèmes, coalescence floue, commande floue, compensation parallèle distribuée, traitement des eaux usées, procédés biologiques.

TAKAGI-SUGENO FUZZY MODELLING AND CONTROL APPLIED TO A WASTEWATER TREATMENT BIOPROCESS

Abstract: This thesis work is placed at the crossroad of Automatic control, Artificial Intelligence and Biotechnology. It aims at developing a modelling and control methodology based on a fuzzy logic approach. The first part of the work presents an introduction to the principles and techniques used in current wastewater treatment plants. It highlights the difficulty in modelling the multiple phenomena involved. From this fact, in the second part, focused on fuzzy modelling and control, we develop initially the identification of affine fuzzy models of Takagi-Sugeno (TS) type from input-output data. We used several fuzzy clustering methods as well as an agglomerative-competitive method which is robust in presence of noise. This type of "grey box" approach allows a rule-based representation which approximates the nonlinear dynamics as a concatenation of locally linear sub-models in the nonlinear autoregressive form (NARX). Moreover, we have developed a graphical version of the FMID toolbox for the fuzzy modelling of systems. Then, we propose a suboptimal linear-quadratic TS fuzzy control relevant to the structure of the identified fuzzy model, by using the parallel distributed compensation (PDC) technique. Finally, the whole methodology is tested and validated in simulation on an aerobic wastewater treatment bioprocess.

Keywords: Takagi-Sugeno fuzzy model, system identification, fuzzy clustering, fuzzy control, parallel distributed compensation, wastewater treatment, biological processes.

MODELADO Y CONTROL DIFUSOS DE TIPO TAKAGI-SUGENO APLICADOS A UN BIOPROCESO DE TRATAMIENTO DE AGUAS RESIDUALES

Resumen: Este trabajo de tesis se enmarca en un punto de encuentro de la Automática, la Inteligencia Artificial y la Biotecnología. El propósito es el desarrollo de una metodología de modelado y control basados en el enfoque de la lógica difusa. La primera parte del trabajo presenta una introducción a los principios y técnicas usados en las plantas actuales de tratamiento de aguas residuales y resalta la dificultad del modelado de los múltiples fenómenos presentes. A partir de este hecho, la segunda parte centrada en el modelado y control difusos, desarrolla inicialmente la identificación a partir de datos de entrada-salida de modelos difusos afines de tipo Takagi-Sugeno (TS). Se consideran distintos métodos de agrupamiento difuso (clustering) así como un método de aglomeración competitiva, robusto en presencia de ruido. Este tipo de enfoque “caja gris” permite una representación basada en reglas que aproxima la dinámica no lineal como la concatenación de submodelos localmente lineales bajo la forma de autoregresión no lineal (NARX). También se ha desarrollado una versión gráfica de la caja de herramientas FMID para el modelado difuso de sistemas. Así mismo, utilizando la filosofía de control de compensación paralela distribuida (PDC), se ha propuesto un control difuso subóptimo lineal cuadrático adaptado a la estructura del modelo difuso identificado. Finalmente la metodología general es probada y validada en simulación en un bioproceso aerobio de tratamiento de aguas residuales.

Palabras clave: Modelo difuso Takagi-Sugeno, identificación de sistemas, agrupamiento difuso, control difuso, compensación paralela distribuida, tratamiento de aguas residuales, procesos biológicos.