



AVERTISSEMENT

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

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

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

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

LIENS

Code de la Propriété Intellectuelle. articles L 122. 4

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

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

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

Ecole Doctorale IAEM Lorraine
DFD Automatique et Production Automatisée

THÈSE

Présentée pour l'obtention du grade de

Docteur de l'Université de Lorraine

(Spécialité Automatique, Traitement du Signal et Génie Informatique)

par

Mohamed Ghassane KABADI

Contribution à la Tolérance aux Défauts des Systèmes Complexes basée sur la Génération de Graphes Causaux

<i>Rapporteurs :</i>	Mireille BAYART	Professeur à l'Université de Lille
	Kamal MEDJAHÉ	Professeur à l'Ecole Nationale d'Ingénieurs de Tarbes
<i>Examineurs :</i>	Taha BOUKHOBZA	Professeur à l'Université de Lorraine
	Ghaleb HOBLOS	MdC, HDR à ESIGELEC, Rouen
<i>Invité :</i>	Maxime MONNIN	Chef de Projets R&D chez Vize Software
<i>Directeurs de thèse :</i>	Frédéric HAMELIN	Professeur à l'Université de Lorraine
	Christophe AUBRUN	Professeur à l'Université de Lorraine

À ceux que j'aime,
À mes parents,
Source de ma joie, secret de ma force et essence de mon existence ;
À Nadia et Nissane,
Exemples de persévérance, de courage et de générosité ;
À ma femme,
Sans parler, elle me comprend.

Remerciements

J'aimerais tout d'abord remercier les membres du jury qui m'ont fait l'honneur de participer à l'examen de ce travail.

Je remercie Mme. **Mireille Bayart**, Professeur à l'Université de Lille, pour avoir accepté de rapporter sur ce travail de thèse et de l'avoir examiné minutieusement.

Je remercie M. **Kamal Medjaher**, Professeur à l'Ecole Nationale d'Ingénieurs de Tarbes, pour avoir accepté de rapporter sur ce manuscrit et pour ses remarques constructives.

J'adresse mes remerciements à M. **Ghaled Hoblos**, MdC, HDR à ESIGELEC, Rouen pour sa participation au jury, ses remarques pertinentes et son regard critique pour ce travail.

Je remercie vivement M. **Taha Boukhobza**, Professeur à l'Université de Lorraine, pour sa disponibilité et pour avoir été un excellent tuteur pendant ma période d'ATER à l'IUT Nancy-Brabois.

Je tiens à exprimer ma gratitude à mes directeurs de thèse, M. **Frédéric Hamelin** et M. **Christophe Aubrun**, Professeurs à l'Université de Lorraine, pour leurs directives scientifiques, leurs implication émotionnelle et relationnelle, de m'avoir également supporté tout au long de cette thèse.

Je tiens également à remercier **Armand Schauder**, **Thomas Desclozeaux**, **Erdogan Tantan**, **Frédéric Schnee** et **Anne-Marie Bresch** pour leur accueil chaleureux chez Lilly France, leur professionnalisme et leur soutien indéfectible. Je remercie également **Fulvio Rovetta** pour les bons moments qu'on a passés ensemble à travailler, à échanger et surtout à rigoler.

Merci à **Frédéric** pour sa bonne humeur en toute circonstance (Restes comme ça)!! Un grand merci à l'aimable **Sabine** pour sa disponibilité à résoudre tous nos problèmes administratifs et bien sûr pour son esprit joyeux et bienveillant. Merci à mes amis, **Abderrahmane**, **Hassan**, **Yacine**, **Krishnan**, **Florian** et sa femme.

Je tiens à remercier vivement mes parents **Abdellah** et **Amina**, ma grande et adorable sœur **Nadia** et, également mon grand et admirable frère **Nissane** pour leurs amour, encouragement, soutien, etc. Je remercie toute ma famille pour m'avoir soutenue depuis mes premiers pas jusqu'à l'obtention du doctorat.

Je remercie également M. **Mounir Rifi**, Professeur à l'Université Hassan II de Casablanca pour son professionnalisme, son amabilité et surtout, pour sa grande humilité et un savoir-être extraordinaire.

Enfin, la dernière mais pas la moindre, **Manal**. Un grand merci pour tout ce qu'elle a fait pour moi, ses encouragements, son soutien et ses interminables relectures du manuscrit tout au long de la thèse et surtout d'avoir supporté mes moments difficiles.

Table des matières

Abréviations	7
Introduction générale	9
1 Projet PAPYRUS	15
1.1 Introduction	15
1.2 Architecture et concept de la plateforme PAPYRUS	18
1.3 Description de la plateforme PAPYRUS	20
1.4 Objectifs du projet et description des lots de travail	20
1.5 Travaux de thèse dans le cadre du projet PAPYRUS	24
1.6 Conclusion	25
2 Diagnostic des systèmes complexes par approche graphique	27
2.1 Introduction	27
2.2 Diagnosticabilité du système à base de modèles structurés	28
2.2.1 Modélisation par graphe orienté	28
2.2.2 Diagnosticabilité structurelle	33
2.2.3 Conclusion	37
2.3 Diagnostic des systèmes sans connaissance à priori d'un modèle	38
2.3.1 Modélisation graphique des relations de causalité	38
2.3.1.1 Inter-corrélation	40
2.3.1.2 Causalité de Granger	51
2.3.1.3 Transfert d'Entropie	59
2.3.2 Matrice de causalité et représentation graphique	68
2.4 Représentation des liens de causalité :	69
2.4.1 Matrice d'adjacence	71
2.4.2 Graphe orienté	75
2.5 Diagnostic et analyse des causes basés sur un modèle de causalité	77
2.5.1 Introduction à l'approche Top-Down pour le diagnostic	77
2.5.2 Caractérisation des dégradations statistiques des performances du système	79
2.5.3 Test statistique de détection de dégradation	81
2.5.3.1 Introduction	81
2.5.3.2 Test de Neyman-Pearson	82

2.5.3.3	Test séquentiel de vraisemblance (Test de Wald)	85
2.5.4	Analyse des causes	89
2.5.4.1	Procédure d'analyse des causes	89
2.5.4.2	Application au procédé de fabrication de papier	93
2.6	Conclusion	107
3	Système tolérant aux défauts : Approche conjointe par gestion de références et analyse graphique	109
3.1	Introduction	109
3.2	Représentation topologique du système	111
3.2.1	Matrice d'adjacence du système	111
3.2.2	Graphe de causalité	113
3.3	Accommodation aux défauts par apprentissage itératif	115
3.4	Synthèse d'un gouverneur de références	118
3.4.1	Définition de l'indice de performance	118
3.4.2	Méthode du gradient pour l'ajustement d'une consigne	119
3.5	Analyse graphique pour la tolérance au défaut	120
3.5.1	Graphe associé des composantes	120
3.5.2	Analyse d'atteignabilité	122
3.5.2.1	Classification des consignes	123
3.5.2.2	Composantes impactées par une sélection de consignes	124
3.5.3	Diagramme de Hasse	126
3.5.3.1	Théorie des ensembles : Définitions	126
3.5.3.2	Diagramme de Hasse	126
3.5.3.3	Diagramme de Hasse pondéré	128
3.6	Stratégie de la tolérance au défaut adoptée	132
3.7	Cas d'étude : Application à un système de chauffage-mixage	135
3.7.1	Présentation du système	135
3.7.2	Représentation graphique du système	136
3.7.3	Identification du diagramme de Hasse pondéré	137
3.7.4	Définitions des indicateurs de performance	139
3.7.5	Modes de fonctionnement	140
3.7.6	Sélection et Ajustement des consignes	142
3.8	Conclusion	146
	Conclusion et perspectives	149

Abréviations

<i>PAPYRUS</i>	<i>Plug And Play monitoring and control architecture for optimization of large scale production processes</i>
<i>COD</i>	Coût-Qualité-Délai
<i>NTIC</i>	Nouvelles Technologies de l'Informatique et de la Communication
<i>SOA</i>	Architecture Orientée Service (<i>Service Oriented Architecture</i>)
<i>CIM</i>	<i>Computer Integrated Manufacturing</i>
<i>PAM</i>	Gestion d'Actifs du Procédé (<i>Plant Asset Management</i>)
<i>R&D</i>	Recherche et Développement
<i>WP_i</i>	<i>Work-Package i</i> (lot de travail <i>i</i>)
<i>PPI</i>	Indicateur de Performance Procédé (<i>Process Performance Indicator</i>)
<i>KPI</i>	Indicateur de Performance Clé (<i>Key Performance Indicator</i>)
<i>LPI</i>	Indicateur local de performance (<i>Local Performance Indicator</i>)
<i>CG</i>	Causalité de Granger
<i>AR</i>	<i>AutoRegressive model</i>

<i>ARX</i>	<i>AutoRegressive model with eXogenous Input</i>
<i>CIA</i>	Critère d'Information Akaike
<i>CIB</i>	Critère d'Information Bayésien
<i>TE</i>	Transfert d'Entropie
<i>KBS</i>	<i>Knowledge-Based System</i>
<i>PFD</i>	<i>Process Flow Diagram</i>
<i>P&ID</i>	<i>Process and Instrumentation Diagram</i>
<i>LSL</i>	<i>Lower Specification Limits</i>
<i>USL</i>	<i>Upper Specification Limits</i>
<i>SPRT</i>	Test Séquentiel du Rapport de Vraisemblance
<i>TD</i>	Temps de Détection
<i>POSET</i>	<i>Partially Ordred Set</i>

Introduction générale

De nos jours, l'industrie joue un rôle majeur dans l'économie européenne nécessitant de ce fait un renouveau constant de ses outils de production. Cependant, cette modernisation liée à l'utilisation de nouvelles technologies, qui a certes pour objectif d'améliorer durablement la qualité des produits, rend les systèmes de plus en plus complexes et sophistiqués.

De plus, un système de production est contraint de produire différents types de recettes (produits) afin de satisfaire les besoins des différents clients. Ceci mène bien souvent vers un système de production de grande dimension induisant un niveau de complexité supplémentaire. Ces aspects rendent le système vulnérable aux défauts ce qui nécessite inéluctablement la mise en œuvre d'outils de diagnostic dont le rôle est de satisfaire une demande accrue de fiabilité, de disponibilité et de sécurité du système en question.

Ce travail de thèse s'inscrit dans le cadre d'un projet européen : projet PAPYRUS développé afin de répondre aux problématiques liées au développement de nouvelles méthodes de modélisation, de diagnostic et de commandes tolérantes aux défauts.

Les notions de diagnostic et de tolérance aux défauts développées dans le projet PAPYRUS, que nous détaillerons plus tard, se basent sur une étape préliminaire et nécessaire qui est la modélisation du système, c'est à dire l'obtention d'une représentation mathématique de la structure et du comportement du système. Cependant, l'obtention d'un modèle mathématique précis est une étape fastidieuse compte tenu des phénomènes physiques souvent complexes intervenant au sein des systèmes. Le problème principal peut être relaxé

en envisageant l'utilisation d'un modèle simple du système mais suffisamment précis pour décrire sa structure et son comportement.

L'obtention d'un modèle graphique exhibant les liens de causalité entre les variables du système à partir des grandeurs mesurées est l'une des tâches du projet PAPYRUS. Le modèle de causalité ainsi obtenu représente la topologie du système.

De nombreuses méthodes ont été développées pour établir les relations entre les variables du système. Granger (1988) s'appuie sur la définition des relations de causalité entre les variables. Dans Faghraoui et al. (2012), les liens sont déterminés à partir du calcul d'un coefficient d'inter-covariance. Enfin, nous citons les travaux majeurs de Schreiber (2000) qui proposent de mettre en évidence les liens structurels à partir du transfert d'entropie entre les différentes variables.

Deux types de modélisation d'un système sont considérés dans notre étude. Le premier type de modélisation est basé sur un modèle de causalité représentant les liens de causalité entre des indicateurs de performances du système de production. Ces indicateurs sont hiérarchisés principalement en deux niveaux : le haut et le bas niveau. Le haut niveau est constitué d'indicateurs de performances clés économiques, de qualité et de disponibilité du système. Le bas niveau est constitué d'indicateurs de performances locaux caractérisant le fonctionnement de chaque sous-système. Le modèle obtenu représente ainsi les interconnexions entre les différentes entités du système qui ont un impact sur ces indicateurs de performances clés. Ce modèle peut donc être dédié à la phase de diagnostic du système.

Le second type de modélisation exhibe les interconnexions entre les différentes variables mesurées du système. Elles correspondent aux variables de commande (consignes) et aux variables de sortie régulées. Ce modèle ainsi obtenu est utilisé à des fins de reconfiguration.

Le but du diagnostic de défauts dans un système, à l'instar du domaine médical, consiste à évaluer le système en comparant ses données courantes mesurées aux données pro-

venant d'un fonctionnement nominal. Ces indicateurs permettent en général, de déterminer des symptômes amenant alors la détection et l'isolation des défauts du système. Le modèle de causalité utilisé pour le diagnostic permet de détecter le défaut ainsi que son impact sur les sous-systèmes dont les performances sont associées aux indicateurs de performances.

Un des avantages majeurs de l'utilisation d'un tel modèle, combiné aux tests statistiques appropriés, est la possibilité d'identifier la ou les cause(s) candidate(s) du défaut en ligne, d'isoler le ou les sous-système(s) en défaut et en déduire une stratégie de tolérance aux pannes ou le cas échéant, établir un plan de maintenance en concordance avec la stratégie d'optimisation de la production.

La phase de détection et de localisation des défauts est certes nécessaire mais insuffisante pour garantir la sûreté de fonctionnement du système. Il est donc nécessaire de considérer une seconde phase de tolérance aux défauts dans le but de garantir un fonctionnement acceptable en cas d'occurrence de défauts. Le modèle de causalité permettra, suivant une stratégie d'accommodation ou de reconfiguration par gestion de références, d'évaluer d'abord les possibilités d'accommodation aux défauts et ensuite d'optimiser les consignes des sous-systèmes afin d'atteindre les performances spécifiées en cas d'occurrence d'un ou plusieurs défauts.

Ce manuscrit s'articule autour de 3 principaux chapitres :

Le premier chapitre est dédié à la présentation du projet PAPYRUS. Les différents membres du consortium se sont fixés pour objectif de développer des approches de modélisation, de diagnostic, de reconfiguration et de stratégies de maintenance adaptées aux types de systèmes étudiés.

Ce premier chapitre présente également les différentes implications de chaque partenaire dans ce projet de façon non exhaustive, ceci afin de présenter de façon cohérente les différents lots de travail et leurs interconnexions.

Le deuxième chapitre est composé de deux parties principales. La première partie est consacrée aux méthodes de modélisation de causalité des systèmes. Plus précisément, les méthodes basées sur le modèle et celles basées sur les données. Pour les méthodes basées sur le modèle, un état de l'art a été établi ainsi qu'une présentation des systèmes structurés usuellement utilisés dans l'étude des propriétés structurelles comme l'observabilité, la commandabilité, la diagnosticabilité, etc. Un exemple d'illustration est fourni afin de comprendre l'outil d'analyse graphique utilisé. Pour les méthodes basées sur les données, trois principales méthodes sont présentées. Il s'agit de l'inter-corrélation, la causalité de Granger et le transfert d'entropie. Un exemple d'illustration est également présenté pour chacune de ces méthodes.

Le modèle de causalité établi à partir de ces méthodes permet la détection des défauts pour les systèmes étudiés à partir d'un test d'hypothèses séquentiel. Ainsi, en l'occurrence d'un défaut, une procédure d'analyse des causes est établie afin de définir l'ensemble des chemins de propagation du défaut à travers le système.

Dans le troisième chapitre, nous nous intéressons aux systèmes tolérants aux défauts. Nous présentons, en premier lieu, la topologie du système qui est basée sur une matrice d'adjacence associée à un graphe de causalité orienté. Ce graphe est ensuite réorganisé pour représenter la structure des différentes boucles de régulation du système. Après avoir représenté graphiquement cette structure des boucles de régulation, une stratégie de sélection des consignes exogènes est effectuée en utilisant le diagramme de Hasse. L'ajustement des consignes exogènes sélectionnées doit permettre de compenser l'effet du défaut survenu au niveau d'une des boucles de régulation. Cette stratégie d'ajustement des consignes repose sur une méthode d'apprentissage itératif permettant de s'affranchir de la méconnaissance de l'amplitude du défaut. Un cas d'étude d'un système de chauffage-mixage est présenté en dernière partie de ce chapitre.

Enfin, un ensemble de conclusions permettent de faire la synthèse des différentes méthodes présentées dans ce manuscrit. L'identification des limites de ce travail permet de proposer des perspectives et de nouvelles idées pour les futurs travaux.

Publications scientifiques :

Ahmed Faghraoui, **Mohamed Ghassane Kabadi**, Taha Boukhobza, Dominique Sauter et Christophe Aubrun : "Data-Driven Causality Digraph Modeling of Large-Scale Complex System Based on Transfer Entropy", IEEE Multi-conference on Systems and Control MSC, Nice/Antibes, France, 2014

Mohamed Ghassane Kabadi, Frédéric Hamelin, Christophe Aubrun et Taha Boukhobza : "Discrete mode observability analysis of switching linear systems with unknown inputs. A graph-theoretic approach", 10^{ème} European Workshop on Advanced Control and Diagnosis ACD, Copenhagen, Denmark, 2012.

Ahmed Faghraoui, **Mohamed Ghassane Kabadi**, Naim Kosayyer, David Morel, Dominique Sauter et Christophe Aubrun : " SOA-based platform implementing a structural modelling for largescale system fault detection : application to a board machine", IEEE Multi-conference on Systems and Control MSC, Dubrovnik, Croatia, 2012.

Taha Boukhobza, Frédéric Hamelin, **Mohamed Ghassane Kabadi** et Samir Aberkane : " Discrete mode Observability of switching linear systems with unknown inputs. A graph-theoretic approach", 18^{ème} IFAC World Congress, Milano, Italie, 2011 .

Mohamed Ghassane Kabadi, Taha Boukhobza et Frédéric Hamelin : "Discrete mode observability analysis of switching linear systems with unknown inputs. A graph-theoretic approach", Conférence Méditerranéenne sur l'Ingénierie Sûre des Systèmes Complexes MISC, Agadir, Maroc, 2011.

Fabien Clanché, **Mohamed Ghassane Kabadi** et Frédéric Hamelin : " Plate-forme pour l'optimisation énergétique des habitats intelligents", CETSIS'2011, Trois-Rivières, Canada, 2011.

Chapitre 1

Projet PAPYRUS

1.1 Introduction

Actuellement, l'industrie joue un rôle essentiel dans l'économie européenne dont la conjoncture financière est difficile. Cette situation pousse les industriels à investir dans de nombreux produits et services innovants et dans des projets porteurs de progrès techniques et technologiques afin d'améliorer davantage la compétitivité industrielle. Cependant, différents facteurs socio-économiques et environnementaux durant la récession économique actuelle, notamment la mondialisation et la concurrence, imposent plus d'exigences en termes d'efficacité et de sûreté de fonctionnement des procédés industriels.

Afin de limiter l'impact de cette récession, les industriels ont concentré leurs efforts sur l'amélioration continue de la rentabilité "économique" en optimisant la production qui est contrainte par divers besoins et exigences en termes de Coût-Qualité-Délai (CQD). La réponse aux exigences des utilisateurs finaux pousse, par le biais du triplet CQD, les industriels à améliorer et à rechercher de nouveaux concepts de gestion globale et à intégrer des procédés de production souvent complexes et de grande dimension. Cela, ouvre ainsi de nouvelles opportunités de recherche et de développement dans le domaine de l'automatique dont l'objectif est d'augmenter la sûreté de fonctionnement des systèmes et ainsi accroître son rendement économique.

Les systèmes complexes sont composés de multiples briques techniques et technologiques en interaction éventuelle avec des opérateurs ou des utilisateurs. Ces systèmes prennent en compte de nombreuses informations pour réaliser une opération complexe de façon plus ou moins automatique. La complexité des systèmes est généralement transparente pour l'utilisateur final, mais ses défaillances peuvent avoir de lourdes conséquences humaines, sociales ou économiques. La conception et le fonctionnement des sous-systèmes qui composent le système complexe font intervenir différents corps de métiers qu'aucun d'entre eux ne pourrait appréhender dans son ensemble Murthy and Krishnamurthy (2009).

Lors de la phase de conception d'un système complexe, son architecture doit être adaptée à la fonction recherchée et au degré de fiabilité et de sûreté nécessaire (y compris par l'introduction de redondances). De plus, la diversité des technologies et des compétences mises en œuvre doit être prise en compte dès la phase de conception, ce qui suppose des outils génériques assurant les liens entre plusieurs de ces technologies et/ou de ces compétences ainsi que des modélisations multi-domaines du même système.

Un système complexe de grande dimension peut être aussi caractérisé par un large nombre de variables (latentes, mesurables), de non linéarités et d'incertitudes. La décomposition en sous-systèmes plus petits, qui peut être organisé sous une forme hiérarchique, rend la gestion du système plus facile. Cependant, il est nécessaire de mettre en œuvre les mécanismes de décentralisation et de coordination efficaces assurant un fonctionnement homogène de l'ensemble des sous-systèmes en interaction et donc assurer le bon fonctionnement du système. La décomposition de ce type de systèmes nécessite souvent l'utilisation de nouvelles approches adaptées. De nouvelles méthodes de modélisation, de diagnostic et de tolérance aux défauts, dont quelques travaux sont cités dans les prochains chapitres, sont présentées en cohérence avec les objectifs du projet PAPYRUS.

L'arrivée des Nouvelles Technologies de l'Information et de la Communication (NTIC) a permis l'implémentation d'architectures ouvertes (par exemple SOA - Service Oriented Architecture) rendant possible l'intégration des méthodes et outils adaptés aux systèmes

complexes de grande dimension au sein des infrastructures technologiques adéquates.

Les exigences basées sur la disponibilité et la fiabilité sont d'une importance particulière pour les systèmes distribués de grande dimension, où un défaut local peut provoquer une dégradation de la performance globale du système, voire conduire à des pertes économiques considérables. Par conséquent, le développement et la synthèse de système tolérant aux défauts est essentiel pour l'industriel.

Le Projet PAPYRUS (*Plug and Play monitoring and control architecture for optimization of large scale production processes*) traite des nouvelles applications destinées aux systèmes d'automation. L'objectif est de développer un système complet de gestion des actifs industriels permettant d'améliorer les performances d'une usine du point de vue économique. Ce projet multipartite est sponsorisé par l'Union Européenne dans le cadre du 7^{ème} programme-cadre européen. Le consortium à l'origine du projet comprend sept partenaires issus de l'industrie : ABB AG en Allemagne, ABB Oy en Finlande, Stora Enso en Finlande, Predict en France, et du milieu universitaire : l'Université de Lorraine en France, l'Université Duisburg-Essen en Allemagne et AALTO University en Finlande. D'une durée de 36 mois, le projet a débuté le 1^{er} Septembre 2010.

Le projet a pour objectif d'améliorer et d'optimiser les activités de l'usine sans augmenter les coûts, ou en d'autres termes, de réduire les coûts d'exploitation tout en conservant le même niveau de performance. Le plan consiste à développer de nouvelles méthodes, algorithmes et outils informatiques tout en tenant compte de l'efficacité globale de l'usine.

Les applications logicielles à développer doivent fonder les diagnostics et traitements futurs sur des considérations de gestion d'actifs globale. Par exemple, sous forme d'un système de contrôle prévisionnel appliqué à toute l'usine. Les directeurs d'usine pourront alors évaluer l'état de fonctionnement de l'usine d'après les indicateurs définis, et éliminer les problèmes en se basant sur ces indicateurs et sur des systèmes d'aide à la décision. Les

actions à mettre en œuvre vont de simples opérations de maintenance à des reconfigurations des systèmes de contrôle. Il sera alors possible d'améliorer ou de stabiliser l'état d'une usine pour un surcoût minime voire nul. Ce nouveau système de gestion des actifs du procédé pour les usines industrielles (*PAM : Plant Asset Management*) est testé sur une machine à carton dans une papèterie Finlandaise détenue par Stora Enso. Ce nouveau système devrait permettre de faire fonctionner l'usine en utilisant une maintenance plus efficace.

La structure du système de gestion des actifs du procédé est détaillée dans la section suivante.

1.2 Architecture et concept de la plateforme PAPYRUS

L'évolution des technologies de communication a eu une forte influence sur l'évolution de la structure des systèmes industriels automatisés. Jusqu'à maintenant, le concept *Computer Integrated Manufacturing* (CIM) permet de décrire l'architecture des systèmes d'automatisation industriels. Dans cette structure pyramidale (hiérarchique), les niveaux de fonctionnalité sont identifiés de façon à ce que chaque dispositif soit conçu pour une tâche spécifique. Des réseaux spécifiques à chaque niveau de la pyramide sont utilisés pour connecter entre eux des composants du même niveau. Cependant, les dispositifs ont récemment évolué et ils commencent à inclure plus d'une fonction, ou d'un module, ce qui augmente le niveau d'intelligence de l'équipement.

Dans le cadre de la surveillance et du diagnostic des systèmes, les PAM (*Plant Asset Management*) permettent d'assurer des tâches élémentaires de diagnostic et d'aide à la maintenance au niveau des équipements de terrain. Ils restent insuffisants pour assurer un niveau de sûreté de fonctionnement acceptable. Ainsi, une dégradation de la qualité, pendant une longue période de production ou un arrêt de la production, peut entraîner des pertes économiques importantes. Il est donc nécessaire de mettre en œuvre une architecture tenant en compte l'ensemble des informations issues des modules de diagnostic, des objectifs économiques et des exigences des parties prenantes afin d'entreprendre le plus rapidement possible, des actions correctives adéquates pouvant accroître la disponibilité du système, et

assurer une qualité de production acceptable en l'occurrence d'un défaut durant la mission du procédé.

Le développement de ce type de système suscite un intérêt majeur dans le domaine de l'automatique et de la sûreté de fonctionnement. Nous citons notamment les principaux travaux Amadi-Echendu et al. (2010); Campbell et al. (2010); Lee et al. (2012) dans le cadre du développement des PAM. Ces travaux sont bien évidemment non-exhaustifs mais ils permettent de faire un tour d'horizon sur les approches utilisées en R&D dans le domaine de l'ingénierie, la maintenance, le diagnostic, le pronostic et la conception des systèmes d'aide à la décision.

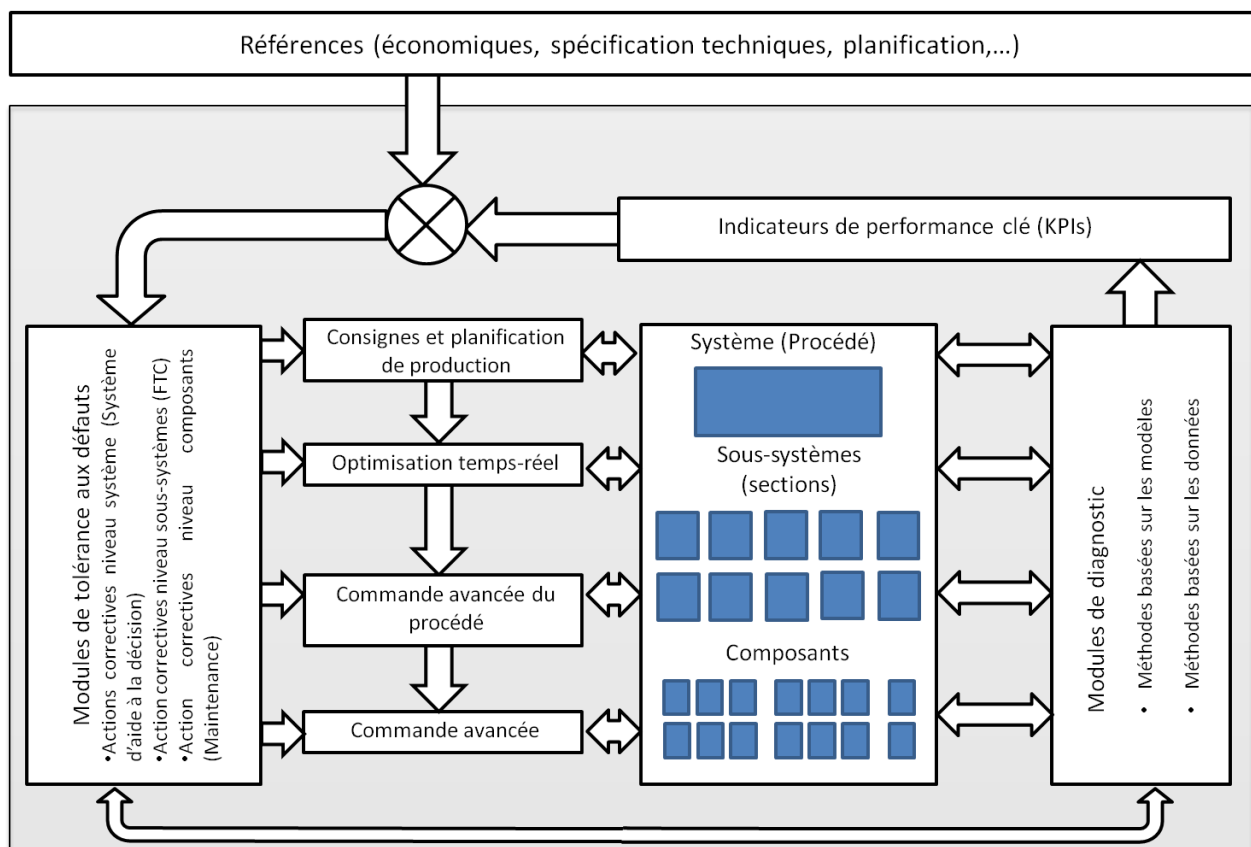


FIGURE 1.1 – Système de Gestion d'Actifs du Procédé (*Plant Asset Management*)

1.3 Description de la plateforme PAPYRUS

1.4 Objectifs du projet et description des lots de travail

Les travaux dans le cadre de PAPYRUS sont décomposés en plusieurs lots de travaux (*Workpackages*), de mise en œuvre et d'application industrielle (WP2-WP11) et de gestion de projet (WP1). Les Figures 1.1 et 1.2 présentent l'organisation de PAPYRUS avec les différentes activités réalisées au niveau de chaque tâche.

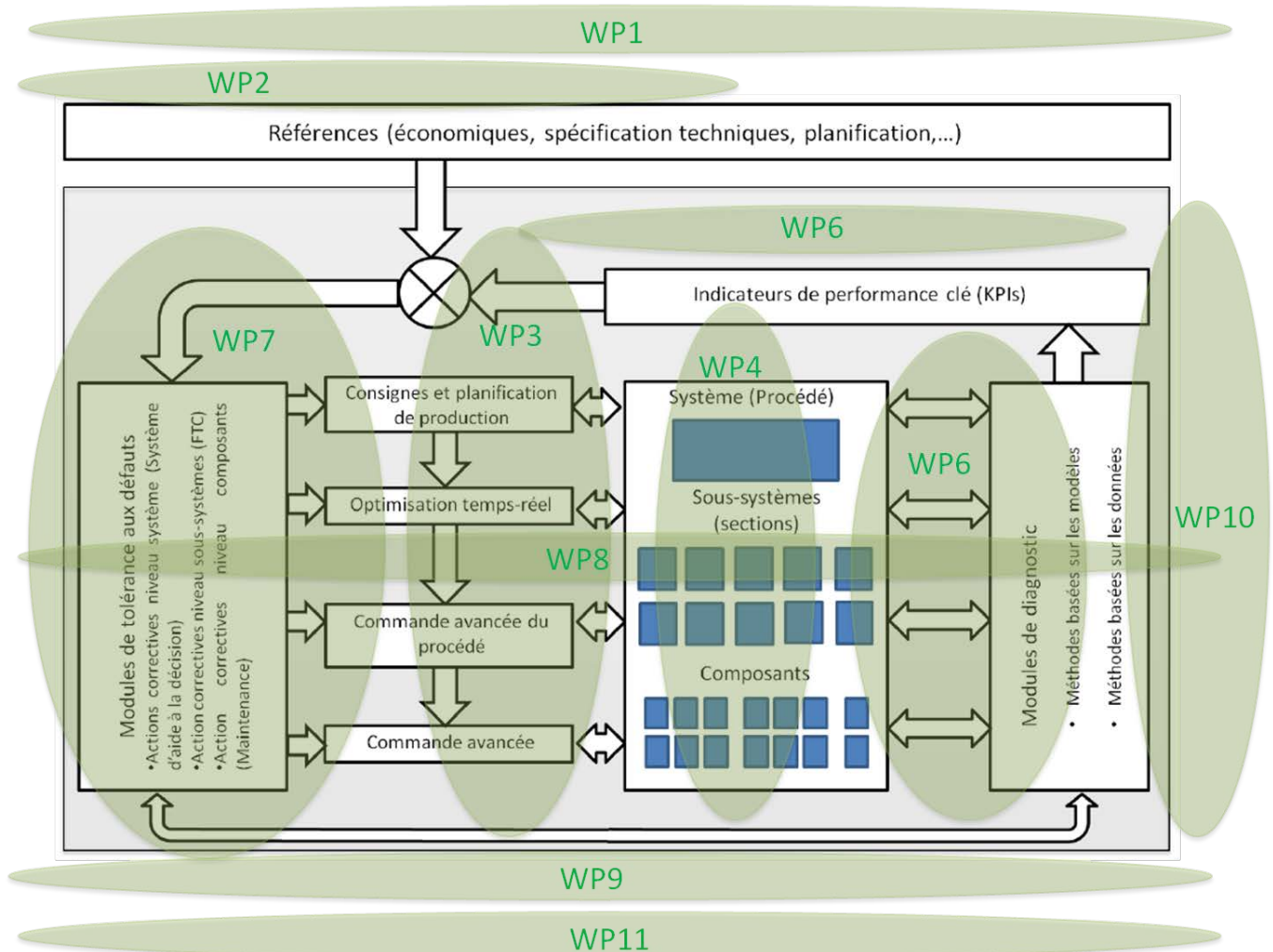


FIGURE 1.2 – Lots de travail

- WP1 (*Project management*) :

L'objectif de ce lot de travail est de surveiller, gérer, coordonner et prioriser les activités des partenaires universitaires et industriels, ainsi qu'initier des actions correctives en fonction du déroulement du projet dans le cas échéant. Il inclut aussi la gestion des risques, la surveillance et la conduite du projet. Un des objectifs majeurs de ce lot de travail est de fixer les résultats obtenus lors de la phase de développement et de recherche par l'ensemble des partenaires.

- WP2 (*Requirements*) :

L'analyse des exigences est la clé du développement et de l'amélioration des systèmes complexes à étudier. La définition des référentiels d'exigences est considérée comme un facteur de succès d'un processus. Le but majeur du projet est de pouvoir utiliser et adapter les résultats obtenus sur différents procédés industriels. L'application industrielle fournie par Stora Enso n'est qu'un système de production parmi d'autres (production de papier).

Ce lot de travail définit l'ensemble des exigences structurelles et fonctionnelles qui concernent tous les aspects du développement des systèmes de diagnostic, de contrôle/commande et de tolérance aux défauts.

- WP3 (*Architecture and platform*) :

WP3 spécifie la plate-forme logicielle de base qui sert à l'implémentation de l'ensemble des méthodes et outils développés par les partenaires durant le projet. La plate-forme logicielle servira pour les simulations et les tests ainsi que pour le procédé de fabrication de Stora Enso.

- **WP4 (*Functionnal and Structural Modeling*)**

Ce lot de travail joue un rôle majeur dans le choix des techniques et des méthodes de diagnostic et de tolérance aux défauts ultérieurs. Il concerne la modélisation fonctionnelle et structurelle du procédé industriel. Il définit le point de départ de la conception d'outils d'analyse hors-ligne de la plate-forme PAPYRUS. Ainsi, les méthodes permettant de modéliser l'architecture et la topologie du procédé, mettent en œuvre des modules de diagnostic de défauts sur différents niveaux du procédé et permettent de générer des actions correctives (système de tolérance aux défauts et/ou maintenance). Le but de la modélisation "structurelle" est de permettre une représentation, d'une part, qualitative exhibant les interdépendances entre les variables d'intérêt du procédé et d'autre part, quantitative visant à une amélioration continue du modèle en fonction des données disponibles du procédé. La modélisation fonctionnelle permet aussi de consolider le modèle. Cette étape permet de faire une validation croisée entre les différents modèles identifiés et de vérifier leur cohérence et leur intégrité. Le modèle fonctionnel est souvent utilisé afin de décomposer le système en plusieurs sous-systèmes. Ces derniers peuvent être ainsi hiérarchisés et facilement analysables. L'analyse fonctionnelle est souvent une première étape vers l'établissement d'un plan de maintenance préventive ou curative.

Les deux approches sont les points de départ du WP5, 6 et 7.

L'Université de Lorraine est principalement responsable de ce lot de travail WP4. Il est subdivisé en 7 tâches :

1. T4.1 : Définition des indicateurs de performance clé (Université de Duisburg-Essen) ;
2. T4.2 et T4.3 : Modélisation fonctionnelle et structurelle du système (Université de Lorraine) ;
3. T4.4 : Validation du modèle (Stora Enso) ;
4. T4.5 : Analyse d'impact des défauts (ABB) ;

5. T4.6 : Conceptualisation des actions correctives (Université de Lorraine et ABB) ;
6. T4.7 : Généralisation des modèles (Université de Lorraine).

- **WP5 (*Fault Detection and Diagnosis*)**

L'objectif principal du WP5 est le développement de méthodes de détection et de diagnostic des défauts (Modules FDD) adaptées aux systèmes complexes de grande dimension. Les méthodes FDD principalement développées sont basées sur les modèles. Ainsi, ces méthodes permettent de fournir des informations sur la nature du défaut ainsi que sur son impact sur les performances du système.

- **WP6 (*Condition and Performance Prognosis*)**

Les objectifs majeurs du WP6 sont :

- Le développement des algorithmes d'estimation/prédiction des KPI ;
- Le développement des méthodes de pronostic basées sur KPI estimé/prédit ;
- L'intégration des algorithmes au niveau de la plate-forme PAM.

- **WP7 (*Corrective Action*)**

WP7 concerne le développement d'un système d'aide à la décision pour le système tolérant aux défauts basé sur les résultats de diagnostic/pronostic de l'état courant/prédit du système. La prise de décision pour le choix d'une action corrective spécifique dépend de chaque niveau du procédé (niveau composant, niveau section ou/et niveau procédé). Ainsi, les actions correctives possibles sont analysées, et des stratégies de prise de décisions sont développées séparément en fonction de chaque niveau du procédé.

- WP8 (*Software Platform*)

L'objectif du WP8 est de mettre en place une plateforme logicielle où les principaux résultats des lots de travail précédent sont implémentés. Cette plateforme permet de vérifier l'étude et de valider l'ensemble des résultats obtenus par les différents partenaires du projet.

- WP9 (*Integration*)

L'objectif du WP9 est l'intégration de l'ensemble des concepts et méthodes développés dans WP4, 5, 6 et 7 dans la plateforme PAPYRUS.

- WP10 (*Application to paper board machine*)

L'objectif du WP10 est la mise en œuvre des essais sur la machine de production de papier de Stora Enso.

Les méthodes, développées durant le projet, sont implémentées et validées en temps réel sur le système industriel. Les résultats des tests et des évaluations sont reportés afin de permettre un raffinement des méthodes implémentées.

- WP11 (*Dessimination and exploitation*)

L'objectif du WP11 est de vérifier l'adéquation et l'homogénéité de la solution globale fournie par les partenaires par rapport aux besoins et exigences définis au début du projet.

1.5 Travaux de thèse dans le cadre du projet PAPYRUS

Le travail de thèse réalisé dans le cadre du projet PAPYRUS concerne principalement le développement des approches permettant de générer un modèle de causalité décrivant la

structure du système étudié (WP4). Ce modèle doit respecter les exigences formulées dans le projet PAPYRUS, notamment :

- ▷ Le modèle doit permettre une simple visualisation (graphique) de la structure, et de l'architecture du système ;
- ▷ Le modèle doit permettre de manipuler des modèles incertains avec des paramètres inconnus pour des systèmes complexes de grande dimension ;
- ▷ La génération du modèle doit être conforme au principe du Plug & Play, et à l'architecture SOA (architecture orientée services).

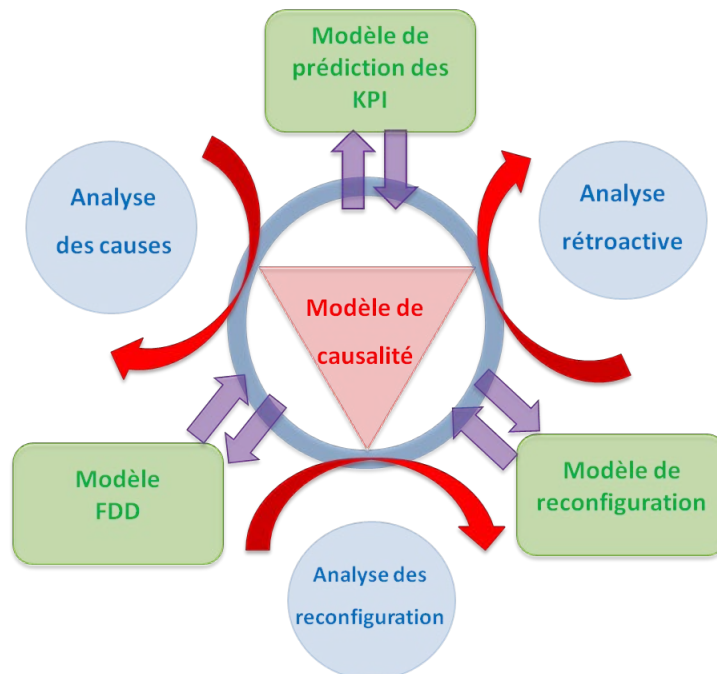


FIGURE 1.3 – Le modèle structurel : modèle pivot

Le modèle de causalité est considéré comme étant un modèle pivot (Figure 1.3) supportant l'ensemble des lots de travail WP5, WP6 et WP7. Ce modèle est adapté ensuite aux contraintes de chaque lot de travail qui sont détaillées dans la section précédente.

1.6 Conclusion

La modélisation, la détection des défauts, le diagnostic et la pronostic sont des éléments clés pour le développement des méthodes dans le projet PAPYRUS pour un PAM dédié aux

systèmes de grande dimension. Étant donné que ces systèmes sont trop complexes pour le premier principe de modélisation, d'autres façons d'obtenir la structure du système doivent être élaborées pour créer les liens entre les PPI et les composants du procédé. Dans ce projet PAPYRUS, l'utilisation des modèles structurels est explorée. Un des objectifs clés de ce projet est de savoir comment intégrer ces modèles dans le concept de PAM et comment les exploiter dans le diagnostic et la détection de défauts pour les systèmes à grande dimension.

Chaque défaillance d'un composant affecte les KPI/PPI d'une manière différente par rapport à une autre. Connaître l'impact "futur" de la défaillance d'un composant est une base importante pour la planification des futurs efforts d'optimisation. Dans le but de minimiser la dégradation des KPI/PPI, il n'est possible de prendre une décision par rapport à des actions correctives possibles que si l'impact de ces actions est connu.

PAPYRUS développe des boîtes à outils pour la surveillance, le diagnostic, le pronostic et le traitement visant à maintenir les niveaux de performances KPI/PPI souhaités. Pour démontrer et valider les possibilités de cette approche, les méthodes développées seront appliquées à un exemple de système de grande dimension. Ce système est connecté à une plate-forme de logiciels où toutes les données sont stockées et accessibles pour les boîtes à outils. Cette mise en œuvre se fera suivant une architecture orientée services de telle façon que la solution sera réutilisable pour d'autres applications et d'autres systèmes.

Chapitre 2

Diagnostic des systèmes complexes par approche graphique

2.1 Introduction

La modélisation statistique et les méthodes d'extraction de données jouent un rôle de plus en plus important, particulièrement, dans les applications qui impliquent prévision et prédiction. Cependant, il est rarement suffisant de simplement découvrir les corrélations statistiques qui existent dans les données. Il est en effet très important de déterminer les relations causales entre les variables du système étudié. La modélisation qui en résulte représente le système et le comportement de ses variables en termes de causalité. Cette causalité peut être représentée graphiquement où les arcs correspondent au sens de la causalité. Le graphe dit “de causalité” est utilisé dans plusieurs applications telles que le diagnostic, la reconfiguration, etc.

Cette étape de la modélisation causale fait l'objet de nombreuses recherches. Plus particulièrement, dans le cadre des réseaux Bayésiens, on peut citer les travaux de Heckerman (1995); Meek (1996); Pearl (2000); Friedman et al. (1999); en plus des graphes de causalité de Silva et al. (2006); Kalisch et al. (2005); Chu and Glymour (2006); Moneta and Spirtes (2006). Les graphes orientés signés ont également été reconnus comme des outils

appropriés afin d'étudier des problèmes de sûreté de fonctionnement des systèmes complexes (Lü and Wang (2007); Pearl (2000); Ram Maurya et al. (2004); Maurya et al. (2007); GAO et al. (2010); Maurya et al. (2006)).

Dans de nombreux systèmes industriels, la plupart des données disponibles sont des mesures temporelles. Ces informations temporelles sont souvent descriptives de la structure causale entre les variables du système. Différents travaux sont dédiés à la modélisation causale utilisant ces données temporelles (Chu and Glymour (2006); Eichler (2013)).

La causalité est une notion difficile à définir. Elle a été étudiée dans de nombreuses et différentes disciplines, y compris la physique, la philosophie, les statistiques et l'informatique. L'aspect temporel est généralement une partie essentielle de notre compréhension intuitive de la causalité, car on juge souvent nécessaire que la cause doit avoir existé dans le temps avant l'effet.

2.2 Diagnosticabilité du système à base de modèles structurés

Dans cette partie, nous allons tout d'abord présenter un outil graphique permettant de matérialiser les liens entre les variables d'un système. Des approches d'analyse de propriétés structurelles sont ensuite présentées. L'accent est mis sur les propriétés structurelles utiles au diagnostic de défauts.

2.2.1 Modélisation par graphe orienté

La modélisation d'un système est, de manière classique, basée sur une série d'équations mathématiques obtenue grâce à l'étude de son comportement physique. Cette modélisation permet d'appréhender le comportement de ce système en fonction d'une certaine commande ou de déterminer la commande permettant d'obtenir un comportement désiré pour ce système.

Dans le cas d'une représentation d'état, le modèle lie les variables d'état, d'entrée, de sortie et de défaut du système. Ces variables sont respectivement groupées dans les vecteurs $x(t) = (x_1(t), \dots, x_n(t)) \in \mathbb{R}^n$, $u(t) = (u_1(t), \dots, u_m(t)) \in \mathbb{R}^m$, $y(t) = (y_1(t), \dots, y_p(t)) \in \mathbb{R}^p$ et $f(t) = (f_1(t), \dots, f_r(t)) \in \mathbb{R}^r$. Le modèle correspond à la forme suivante :

$$\begin{cases} \dot{x}(t) = Ax(t) + Bu(t) + Ef(t) \\ y(t) = Cx(t) + Du(t) \end{cases} \quad (2.1)$$

Les relations entre les variables d'état sont représentées par la matrice $A \in \mathbb{R}^{n \times n}$, celles entre les variables d'état et les entrées sont représentées par la matrice $B \in \mathbb{R}^{n \times m}$, celles entre les variables d'état et les défauts sont représentées par la matrice $E \in \mathbb{R}^{n \times r}$, celles entre les sorties et les variables d'état sont représentées par la matrice $C \in \mathbb{R}^{p \times n}$ et celles entre les sorties et les entrées sont représentées par la matrice $D \in \mathbb{R}^{p \times m}$. Nous supposons dans ce paragraphe que les défauts peuvent se présenter de manière additive.

Les valeurs numériques des matrices A , B , C , D et E ne sont que rarement connues surtout lors de la phase de conception du système. En revanche, la structure du système, c'est-à-dire la nature des variables et les relations de causalité qui les lient peuvent être plus facilement obtenus. Ainsi, nous étudions par la suite des systèmes structurés où les paramètres des matrices d'état sont donnés par des paramètres indépendants non nuls notés λ_i . Dans la représentation d'état d'un système structuré, les matrices d'état contiennent donc des paramètres λ_i correspondant aux liens de causalité. La présence de "0" traduit l'absence de lien de causalité. Le modèle de l'équation (2.1) est alors récrit comme suit :

$$\begin{cases} \dot{x}(t) = A^\lambda x(t) + B^\lambda u(t) + E^\lambda f(t) \\ y(t) = C^\lambda x(t) + D^\lambda u(t) \end{cases} \quad (2.2)$$

Ce type de modèle peut être représenté graphiquement par des graphes orientés représentant les variables du système modélisé et les relations les liant. Ces graphes contiennent toute l'information du modèle et permettent une analyse graphique efficace et de moindre

complexité qu'en utilisant le modèle d'état. De plus, cette représentation graphique permet la visualisation simple et intuitive de certaines propriétés structurelles du système représenté.

Un graphe orienté, noté $\mathcal{G}(\mathcal{V}, \mathcal{E})$, est composé de deux ensembles. Le premier ensemble, noté $\mathcal{V}$, contient les sommets du graphe. Le deuxième ensemble, noté $\mathcal{E}$, contient les couples ordonnés de sommets appelés arcs. L'ensemble des sommets $\mathcal{V}$ est associé aux variables du système. L'ensemble des arcs $\mathcal{E}$ représente l'existence de liens de causalité entre ces variables. Plus précisément, l'existence d'un arc liant deux sommets de $\mathcal{V}$ signifie que la variation d'une variable affecte la variation de l'autre en fonction du sens de cet arc.

Pour un système linéaire, l'ensemble de sommets $\mathcal{V}$ est donné par $\mathcal{V} = \mathbf{X} \cup \mathbf{U} \cup \mathbf{Y} \cup \mathbf{F}$ avec $\mathbf{X}$, $\mathbf{U}$, $\mathbf{Y}$ et $\mathbf{F}$ représentant respectivement les sommets associés aux variables d'état, d'entrée, de sortie et de défaut, et de l'ensemble d'arcs $\mathcal{E} = A_{arcs} \cup B_{arcs} \cup E_{arcs} \cup C_{arcs} \cup D_{arcs}$, avec $A_{arcs} = \{(\mathbf{x}_j, \mathbf{x}_i) | A(i, j) \neq 0\}$, $B_{arcs} = \{(\mathbf{u}_j, \mathbf{x}_i) | B(i, j) \neq 0\}$, $E_{arcs} = \{(\mathbf{f}_j, \mathbf{x}_i) | E(i, j) \neq 0\}$, $C_{arcs} = \{(\mathbf{x}_j, \mathbf{y}_i) | C(i, j) \neq 0\}$ et $D_{arcs} = \{(\mathbf{u}_j, \mathbf{y}_i) | D(i, j) \neq 0\}$.

À partir de cette représentation graphique, il est possible d'évaluer la validité de certaines propriétés dites structurelles telles que l'observabilité, la commandabilité ou la diagnosticabilité/tolérance aux défauts structurelles. Cette évaluation dépend de certains critères graphiques basés sur l'existence ou pas de chemins, de cycles, etc.

Exemple 1 Soit le système hydraulique représenté à la Figure 2.1.

Ce système est composé de deux cuves. Ces cuves sont alimentés par deux débits d'eau extérieurs q_1 et q_2 à une température T_0 . T_i , N_i et A_i correspondent respectivement à la température, le niveau et la section des deux cuves. La température de la première cuve est commandée par une résistance de chauffe H .

Ce système correspond aux équations (2.3) où C_1 et C_2 correspondent aux capacités calorifiques des deux cuves.

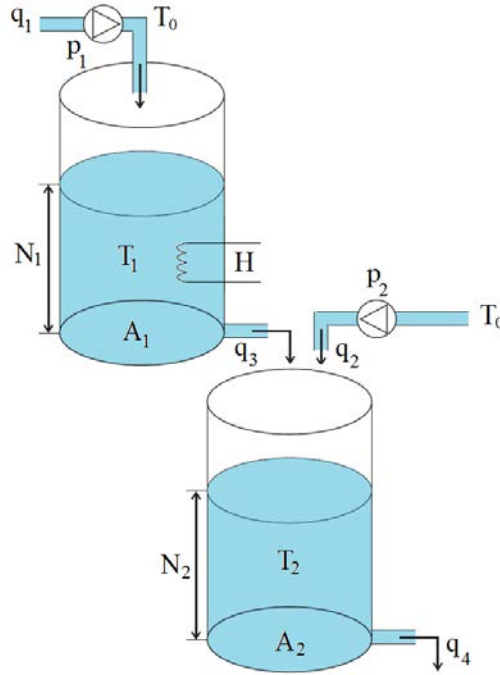


FIGURE 2.1 – Système hydraulique

$$\begin{cases} A_1 \dot{N}_1(t) &= -q_3(t) + q_1(t) \\ A_2 \dot{N}_2(t) &= q_3(t) - q_4(t) + q_2(t) \\ C_1 \dot{T}_1(t) &= q_1(t)T_0 - q_3(t)T_1(t) + H(t) \\ C_2 \dot{T}_2(t) &= q_3(t)T_1(t) + q_2(t)T_0 - q_4(t)T_2(t) \end{cases} \quad (2.3)$$

Soient deux défauts actionneurs f_1 et f_2 au niveau de la pompe p_1 et de la résistance de chauffe H . Pour ce système, nous considérons que le défaut f_1 affecte le niveau N_1 de la cuve 1 et, f_2 affecte la température T_1 de la cuve 1.

Après la linéarisation de ce système autour d'un point de fonctionnement, nous obtenons le modèle linéaire de la forme (2.2) avec $x = (x_1, x_2, x_3, x_4)^T = (N_1, N_2, T_1, T_2)^T$, $u = (u_1, u_2, u_3)^T = (q_1, q_2, H)^T$, $f = (f_1, f_2)^T$ et $y = (y_1, y_2)^T = (N_2, T_2)^T$. Les matrices correspondant à ce système structuré sont données par :

$$A^\lambda = \begin{pmatrix} \lambda_1 & 0 & 0 & 0 \\ \lambda_2 & \lambda_3 & 0 & 0 \\ 0 & 0 & \lambda_4 & 0 \\ 0 & 0 & \lambda_5 & \lambda_6 \end{pmatrix}; \quad B^\lambda = \begin{pmatrix} \lambda_7 & 0 & 0 \\ 0 & \lambda_8 & 0 \\ \lambda_9 & 0 & \lambda_{10} \\ 0 & \lambda_{11} & 0 \end{pmatrix}; \quad C^\lambda = \begin{pmatrix} 0 & \lambda_{12} & 0 & 0 \\ 0 & 0 & 0 & \lambda_{13} \end{pmatrix};$$

$$E^\lambda = \begin{pmatrix} \lambda_{14} & 0 \\ 0 & 0 \\ 0 & \lambda_{15} \\ 0 & 0 \end{pmatrix} \text{ et } D^\lambda = \begin{pmatrix} 0 & 0 & 0 \\ 0 & 0 & 0 \end{pmatrix}.$$

Le graphe orienté représentant ce système est donné à la Figure 2.2.

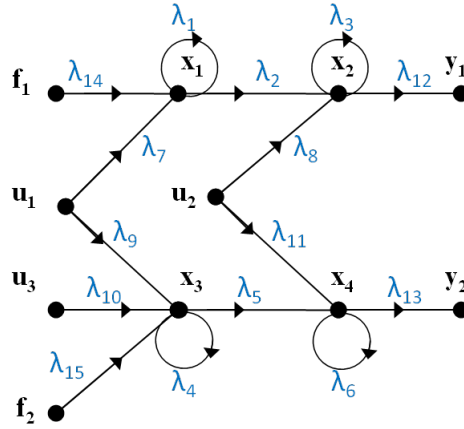


FIGURE 2.2 – Graphe orienté correspondant à l'Exemple 1

Quelques notations et définitions Dion et al. (2003); Boukhobza (2010) :

Dans un graphe orienté $\mathcal{G}(\mathcal{V}, \mathcal{E})$:

- Un chemin p qui couvre les sommets $\mathbf{v}_{r_0}, \dots, \mathbf{v}_{r_i}$ est noté $p = \mathbf{v}_{r_0} \rightarrow \mathbf{v}_{r_1} \dots \rightarrow \mathbf{v}_{r_i}$ où pour $j = 0, 1, \dots, i-1$, $(\mathbf{v}_{r_j}, \mathbf{v}_{r_{j+1}}) \in \mathcal{E}$.
- Des chemins sont disjoints s'ils n'ont aucun sommet commun.
- Un ensemble de k chemins disjoints $\mathcal{V}_1 - \mathcal{V}_2$ forme un lien $\mathcal{V}_1 - \mathcal{V}_2$. Les liens consistant en un nombre maximal de chemins sont dits liens maximaux. La taille d'un lien est le

nombre de ses chemins.

- $\rho(\mathcal{V}_1, \mathcal{V}_2)$ représente le nombre de chemins dans un lien $\mathcal{V}_1$ - $\mathcal{V}_2$ maximal ou autrement dit le nombre maximal de chemins disjoints partant de $\mathcal{V}_1$ et arrivant à $\mathcal{V}_2$. $\rho(\mathcal{V}_1, \mathcal{V}_2)$ est appelé la taille du lien $\mathcal{V}_1$ - $\mathcal{V}_2$.
- Soit un ensemble $\mathcal{V}_1$. $\text{card}(\mathcal{V}_1)$ représente la cardinalité de $\mathcal{V}_1$ *i.e.* le nombre d'éléments dans cet ensemble.

2.2.2 Diagnosticabilité structurelle

Afin de maximiser les fonctionnalités opérationnelles, l'efficacité de la mission ainsi que d'assurer un niveau de sécurité et de sûreté de fonctionnement de tout système, il est nécessaire que les défauts soient détectés aussi vite que possible après leur apparition. Ainsi, la capacité d'un système à détecter et à localiser un défaut, dite diagnosticabilité, est une étape fondamentale dans l'étude d'un système (Cocquempot (2004); Franck (1990); Dustegor et al. (2004); Zhang and Jiang (2006)).

Afin de présenter différents travaux relatifs à la diagnosticabilité des systèmes structurés, nous considérons un résidu conçu en utilisant un ensemble de r observateurs selon le schéma d'observateur présenté dans Chen and Patton (1999). Chaque résidu est conçu pour être sensible à un défaut unique tout en restant insensible aux autres défauts.

Pour un système de la forme 2.1 et dans le même esprit que Commault et al. (2008), le $i^{\text{ème}}$ des r observateurs est défini comme suit :

$$\dot{\hat{x}}^i(t) = A\hat{x}^i(t) + K^i(y(t) - C\hat{x}^i(t))$$

avec $\hat{x}^i(t) \in \mathbb{R}^n$ l'état du $i^{\text{ème}}$ observateur et K^i le gain de cet observateur tel que $\hat{x}^i(t)$ converge asymptotiquement vers $x(t)$ quand $f(t) = 0$. Ainsi, un vecteur de résidu noté $\rho(t)$ est généré par :

$$\rho_i(t) = Q^i(y(t) - C\hat{x}^i(t)), \text{ pour } i = 1, \dots, r$$

où Q^i est un vecteur de dimension p (p est la cardinalité de $y(t)$).

Le problème de diagnosticabilité structurelle de défauts Commault et al. (2008) consiste à vérifier l'existence des matrices K^i et Q^i pour $i = 1, 2, \dots, r$, de telle sorte que, $A - K^i C$ soit stable, et la matrice de transfert des résidus soit non nulle et diagonale *i.e.* le transfert des défauts $f(t)$ vers les résidus $\rho(t)$ est de la forme :

$$\begin{pmatrix} t_{11}(s) & 0 & \dots & 0 \\ 0 & t_{22}(s) & \dots & 0 \\ \dots & \dots & \ddots & \dots \\ 0 & 0 & \dots & t_{rr}(s) \end{pmatrix}$$

où $t_{ii}(s) \neq 0$ pour $i = 1, 2, \dots, r$.

Cette structure diagonale nous permet de détecter les défauts simultanément. Les conditions graphiques de solvabilité de ce problème sont détaillées dans le paragraphe suivant. Ces conditions expriment notamment la nécessité d'avoir un nombre suffisant de sorties y_i pour permettre la détection et la localisation des défauts.

En utilisant l'approche graphique, le problème de détection et de localisation des défauts peut être résolu selon le Théorème 1 proposé dans Commault et al. (2002).

Théorème 1 *Soit un système linéaire structuré observable soumis à r défauts et associé à un graphe orienté $\mathcal{G}(\mathcal{V}, \mathcal{E})$. Le problème de détection et de localisation des r défauts est soluble, pour presque toutes les valeurs des paramètres λ_i , si et seulement si : $\rho(F - Y) = r$ *i.e.* $\rho(F - Y) = \text{card}(F)$.*

Ainsi, la solvabilité du problème de détection et de localisation des défauts est associée à la condition graphique du Théorème 1. Cette condition consiste à avoir des chemins disjoints entre les sommets de défauts dans F et les sommets de sorties dans Y dont le nombre est égal à r .

Exemple 2 *Considérons le système linéaire de l'Exemple 1 représenté par son graphe orienté de la Figure 2.2.*

Pour ce système, nous avons deux défauts f_1 et f_2 . Ce qui correspond à $r = 2$. Nous avons également un lien maximal entre F et Y qui correspond aux deux chemins $\mathbf{f}_1 \longrightarrow \mathbf{x}_1 \longrightarrow \mathbf{x}_2 \longrightarrow \mathbf{y}_1$ et $\mathbf{f}_2 \longrightarrow \mathbf{x}_3 \longrightarrow \mathbf{x}_4 \longrightarrow \mathbf{y}_2$. Puisque ce lien maximal contient deux chemins, alors $\rho(F - Y) = 2$.

Ainsi, la condition graphique $\rho(F - Y) = \text{card}(F) = 2$ est vérifiée et le problème de détection et de localisation des défauts est soluble pour ce système linéaire.

Observabilité du mode discret pour les systèmes linéaires à commutations

Un système linéaire à commutations est un système qui peut admettre plusieurs modes de fonctionnement. Suite à des événements donnés, le système commute d'un mode à un autre où chaque mode est associé à un modèle. Un système linéaire à commutations combine donc une dynamique événementielle et temporelle.

Pour la détection et la localisation des défauts, les événements provoquant la commutation du système d'un mode à un autre peuvent être liés à l'occurrence de défauts.

C'est l'hypothèse qui est faite dans ce paragraphe, à savoir qu'un défaut consiste à un changement au niveau des paramètres des matrices du modèle du système. En considérant ces changements, plusieurs modèles du système sont considérés. Ainsi, chaque défaut correspond à un mode distinct du système représenté par un graphe.

La Figure 2.3 donne un exemple de système linéaire à commutations avec 3 modes. Le premier mode ne correspond à aucun défaut (structure relative au fonctionnement nominal). Le deuxième mode correspond à un défaut paramétrique au niveau de λ_2 qui change en λ_8 . Le troisième mode correspond à un défaut paramétrique au niveau de λ_5 qui change en λ_9 .

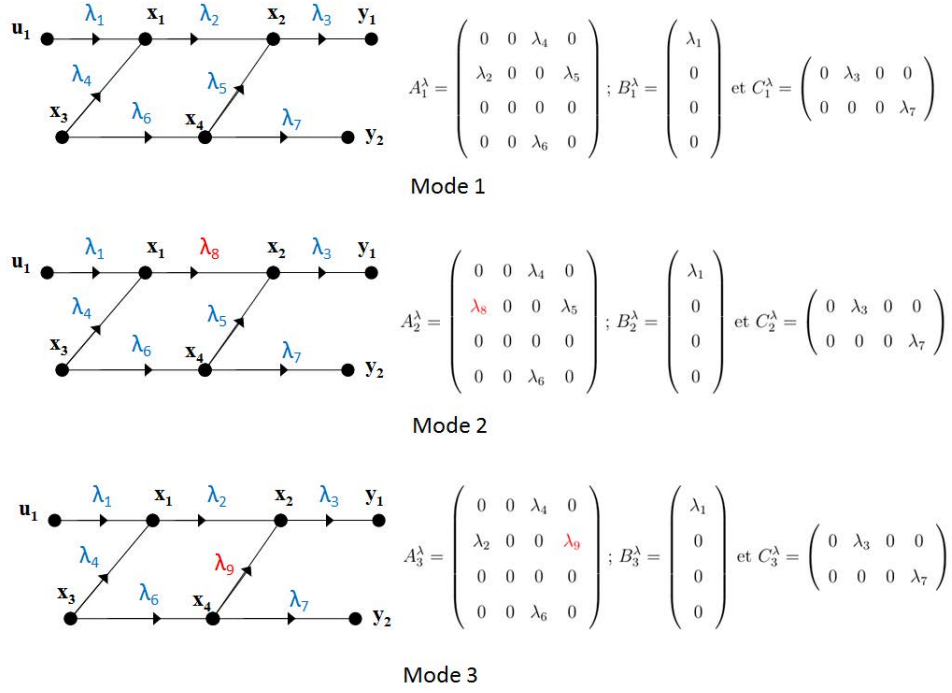


FIGURE 2.3 – Exemple de défauts paramétrique

Plusieurs travaux dans la littérature s'intéressent à l'observabilité du mode discret de systèmes à commutations. Dans Kabadi et al. (2012), nous nous intéressons à la formalisation des conditions suffisantes pour l'observabilité du mode discret en supposant que la structure du système est connue (graphe orienté du système). Cette observabilité permet de préciser le mode dans lequel le système se trouve suite à une ou plusieurs commutations. En utilisant les propriétés du graphe orienté associé au système, les conditions suffisantes permettent d'obtenir des critères de placement de capteurs afin de récupérer l'observabilité d'un mode discret. En considérant que chaque mode est supposé être associé à l'occurrence d'un défaut donné, la détection et la localisation des défauts peut être effectuée à travers l'observabilité du mode discret du système. Ainsi, le problème de détection et de localisation du défaut f_i revient à observer le mode discret correspondant à ce défaut.

Remarque 1 *Un défaut peut également être associé au changement d'un paramètre nul en un paramètre non nul. Dans ce cas, de nouveaux arcs apparaissent sur le graphe représentant le mode correspondant à ce défaut. La Figure 2.4 montre un exemple de système linéaire à commutations où les défauts provoquent le changement de la structure du système.*

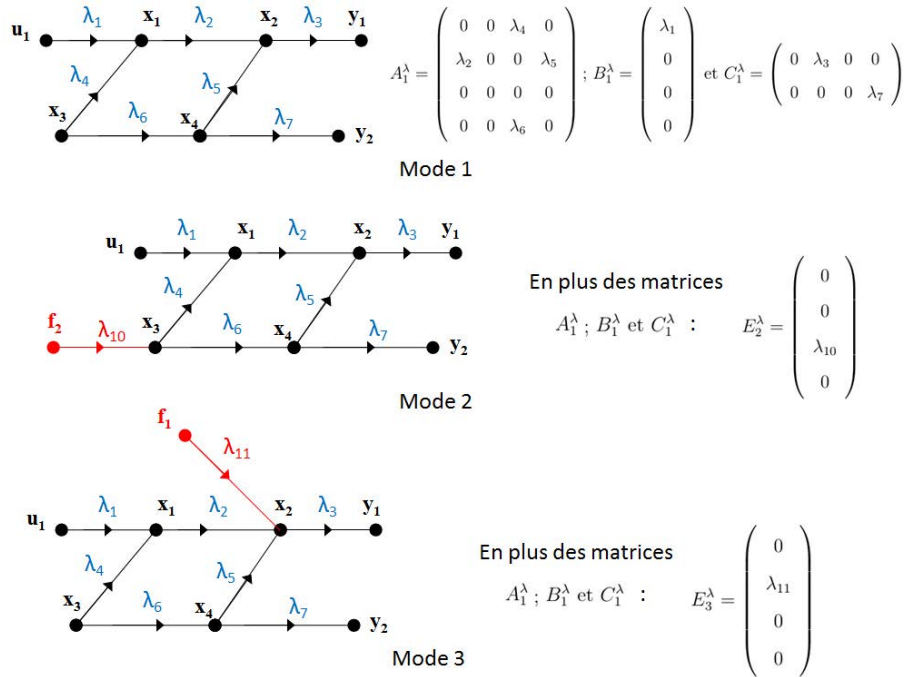


FIGURE 2.4 – Exemple de défauts structurels

Pour ces défauts dits structurels, l'analyse de la détection et de la localisation des défauts peut être effectuée de la même façon que dans l'exemple 2.

2.2.3 Conclusion

Les méthodes de diagnosticabilité présentées dans cette section sont basées sur un modèle structurel du système. Ce modèle correspond à la structure du système indépendamment de la connaissance des valeurs numériques des paramètres du modèle. Ceci représente un avantage pour l'analyse d'un système puisque les propriétés structurelles de ce dernier peuvent être traitées dès la phase de conception du système. Cette analyse en phase de conception permet une étude des propriétés du système en boucle ouverte.

Le modèle structurel considéré résulte de la connaissance physique que l'on a d'un système. Ce type de modèle est difficile à obtenir sachant que le graphe est généré à partir de la représentation d'état (2.2). Dans le cadre des systèmes complexes de grande dimension exhibant souvent des phénomènes physiques, la représentation par modèle structurel requiert une expertise poussée.

2.3 Diagnostic des systèmes sans connaissance à priori d'un modèle

L'approche proposée nécessite l'acquisition de différentes séries temporelles représentant les mesures des différentes grandeurs du système en fonction du temps. De ces séries sont déduites des informations temporelles de causalité du fait de la dynamique du système et des éventuels retards purs qui peuvent exister entre les différentes variables mesurées.

Nous présentons par la suite, des approches statistiques traduisant ces notions de causalité au travers de l'analyse des séries temporelles. Principalement, trois types de méthodes sont présentés : l'inter-corrélation, la causalité de Granger et le transfert d'entropie.

2.3.1 Modélisation graphique des relations de causalité

L'établissement des relations de causalité entre les différentes entités d'un système est une étape importante. Ces relations constituent une information nécessaire à la détermination des chemins de propagation de défauts au travers des entités du système. Ainsi, il est intéressant d'établir, pour des systèmes complexes de grande dimension, un modèle permettant de représenter intuitivement les causalités entre les variables clés du système.

Dans la littérature, ces modèles de causalité sont souvent associés à une représentation graphique (graphe orienté Boukhobza (2010), graphe de liaisons Thoma (1976), Réseaux Bayésiens Jensen (1996), etc). Le choix de la représentation dépend de la granularité du modèle du système (boite noire, blanche ou grise) et dépend également de l'objectif de son utilisation (diagnostic, reconfiguration, maintenance, etc).

Dans la section précédente, nous avons illustré l'utilisation des graphes orientés dans le cadre des systèmes structurés. L'analyse graphique sur ce type de système permet génériquement

de statuer sur la satisfaction de propriétés structurelles (observabilité, commandabilité, tolérance aux défauts, etc) ayant une importance majeure lors de l'étape de conception du système pour le contrôle/commande. Cependant, dans la majorité des cas, cette analyse requiert un modèle basé sur les équations d'état structurées. Ces équations étant souvent difficiles à obtenir, nous proposons une méthode par la suite analysant la causalité entre les variables du système lorsque suffisamment de données sont disponibles. Elle se base sur le principe de la causalité qui précise qu'un effet est toujours une conséquence d'une cause. Ce principe permet d'affirmer qu'il existe une dépendance dans l'espace et dans le temps entre la cause et l'effet. Nous pouvons citer plusieurs références portant sur la définition de la causalité sous différents points de vue : (Lear (1988); Wallace (1974); May (1970)).

Nous nous focalisons principalement sur la définition de causalité en cohérence avec le temps. Une définition intuitive de la causalité est liée à une dépendance temporelle entre la cause et l'effet. Formellement, l'occurrence d'un événement à l'instant t peut avoir un impact sur un autre événement uniquement à l'instant $t + dt$ et non vice-versa.

Ceci est en lien avec la notion de système causal défini notamment dans Oppenheim (1997) en théorie de la commande. Une définition supplémentaire permet de stipuler que les valeurs des entrées d'un système causal à l'instant $t > t_1$ n'ont aucune influence sur les réponses du système à l'instant t_1 .

D'après les précédentes définitions, dans le cas où les entrées sont directement liées aux sorties du système par des équations algébriques, le sens de la causalité ne peut être établi. Il peut être identifié en utilisant des outils probabilistes de mesure de causalité (transfert d'entropie par exemple) et/ou l'exploitation des connaissances à priori (équations mathématiques, expert, etc). Cette dernière est souvent utilisée en industrie lors la phase de validation d'un modèle du système.

2.3.1.1 Inter-corrélation

Principe et formalisation mathématique

L'inter-corrélation, présentée en détail dans Johansson (1993), mesure la similitude linéaire entre deux séries temporelles (vecteurs d'observations) que nous noterons $x(k) \in \mathbb{R}$ et $y(k) \in \mathbb{R}$ avec $k = 1, \dots, n$. Ces séries peuvent être retardées l'une de l'autre par un retard constant h .

En déterminant la valeur maximale de l'inter-corrélation entre les deux séries temporelles qui se manifeste théoriquement en h , il est possible de déceler la dépendance de causalité en termes de corrélation décalée dans le temps des deux séries temporelles. Ainsi, le retard h maximisant la fonction d'inter-corrélation est utilisé pour calculer une matrice appelée "la matrice de causalité" entre les variables mesurées du système.

Nous présentons ci-dessous la méthode basée sur la corrélation de Pearson. Cette méthode présente l'avantage de fournir un résultat facile à appréhender et, comme indiqué dans Burnham and Anderson (2002a), l'inter-corrélation est assez "robuste" au bruit (additif) affectant les données. L'inter-corrélation entre la paire de variables x et y est notée $\hat{C}_{xy}(h)$. Cette fonction est donnée par :

$$\hat{C}_{xy}(h) = \frac{1}{n} \cdot \frac{\sum_{k=1}^{n-h} (x(k) - \hat{\mu}_x) \cdot (y(k+h) - \hat{\mu}_y)}{\sqrt{\sum_{k=1}^{n-h} (x(k) - \hat{\mu}_x)^2 \cdot \sum_{k=1}^{n-h} (y(k+h) - \hat{\mu}_y)^2}} \quad (2.4)$$

avec $h \in \{1 - n, 2 - n, \dots, n - 2, n - 1\}$ et

$$\hat{\mu}_x = \frac{1}{n} \cdot \sum_{k=1}^n x(k), \quad \hat{\mu}_y = \frac{1}{n} \cdot \sum_{k=1}^n y(k) \quad (2.5)$$

Le dénominateur de $\hat{C}_{xy}$ décrit les écarts-types de $x(k)$ et $y(k)$. Le résultat de l'inter-corrélation est donc normalisé c.à.d. $\hat{C}_{xy}(h) \in [-1, 1]$. Si le maximum de $|\hat{C}_{xy}(h)| = 1$, cela signifie que $x(k)$ et $y(k)$ sont parfaitement corrélés à un retard h donné.

Au contraire, les valeurs proches de zéro indiquent qu'aucune corrélation n'existe entre les séries temporelles. Formellement, il est indiqué que, sous l'hypothèse que $x(k)$ soit un bruit

blanc stationnaire et que $y(k)$ dépend linéairement de $x(k)$, il est possible de mettre en avant l'existence d'une causalité temporelle par l'analyse de cette inter-corrélation entre $x(k)$ et $y(k)$.

Mise en œuvre sur un système 1^{er} ordre :

Afin d'illustrer le calcul de l'inter-corrélation, nous considérons le système de premier ordre représenté par l'équation 2.6.

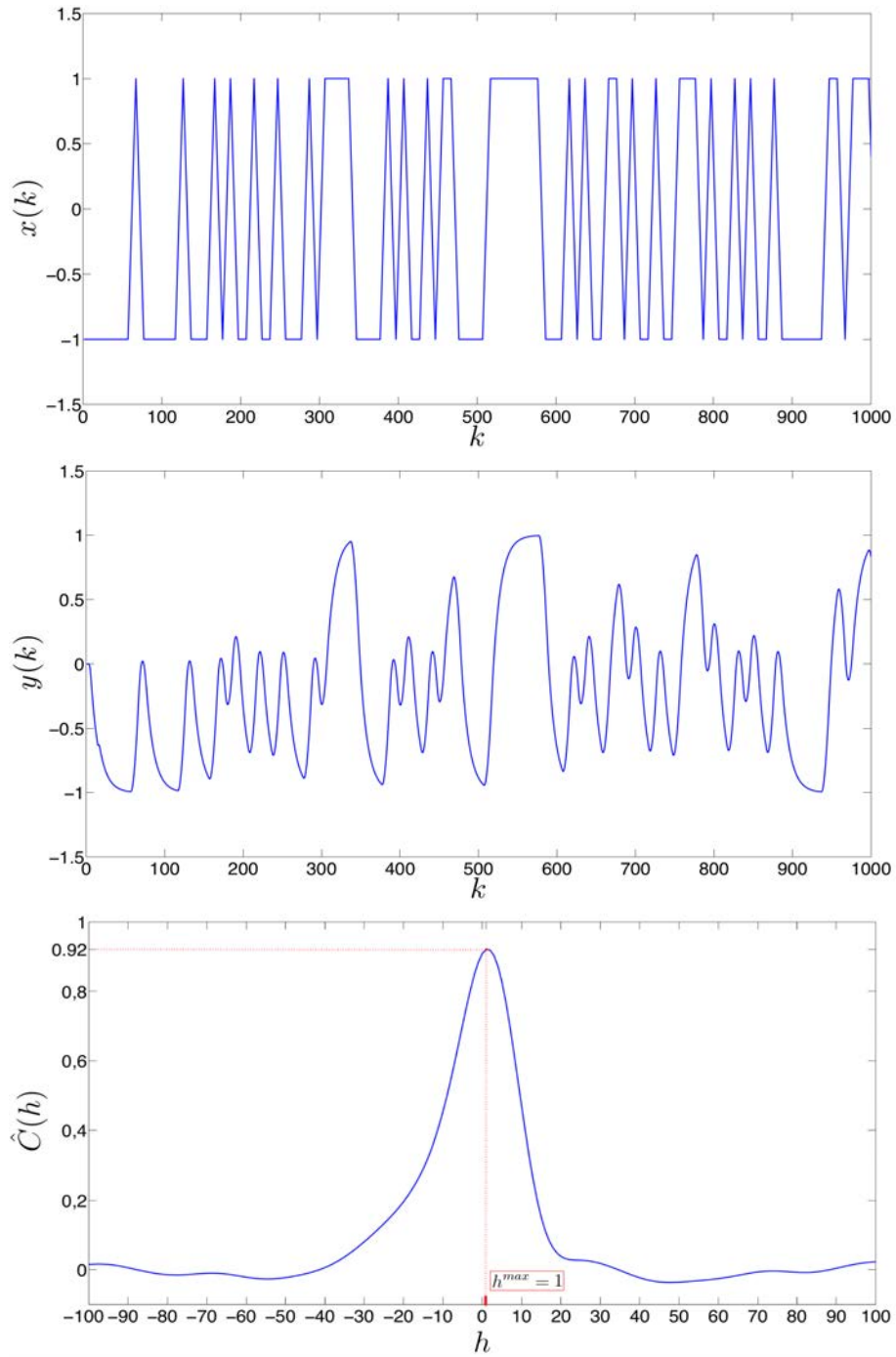
$$\frac{y(s)}{x(s)} = \frac{1}{1 + s} \quad (2.6)$$

La Figure 2.5 montre le signal d'excitation $x(t)$, la réponse du système $y(t)$ et la fonction d'inter-corrélation liée au décalage $h \in \{-100, \dots, 100\}$.

Le choix de l'intervalle réduit permet de visualiser le coefficient d'inter-corrélation en fonction du décalage temporel h sur un horizon fini réaliste par rapport au type du système étudié. Ainsi, les décalages maximisant la fonction d'inter-corrélation en dehors de cet intervalle ne représentent qu'une "coïncidence statistique" et ne reflète véritablement pas l'existence d'une causalité entre la paire de variables x et y .

Dans cet exemple, $x(k)$ représente une séquence binaire pseudo aléatoire et $y(k)$ représente la réponse du système. La fréquence d'échantillonnage de ce système est $T_s = 0,1s$. Pour $h > 0$, l'inter-corrélation est maximale à $h = 1$. Tandis que pour $h \leq 0$, l'inter-corrélation a des valeurs proches de zéro. Cela conduit à la conclusion qu'il existe une dépendance causale de $x \rightarrow y$ pour $h > 0$ mais pas de $y \rightarrow x$.

L'utilisation de l'inter-corrélation pour évaluer la présence d'une dépendance causale entre plusieurs variables a fait l'objet de plusieurs travaux tels que Horch (2000). L'algorithme présenté dans Horch (2000) est utilisé pour tester si l'amplitude maximale de l'inter-corrélation déterminée par un ensemble de séries temporelles est significative

FIGURE 2.5 – $x(k)$, $y(k)$ et $\hat{C}_{xy}(h)$

par rapport à l'amplitude maximale de l'inter-corrélation entre deux variables aléatoires indépendantes.

Dans Bauer (2005), un critère contenant les amplitudes maximales de l'inter-corrélation pour $h > 0$ et $h < 0$ est proposé dans le but de prendre une décision. La valeur de ce critère

est testée statistiquement en termes de la règle de $68 - 95 - 99.7 - \sigma$ (règle de $1 - 2 - 3 - \sigma$) contre la valeur du critère calculée à partir de deux variables aléatoires indépendantes (absence de causalité), c.à.d. quand la valeur de l'inter-corrélation est supérieure à 99.7% ($3 - \sigma$) des valeurs probables sous l'hypothèse $\mathcal{H}_0$, alors on en déduit l'existence d'une causalité entre les deux séries temporelles. L'idée de la règle de $3 - \sigma$ est de vérifier si la valeur du critère pour une série temporelle non-permutée dépasse 3 fois la valeur de l'écart-type généré à partir de la valeur du critère relatif aux permutations aléatoires du signal d'entrée $x(k)$, ainsi pour chaque permutation i , $x_{pi}(k)$ correspond au signal dont les observations sont permutées aléatoirement.

L'utilisation des permutations de la série originale $x(k)$ notés $x_{pi}(k)$, pour le test de présence de causalité, présente l'avantage de conservation des mêmes caractéristiques en terme d'amplitude et de distribution de probabilité de la série temporelle. Tandis que les dépendances de causalité d'une série temporelle par rapport à l'autre sont détruites grâce à la procédure de permutation aléatoire. Le principe de permutation et le test de signifiante associé sont illustré à la Figure 2.6. Ce principe de test de signifiante est adopté pour l'ensemble des méthodes basées sur les données citées dans ce chapitre.

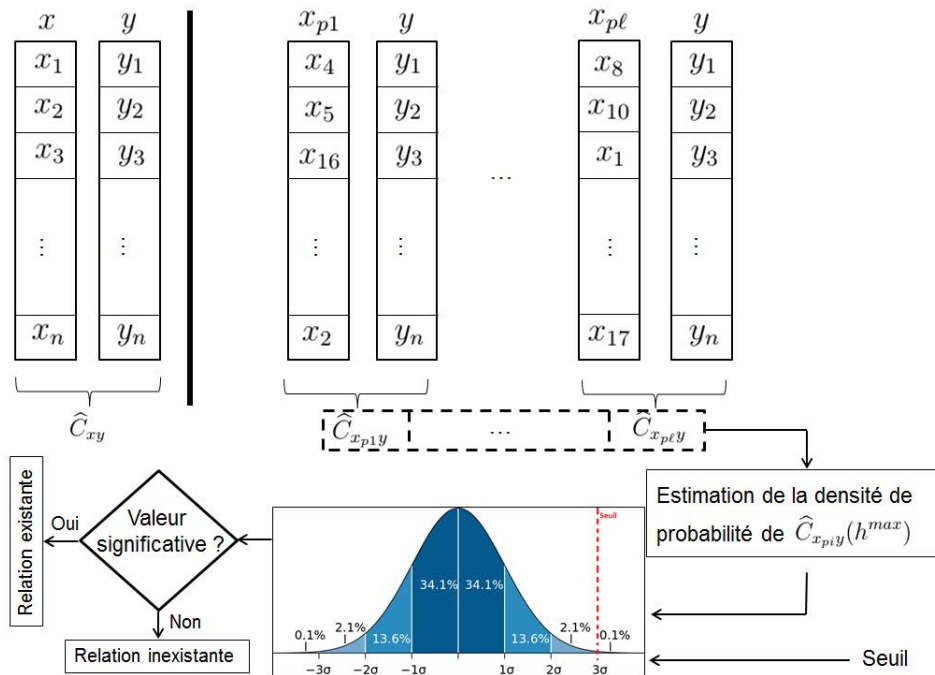


FIGURE 2.6 – Test de signifiante basé sur le principe de permutations

Test de présence de causalité

L'utilisation de l'inter-corrélation pour la détection d'une éventuelle relation de cause à effet entre x et y présente l'inconvénient de toujours présenter un maximum lorsqu'elle est tracée en valeur absolue. Par conséquent, il est important d'établir deux tests. Le premier test vérifie si la valeur maximale trouvée de la fonction d'inter-corrélation diffère significativement de zéro ou si elle peut être le résultat de deux séries temporelles non corrélées. Dans le second test, nous vérifions si l'éventuelle causalité trouvée $x \rightarrow y$ diffère sensiblement de $y \rightarrow x$. En outre, la valeur obtenue lors du deuxième test est utilisée pour définir la "force" de la relation de cause à effet trouvée.

1^{er} Test : Signification statistique du coefficient de corrélation

Pour tester si l'amplitude maximale absolue de l'inter-corrélation est significativement différente de zéro et qu'elle n'est pas due à deux variables aléatoires indépendantes, un test d'hypothèse basé sur le coefficient de corrélation de Pearson est effectué. La première étape est le choix de l'amplitude maximale de l'inter-corrélation. Comme il est possible que le coefficient d'inter-corrélation soit négatif, la valeur absolue est utilisée. Ceci signifie que le décalage h^{max} entre les deux séries x et y s'associant à la position l'amplitude maximale de l'inter-corrélation peut être définie comme suit :

$$h^{max} = \operatorname{argmax}_h |\hat{C}_{xy}(h)| \quad \text{avec } h = 1 - n, \dots, -1, 1, \dots, n - 1. \quad (2.7)$$

Si $h^{max} > 0$, ceci implique qu'il existe une relation de causalité $x \rightarrow y$ et le test d'hypothèse peut être effectué. Considérons la définition du coefficient de corrélation $\hat{\rho}_{xy}$ associé au décalage temporel h^{max} :

$$\hat{\rho}_{xy} := \hat{C}_{xy}(h^{max}) \quad (2.8)$$

Cette valeur sera utilisée pour ce 1^{er} test. Puisque le nombre d'observations est limité pour chaque série temporelle étudiée, l'équation 2.4 ne donne qu'une estimation de la fonction d'inter-corrélation. Ceci signifie que $\hat{\rho}_{xy}$ différera toujours de zéro lors du calcul de

l'inter-corrélation pour deux séries temporelles non corrélées.

Par conséquent, un test statistique est nécessaire pour vérifier si un lien important entre les signaux est présent. Selon Heiman (2010), l'estimation de $\hat{\rho}_{xy}$ à partir de deux signaux non corrélés suit une distribution de Student (distribution t) avec un degré de liberté $n - 2$, ce qui signifie qu'un t-test classique défini par l'équation 2.9 peut être effectué pour vérifier si une corrélation significative existe. Pour ce test, l'hypothèse nulle $\mathcal{H}_0$ est définie de telle façon qu'il n'y a pas de corrélation entre les deux séries, ce qui signifie que si $t > t_{(n-2;1-\alpha/2)}$, il est possible de supposer que les deux signaux sont corrélés.

$$t = \hat{\rho}_{xy} \sqrt{\frac{n-2}{1-\hat{\rho}_{xy}^2}} \quad (2.9)$$

α est défini comme un seuil de signification statistique. Nous considérons qu'il est fixé à $\alpha = 0,05$.

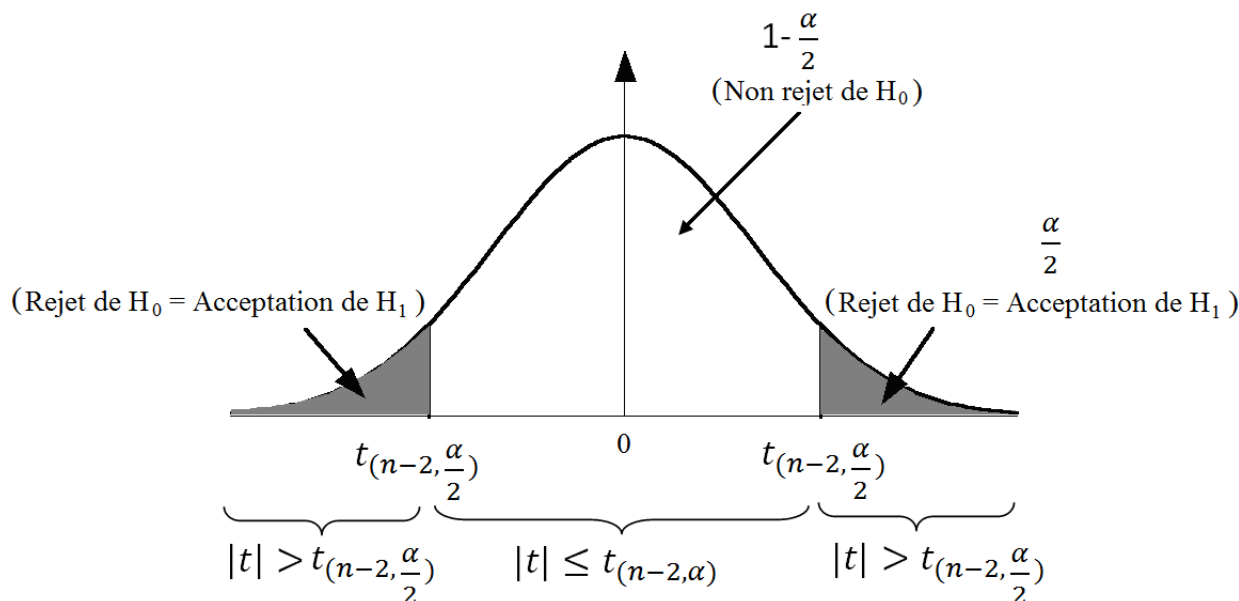
Si le test indique qu'à h^{max} l'hypothèse nulle peut être rejetée signifiant qu'il existe une corrélation significative, un deuxième test basé sur les amplitudes maximales pour $h < 0$ et $h > 0$ est effectué. Si ce test échoue ou $h^{max} < 0$, nous concluons donc qu'aucune relation de cause à effet n'existe de x vers y ($x \rightarrow y$) et la causalité de y vers x ($y \rightarrow x$) peut ainsi être testée à son tour.

Ce test peut être formulé en considérant deux hypothèses :

- Hypothèse nulle $\mathcal{H}_0 : \hat{\rho}_{xy} = 0$ (absence de corrélation entre les deux séries temporelles $x(k)$ et $y(k)$)
- Hypothèse alternative $\mathcal{H}_1 : \hat{\rho}_{xy} \neq 0$ (la corrélation est significative entre les deux séries temporelles $x(k)$ et $y(k)$)

La règle de décision consiste à accepter :

- $\mathcal{H}_0$ si $|t| < t_{(n-2; \frac{\alpha}{2})}$
- $\mathcal{H}_1$ si $|t| > t_{(n-2; \frac{\alpha}{2})}$

FIGURE 2.7 – Distribution de Student et prise de décision au risque α

La prise de décision, en considérant ce test bilatéral peut être illustrée par la Figure 2.7.

Il est aussi évident que si le degré de liberté $(n - 2) \rightarrow \infty$ c.à.d si le nombre de mesures tend vers l'infini, alors cette distribution tend vers une distribution normale (Théorème Central Limite Rokhlin (2015)). Ainsi, la prise de décision est semblable à celle décrite dans le paragraphe précédent et basée sur une distribution usuellement utilisée dans le domaine de la statistique.

En réalité, les caractéristiques de cette distribution ainsi que le risque α doivent être ajustés en fonction des caractéristiques statistiques des mesures issues du procédé. En effet, le bruit peut affecter considérablement les coefficients d'inter-corrélation rendant ce test peu efficace. Dans Yang et al. (2012); Bauer et al. (2007); Bauer and Thornhill (2008), des améliorations sont proposées en termes de prise de décision basée sur des méthodes empiriques.

2^{ème} Test : Signifiante statistique du sens de la causalité

L'objectif du premier test est d'évaluer si les deux séries temporelles sont corrélées par le choix de la valeur maximale de leur fonction d'inter-corrélation. Celui-ci ne permet

pas de prendre en compte le cas où les valeurs optimales de la fonction d'inter-corrélation correspondent à un maximum global pour $h > 0$ à un minimum local $h < 0$. Dans ce cas, il est difficile de statuer sur le sens de la relation de causalité entre une paire de variables.

L'objectif du second test est d'évaluer si la valeur du maximum global pour $h > 0$ est significativement différente de celle qui correspond à $h < 0$.

Ainsi, le critère de comparaison est un ratio qui peut être défini par :

$$\Psi^C = \frac{\max_{h>0} |\hat{C}_{xy}(h)| - \max_{h<0} |\hat{C}_{xy}(h)|}{\max_{h>0} |\hat{C}_{xy}(h)| + \max_{h<0} |\hat{C}_{xy}(h)|} \quad (2.10)$$

avec $\Psi^C \in [-1 \ 1]$.

La valeur de $\Psi^C > 0$ implique la présence d'une relation causale de x vers y ($x \rightarrow y$). Cette valeur dépend fortement des caractéristiques statistiques des deux séries temporelles évaluées et un test basé sur la règle de $68-95-99.7-\sigma$ (règle de 1-2-3- σ) est nécessaire.

Comme les permutations aléatoires de la série temporelle $x(k)$ ne permettent pas de mettre en valeur un lien de causalité de $x \rightarrow y$, les valeurs de Ψ^C doivent être proches de zéro. Cette idée peut être exploitée en calculant, pour des permutations aléatoires de l'entrée $x(k)$, la valeur du critère (2.10) associé à ces permutations. La valeur du critère ainsi obtenue, notée Ψ_p^C , permet de représenter la valeur du critère quand la causalité entre une paire de variables n'existe pas.

Une prise de décision sur le sens de causalité dépend du seuil défini par :

$$\Psi_{seuil}^C = \mu_{\Psi_p^C} + 3\sigma_{\Psi_p^C} \quad (2.11)$$

avec,

$$\mu_{\Psi_p^C} = \frac{1}{N} \sum_{k=1}^N \Psi_p^C(k), \quad \sigma_{\Psi_p^C} = \sqrt{\frac{1}{N} \sum_{k=1}^N (\Psi_p^C(k) - \mu_{\Psi_p^C})^2} \quad (2.12)$$

$\mu_{\Psi_p^C}$ et $\sigma_{\Psi_p^C}$ représentent respectivement la moyenne empirique et l'écart-type empirique de Ψ_p^C .

Si $\Psi^C > \Psi_{seuil}^C$, la relation de causalité est considérée comme significative. Notons que le choix d'un test de $99.7\text{-}\sigma$ ($3\text{-}\sigma$) dépend de l'appréciation de l'utilisateur et de la procédure de modélisation. Cependant, ce test est largement utilisé en industrie grâce la simplicité de sa mise en œuvre. Il est considéré comme étant une règle heuristique qui stipule que la plupart des observations (98 %) sont à l'intérieur de l'intervalle de confiance $[\mu - 3\sigma, \mu + 3\sigma]$. Nous citons Pukelsheim (1994); Wheeler et al. (1992) décrivant l'utilisation du test dans le cadre de la maîtrise statistique du procédé. Cet intervalle peut bien évidemment être ajusté en fonction des cas à étudier.

Considérons le système de premier ordre décrit précédemment, la Figure 2.8 illustre l'estimation de la densité de probabilité du critère Ψ_p^C . Ce ratio est donné par l'Equation 2.10.

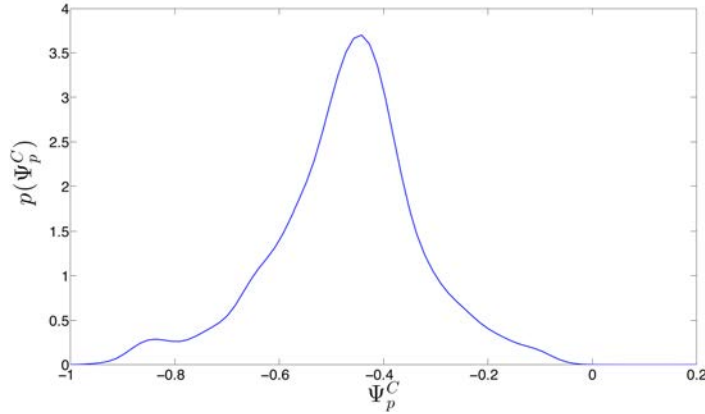


FIGURE 2.8 – Résultat de l'estimation de la densité de probabilité de Ψ_p^C pour $n = 1000$

A partir de cette densité, la moyenne et la variance empiriques peuvent être estimées. Ainsi, $\mu_{\Psi_p^C} \approx -0.47$ et $\sigma_{\Psi_p^C} \approx 0.14$. De ces deux caractéristiques statistiques, le seuil $\Psi_{seuil}^C = \mu_{\Psi_p^C} + 3\sigma_{\Psi_p^C}$ est calculé et $\Psi_{seuil}^C = -0.05 \Rightarrow \Psi_{seuil}^C \approx 0$. En fonction du seuil ainsi calculé, la valeur de $\Psi^C = 0.39$ est significativement plus grande que Ψ_{seuil}^C . Ceci mène vers la conclusion que la variable x est la cause de y ($x \rightarrow y$).

Exemple d'illustration :

Considérons le système décrit par les équations récurrentes suivantes dont la simulation permet de générer des mesures bruitées :

$$\begin{cases} x_1(k) = 0.8x_1(k-1) - 0.9x_1(k-2) + b_1(k) \\ x_2(k) = -0.5x_1(k-2) + 0.25x_2(k-1) + b_2(k) \\ x_3(k) = -0.5x_2(k-4) - 0.25x_3(k-1) + b_3(k) \end{cases} \quad (2.13)$$

avec b_i des bruits blancs gaussiens indépendants de moyenne nulle pour $i \in \{1, 2, 3\}$. Si nous considérons le modèle décrit par les Équations (2.13), le graphe orienté réel associé et les séries temporelles du système sont respectivement représentés à la Figure 2.9 et 2.10.

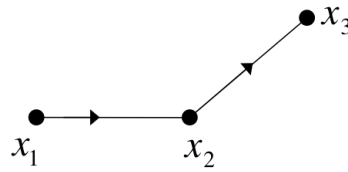


FIGURE 2.9 – Graphe de causalité réel associé au modèle (2.13)

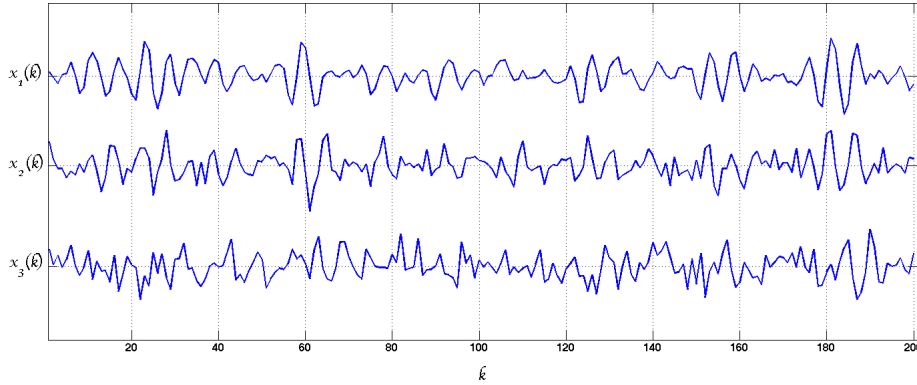


FIGURE 2.10 – Séries temporelles x_1 , x_2 et x_3 du modèle autorégressif (2.13)

A partir des résultats d'analyse par inter-corrélation montrés à la Figure 2.11, nous pouvons constater différents maxima locaux de la fonction d'inter-corrélation pour chaque paire de variables (x_i, x_j) . En effet, le coefficient d'inter-corrélation $\hat{C}_{x_1x_2}$ de $x_1 \rightarrow x_2$ est maximal à $h = 2$ et $\hat{C}_{x_1x_2}(2) \approx 0.77$; le coefficient $\hat{C}_{x_2x_3}$ de $x_2 \rightarrow x_3$ est maximal à $h = 4$ et $\hat{C}_{x_2x_3}(4) \approx 0.65$.

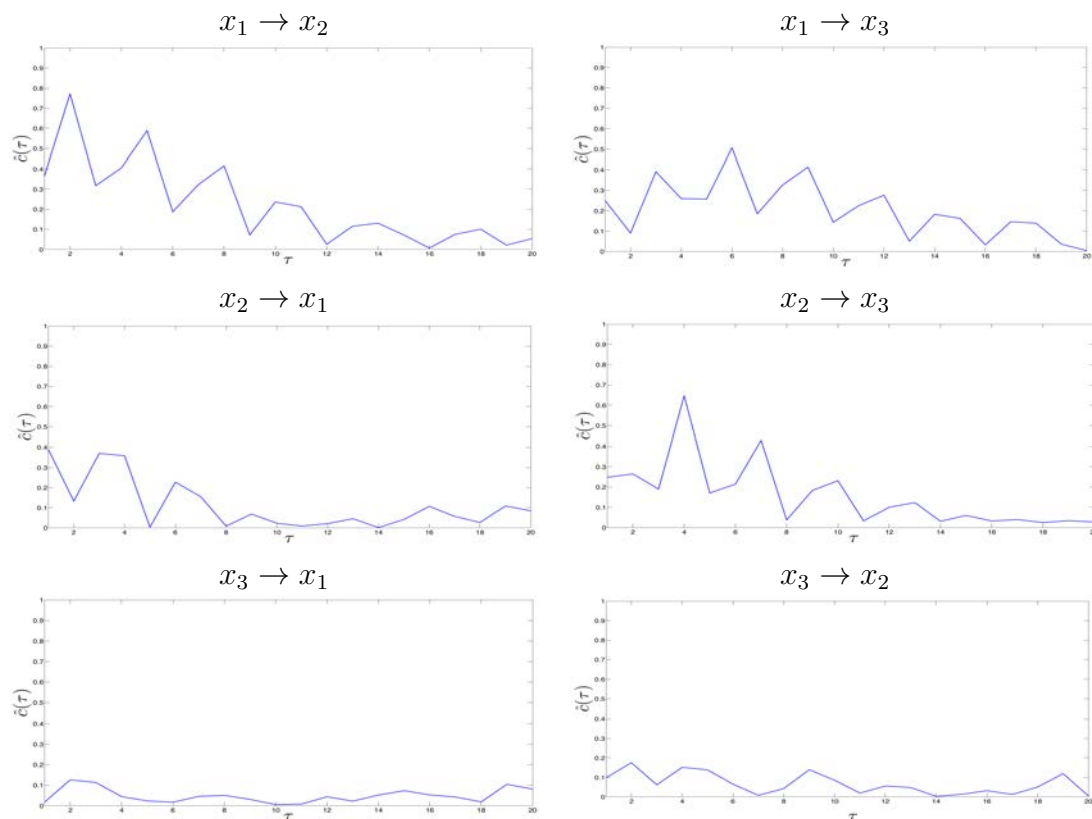


FIGURE 2.11 – Fonctions d’inter-corrélation

Il existe aussi une relation de causalité indirecte entre x_1 et x_3 due à l’existence d’un chemin $x_1 \rightarrow x_2 \rightarrow x_3$. Cette relation est liée à l’existence d’une variable intermédiaire x_2 avec un retard de 6 échantillons, le coefficient $\hat{C}_{x_1x_3}(6) \approx 0.51$. Les causalités $x_3 \rightarrow x_1$ et $x_3 \rightarrow x_2$ sont inexistantes. Les coefficients d’inter-corrélation respectifs des deux relations de causalité $x_3 \rightarrow x_1$ et $x_3 \rightarrow x_2$ sont très faibles et éliminés par le test statistique de Student énoncé précédemment puisque $t_{x_3 \rightarrow x_1} = 1.69$ et $t_{x_3 \rightarrow x_2} = 2.40$ sachant que le seuil est $t_{(n-2;\alpha/2)} = 2.75$.

Afin de retrouver le graphe de causalité liant x_1 , x_2 , et x_3 , nous utilisons la méthode empirique basée sur la permutation aléatoire de la série temporelle x_i . A chaque permutation, nous calculons le coefficient de corrélation $\hat{\rho}_p = \hat{\rho}_{x_i y}$ (coefficient associé à la permutation p). Nous considérons pour cet exemple 100 permutations. Nous pouvons tracer la densité de probabilité des coefficients de corrélation relative à chaque permutation. Ceci est présenté à la Figure 2.12.

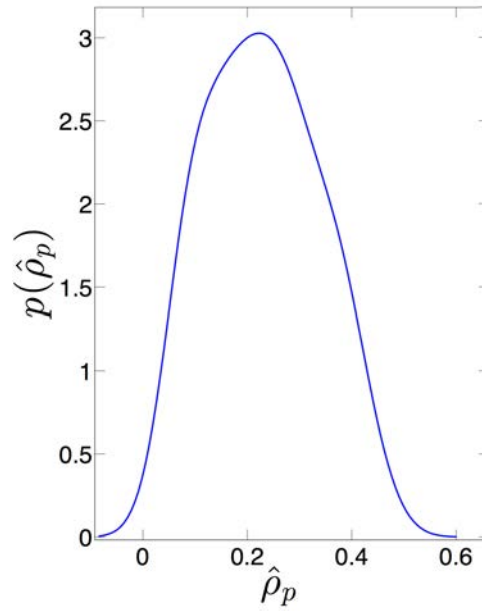


FIGURE 2.12 – Estimation de la densité de probabilité des coefficients de corrélation $\hat{\rho}_p$ sous l'hypothèse $\mathcal{H}_0$

Comme expliqué précédemment, lors de la partie relative au deuxième test, nous pouvons calculer le seuil par l'intermédiaire de l'équation suivante :

$$\hat{\rho}_{seuil} = \mu_{\hat{\rho}_p} + 3\sigma_{\hat{\rho}_p} \quad (2.14)$$

La valeur du seuil est $\hat{\rho}_{seuil} \approx 0.45$ (Figure 2.13).

Le graphe orienté pondéré obtenu par rapport au seuil prédéfini est illustré à la Figure 2.14. Les pondérations apparaissent sur les arcs traduisent le décalage temporel (>0) par lequel l'inter-corrélation est maximale.

Nous pouvons à travers ce type de structure éliminer la relation indirecte $x_1 \rightarrow x_3$ induite par l'existence du chemin $x_1 \rightarrow x_2 \rightarrow x_3$.

En effet, nous partons de l'hypothèse que si la somme des retards sur le chemin $x_1 \rightarrow x_2 \rightarrow x_3$, notée $h_{x_1 \rightarrow x_2} + h_{x_2 \rightarrow x_3} = 2 + 4$ est égale à $h_{x_1 \rightarrow x_3} = 6$ alors l'arc (x_1, x_3) peut être éliminé.

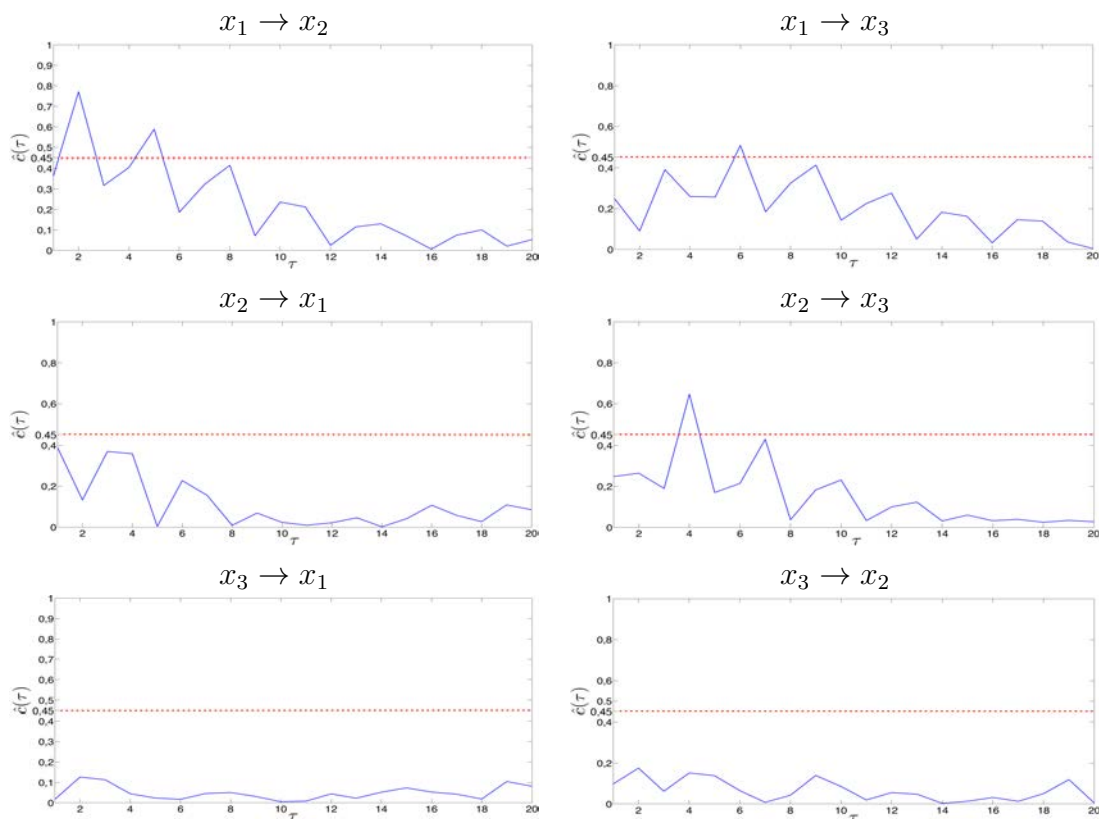
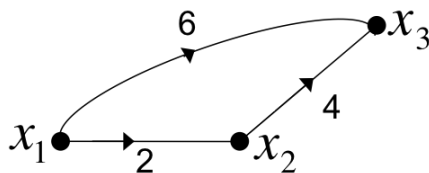
FIGURE 2.13 – La fonction d’inter-corrélation vs retard h avec seuillage $\hat{\rho}_{seuil} \approx 0.45$ 

FIGURE 2.14 – Graphe orienté pondéré identifié par la méthode d’inter-corrélation

2.3.1.2 Causalité de Granger

Le concept de causalité de Granger (CG) a été introduit à l’origine dans le domaine de l’économie par Clive Granger en 1969 (Granger (1969)) qui l’a utilisé pour déterminer les relations entre différents modèles économétriques. La causalité de Granger est basée sur un test de causalité multivariée. La série temporelle $x(k)$ est étendue pour cette mesure de causalité à un ensemble de variables d’entrée qui consiste en r séries temporelles $x_i(k)$ $i \in \{1, \dots, r\}$.

Le concept de causalité de Granger peut s’expliquer aisément en considérant le cas bivariable où deux séries $x(k)$ et $y(k)$ temporelles sont considérées. La relation de dépendance causale

de x vers y ($x \rightarrow y$), au sens de Granger, n'est existante que si les valeurs de $x_i(k)$ passées permettent une meilleure prédiction de la série temporelle $y(k)$ plutôt que de considérer les valeurs passées de $y(k)$ seules.

L'évaluation de la causalité se base principalement sur la comparaison de la variance de l'erreur de prédiction d'un modèle autorégressif (modèle restreint AR) et d'un modèle autorégressif avec entrées exogènes (modèle non restreint ARX) en utilisant des tests d'hypothèses permettant de statuer sur la présence ou non de causalité entre les variables étudiées.

La Figure 2.15 permet d'illustrer globalement le concept de la causalité de Granger.

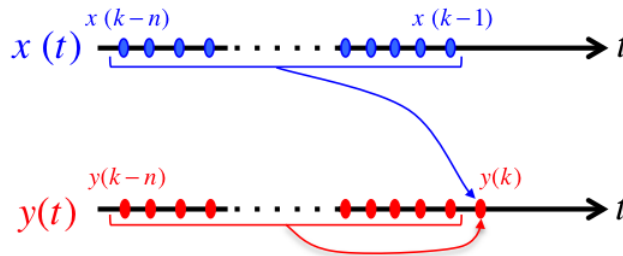


FIGURE 2.15 – Concept de la causalité de Granger basé sur la comparaison de modèles AR-ARX

Les travaux de Friston et al. (2014, 2013) permettent de rapporter sur l'utilisation de la causalité de Granger principalement dans le domaine fréquentiel. Des cas d'études basés sur l'analyse de séries temporelles sont également présentés afin d'illustrer la mesure de causalité proposée. Nous citons également Eichler (2013) dressant un état de l'art sur le concept de la causalité de Granger et mettant en relief le problème d'existence d'une fausse causalité. Cette fausse causalité peut se manifester par l'existence d'un arc entre deux variables $x \rightarrow y$ induit par l'existence d'une variable intermédiaire z sur le chemin $x \rightarrow z \rightarrow y$.

Principe et formalisation mathématique

Comme introduit précédemment, la causalité de Granger est souvent utilisée comme étant une mesure de causalité multivariée. Cela signifie que pendant l'évaluation de la causalité entre $x_i \rightarrow y$, la série temporelle $y(k)$ peut dépendre aussi des autres variables mesurées du

système (entrées) $x_j(k)$ $j \neq i$. Ainsi, il est utile de redéfinir les modèles restreints (AR) et non-restreints (ARX) en tenant compte de ces variables $x_j(k)$ avec $j \neq i$.

La causalité de Granger est évaluée en comparant les deux modèles AR et ARX. La méthode se base sur la prédiction de $y(k)$ en supposant connues les observations antérieures (principe de causalité temporelle). Si la norme quadratique de l'erreur de prédiction en utilisant le modèle non-restreint $\{x_l(k), \dots, x_r(k)\}$ est sensiblement plus petite que celle du modèle restreint ne considérant que les seules séries temporelles $\{x_l(k), \dots, x_r(k)\} \setminus x_i(k)$ c.à.d, n'intégrant pas la variable x_i du modèle, alors on peut conclure que x_i explique par partie y .

Considérons une variable dépendante y et deux ensembles de variables d'entrée $X = \{x_l(k), \dots, x_r(k)\}$ et $X \setminus x_i = \{x_l(k), \dots, x_r(k)\} \setminus x_i$.

La norme quadratique de l'erreur de prédiction, en considérant un modèle d'ordre m et un nombre d'observations n , est donnée par :

$$\left\{ \begin{array}{l} \xi_X^2 = \sum_{k=m+1}^n \left(y(k) - \hat{a}_0 - \sum_{j=1}^m \hat{a}_j y(k-j) - \sum_{l=1}^r \sum_{j=1}^m \hat{b}_{lj} x_l(k-j) \right)^2 \\ \xi_{X \setminus x_i}^2 = \sum_{k=m+1}^n \left(y(k) - \hat{a}'_0 - \sum_{j=1}^m \hat{a}'_j y(k-j) - \sum_{\substack{l=1 \\ l \neq i}}^r \sum_{j=1}^m \hat{b}'_{lj} x_l(k-j) \right)^2 \end{array} \right. \quad (2.15)$$

Les paramètres $\hat{a}_i$, $\hat{b}_{lj}$, $\hat{a}'_i$ et $\hat{b}'_{lj}$ des modèles restreints et non-restreints sont obtenus par l'utilisation de certaines méthodes d'identification.

Les valeurs de ξ_X et $\xi_{X \setminus x_i}$ sont ensuite utilisées pour évaluer l'existence d'une relation de causalité et en estimer éventuellement la force. Dans le paragraphe suivant, nous présentons différents critères de sélection de l'ordre du modèle et de l'évaluation de sa consistance en termes d'erreur de prédiction.

La causalité au sens de Granger a été définie tant dans le domaine temporel que fréquentiel Barnett and Seth (2014). Dans ce travail, nous nous intéressons exclusivement au domaine temporel. Dans ce cas, si l'on considère les variables x et y , la causalité est définie

par le ratio suivant :

$$\mathcal{GC}_{x_i \rightarrow y} = \ln \left(\frac{\text{var}(\xi_{X \setminus x_i})}{\text{var}(\xi_X)} \right) \quad (2.16)$$

où $\text{var}(\cdot)$ exprime la variance de ξ_X et $\xi_{X \setminus x_i}$. L'expression $\mathcal{GC}_{x_i \rightarrow y}$ permet d'évaluer l'apport de la variable x en considérant la réduction de l'erreur de prédiction.

2.3.1.2.1 Estimation de l'ordre du modèle

Comme expliqué précédemment, la causalité de Granger se base sur l'estimation des paramètres d'un modèle autorégressif. Il est clair que les performances du modèle dépendent de l'ordre choisi afin de permettre une meilleure prédiction de la sortie $y(k)$.

Dans la littérature, deux approches sont utilisées pour estimer l'ordre d'un modèle : Critère d'Information Akaike (CIA) (Seth (2010)) et Critère d'Information Bayésien (CIB) (Rissanen (1978)).

Ces critères reposent sur un compromis entre la qualité de l'ajustement et la complexité du modèle, en pénalisant les modèles ayant un grand nombre de paramètres limitant les effets du sur-ajustement. Ces critères tiennent compte de l'erreur de prédiction ξ_X ou $\xi_{X \setminus x_i}$, la taille n de la série temporelle (vecteur des observations), le nombre des variables m considérées et l'ordre du modèle choisi.

Ainsi, l'estimation se base sur la fonction coût $\mathcal{F} = \xi_X$ pour le modèle non-restreint et $\mathcal{F} = \xi_{X \setminus x_i}$ pour le modèle restreint. Nous ne présenterons dans cette partie que le critère CIA. Afin d'avoir plus de détail concernant ces métriques, une étude comparative exhaustive est conduite par Yang (2005); Burnham and Anderson (2004, 2002b). Elle montre l'équivalence des deux critères :

$$CIA(m) = \log(\mathcal{F}) + \frac{2mp^2}{n} \quad (2.17)$$

Le 2^{ème} terme $\frac{2mp^2}{n}$ de l'Équation (2.17) représente la pénalisation du critère $CIA(m)$ par rapport à l'ordre du modèle m et le nombre de variables explicatives p ; c.à.d. le nombre de paramètres à estimer.

Test de signification statistique

Le test de signification statistique consiste à évaluer si les erreurs de prédiction quadratiques du modèle restreint ($\xi_{X \setminus x_i}$) et non restreint (ξ_X) diffèrent significativement. En considérant l'hypothèse que les deux erreurs de prédiction quadratiques suivent une distribution du χ^2 , les travaux de Seth (2010); Barnett and Seth (2014) se basent sur le test de Fisher pour évaluer le lien de causalité éventuel entre $x_i(k)$ et $y(k)$.

La statistique associée au test de Fisher correspond, dans notre cas, à l'équation suivante :

$$F(m, n - m - p) = \frac{\xi_{X \setminus x_i} - \xi_X}{\xi_X} \cdot \frac{n - m(1 + p)}{m} \quad (2.18)$$

La valeur de la statistique $F(m, n - m - p)$ est à comparer à la valeur du test de Fisher F_α correspondant à un seuil de significativité (au risque) de α . Si $F(m, n - m - p) > F_\alpha$, alors x_i peut être interprétée comme une cause probable de y . Ceci peut être formulé par un test d'hypothèse unilatéral (à droite) sur les variances $\sigma_{\xi_{X \setminus x_i}}$ et σ_{ξ_X} des erreurs de prédiction quadratiques du modèle restreint et non restreint.

Nous considérons pour ce test de Fischer, les deux hypothèses suivantes :

- $\mathcal{H}_0 : \sigma_{\xi_{X \setminus x_i}} = \sigma_{\xi_X}$
- $\mathcal{H}_1 : \sigma_{\xi_{X \setminus x_i}} > \sigma_{\xi_X}$

L'hypothèse $\mathcal{H}_0$ est équivalente à formuler le fait que x_i n'est pas une cause de y en considérant que les variances des erreurs de prédiction quadratiques des deux modèles en

tenant compte de la variable x_i dans le modèle non restreint et sans x_i dans le modèle restreint n'affecte pas la prédictibilité des modèles.

A contrario, l'hypothèse $\mathcal{H}_1$ permet de formuler l'existence d'une causalité entre x_i et y en évaluant les deux variances des erreurs de prédiction quadratiques. Ceci correspond à évaluer si l'introduction de la variable x_i dans le modèle non restreint, réduit significativement l'erreur de prédiction quadratique par rapport à celle du modèle restreint. Ainsi, la prise de décision est simple. En effet,

- si $F > F_\alpha$ alors $\mathcal{H}_1$ est vraie
- sinon acceptation de l'hypothèse $\mathcal{H}_0$

Ce test peut être ré-exprimé autrement en évaluant la valeur du risque α par rapport à une p-valeur. Cette p-valeur est associée à l'observation d'une statistique de test F et correspond au seuil auquel on rejeterait l'hypothèse nulle $\mathcal{H}_0$ compte tenu de cette observation.

Formellement, sous l'hypothèse que la statistique F suit une loi de Fischer, la p-valeur notée $p - val$ s'exprime par :

$$p - val = 1 - \Phi(F, dl1, dl2) \quad (2.19)$$

La fonction Φ est la densité de probabilité cumulée de Fisher. $dl1$ et $dl2$ sont respectivement les degrés de liberté du numérateur et du dénominateur de la statistique F .

Ainsi, l'acceptation de l'hypothèse $\mathcal{H}_1$ dépend de la validité de l'inégalité $p - val < \alpha$, sinon l'hypothèse nulle $\mathcal{H}_0$ est acceptée.

Exemple d'illustration :

Considérons le même modèle décrit lors de l'exemple d'illustration de la section 2.3.1.1. Le choix de l'ordre du modèle s'établit par le critère CIA comme illustré à la Figure 2.16

L'ordre du modèle choisi est $m=4$. Ainsi, les modèles restreint et non restreint pour l'estimation de la causalité de Granger sont d'ordre 4.

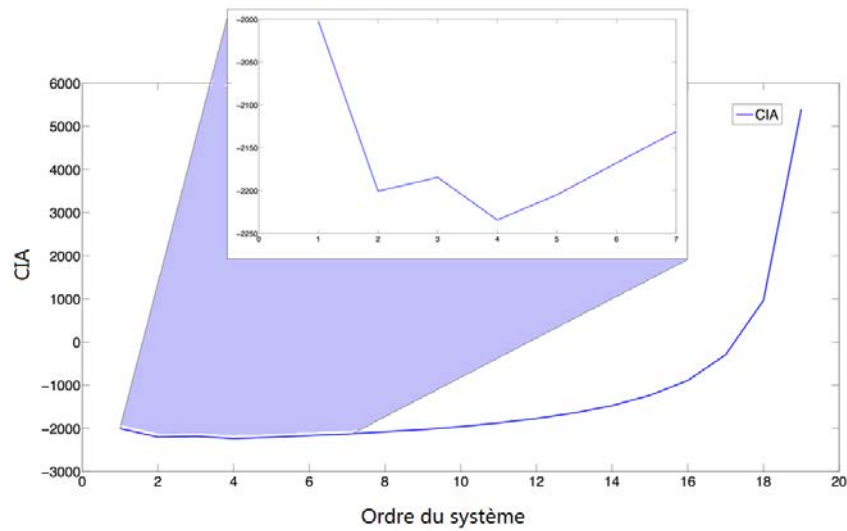


FIGURE 2.16 – Critère CIA du modèle autorégressif du système

La Figure 2.17 représente les causalités qui peuvent exister entre les variables du système décrit par l'Équations (2.13).

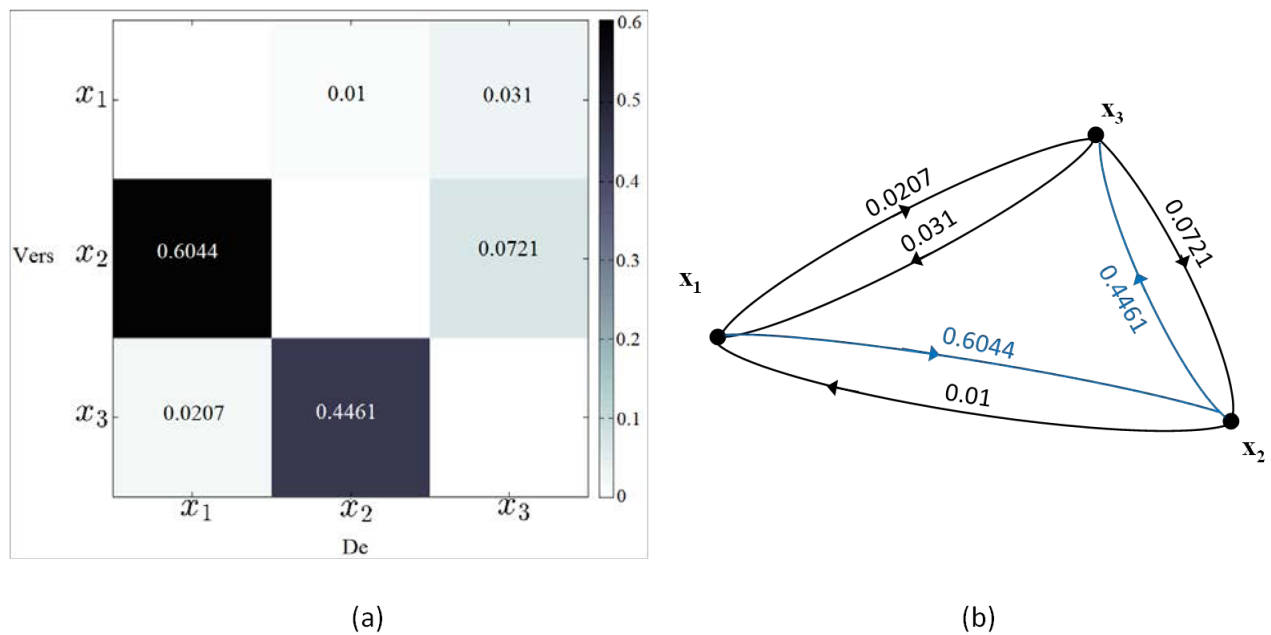


FIGURE 2.17 – Causalité de Granger : (a) Carte thermique, (b) Graphe de causalité pondéré

Interprétation

A partir de la Figure 2.17, nous pouvons statuer sur l'existence des différentes rela-

tions de causalité entre les variables x_i en évaluant les coefficients correspondant aux ratios de vraisemblance exprimés dans l'Équation 2.16. Ainsi, pour cet exemple, seuls les coefficients 0.6044 et 0.4461 représentant respectivement la présence de causalité $x_1 \longrightarrow x_2$ et $x_2 \longrightarrow x_3$ sont retenus puisque les autres coefficients sont très proches de 0 exprimant ainsi l'absence d'une causalité.

Afin de statuer systématiquement sur la présence de causalité, nous appliquons le test statistique de Fisher en calculant la p-valeur. Nous adoptons la notation suivante de p-valeur associée à une relation de causalité (x_i, x_j) notée $p_{(x_i, x_j)\text{-val}}$.

Pour un seuil de $\alpha = 5\%$, nous avons $p_{(x_1, x_2)\text{-val}} = 0$, cette valeur correspond la probabilité de rejeter l'hypothèse $\mathcal{H}_0$ sachant que cette hypothèse est vraie c.à.d. $p_{(x_1, x_2)\text{-val}} < \alpha$. Ainsi, la relation de causalité entre $x_1 \rightarrow x_2$ existe.

Nous résumons les autres résultats dans la Table 2.1 suivante :

(x_i, x_j)	p-val	$\mathcal{H}_i$ pour $i \in \{0, 1\}$
(x_1, x_2)	0	$\mathcal{H}_1$
(x_1, x_3)	0.093 (9.3 %)	$\mathcal{H}_0$
(x_2, x_1)	0.42 (42 %)	$\mathcal{H}_0$
(x_2, x_3)	0	$\mathcal{H}_1$
(x_3, x_1)	0.017 (17 %)	$\mathcal{H}_0$
(x_3, x_2)	$1.310^{-05} \approx 0$	$\mathcal{H}_1$

TABLE 2.1 – Existence de relation de causalité (x_i, x_j)

La présence du faux lien de causalité (x_3, x_2) est due à la présence de colinéarité entre x_3 et la paire de variables x_1 et x_2 .

Différents travaux (Gilson and Van den Hof (2005); Söderström and Stoica (2002); Forssell and Ljung (1999)) relatifs à l'identification des systèmes permettent d'améliorer les résultats de mesure de causalité de Granger. Ces travaux proposent des méthodes basées sur l'exploitation de variables instrumentales. Ces méthodes permettent également de considérer l'endogénéité de certaines variables dans un problème de régression contrairement à la méthode des moindres carrés ordinaire utilisée dans ce manuscrit pour l'estimation des modèles restreint et non restreint.

2.3.1.3 Transfert d'Entropie

Une approche non paramétrique de la mesure de causalité consiste en celle de la théorie de l'information dans laquelle on peut introduire différents indicateurs, comme l'information dirigée. Une autre caractérisation de nature entropique est le transfert d'entropie (TE). Ce concept récent a été introduit par Schreiber (Schreiber (2000)) pour décrire le flux d'informations entre deux séries temporelles.

D'après sa définition, cette mesure de causalité est asymétrique. En effet, TE mesure "la quantité d'information" de la série temporelle $x(k)$ vers la série $y(k)$ et mesure de nouveau "la quantité d'information" de $y(k)$ vers $x(k)$.

Contrairement à l'inter-corrélation qui mesure la causalité entre deux séries temporelles en fonction d'un retard, TE mesure la causalité par l'évaluation des probabilités de transition. Cette approche permet de détecter les relations de causalité non linéaires qui peuvent exister entre les séries temporelles d'un jeu de données.

TE est utilisée dans divers domaines (Staniek and Lehnertz (2009); Chávez et al. (2003); Sensoy et al. (2014)) tels qu'en neurosciences, en économie et analyse financière. Nous nous intéresserons plus particulièrement aux travaux pour la modélisation de causalité des procédés et notamment ceux de Shu and Zhao (2013); Naghoosi et al. (2013) qui concernent la modélisation des procédés chimiques. Dans (Faghraoui et al. (2012, 2014)), nous nous sommes intéressés à la modélisation des systèmes à grande dimension à travers les liens de causalité entre les variables du système.

Formalisation mathématique

Comme introduit précédemment, TE se base sur le calcul de probabilités de transition pouvant montrer l'existence de relations de causalité.

Considérons les séries temporelles $x(k)$ et $y(k)$ écrites respectivement sous forme vectorielle :

$$x_k^{(l)} = (x(k), x(k-1), \dots, x(k-l+1)) \quad (2.20)$$

$$y_k^{(m)} = (y(k), y(k-1), \dots, y(k-m+1)) \quad (2.21)$$

$x_k^{(l)}$ et $y_k^{(m)}$ sont deux vecteurs d'observation de taille l et m .

Considérons le vecteur d'observations $y_k^{(m)}$, la probabilité de transition, définie par $p(y_{k+1}|y_k)$, s'écrit comme suit :

$$p(y_{k+1}|y_k^{(m)}) = p(y_{k+1}|y_k, \dots, y_{k-m+1}) \quad (2.22)$$

Si la valeur future y_{k+1} ne dépend que de y_k , alors la probabilité de transition (2.22) s'écrit sous la forme :

$$p(y_{k+1}|y_k^{(m)}) = p(y_{k+1}|y_k) \quad (2.23)$$

Cette probabilité est appelée "probabilité Markovienne de premier ordre" .

Le calcul de TE nécessite aussi le calcul de la probabilité de transition $p(y_{k+1}|y_k^{(m)}, x_k^{(l)})$.

Cette dernière peut s'écrire sous la forme :

$$p(y_{k+1}|y_k^{(m)}, x_k^{(l)}) = p(y_{k+1}|y_k, \dots, y_{k-m+1}, x_k, \dots, x_{k-l+1}) \quad (2.24)$$

Avant de présenter l'équation permettant le calcul du TE, quelques définitions sont nécessaires :

Définition 1

L'ensemble probabiliste Ω est un ensemble non vide dont les éléments représentent tous les résultats possibles d'une expérience aléatoire donnée.

Définition 2

Soit X une variable aléatoire sur un espace probabiliste Ω . La loi de X est définie par :

$$p(x) = \mathbb{P}(X = x) = \sum_{X(\omega)=x} \mathbb{P}[\omega] \text{ pour } x \in \mathbb{R}$$

En d'autres termes, $p(x)$ est la probabilité des événements ω vérifiant $X(\omega) = x$, donc la probabilité que X prenne la valeur x à une éventualité donnée (un échantillon donné).

Définition 3

Soient X et Y deux variables aléatoires, la loi de probabilité conjointe est définie par :

$$p(x, y) = \mathbb{P}(X = x \text{ ET } Y = y) \text{ pour } (x, y) \in \mathbb{R}^2$$

$p(x, y)$ est la probabilité des événements ω vérifiant simultanément $X(\omega) = x$ et $Y(\omega) = y$.

L'égalité $\sum_{(x,y) \in \mathbb{R}^2} p(x, y) = 1$ est vérifiée.

Définition 4

La probabilité conditionnelle notée $p(y | x)$ est la probabilité que la variable aléatoire $Y(\omega) = y$ sachant la valeur de la variable aléatoire $X(\omega) = x$. Cette probabilité est définie par :

$$p(y | x) = \frac{p(y, x)}{p(x)}$$

Cette définition est issue du théorème de Bayes Jensen (1996).

Sachant que les densités de probabilité sont souvent inconnues sur un système réel, leur estimation s'avère nécessaire. La méthode par histogramme est souvent utilisée car cette méthode présente l'avantage d'être simple à mettre en œuvre.

Le transfert d'entropie représente la différence entre l'information de l'entropie de y_{k+1} quand les passés de y_k et x_k sont connus. La connaissance de y_{k+1} dépend de x_k . Cette définition est similaire à celle de la causalité de Granger.

Ainsi, TE décrit l'incertitude sur la valeur de y_{k+1} déduite à partir de $x_k^{(l)}$ et de $y_k^{(m)}$.

TE s'écrit :

$$TE_{x \rightarrow y} = \sum_{y_{k+1}, y_k^{(m)}, x_k^{(l)}} p(y_{k+1}, y_k^{(m)}, x_k^{(l)}) \log \frac{p(y_{k+1} | y_k^{(m)}, x_k^{(l)})}{p(y_{k+1} | y_k^{(m)})} \quad (2.25)$$

L'intervalle des valeurs de TE est $[0 H_y]$. Cet intervalle est défini dans Marschinski and Kantz (2002).

Afin de prendre en compte la présence d'un retard pur, la variable τ est introduite dans l'expression de TE . Ceci consiste à décaler la série temporelle x_k par rapport à y_k . L'équation (2.25) est calculée ainsi pour différents retards τ de la série temporelle $x(k - \tau)$. Nous introduisons la notation $TE_{x \rightarrow y}(\tau_{max}) = \max_{\tau} (TE_{x \rightarrow y}(\tau))$ signifiant le choix d'une valeur maximale en fonction de τ .

L'équation (2.25) en fonction de τ peut être réécrite :

$$TE_{x \rightarrow y}(\tau) = \sum_{y_{k+1}, y_k^{(m)}, x_{k-\tau}^{(l)}} p(y_{k+1}, y_k^{(m)}, x_{k-\tau}^{(l)}) \log \frac{p(y_{k+1} | y_k^{(m)}, x_{k-\tau}^{(l)})}{p(y_{k+1} | y_k^{(m)})} \quad (2.26)$$

Ceci présente un inconvénient majeur en pratique. En effet, le calcul numérique peut devenir conséquent induit par le nombre de variables du système étudié et aussi par l'ordre des probabilités de transition qui dépendent de m et l .

Exemple d'illustration :

Avant d'illustrer les résultats obtenus, rappelons le système d'équations suivant :

$$\begin{cases} x_1(k) = 0.8x_1(k-1) - 0.9x_1(k-2) + b_1(k) \\ x_2(k) = -0.5x_1(k-2) + 0.25x_2(k-1) + b_2(k) \\ x_3(k) = -0.5x_2(k-4) - 0.25x_3(k-1) + b_3(k) \end{cases} \quad (2.27)$$

avec b_i des bruits blancs gaussiens indépendants de moyenne nulle pour $i \in \{1, 2, 3\}$. Si nous considérons le modèle décrit par les équations (2.27), le graphe orienté réel associé et les séries temporelles du système sont respectivement représentés à la Figure 2.18 et 2.19.

Cet exemple d'illustration est semblable à celui utilisé lors des méthodes précédemment présentées (l'inter-corrélation et la causalité de Granger).

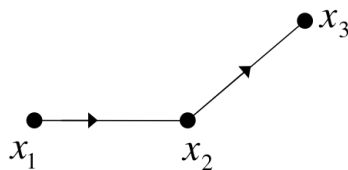
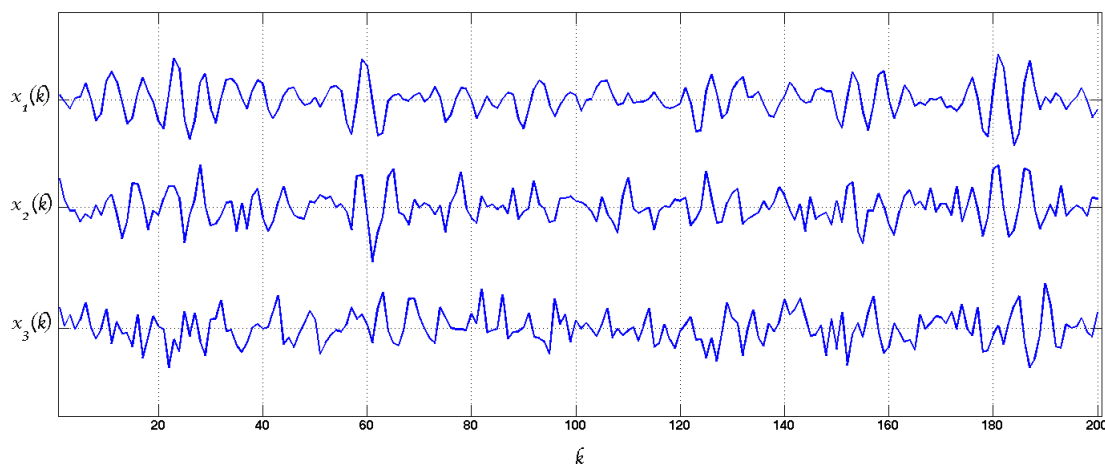


FIGURE 2.18 – Graphe de causalité réel associé au modèle (2.27)

FIGURE 2.19 – Séries temporelles x_1 , x_2 et x_3 du modèle autorégressif (2.27)

Afin de construire le graphe de causalité basé sur l'estimation par la méthode TE entre les séries temporelles étudiées, nous considérons les étapes explicitées ci-dessous :

a) Choix des dimensions l et m des séries temporelles et de l'horizon de prédiction

Avant le calcul de TE pour chaque paire de variables (x_i, x_j) , il est nécessaire de définir deux paramètres liés à la dimension l et m des séries temporelles $x_i^{(l)}$ et $x_j^{(m)}$. Si les valeurs des paramètres l et m choisies sont petites, la capture de la dynamique du système est omise. Ceci se traduit par des résultats d'estimation de TE mitigés.

À contrario, si les paramètres l et m sont grands, le calcul de TE peut devenir chronophage. En effet, le calcul dépend principalement de l'estimation des combinaisons des probabilités conditionnelles $p(y_{k+1}|y_k^{(m)}, x_k^{(l)})$ et $p(y_{k+1}|y_k^{(m)})$.

Le calcul de TE proposé par Schreiber (2000) et Bauer (2005) se base sur des valeurs de l et m prédéfinis ($l=m=1$) en se basant sur l'hypothèse que les séries temporelles peuvent être représentées par des processus de Markov de premier ordre 2.24. Plusieurs méthodes ont été proposées. On cite Cao (1997) et Ragwitz and Kantz (2002) où la méthode présentée dans Ragwitz and Kantz (2002) est utilisée dans ce travail et il en résulte $l=m=5$.

Le choix de l'horizon de prédiction doit être réaliste par rapport au type du système étudié. Ainsi, les décalages maximisant TE en dehors de cet intervalle ne représentent qu'une "coïncidence statistique" et ne reflète pas véritablement l'existence d'une causalité entre la paire de variables x et y . Un horizon très grand engendre une complexité algorithmique temporelle prohibitive.

Dans notre exemple, l'horizon de prédiction est fixé à $[0 \ 10]$.

b) Test de signifiante statistique

Le principe du test statistique présenté dans cette section est similaire à celui présenté dans la section 2.3.1.1.

Nous utilisons la méthode empirique basée sur la permutation aléatoire de la série temporelle x_i . À chaque permutation de la série temporelle x_i , nous calculons la valeur TE c.à.d. $\hat{TE}_p = TE_{x_i \rightarrow y}$ (valeur TE associée à la permutation p).

Nous considérons pour cet exemple 100 permutations. Nous pouvons ainsi tracer la densité de probabilité des valeurs TE relative à chaque permutation p . La densité de probabilité estimée est illustrée à la Figure 2.20.

Le seuil $\hat{TE}_{seuil}$ est donné à l'Équation 2.28.

$$\hat{TE}_{seuil} = \mu_{\hat{TE}_p} + 3\sigma_{\hat{TE}_p} \quad (2.28)$$

La valeur du seuil est $\hat{TE}_{seuil} \approx 0.0034$ (Figure 2.21).

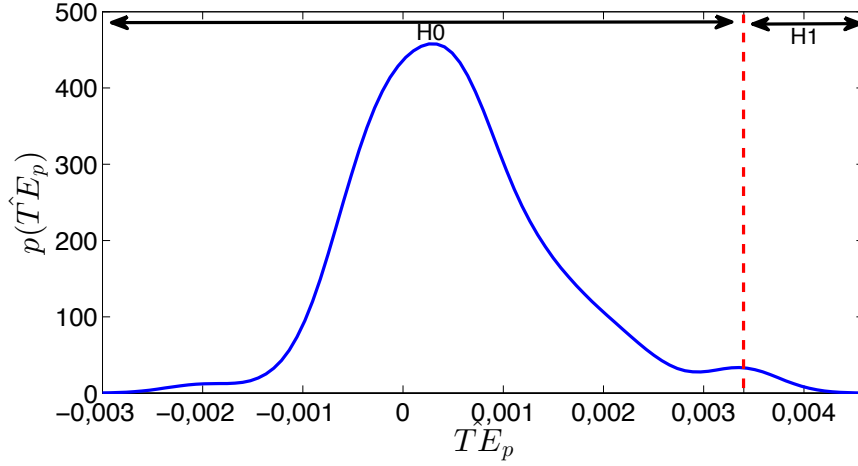


FIGURE 2.20 – Estimation de la densité de probabilité des valeurs TE sous l’hypothèse $\mathcal{H}_0$ vs $\mathcal{H}_1$ ($\mathcal{H}_0$: inexistence de causalité)

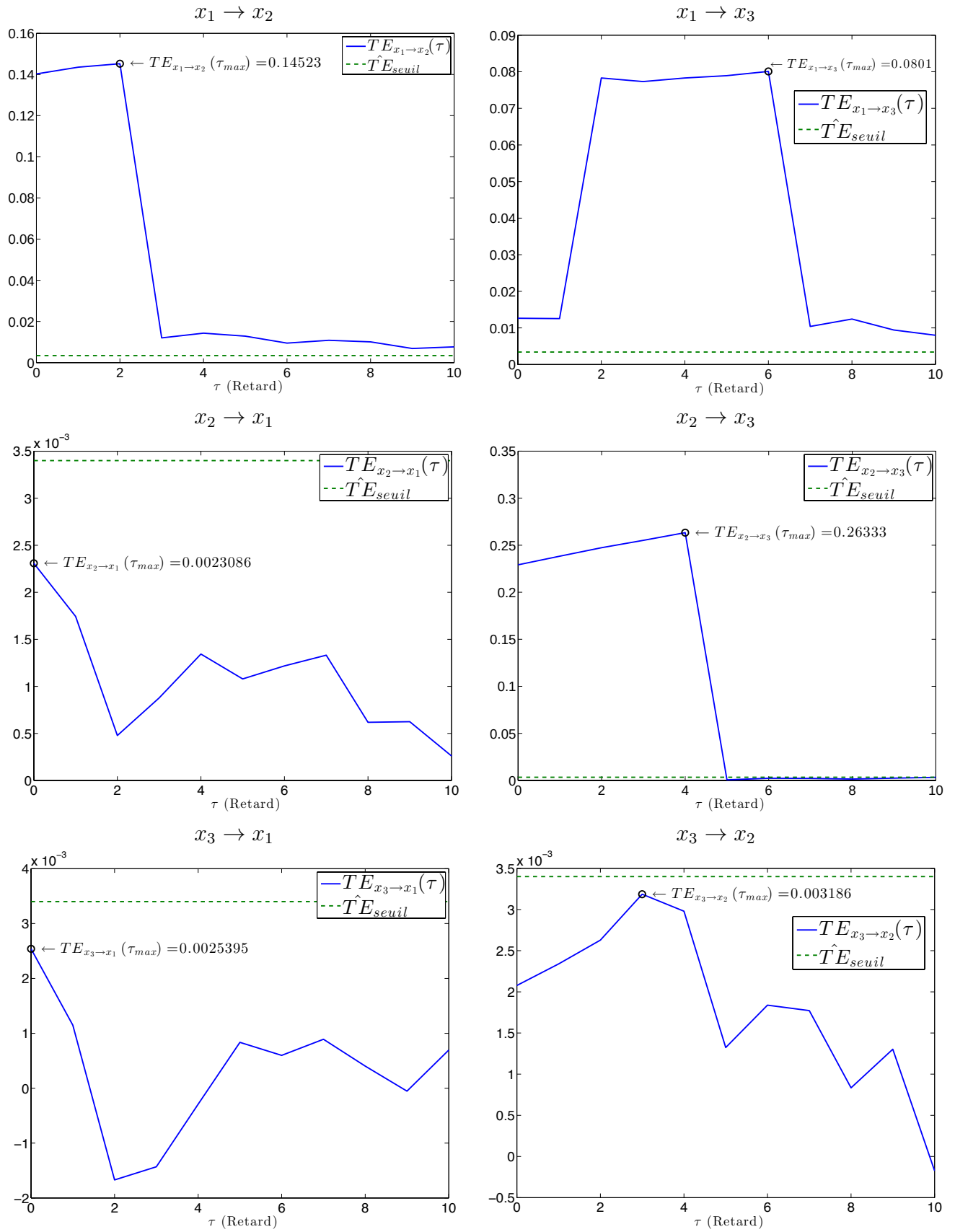
Les résultats du calcul de TE, appliqué à l’exemple 2.27, sont illustrés par la Figure 2.21.

Nous pouvons constater différents maxima locaux de la mesure TE pour chaque paire de variables (x_i, x_j) . En effet, la valeur $TE_{x_1 \rightarrow x_2}$ relative au lien $x_1 \rightarrow x_2$ est maximal à $\tau = 2$ et $TE_{x_1 \rightarrow x_2}(2) \approx 0.14$. La valeur $TE_{x_2 \rightarrow x_3}$ relative au lien $x_2 \rightarrow x_3$ est maximal à $h = 4$ et $TE_{x_2 \rightarrow x_3}(2) \approx 0.26$.

Il existe aussi une relation de causalité indirecte entre x_1 et x_3 due à l’existence d’un chemin $x_1 \rightarrow x_2 \rightarrow x_3$. Cette relation est liée à l’existence d’une variable intermédiaire x_2 avec un retard de 6 échantillons. La valeur $TE_{x_1 \rightarrow x_3}(6) \approx 0.08$.

Les causalités $x_3 \rightarrow x_1$ et $x_3 \rightarrow x_2$ sont inexistantes. Les valeurs TE respectives des deux relations de causalité $x_3 \rightarrow x_1$ et $x_3 \rightarrow x_2$ sont très faibles (en dessous du seuil $\hat{TE}_{seuil}$) et éliminées par le test statistique non-paramétrique précédemment énoncé et illustré sur les différents graphiques de la Figure 2.21.

Le graphe orienté pondéré obtenu par rapport au seuil $\hat{TE}_{seuil}$ est illustré à la Figure 2.22. Les pondérations des arcs traduisent le décalage temporel (>0) par lequel TE est maximale sur l’horizon de prédiction $[0 \ 10]$.

FIGURE 2.21 – TE entre paires de variables (x_i, x_j) relatif à l'exemple 2.27

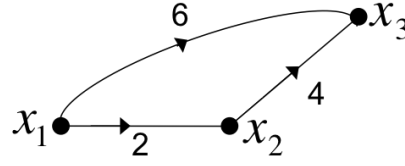


FIGURE 2.22 – Graphe orienté pondéré identifié par la méthode transfert d'entropie

Nous pouvons à travers ce type de structure éliminer la relation indirecte $x_1 \rightarrow x_3$ induite par l'existence du chemin $x_1 \rightarrow x_2 \rightarrow x_3$.

En effet, nous partons de l'hypothèse que si la somme des retards maximisant TE pour chaque paire de variables (x_i, x_j) sur le chemin $x_1 \rightarrow x_2 \rightarrow x_3$ donnée par $\tau_{x_1 \rightarrow x_2} + \tau_{x_2 \rightarrow x_3} = 2 + 4$ est égale à $\tau_{x_1 \rightarrow x_3} = 6$ alors l'arc (x_1, x_3) peut être éliminé.

2.3.2 Matrice de causalité et représentation graphique

La matrice de causalité permet de synthétiser l'ensemble des informations concernant la présence de causalités qui peuvent exister entre les variables d'un procédé. Afin de représenter l'ensemble des liens causaux des variables du procédé étudié, une matrice de causalité, notée C comme illustrée par l'équation 2.29 peut être générée.

$$C = \begin{pmatrix} * & c_{x_2 \rightarrow x_1} & \cdots & c_{x_n \rightarrow x_1} \\ c_{x_1 \rightarrow x_2} & * & \cdots & c_{x_n \rightarrow x_2} \\ * \vdots & \vdots & \ddots & \vdots \\ c_{x_1 \rightarrow x_n} & c_{x_2 \rightarrow x_n} & \cdots & * \end{pmatrix} \quad (2.29)$$

Cette matrice de dimension $n \times n$ est associée à n variables notées x_i , $i \in \{1 \dots n\}$ et les éléments de cette matrice notés $c_{x_i \rightarrow x_j}$ représente une mesure de causalité entre la paire de variables x_i et x_j avec $i \neq j$. La variable $c_{x_i \rightarrow x_j}$ est une mesure heuristique qui dépend de la méthode choisie (inter-corrélation, Causalité de Granger, etc) et représente la force de la causalité. Cette mesure est souvent standardisée ($0 \leq c_{x_i \rightarrow x_j} \leq 1$). Ainsi une valeur proche de 0 représente une relation causale faible. En revanche, une valeur proche de 1 représente une forte relation causale.

L'un des problèmes induits par la considération des mesures réelles est la présence du bruit. Ce dernier influe souvent l'estimateur de causalité, nécessitant ainsi une étape de seuillage fondée sur des méthodes statistiques et/ou un retour d'expérience d'expert ayant des connaissances a priori sur la présence de causalités entre les variables du procédé (Bauer (2005); Bauer et al. (2007)). Ainsi, estimer les dépendances causales entre n variables nécessite le calcul de $n(n - 1)$ valeurs $c_{x_i \rightarrow x_j}$. Au moyen des différents types de méthodes présentées précédemment dans ce manuscrit, cette matrice peut être aisément générée. Nous insistons sur le fait que cette matrice peut être également construite à partir de connaissances d'expertises ou par des ressources matérielles et notamment des schémas du procédé.

A cet effet, nous allons aborder à la section 2.4, les principaux travaux menés dans le cadre du développement de méthodes combinées pour la modélisation de causalité des procédés industriels.

Une interprétation plus intuitive de la matrice de causalité a fait l'objet de nombreux travaux. L'objectif est de permettre une compréhension plus facile et rapide de la topologie du système. Dans le cadre du diagnostic, l'analyse des causes doit être explicite dans le sens où la partie en défaut doit être aisément identifiable et les chemins de propagation du défaut facilement visualisables.

Nous citons Bauer (2005) et Seth (2010) qui proposent un diagramme à bulles. La largeur de chaque bulle représente la force de la dépendance causale entre une paire de variables. Une représentation par carte thermique est également souvent utilisée dans le domaine du traitement d'image où les valeurs sont associées à différentes couleurs. Le problème de cette représentation est qu'il est parfois difficile de comprendre la structure du système en se basant sur des couleurs dont les nuances peuvent se rapprocher.

Dans ce manuscrit, nous nous intéressons particulièrement à la représentation de la

topologie du système par un graphe orienté, noté $\mathcal{G}(\mathcal{S}, \mathcal{A})$, composé de deux ensembles finis :

- L'ensemble $\mathcal{S}$ qui représente l'ensemble des sommets associés aux variables mesurées du système ;
- L'ensemble $\mathcal{A}$ qui représente l'ensemble des arcs associés à l'existence de relations de causalité entre les variables mesurées du système.

2.4 Représentation des liens de causalité :

La complexité et la taille de certains systèmes rendent délicat leur diagnostic. Ce dernier requiert en effet l'utilisation de différentes approches transverses basées sur des connaissances graphiques et / ou directement sur des données. Divers travaux ont reconnu la nécessité et le bénéfice de mixer plusieurs approches complémentaires. Nous pouvons citer Vedam and Venkatasubramanian (1999); Chiang and Braatz (2003) qui proposent un diagnostic basé sur les données en utilisant un modèle qualitatif du système. L'intérêt de ce modèle est de consolider les résultats d'analyse des causes à travers l'analyse d'un graphe orienté signé représentant le procédé. Norvilas et al. (2000) proposent une approche de diagnostic basée conjointement sur l'utilisation de techniques statistiques multivariées et l'implémentation de règles d'expertises a priori (KBS¹). Les KBS sont des plateformes informatiques permettant d'implémenter des descriptions liées au type du procédé.

Leung and Romagnoli (2002) proposent également une méthode basée conjointement sur l'analyse statistique multivariée des mesures du procédé et sur une carte causale (modèle graphique) du procédé identifiée manuellement en utilisant la connaissance d'experts. Cette carte causale permet d'affiner les résultats de la localisation des défauts.

Lee et al. (2003) combinent des graphes orientés signés et des méthodes statistiques pour améliorer l'analyse des causes de défauts. Nous citons enfin Maurya et al. (2007) qui montrent les liens pouvant exister entre les graphes orientés signés et l'analyse des tendances pour le diagnostic. Les travaux de Maurya ont mis l'accent sur l'identification du modèle

1. Knowledge-based system

de causalité à partir des schémas P&ID ou PFD du système. Cela permet d'améliorer considérablement le diagnostic basé sur les données car le graphe est enrichi d'informations permettant d'augmenter la granularité du modèle descriptif. La combinaison de deux techniques issues d'approches basées sur les modèles / données permet de réduire le nombre de solutions inadéquates ou "parasites" .

Dans cette section, nous nous intéressons plus particulièrement aux modèles de connectivité qui permettent de représenter des liens physiques / informationnels (liens de causalité liés à la mise en place d'une structure de régulation, etc) entre les différentes parties du système (composants par exemple). Ces liens représentent une connaissance qualitative du procédé qui ne nécessite pas un modèle mathématique a priori. La ressource nécessaire à l'obtention d'un modèle de connectivité d'un système est un ensemble de paradigmes de modélisation d'un point de vue organique classiquement utilisés en industrie. Nous pouvons citer notamment, les diagrammes de flux du procédé (PFDs²) ou encore les schémas de tuyauterie et d'instrumentation (P&IDs³). Ces ressources permettent de décrire, de façon exhaustive, les interconnexions et les flux qui existent entre les composants du système. Ainsi, nous pouvons agréger ces ressources et les convertir vers des formats standards facilement utilisables par les opérateurs. Le format de conversation dépend de l'objectif de modélisation, par exemple le diagnostic, la tolérance aux défauts, la maintenance ou simplement la visualisation de la structure du système pour la compréhension et la consolidation d'une connaissance pour des fins d'expertise.

Dans cette section, nous présentons deux méthodes classiques de représentation de la connectivité/causalité. Un exemple est fourni pour les illustrer.

2.4.1 Matrice d'adjacence

La matrice d'adjacence permet de représenter les liens de causalité entre différentes entités du système (composants, variables, tâches, etc). La notion introduite par cette matrice est la

2. Process Flow Diagram

3. Process and Instrumentation Diagram

contiguïté. En effet, les relations de causalité représentées sont des relations de connectivité directe (notion qui peut être associée à un arc). Ainsi, les liens indirects (associés à la présence de chemins) ne peuvent être définis qu'après inférence de la matrice c.à.d. calcul de la matrice d'atteignabilité introduite dans cette section.

Avant de définir formellement la matrice d'adjacence, deux définitions sont introduites.

Définition 5 *Une matrice ne comportant que des éléments de valeur 0 ou 1 est appelée matrice booléenne.*

Définition 6 *L'arithmétique booléenne est basée sur l'utilisation des opérateurs logiques $\wedge$ (ET) et $\vee$ (OU) définis entre une paire de variables booléennes a et b par :*

$$a \wedge b = \begin{cases} 1 & \text{si } a = b = 1; \\ 0 & \text{sinon.} \end{cases} \quad (2.30)$$

$$a \vee b = \begin{cases} 1 & \text{si } a = 1 \text{ ou } b = 1; \\ 0 & \text{sinon.} \end{cases} \quad (2.31)$$

Considérons un système composé de n entités notées s_i avec $i \in \{1 \dots n\}$. Une matrice d'adjacence A de dimension $n \times n$ peut donc être définie de la manière suivante. Chaque élément de la matrice A noté a_{ij} est binaire (0 ou 1) : si un élément s_i est adjacent ou directement connecté à un élément s_j , alors $a_{ij} = 1$; sinon $a_{ij} = 0$.

La matrice d'atteignabilité, calculée à partir de la matrice d'adjacence A , est notée R . Elle permet de visualiser les liens directs et indirects entre les variables. La taille des liens indirects dépend de la puissance de la matrice A . Si une entité s_i n'est pas directement connectée à s_j , l'entité s_j peut éventuellement être connectée à s_i à travers une troisième entité s_l . L'atteignabilité de s_j s'intitule atteignabilité à 2-pas. Cette dernière notion est introduite afin de la différencier de celle de la matrice d'adjacence (atteignabilité à 1-pas).

Cette notion d'atteignabilité à 2-pas peut être aisément étendue à k-pas. Avant de

formellement définir la matrice d'atteignabilité R , quelques définitions sont nécessaires.

Définition 7 *La somme booléenne de deux matrices A et B de même dimension est la matrice dont les éléments correspondent à la somme booléenne des éléments de A et B . Cette somme est notée $A \oplus B$.*

La somme booléenne consiste à remplacer dans la somme classique tous les réels supérieurs ou égaux à 1 par 1.

Définition 8 *Soient A et B deux matrices booléennes respectivement de dimensions $m \times k$ et $k \times n$. Alors le produit booléen de A par B noté $A \odot B$ est une matrice C de dimension $m \times n$ avec un $(i,j)^{\text{ème}}$ élément noté c_{ij} et donné par :*

$$c_{ij} = (a_{i1} \wedge b_{1j}) \vee (a_{i2} \wedge b_{2j}) \vee \dots \vee (a_{ik} \wedge b_{kj}) \quad (2.32)$$

Les opérateurs $\wedge$ et $\vee$ représentent respectivement l'opérateur "ET" logique (conjonction booléenne) et "OU" logique (disjonction booléenne).

D'après cette définition, la $n^{\text{ème}}$ puissance booléenne d'une matrice A correspond à n produits de la matrice A :

$$A^n = \underbrace{A \odot A \odot \dots \odot A}_{n \text{ fois}} \quad (2.33)$$

Si une entité s_j n'est pas atteignable par s_i au bout de n -pas, alors s_j ne sera jamais atteignable par s_i . Ceci permet d'écrire la matrice d'atteignabilité comme une somme booléenne de puissance de A^i avec $i \in \{1 \dots n\}$:

$$R^{(n)} = A^1 \oplus \dots \oplus A^n \quad (2.34)$$

Les éléments de $R^{(n)}$ notés r_{ij} représentent la présence d'un lien directs ou indirects (chemin) entre s_i et s_j de longueur inférieure ou égale à n .

Exemple d'illustration :

Considérons le système à cuve dont le schéma P&ID est représenté à la Figure 2.23 :

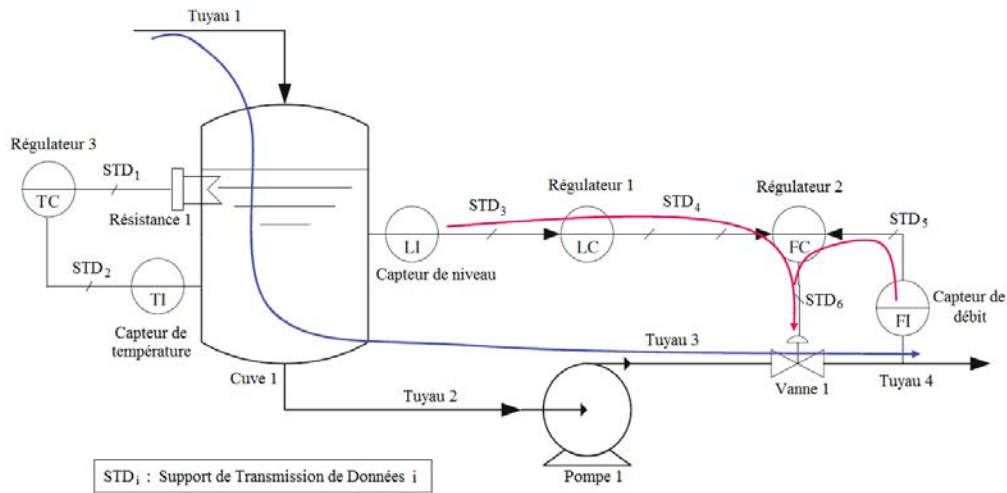


FIGURE 2.23 – Schéma P&ID du système à cuve

Le système de commande se compose de deux structures de boucle de régulation :

- Boucle de régulation simple pour la régulation de la température du liquide au niveau de la cuve ;
- Boucle de régulation en cascade pour la régulation de niveau de la cuve.

Pour représenter la connectivité entre les entités telles que la cuve, les capteurs (de débit, de niveau et de température), les actionneurs, les liens d'information (connexions entre les organes de contrôle/commande) et les composants élémentaires (tuyauterie), une matrice d'adjacence peut être construite. À travers cette matrice d'adjacence, l'atteignabilité à k -pas peut être obtenue. Pour cet exemple, la matrice d'adjacence, la matrice d'atteignabilité à (1 ou 2)-pas et la matrice d'atteignabilité à n -pas sont données à titre d'illustration à la Figure 2.24. La Table 2.2 résume les correspondances entre les entités s_i (ou s_j) et les éléments dans le schéma P&ID pour rendre sa lecture plus aisée.

2.4.2 Graphe orienté

Une alternative à la représentation matricielle de la connectivité par la matrice d'adjacence est le graphe orienté. Un graphe orienté offre une intuitivité de visualisation des entités du

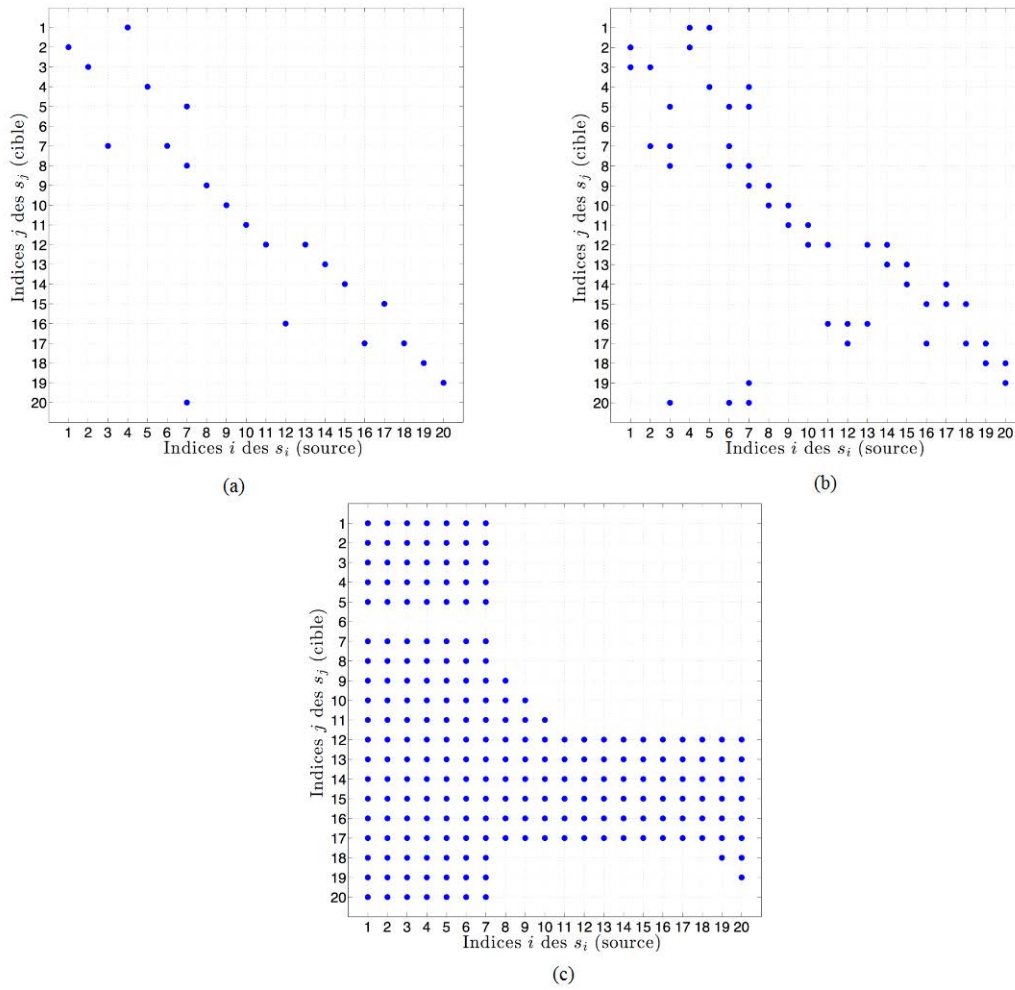


FIGURE 2.24 – Topologie du système : (a) Matrice d'adjacence (b) Matrice d'atteignabilité à (1 ou 2)-pas (c) Matrice d'atteignabilité $R^{(n)}$

<i>Sommet</i>	<i>Correspondance</i>	<i>Sommet</i>	<i>Correspondance</i>	<i>Sommet</i>	<i>Correspondance</i>
s_1	Régulateur 3	s_2	STD_1	s_3	Résistance 1
s_4	STD_2	s_5	Capteur de température	s_6	Tuyau 1
s_7	Cuve 1	s_8	Capteur de niveau	s_9	STD_3
s_{10}	Régulateur 1	s_{11}	STD_4	s_{12}	Régulateur 2
s_{13}	STD_5	s_{14}	Capteur de débit	s_{15}	Tuyau 4
s_{16}	STD_6	s_{17}	Vanne 1	s_{18}	Tuyau 3
s_{19}	Pompe 1	s_{20}	Tuyau 2		

TABLE 2.2 – Sommets et leurs correspondances

système et leurs interconnexions car le graphe orienté représente simplement le PFD ou le P&ID par l'abstraction de chaque entité en un sommet du graphe. Le passage d'une matrice d'adjacence à un graphe orienté est immédiat. En fait, chaque entité est représentée par un sommet et chaque élément à un (1) dans la matrice A signifie la présence d'un arc entre la

paire de sommets du graphe correspondant.

Ainsi, à travers les techniques de la théorie des graphes, des analyses de certaines propriétés du système peuvent être conduites. Les méthodes de parcours de graphe peuvent être utilisées comme une variante par rapport au calcul matriciel. Le résultat est équivalent et peut être visualisé aisément.

Les algorithmes d'analyse graphique ont souvent des complexités algorithmiques polynomiales offrant un avantage majeur lors de l'analyse graphique d'un système de grande dimension.

Pour tester l'atteignabilité d'un sommet par rapport à un autre, un parcours du graphe en profondeur peut être utilisé afin d'identifier l'ensemble des chemins entre ces sommets. Ceci est équivalent à l'évaluation de certains éléments r_{ij} de la matrice R afin de statuer sur la présence des liens directs/indirects entre les entités du système.

Le graphe établi à partir du PFD ou du P&ID peut être utilisé pour différents objectifs. Des paramètres quantitatifs peuvent être ajoutés au graphe orienté afin d'augmenter sa granularité descriptive. Ceci peut s'avérer nécessaire pour certains procédés dont le graphe orienté seul ne permet pas une bonne compréhension, voire induit une ambiguïté ou une fausse analyse d'une problématique donnée.

Le graphe orienté du système associé au schéma P&ID de la Figure 2.23 est donné à la Figure 2.25.

Pour cet exemple, il est possible de déduire du graphe orienté l'arbre de défaillances qui considère l'événement redouté suivant : débit sortant insuffisant. Cet arbre de défaillances est donné à la Figure 2.26. A travers l'arbre de défaillances, on peut associer à chaque cause possible de l'événement redouté une fréquence d'apparition. Cette fréquence peut être calculée en utilisant des algorithmes de détection de défaut basés sur des mo-

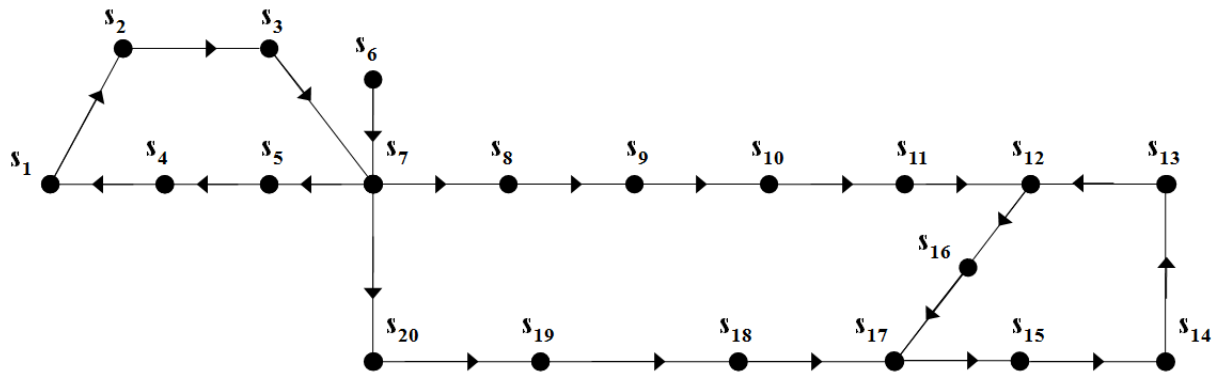


FIGURE 2.25 – Graphe orienté du système à cuve

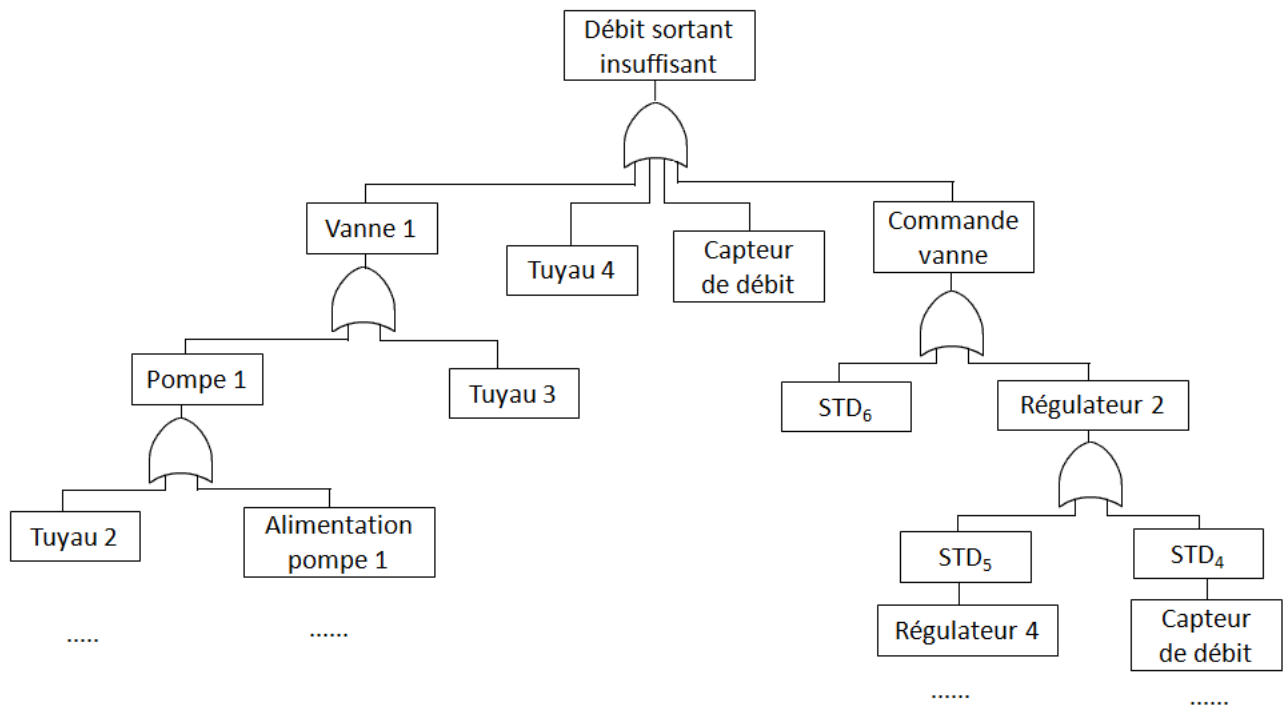


FIGURE 2.26 – Arbre de défaillances associé au système de la Figure 2.23

dèles/données afin d'associer à chaque cause candidate une fréquence permettant ainsi l'aide à la décision lors de la phase de maintenance du procédé. En industrie, il existe différentes sources de capitalisation de ce type d'informations. En effet, des systèmes de maintenance intégrée tels que les SAP permettent l'historisation de données et d'informations à des fins de mise en place de stratégies de maintenance préventive et/ou curative.

2.5 Diagnostic et analyse des causes basés sur un modèle de causalité

2.5.1 Introduction à l'approche Top-Down pour le diagnostic

Dans cette partie, nous présentons l'approche adoptée pour le diagnostic des systèmes de grande dimension.

En l'occurrence d'un défaut dégradant certains indicateurs clés de performance PPI⁴ d'un système, il est nécessaire de déterminer l'ensemble des causes candidates. Il est probable que ce défaut s'est propagé à travers différents sous-systèmes c.à.d. des boucles de régulation ce qui provoque une dégradation des performances du système.

Pour un système complexe de grande dimension, la localisation de la boucle de régulation en défaut par les ressources humaines est souvent une tâche fastidieuse, difficile et coûteuse. Ceci est dû à la complexité des interconnexions et aussi à la taille du système qui devient parfois difficilement gérable par les opérateurs.

Afin de rendre la détection et la localisation des défauts d'actionneurs plus aisées, une approche de diagnostic Top-Down est proposée Blanke et al. (2006). Le principe de cette approche est la surveillance en temps-réel des indicateurs clé de performance du système dont le nombre est souvent restreint. Comme présenté à la Figure 2.27, lors de la détection d'une dégradation d'un indicateur clé, une procédure d'analyse des causes est lancée afin d'identifier "le(s) chemin(s)" de propagation de(s) défaut(s).

Ainsi, l'objectif principal est de remonter à la boucle de régulation en défaut causant la dégradation du PPI. Cette procédure se base conjointement sur le parcours du graphe de causalité et l'utilisation de tests statistiques séquentiels afin de permettre l'identification de l'ensemble des chemins de propagation de(s) défaut(s). Le parcours du graphe répond aussi à une contrainte essentielle liée à la réduction de la dimension du système. En effet, l'analyse

4. Indicateur de Performance Procédé

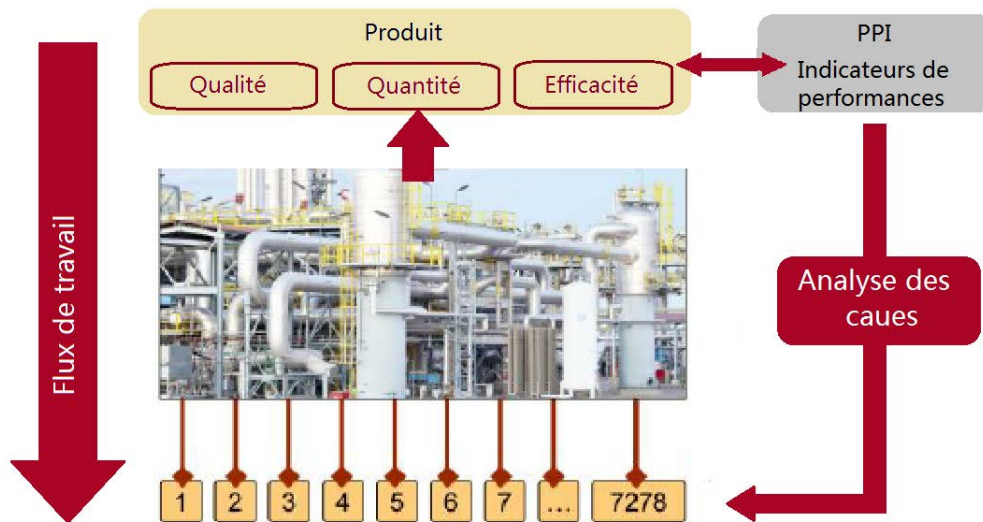


FIGURE 2.27 – Approche Top-Down

des liens de causalité, après la détection d'un défaut sur un PPI, permet d'appliquer les tests statiques sur un nombre limité d'indicateurs précédant chaque indicateur dégradé. Ceci permet d'utiliser le principe de causalité entre les indicateurs de performances afin d'identifier les causes candidates en utilisant le moins possible les ressources humaines/matérielles.

2.5.2 Caractérisation des dégradations statistiques des performances du système

Dans tout procédé, quelle que soit sa conception ou sa maintenance, un certain nombre de facteurs de variabilité naturelle existent. Cette variabilité est l'effet cumulé de plusieurs causes souvent inévitables. Au sens gaussien, cette variabilité est appelée "causes communes" (main d'œuvre, matières premières, outils de fabrication, instruments de mesure, etc) Fleming (1974). Un processus qui fonctionne dans le cadre d'une variabilité naturelle est sous une maîtrise statistique ou considéré en mode nominal. À contrario, nous appelons "causes spéciales" la variabilité exogène due éventuellement à des défauts dans le procédé dont les mesures permettent de les détecter et de les localiser si le modèle utilisé est pertinent.

Cette variabilité exogène peut parfois être présente sur différentes mesures du procédé dont les indices de performance clés. Elle découle habituellement de 4 sources :

commande inadaptée du système, erreurs de manipulation (opérateur), matières premières non conformes ou système en défaut. Cette variabilité est généralement plus importante par rapport à la variabilité naturelle, et elle représente habituellement une dégradation inacceptable des performances du système. Un système qui fonctionne en présence d'une cause exogène de variabilité due à un ou plusieurs défauts est dit "en mode défectueux".

Plusieurs cas de changement des caractéristiques statistiques sont présentés à la Figure 2.28.

Jusqu'à l'instant t_1 , le système est sous maîtrise statistique. Seule la variabilité naturelle est présente. En conséquence, la moyenne et l'écart-type de la série temporelle, relatives à un indicateur de performance étudié, sont à leurs valeurs nominales μ_0 et σ_0 .

- Cas A : à l'instant t_1 , un premier défaut se produit. Comme le montre la Figure 2.28, l'effet de ce défaut se traduit par un changement de la moyenne de la série temporelle en une nouvelle valeur $\mu_1 > \mu_0$ (saut de moyenne).
- Cas B : à l'instant t_2 , un second défaut se produit. Dans ce cas, l'écart-type de la série temporelle a changé en une valeur plus grande $\sigma_1 > \sigma_0$ (saut de variance).
- Cas C : à l'instant t_3 , un dernier défaut impacte simultanément la moyenne et l'écart type de la série temporelle.

Lorsqu'un défaut majeur se produit, le système change de mode de fonctionnement d'un mode nominal à un mode défectueux. Ce changement rend les indices de performance du système non conformes aux spécifications et aux exigences prédéfinies.

Par exemple, à partir de la Figure 2.28, lorsque le système est en mode de fonctionnement nominal, la plupart des échantillons se situent entre les limites des exigences

inférieure et supérieure, respectivement LSL et USL (Lower and Upper Specification Limits). Lorsque le système est en défaut, une proportion plus élevée des échantillons du procédé se trouve en dehors de ces exigences.

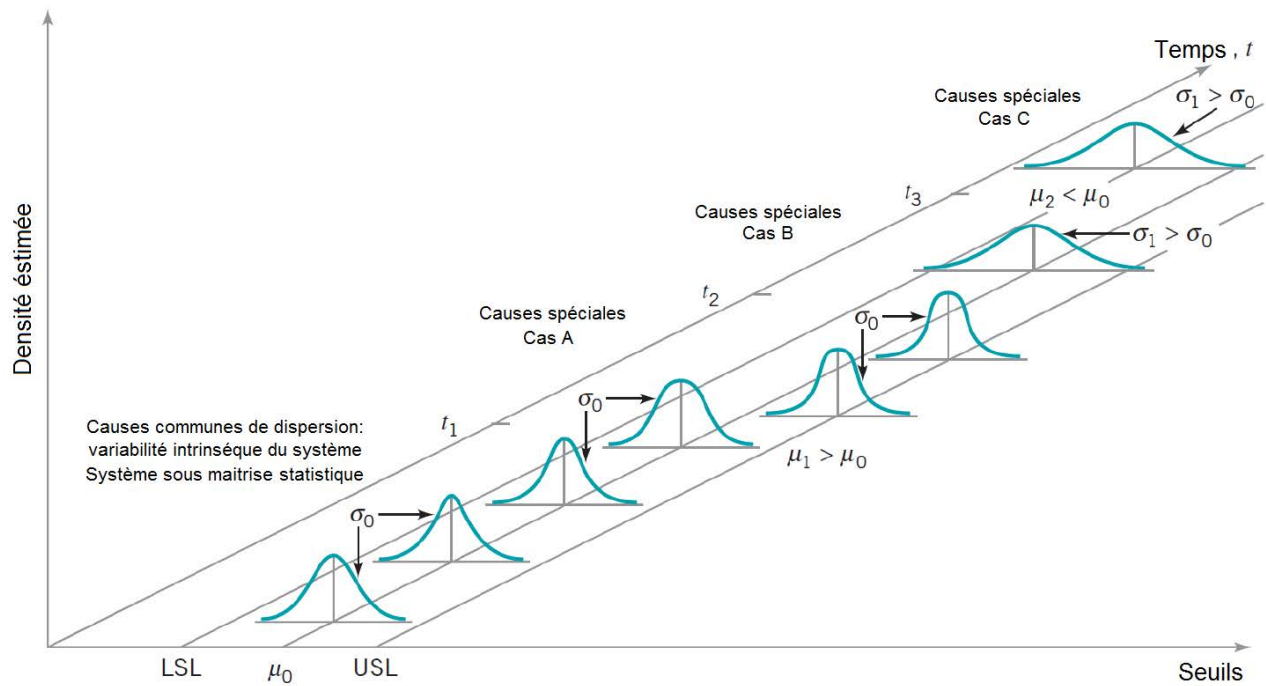


FIGURE 2.28 – Modes de fonctionnement définis par les caractéristiques statistiques des indicateurs de performance

Un objectif important de la maîtrise statistique des procédés est de détecter rapidement l'apparition de changements dans les caractéristiques statistiques d'un nombre limité d'observations KPI⁵, puis de déclencher le plus rapidement possible une mesure corrective adéquate pour rétablir les KPI dans des limites acceptables. La section suivante a pour objectif d'introduire deux types de tests statistiques. Le premier test statistique est un test hors ligne. Ce test consiste à maximiser la probabilité de détection P_D étant donnée la probabilité de fausse alarme P_F . Ce test se base sur un nombre fixe d'observations.

5. Indicateur de Performance Clé

Le second test est un test en ligne. Dans ce type de test (appelé test de Wald) le nombre d'observations n'est pas connu à priori. Si l'information n'est pas suffisante pour prendre une décision sur un risque donné, la décision n'est pas prise en attente d'une autre observation. Ce type de test permet une prise de décision en utilisant un nombre d'observations limité.

2.5.3 Test statistique de détection de dégradation

2.5.3.1 Introduction

Les développements liés aux tests statistiques couvrent un champ thématique très vaste allant de la théorie de la décision aux mathématiques appliquées en passant par le traitement du signal (Brunet et al. (1990); Basseville and Nikiforov (1993); Wald (2004)).

Dans cette partie, un test d'hypothèses séquentiel pour la surveillance des indicateurs de performances du système est présenté. L'objectif visé par ce type d'approche est de détecter le plus rapidement possible la présence d'un défaut. Dans le cadre de notre application, on considère deux modes de fonctionnement du système : l'état nominal et l'état en défaut.

Nous présentons deux tests statistiques d'hypothèses simples $\mathcal{H}_0$ et $\mathcal{H}_1$. Le premier est un test hors ligne et le second est un test en ligne. Nous nous intéressons plus particulièrement au deuxième test pour la surveillance du système et l'identification de son mode de fonctionnement en temps réel.

Pour la construction de ce test, certaines hypothèses sont émises :

Hypothèse 1

Considérons des mesures z_k . La distribution des observations z_k suit une loi normale $\mathcal{N}$ de moyenne μ et d'écart type σ .

Les mesures sont considérées comme des variables aléatoires indépendantes et identiquement distribuées (i.i.d).

	Hypothèse $\mathcal{H}_0$ est vraie	Hypothèse $\mathcal{H}_1$ est vraie
Accepter $\mathcal{H}_0$	Décision correcte (D_{00})	Décision incorrecte (D_{01})
Accepter $\mathcal{H}_1$	Décision incorrecte (D_{10})	Décision correcte (D_{11})

TABLE 2.3 – Cas possibles de la prise de décision

2.5.3.2 Test de Neyman-Pearson

Considérons que les densités de probabilité conditionnelles des observations z_k relativement aux deux hypothèses $\mathcal{H}_0$ et $\mathcal{H}_1$ sont connues. L'objectif du test consiste à décider, à partir des observations effectuées sur le système, selon quelle hypothèse le système fonctionne (test hors ligne). Notons dorénavant Z le vecteur de dimension n des observations z_k .

Supposons que l'espace d'observation est divisé en deux sous-espaces disjoints (E_0 et E_1) tels que $E_0 : z \leq \gamma$ et $E_1 : z > \gamma$. Si l'observation courante appartient à E_0 , l'hypothèse $\mathcal{H}_0$ est choisie, sinon $\mathcal{H}_1$. Ainsi, quatre décisions (D_{00} , D_{01} , D_{10} et D_{11}) peuvent être considérées. Elles sont résumées dans la Table 2.3. Il apparaît deux situations de décision incorrecte, D_{01} qui correspond à une erreur de type I et D_{10} qui correspond à une erreur de type II (Figure 2.29).

Si l'hypothèse $\mathcal{H}_1$ correspond à la présence d'un défaut, la décision D_{11} est relative à une bonne détection dont la probabilité est $P_D = 1 - \beta$. D_{10} correspond à une fausse alarme dont la probabilité est $P_F = \alpha$ et D_{01} à un défaut non-déecté de probabilité $P_{ND} = \beta$.

Les densités de probabilité conditionnelles des observations z_k relativement aux deux hypothèses $\mathcal{H}_0$ et $\mathcal{H}_1$ étant connues, ces trois situations sont caractérisées par les probabilités :

$$\left\{ \begin{array}{l} P_D = 1 - \beta = P(\mathcal{H}_1 \text{ choisie} \mid \mathcal{H}_1 \text{ vraie}) = \int_{\gamma}^{\infty} p(Z \mid \mathcal{H}_1) dZ. \\ P_F = \alpha = P(\mathcal{H}_1 \text{ choisie} \mid \mathcal{H}_0 \text{ vraie}) = \int_{\gamma}^{\infty} p(Z \mid \mathcal{H}_0) dZ. \\ P_{ND} = \beta = P(\mathcal{H}_0 \text{ choisie} \mid \mathcal{H}_1 \text{ vraie}) = \int_{-\infty}^{\gamma} p(Z \mid \mathcal{H}_1) dZ. \end{array} \right. \quad (2.35)$$

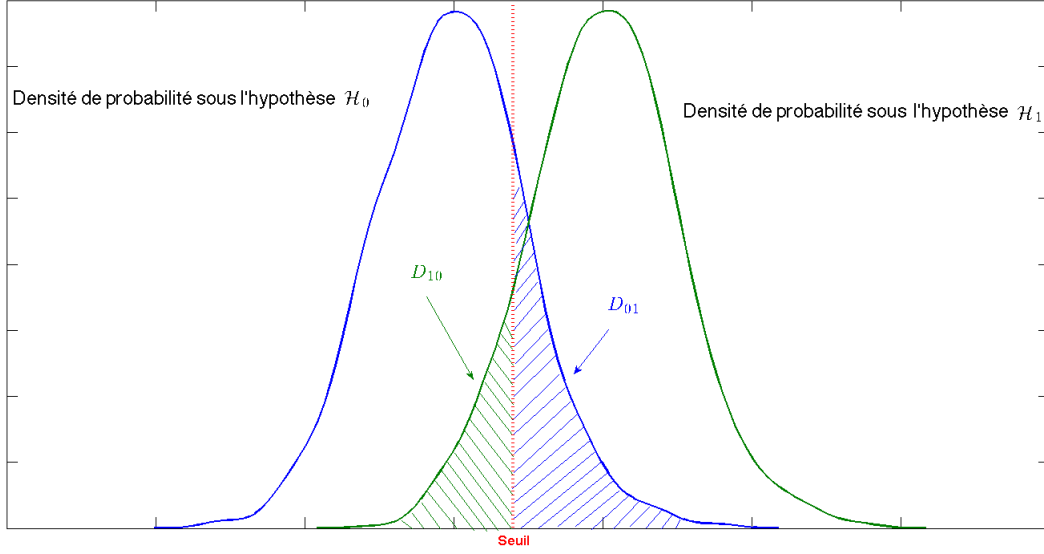


FIGURE 2.29 – Régions de décision

Le test de Neyman-Pearson, à probabilité de fausse alarme fixée $P_F = \alpha$, maximise la puissance du test P_D . Ces deux exigences étant contradictoires, la décision résulte forcément d'un compromis entre les deux exigences.

Ce problème peut être vu comme un problème d'optimisation consistant à minimiser P_{ND} sous contrainte $P_F - \alpha = 0$.

Ainsi la fonction coût peut être mise sous la forme :

$$J = P_{ND} + \lambda(P_F - \alpha) \quad (2.36)$$

En remplaçant P_{ND} et P_F par leur expression donnée dans (2.35) :

$$J = \int_{-\infty}^{\gamma} p(Z | \mathcal{H}_1) dZ + \lambda \left(\int_{\gamma}^{\infty} p(Z | \mathcal{H}_0) dZ - \alpha \right) \quad (2.37)$$

La fonction coût (2.37) peut être mise aussi sous la forme :

$$J = \int_{-\infty}^{\gamma} (p(Z | \mathcal{H}_1) - \lambda p(Z | \mathcal{H}_0)) dZ + \lambda(1 - \alpha) \quad (2.38)$$

Avec : $\int_{\gamma}^{\infty} p(Z | \mathcal{H}_0) dZ = 1 - \int_{-\infty}^{\gamma} p(Z | \mathcal{H}_0) dZ.$

Le critère J est minimal si l'on affecte au sous-espace E_0 (l'espace qui correspond à l'acceptation de $\mathcal{H}_0$) les observations telles que $p(Z | \mathcal{H}_1) - \lambda p(Z | \mathcal{H}_0) < 0$, c'est-à-dire :

$$\frac{p(Z | \mathcal{H}_1)}{p(Z | \mathcal{H}_0)} < \lambda$$

La règle de la décision s'exprime alors de la manière suivante :

$$\Lambda(Z) = \frac{p(Z | \mathcal{H}_1)}{p(Z | \mathcal{H}_0)} \underset{H_0}{\overset{H_1}{\geq}} \lambda \quad (2.39)$$

Le rapport $\Lambda(Z)$ est le rapport de vraisemblance des densités de probabilité conditionnelles et λ est le seuil de détection. La notation employée signifie que si ce rapport est inférieur à λ , on accepte l'hypothèse $\mathcal{H}_0$; dans le cas contraire, on accepte $\mathcal{H}_1$.

La valeur du paramètre λ peut être définie en utilisant la contrainte $P_F - \alpha = 0$, on peut écrire :

$$\int_{\gamma}^{\infty} p(Z | \mathcal{H}_0) dZ = \alpha \quad (2.40)$$

La résolution de cette équation fournit la valeur γ . Les lois de probabilités conditionnelles $p(Z | \mathcal{H}_0)$ et $p(Z | \mathcal{H}_1)$ étant connues, l'inéquation (2.39) fournit la valeur du seuil de détection qui est solution de :

$$\left. \frac{p(Z | \mathcal{H}_1)}{p(Z | \mathcal{H}_0)} \right|_{Z=\gamma} = \lambda \quad (2.41)$$

Nous présentons maintenant un test séquentiel usuellement utilisé et permettant de détecter des sauts en temps réel. Ce test présente l'intérêt d'être facile d'implémentation permettant ainsi de répondre à certaines exigences du projet européen et notamment à celle relative à la proposition d'algorithmes simples (contrainte de réactivité).

2.5.3.3 Test séquentiel de vraisemblance (Test de Wald)

Les tests séquentiels d'hypothèses ont été introduits par Wald (2004, 1945). Ils sont généralement utilisés pour la détection de changements dans un système. Ils se basent sur le principe suivant lequel la loi de probabilité de la variable décrivant le système/sous-système diffère de la loi avant le changement après un changement. Parmi les tests d'hypothèses

séquentiels les plus fréquemment utilisés, nous citons le test séquentiel du rapport de vraisemblance (SPRT) qui s'applique à deux hypothèses simples. Ce test est utilisé afin de détecter des changements abrupts (saut de moyenne/variance) et ainsi apprécier le mode de fonctionnement du système à un instant donné.

L'intégrité des décisions proposées par le test précédent est directement liée au nombre d'observations qui constituent le vecteur de mesures Z . L'avantage majeur des tests statistiques séquentiels est leur implémentation en ligne.

2.5.3.3.1 Principe de fonctionnement du test séquentiel de vraisemblance

Le test séquentiel de vraisemblance a pour objectif de décider entre deux hypothèses simples concernant un paramètre Θ (moyenne ou variance).

$$\begin{cases} \mathcal{H}_0 : \Theta = \Theta_0 \text{ (fonctionnement nominal)} \\ \mathcal{H}_1 : \Theta = \Theta_1 \text{ (fonctionnement en défaut)} \end{cases} \quad (2.42)$$

Dans ce type de test, au fur et à mesure de l'obtention des mesures, une fonction de ces dernières est calculée. Cette fonction est confrontée à deux limites correspondant au refus et à l'acceptation d'une hypothèse.

Après chaque mesure, une des trois décisions suivantes est prise :

- Accepter $\mathcal{H}_0$;
- Accepter $\mathcal{H}_1$;
- Une mesure supplémentaire est nécessaire (pas de prise de décision).

L'acceptation de l'une ou l'autre des hypothèses ou la non-prise de décision s'effectue en fonction de la valeur du rapport de vraisemblance défini précédemment.

Soit $z_1, z_2, \dots, z_k$ la séquence d'observations disponibles à l'instant k . Le rapport de vraisemblance s'écrit :

$$\Lambda(z_1, z_2, \dots, z_k) = \frac{p(z_1, z_2, \dots, z_k \mid \mathcal{H}_1)}{p(z_1, z_2, \dots, z_k \mid \mathcal{H}_0)} \quad (2.43)$$

Les valeurs des risques d'erreurs α et β sont fixées. Pour retenir l'hypothèse $\mathcal{H}_1$, il faut que ce rapport de vraisemblance soit tel que le risque d'erreur de première espèce (erreur de type I) soit inférieur ou égal à α si l'on retient $\mathcal{H}_1$ alors que $\mathcal{H}_0$ est vraie.

Si γ_1 est le seuil de décision correspondant, il est nécessaire d'avoir :

$$\Lambda(z_1, z_2, \dots, z_k) \geq \gamma_1 \quad (2.44)$$

Le seuil γ_1 peut être déterminé en considérant le cas limite où l'on a l'égalité des deux membres, c.à.d. :

$$\int_{\gamma_1}^{\infty} p((z_1, z_2, \dots, z_k | \mathcal{H}_1) dz_1 \dots dz_k = \gamma_1 \int_{\gamma_1}^{\infty} p((z_1, z_2, \dots, z_k | \mathcal{H}_0) dz_1 \dots dz_k \quad (2.45)$$

ou encore : $1 - \beta = \alpha \gamma_1$

Le seuil de décision s'exprime donc en fonction des risques d'erreurs sous la forme :

$$\gamma_1 = \frac{1 - \beta}{\alpha} \text{ (Égalité de Wald)} \quad (2.46)$$

De manière similaire, l'hypothèse $\mathcal{H}_0$ sera retenue si le rapport de vraisemblance est tel que le risque d'erreur de seconde espèce est inférieur ou égal à β si l'on retient $\mathcal{H}_0$ alors que $\mathcal{H}_1$ est vraie. Une démarche tout à fait analogue à la précédente permet de déterminer un second seuil de décision γ_0 qui s'exprime également en fonction des risques d'erreurs :

$$\gamma_0 = \frac{\beta}{1 - \alpha} \text{ (Égalité de Wald)} \quad (2.47)$$

La règle de décision est :

- Si $\Lambda(z_1, z_2, \dots, z_k) \leq \gamma_0$, alors $\mathcal{H}_0$ est acceptée ;
- Si $\Lambda(z_1, z_2, \dots, z_k) \geq \gamma_1$, alors $\mathcal{H}_1$ est acceptée ;
- Si $\gamma_0 < \Lambda(z_1, z_2, \dots, z_k) < \gamma_1$, alors pas de prise de décision et la collecte des mesures doit se poursuivre jusqu'à l'obtention du premier rapport dont la valeur est extérieure aux limites.

L'instant de détection de défaut $T_D(k)$ peut être aisément défini par :

$$T_D(k) = \min\{k > 0; \Lambda(z_1, z_2, \dots, z_k) \geq \gamma_1\} \quad (2.48)$$

De façon similaire, l'instant de retour à un fonctionnement nominal du système peut être défini par :

$$T_{FN}(k) = \min\{k > 0; \Lambda(z_1, z_2, \dots, z_k) \leq \gamma_0\} \quad (2.49)$$

La stratégie que nous mettons en œuvre, et qui est souvent utilisée dans le cadre du test séquentiel de vraisemblance (Baghli (2006)), consiste à poursuivre la construction du test même sous validation des hypothèses $\mathcal{H}_0$ ou $\mathcal{H}_1$. Le test de Wald est efficace pour détecter des ruptures dans le comportement d'un signal au cours du temps. Cette propriété est particulièrement intéressante dans le contexte de la détection de défauts.

Test séquentiel et défauts saturants

Dans le cadre du diagnostic de défauts saturants, le défaut peut être considéré comme étant un saut de moyenne ou de variance. Nous mettons ainsi l'accent sur l'application du test séquentiel dans la détection d'un biais/saut reflétant une dégradation des performances.

Test séquentiel sur la moyenne

Soit Z une variable aléatoire (i.i.d) normalement distribuée de variance σ^2 . Le fonctionnement nominal est caractérisé par une moyenne μ_0 alors qu'en l'occurrence d'un défaut, la moyenne de Z devient μ_1 . Nous distinguons les deux hypothèses suivantes :

$$\begin{cases} \mathcal{H}_0 : \Theta = \mu_0 \text{ (fonctionnement nominal)} \\ \mathcal{H}_1 : \Theta = \mu_1 \text{ (fonctionnement en défaut)} \end{cases} \quad (2.50)$$

Le rapport de vraisemblance peut s'écrire pour la $k^{\text{ème}}$ observation :

$$\begin{aligned}\Lambda(z_k) &= \frac{p(z_k | \mathcal{H}_1)}{p(z_k | \mathcal{H}_0)} \\ \Lambda(z_k) &= \frac{\frac{1}{\sigma\sqrt{2\pi}} \exp(-\frac{1}{2\sigma^2}(z_k - \mu_1)^2)}{\frac{1}{\sigma\sqrt{2\pi}} \exp(-\frac{1}{2\sigma^2}(z_k - \mu_0)^2)} \\ \Lambda(z_k) &= \exp(-\frac{1}{2\sigma^2}(z_k - \mu_1)^2 + \frac{1}{2\sigma^2}(z_k - \mu_0)^2)\end{aligned}\tag{2.51}$$

Considérons une séquence d'observations $z_1, z_2, \dots, z_k$. Le rapport de vraisemblance s'écrit :

$$\begin{aligned}\Lambda(z_1, z_2, \dots, z_k) &= \prod_{i=1}^k \Lambda(z_i) \\ &= \exp\left(-\frac{1}{\sigma^2}(\mu_1 - \mu_0) \sum_{i=1}^k (z_i - \frac{(\mu_0 + \mu_1)}{2})\right)\end{aligned}\tag{2.52}$$

Si les seuils λ_0 et λ_1 sont choisis en fonction de la probabilité de fausse alarme P_F et de la probabilité de non-détection P_{ND} que l'on accepte, le test séquentiel de vraisemblance est le suivant :

- Si $\frac{1}{\sigma^2}(\mu_1 - \mu_0) \sum_{i=1}^k (z_i - \frac{(\mu_0 + \mu_1)}{2}) \leq \log(\lambda_0)$ alors accepter $\mathcal{H}_0$.
- Si $\log(\lambda_0) < \frac{1}{\sigma^2}(\mu_1 - \mu_0) \sum_{i=1}^k (z_i - \frac{(\mu_0 + \mu_1)}{2}) < \log(\lambda_1)$ alors collecter une nouvelle mesure.
- Si $\log(\lambda_1) \leq \frac{1}{\sigma^2}(\mu_1 - \mu_0) \sum_{i=1}^k (z_i - \frac{(\mu_0 + \mu_1)}{2})$ alors accepter $\mathcal{H}_1$.

2.5.4 Analyse des causes

2.5.4.1 Procédure d'analyse des causes

Dans cette section, nous présentons une procédure permettant d'analyser les causes d'un défaut affectant un indicateur de performance procédé (PPI) en utilisant une approche graphique. Cette approche est basée sur le concept Top-Down précédemment décrit.

Les indicateurs de performance d'un système et leurs liens de causalité sont représentés dans un modèle graphique. Ce modèle permet d'exhiber les différentes relations entre un ensemble d'indicateurs locaux de performance LPI et des indicateurs clés de performance procédé PPI. L'analyse des causes proposée dans ce travail de thèse peut être présentée sous formes d'étapes. Ces étapes consistent à déterminer les chemins de propagation candidats à la dégradation d'un indicateur clé de performance (PPI).

Le graphe, noté $\mathcal{G}$, représentant le système est décomposé en plusieurs niveaux. Ce graphe contient l'indicateur de performance PPI en défaut, défini comme étant un sommet puits, et pour chaque niveau, des prédécesseurs LPI de cet indicateur (PPI). A titre d'illustration de la procédure proposée, nous considérons la Figure 2.30 avec un graphe composé de 3 niveaux où le sommet cible en défaut est s_1 . Les sommets en défaut (correspondants aux indicateurs de performance en défaut) sont représentés en rouge.

Hypothèse 1 Nous considérons pour cette analyse des causes que seul un défaut est pris en compte.

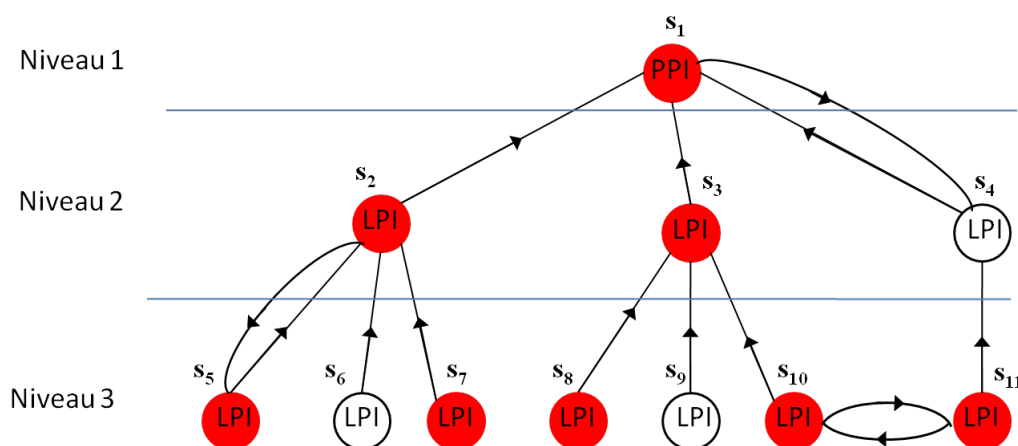


FIGURE 2.30 – Exemple d'un graphe $\mathcal{G}$ de causalité décomposé en niveaux

Chaque sommet s_k en défaut est associé à un temps de détection TD_k . Pour les sommets de la Figure 2.30, les TD_k correspondants sont regroupés dans la Table 2.4.

sommet s_k	Temps de détection TD_k	sommet s_k	Temps de détection TD_k
s_1	8s	s_2	5s
s_3	6s	s_5	3s
s_7	7s	s_8	1s
s_{10}	3s	s_{11}	5s

TABLE 2.4 – Temps de détection pour les sommets en défaut

Étape 1 :

Cette étape consiste à classer les sommets prédécesseurs du sommet en défaut en fonctions de leurs temps de détection TD_k .

Pour chaque niveau i , nous commençons par définir un ensemble de vecteurs Q_{ij} où i correspond aux niveaux du graphe $\mathcal{G}$ et j correspond aux indices des sommets prédécesseurs (direct) en défaut dans le niveau $i - 1$.

Chaque vecteur Q_{ij} contient les temps de détection TD_k inférieurs au temps de détection TD_j i.e. $TD_k < TD_j$ par respect du principe de la causalité temporelle. Ces temps de détection TD_k sont ordonnés dans les vecteurs Q_{ij} par ordre décroissant.

Illustration sur l'exemple de la Figure 2.30.

Niveau 1 :

$$Q_{10} = [TD_1]$$

Dans ce niveau 1, il n'y a qu'un seul sommet en défaut qui correspond à TD_1 c.à.d. le sommet s_1 PPI.

Niveau 2 :

$$Q_{21} = [TD_3, TD_2]$$

Dans ce niveau 2, il y a deux sommets en défaut prédécesseurs de s_1 et correspondant à TD_2 et TD_3 .

Niveau 3 :

$$Q_{32} = [TD_5], Q_{33} = [TD_{10}, TD_8]$$

Dans ce niveau 3, il y a cinq sommets en défaut.

s_5 et s_7 sont prédécesseurs de s_2 . Seul s_5 est considéré dans Q_{32} puisque $TD_5 = 3s < TD_2 = 5s$, ce qui n'est pas le cas de s_7 ($TD_7 = 7s > TD_2 = 5s$). s_8 et s_{10} sont prédécesseurs de s_3 . Les temps de détection correspondants TD_8 et TD_{10} sont représentés dans Q_{33} dans l'ordre croissant puisqu'ils sont tous les deux inférieurs à TD_3 . s_{11} est en défaut et appartient au niveau 3 mais il n'est pas considéré dans Q_{33} puisqu'il n'est pas successeur direct de s_3 .

Étape 2 :

Dans cette étape, nous utilisons les résultats de la première étape pour obtenir un nouveau graphe pondéré permettant la représentation des chemins de propagation du défaut.

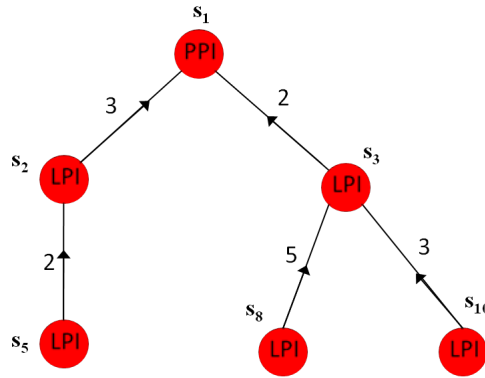
En effet, après avoir défini les vecteurs Q_{ij} en première étape, nous définissons un sous-graphe pondéré $\mathcal{G}_1$ issu du graphe initial $\mathcal{G}$. Ce nouveau graphe ne contient que les sommets associés aux TD_k représentés dans les vecteurs Q_{ij} .

Les arcs (s_t, s_h) de ce sous-graphe $\mathcal{G}_1$ sont pondérés par des indices $\Delta_{ht} = TD_h - TD_t > 0$ (distance temporelle) tels que TD_h correspond au sommet s_h et TD_t correspond au sommet s_t .

Pour l'exemple de la Figure 2.30, le sous-graphe $\mathcal{G}_1$ est donné par la Figure 2.31.

Étape 3 :

Cette troisième étape permet d'associer chaque chemin du sous-graphe $\mathcal{G}_1$ à un indice égal à la somme des indices de ses arcs.

FIGURE 2.31 – Sous-graphe pondéré $\mathcal{G}_1$

Les chemins du sous-graphe $\mathcal{G}_1$ de la Figures 2.31 sont associés aux indices suivants :

$$s_5 \longrightarrow s_2 \longrightarrow s_1 : 5; \quad s_8 \longrightarrow s_3 \longrightarrow s_1 : 7; \quad s_{10} \longrightarrow s_3 \longrightarrow s_1 : 5.$$

Étape 4 :

Après avoir associé chaque chemin du graphe $\mathcal{G}_1$ à un indice, les chemins de propagation du défaut candidats peuvent être classés par ordre de priorité en fonction de la durée de propagation du défaut. La priorité la plus importante est attribuée au chemin ayant le plus petit indice indépendamment du nombre des arcs constituant le chemin.

Dans le cas où le même plus petit indice est attribué à plusieurs chemins, nous privilégions celui couvert par l'arc entrant vers le sommet puits avec le plus petit indice et ainsi de suite si égalité.

Pour le sous-graphe $\mathcal{G}_1$, le chemin de propagation candidats à choisir sont classés par priorité :

Priorité 1 : $s_{10} \longrightarrow s_3 \longrightarrow s_1 : 5$.

Priorité 2 : $s_5 \longrightarrow s_2 \longrightarrow s_1 : 5$.

Priorité 3 : $s_8 \longrightarrow s_3 \longrightarrow s_1 : 7$.

2.5.4.2 Application au procédé de fabrication de papier

L'objectif de cette section est d'illustrer la méthode d'analyse des causes présentée dans les sections précédentes.

Nous considérons le procédé de fabrication de papier décrit en détail dans le paragraphe suivant. Ce procédé est instrumenté et différentes mesures relatives à chaque section de celui-ci sont disponibles. Du fait de la grande dimension du système, le graphe de causalité présenté ne concerne qu'une partie du procédé à savoir la section séchage. Cette section joue un rôle majeur dans la qualité du produit en terme d'hygrométrie.

Nous présentons tout d'abord les résultats concernant la génération du graphe orienté associé au modèle de causalité par l'intermédiaire d'une boîte à outils proposée dans le cadre du projet Papyrus. Le modèle graphique généré décrit les liens entre les indicateurs de performances du procédé.

Ce modèle permet de représenter une grande connexité entre les différents indicateurs de performance LPI (Indicateur Local de Performance) et PPI (Indicateur de Performance Procédé). En l'occurrence d'un défaut impactant la qualité du produit, une procédure d'analyse des cause(s) est appliquée et des chemins candidats de propagation de défauts sont identifiés. Le Schéma P&ID (Figure 2.32) et le graphe de causalité relatif à la section séchage (Figure 2.33) sont représentés ci-après.

Les défauts de grande amplitude ne peuvent pas être compensés par les correcteurs des boucles de régulation locales et peuvent donc se propager à travers les composants du système. L'intérêt du modèle graphique est qu'il ne représente que les variables affectant un indicateur PPI donné. Ainsi, il est possible d'analyser itérativement des prédécesseurs directs et d'appliquer des tests statistiques sur un ensemble réduit d'indicateurs de performance. Cette procédure d'utilisation conjointe du modèle graphique et de méthodes statistiques pour le diagnostic permet de répondre à la problématique liée à la taille du système.

Cette section est organisée comme suit : le procédé de fabrication de papier tout d'abord puis un modèle graphique du procédé réel sont fournis. Un scénario de défaut est considéré et la procédure d'analyse des causes est mise en œuvre. Les résultats de validation des partenaires industriels sont enfin énoncés afin d'apprécier les résultats que nous avons obtenus.

Schémas P&ID et Modèle de causalité du système

Le procédé de fabrication de papier est composé de plusieurs sections comme le montre la Figure 2.32.

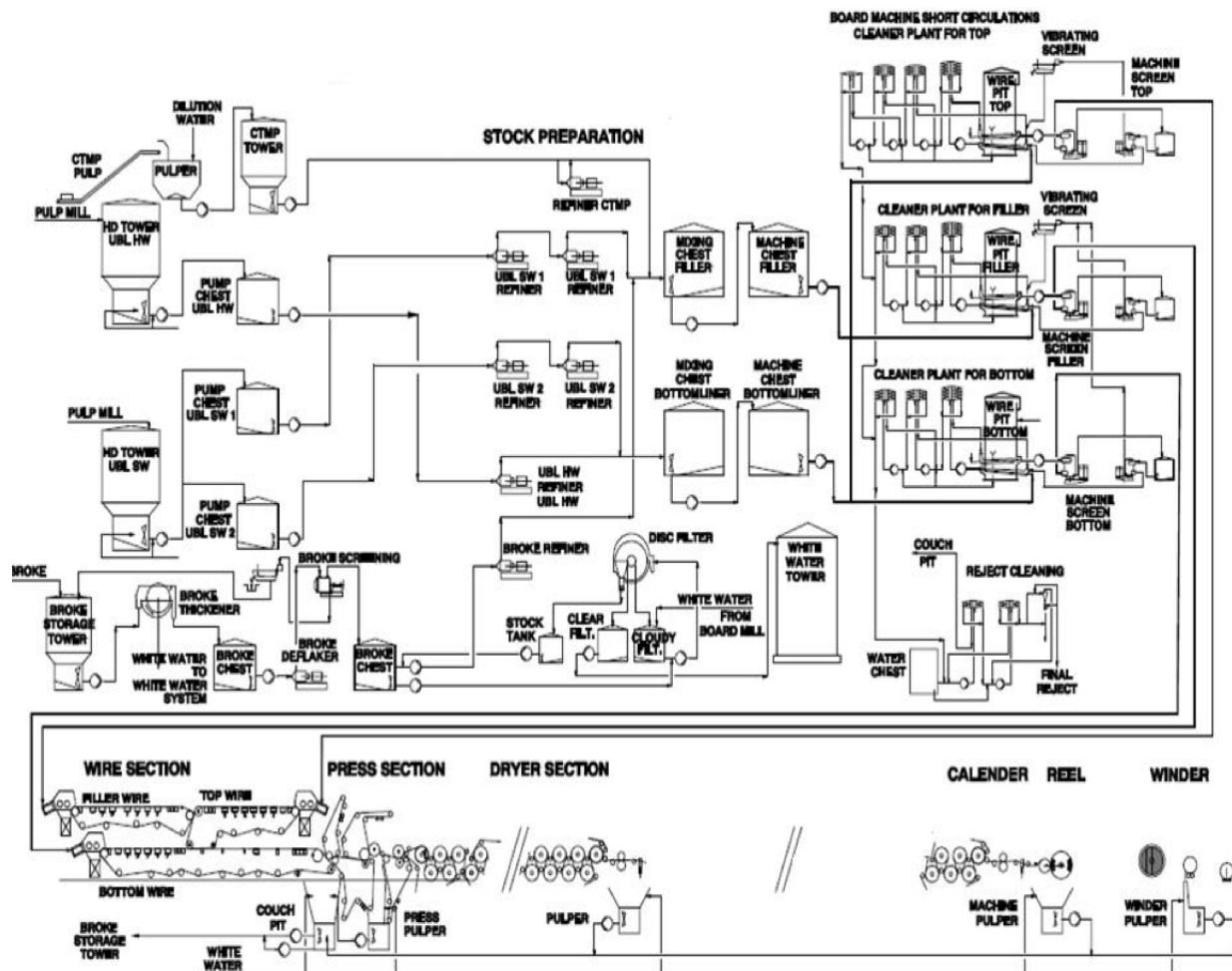


FIGURE 2.32 – Schéma illustrant les différentes sections du procédé de fabrication

Dans cette section, nous nous intéressons à la section de séchage. Pendant le séchage, le

papier est soumis à un processus de réduction de la teneur en eau : l'évaporation.

La section de séchage est composée d'un ensemble de cylindres chauffés à la vapeur (groupes de vapeur de 1 à 8), sur lesquels passe le papier. Comme présenté à la Figure 2.33, les cylindres sont disposés de façon à ce que les deux faces de la feuille de papier entrent l'une après l'autre en contact avec les cylindres. Le produit obtenu doit être conforme aux exigences prédéfinies et peut éventuellement subir d'autres traitements supplémentaires en fonction de l'utilisation future du papier. Ainsi, des mesures liées à la qualité du produit sont effectuées à la sortie de cette section de séchage (mesure du taux d'humidité KA4_XM1009_PPI).

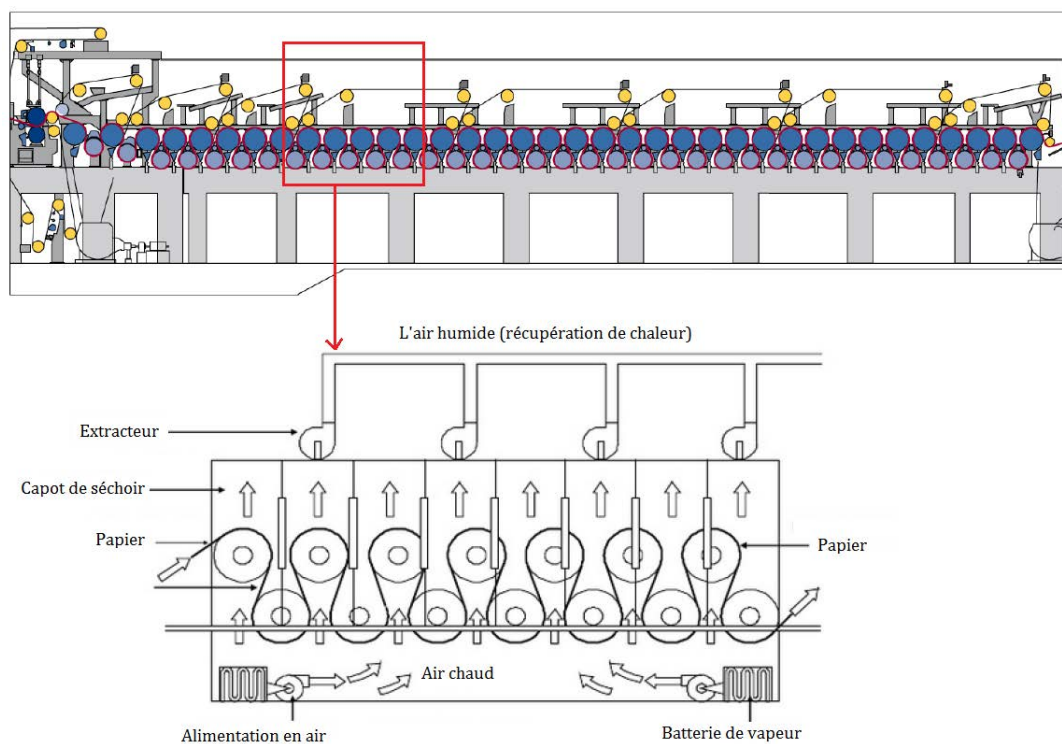


FIGURE 2.33 – Section séchage

Scénario de défaut étudié

Nous considérons un jeu de données d'une journée. Ce jeu est relatif aux mesures liées aux indicateurs de performance des différentes boucles de régulation de la section séchage

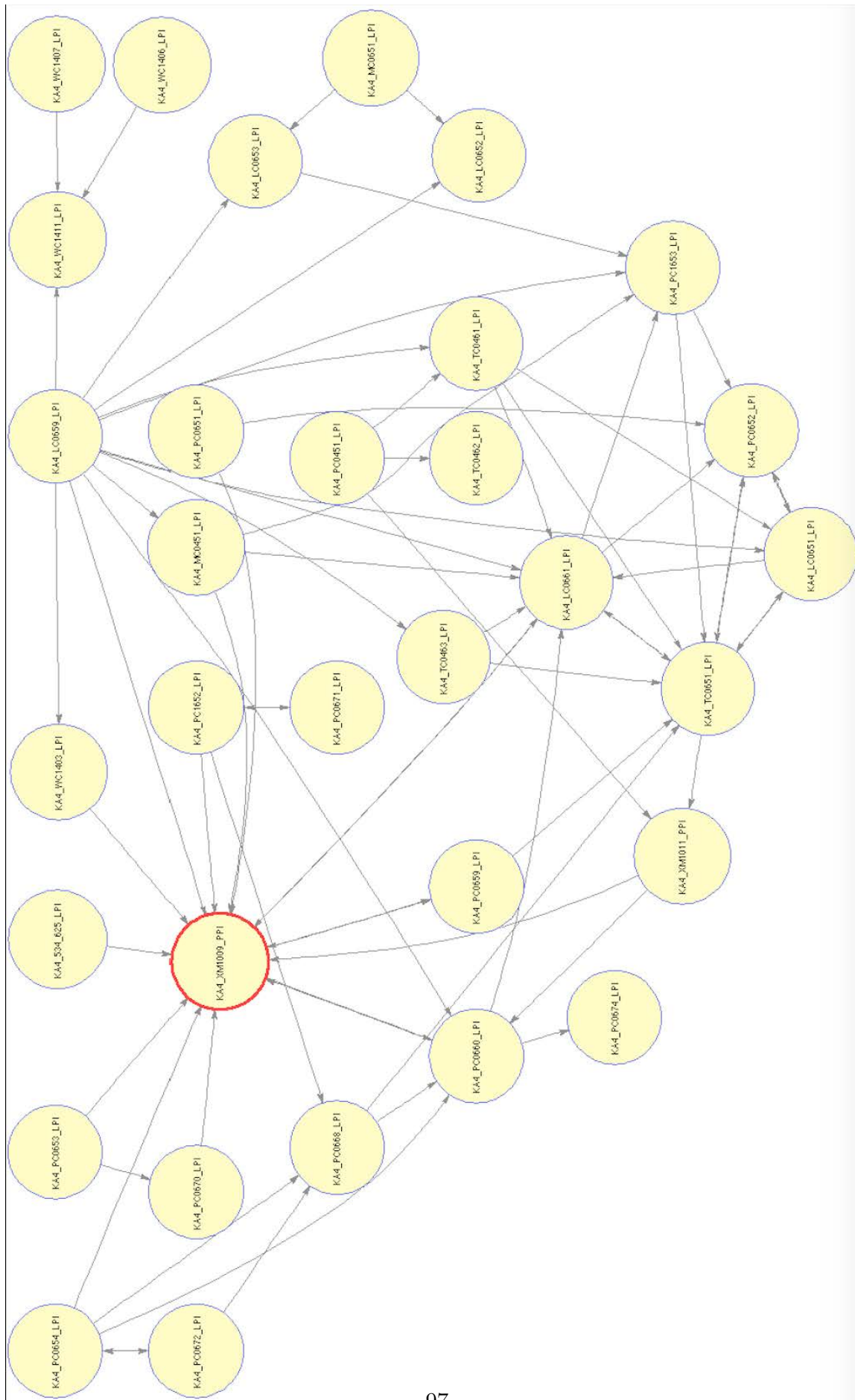


FIGURE 2.34 – Graphe de causalité de la section séchage : le sommet dont le périmètre est rouge est relatif au PPI de l'humidité associé au Tag **KA4_XM1009_PPI**

(42 boucles). Nous considérons la boucle de régulation de “haut niveau” (PPI) relative à l’humidité (en anglais *moisture*) présentée à la Figure 2.34. Les indicateurs de performance des 40 boucles de régulation restantes sont des indicateurs de performance locaux (LPI). La Figure 2.35 illustre le PPI de l’humidité pendant une journée.

Dans ce jeu de mesures, nous nous intéressons aux événements ayant un impact sur la qualité du produit, à savoir un changement des caractéristiques statistiques de l’indicateur PPI d’humidité du papier. Nous pouvons distinguer trois événements ayant dégradé la qualité du produit à différents instants. Cette dégradation se voit à travers des sauts de moyenne sur l’indicateur de performance du procédé relatif à l’humidité du produit.

Description des épisodes en défaut :

- Épisode 1 : Un défaut sur l’humidité (variabilité exogène) aux environs de 9h45.
- Épisode 2 : Un défaut sur l’humidité (variabilité exogène) aux environs de 12h30.
- Épisode 3 : Un défaut sur l’humidité (variabilité exogène) aux environs de 13h15.

Nous nous intéressons à l’identification de(s) cause(s) relatives au dernier événement à savoir le troisième épisode. La procédure d’analyse de(s) cause(s) est identique et, l’illustration pour les autres événements n’est pas nécessaire. Les tests sont conduits, dans un premier temps, à l’aveugle et les résultats sont croisés avec les données horodatées de l’industriel afin de les apprécier et de consolider les paramétrages des tests statistiques pour la mise en production.

La cause du défaut sera explicitée à la fin de la section 2.5.4.2.

La Figure 2.36 montre les différents épisodes :

- L’épisode (A) représente un comportement anormal dû à un changement de produit (épaisseur de papier à produire). Cet épisode représente un régime transitoire engendré

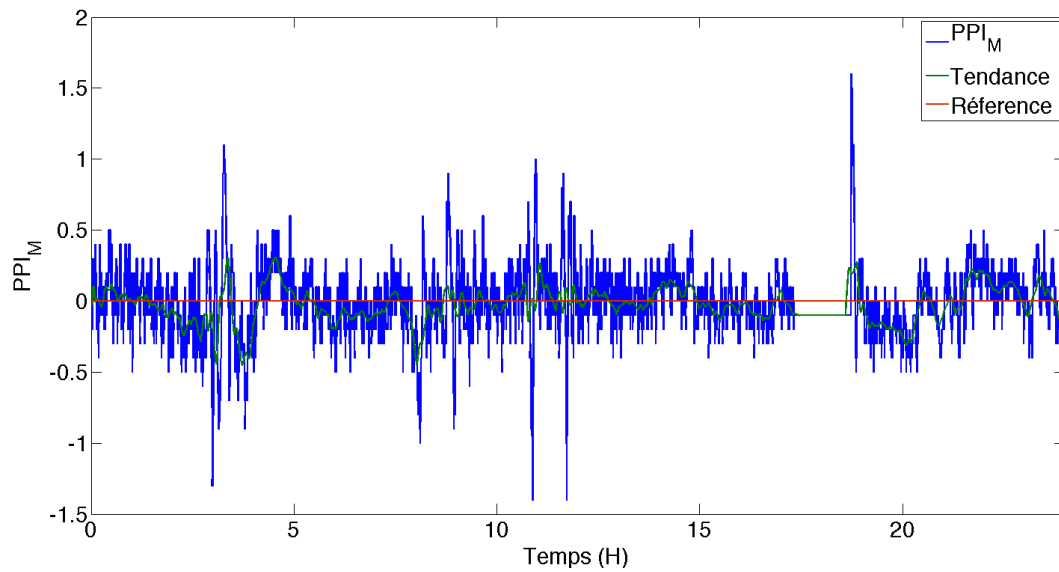


FIGURE 2.35 – PPI (KA4_XM1009_PPI) de l'humidité sur une journée

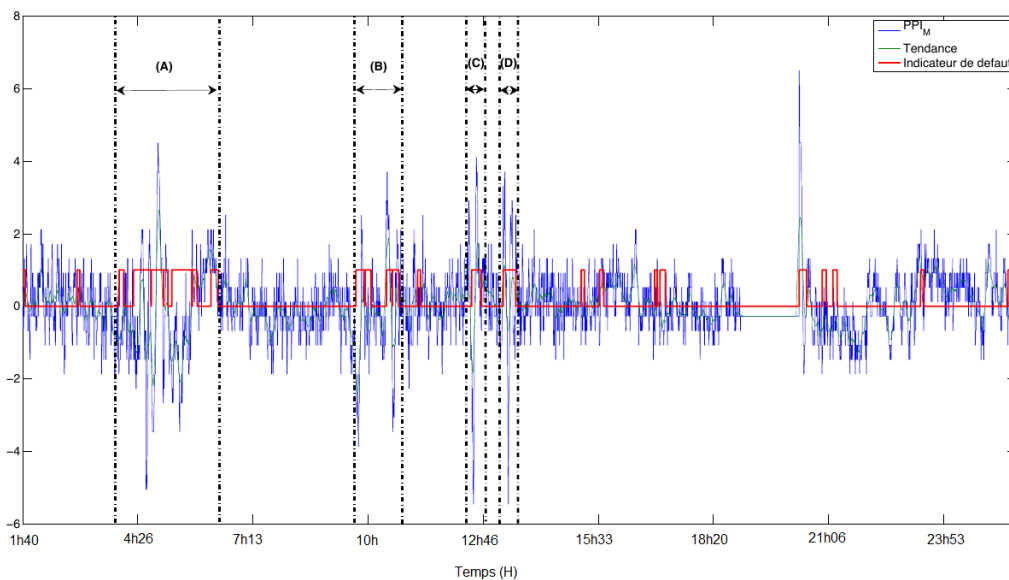


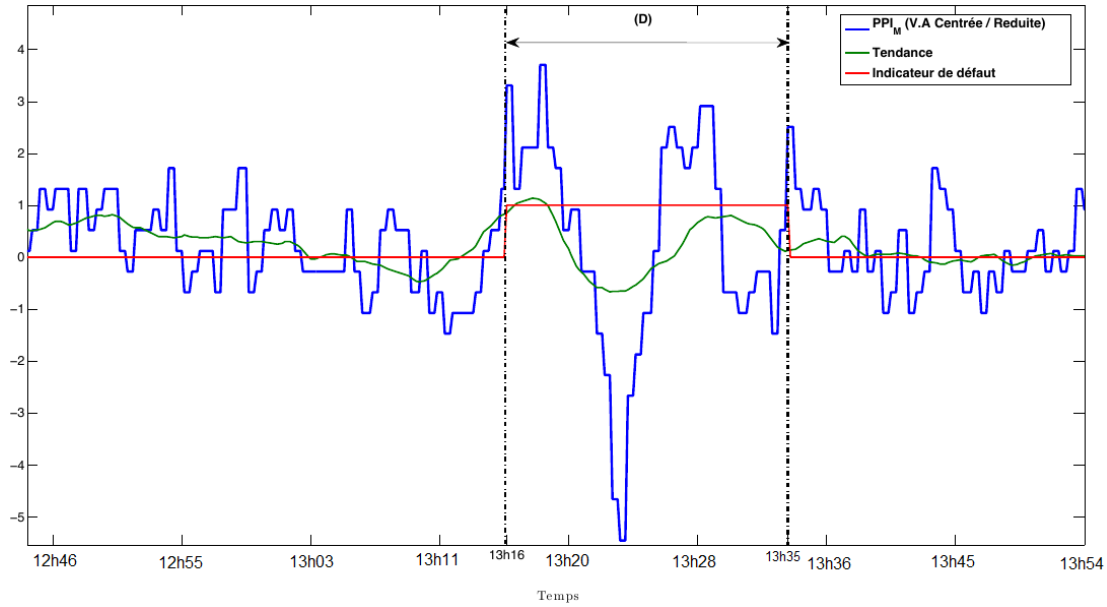
FIGURE 2.36 – Épisodes des dégradations des performances à différentes périodes de la journée. Un indicateur de défaut comme résultat synthétique du test statistique est associé pour la détection des sauts des moyennes aux signes inconnus

par des changements de consignes des différentes boucles de régulation et notamment des boucles de pression. L'épisode se termine approximativement à l'instant 5h52.

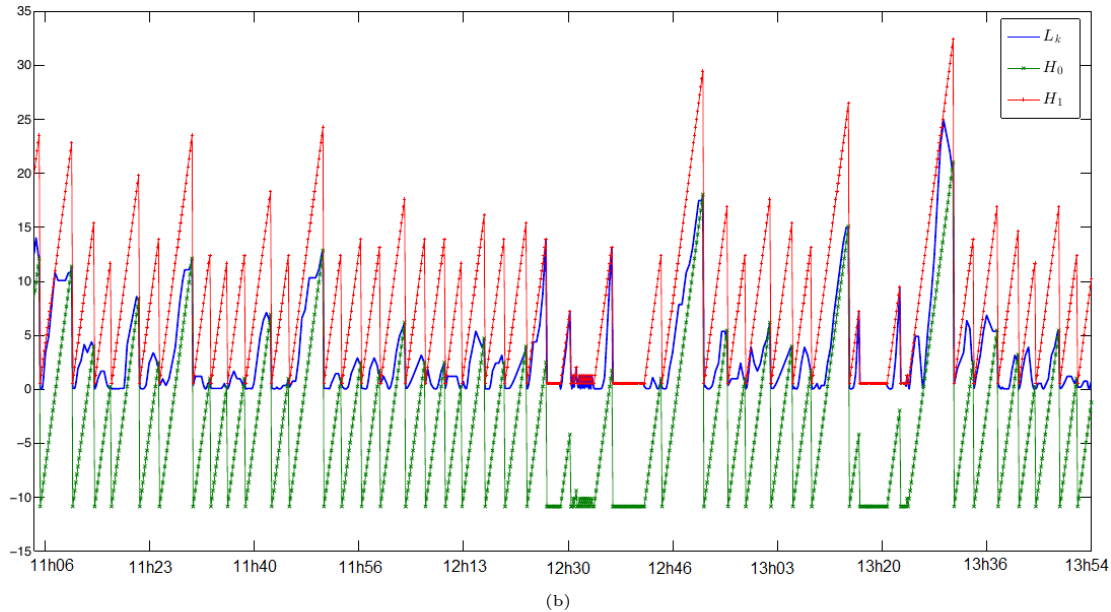
- Les épisodes (B), (C) et (D) représentent des épisodes de défauts. Ces défauts sont causés par l'incapacité de certains groupes de séchage à poursuivre leurs consignes et donc à déshumidifier convenablement le produit.

L'épisode (D) (Figure 2.37) étudié dure environ 20 minutes.

La Table 2.5 résume les principaux paramètres pour l'utilisation du test statistique relatif à l'indicateur de performance PPI_M (PPI correspondant à $KA4_XM1009_PPI$).



(a)



(b)

FIGURE 2.37 – Test statistique SPRT sur l'épisode (D) : (a) Indicateur de défaut, (b) SPRT :

$\Lambda(z_1, z_2, \dots, z_k) = L_k$ (fonction de vraisemblance) vs $\mathcal{H}_0 \vee \mathcal{H}_1$

A partir de l'instant de détection de la dégradation du PPI_M (c.à.d. 13 h 16), la procédure d'analyse de(s) défaut(s) permettant l'identification de(s) cause(s) candidate(s) est lancée.

La fonction de vraisemblance relative au test de saut de moyenne de signe inconnu	
$\Lambda(z_1, z_2, \dots, z_k) = L_k$ avec z_i une observation à l'instant i	$e^{-\frac{k\mu_1^2}{2\sigma_0^2}} \cdot \cosh\left(\frac{\mu_1}{\sigma_0^2} \sum_{i=1}^k (z_i - \mu_0)\right)$
$\mu_0(KA4_XM1009_PPI)$	0.15
$\mu_1(KA4_XM1009_PPI)$	$4^* \mu_0$
$\sigma_0(KA4_XM1009_PPI)$	0.4928
$\alpha(KA4_XM1009_PPI)$	90 %
$\beta(KA4_XM1009_PPI)$	1 %

TABLE 2.5 – Paramétrisation du test statistique KA4_XM1009_PPI

Le paragraphe suivant permet de décrire la procédure d'analyse de(s) cause(s) basée sur les méthodes et algorithmes introduits précédemment.

L'analyse conjointe permet de lister les chemins candidats de propagation de(s) défaut(s) et de les prioriser comme décrit à la section 2.5.4.

Analyse de(s) cause(s) (AdC)

L'analyse de(s) cause(s) se base sur l'information des causalités qui peuvent exister entre les variables explicatives du procédé c.à.d., dans notre contexte, les indicateurs de performances. A travers ces informations et le(s) temps de(s) détection(s) de(s) défaut(s), nous pouvons identifier le(s) indicateur(s) candidat(s) en ne s'intéressant exclusivement qu'aux prédécesseurs directs d'un indicateur de performance dégradé.

Ainsi, la procédure AdC s'achève lorsque les prédécesseurs d'un indicateur de performance ne se sont pas dégradés. L'AdC proposée dans ce manuscrit se base sur la hiérarchisation à plusieurs niveaux, celle-ci se base sur l'approche Top-Down introduite précédemment. L'objectif de cette approche est de générer un modèle utile pour le diagnostic et d'identifier la(s) cause(s) candidate(s).

Ainsi, les niveaux 1 et 2 (distance de 1 du PPI en défaut) permettent de représenter les relations $PPI_i \leftarrow LPI_i$ et les niveaux de 2 à n (distance de 2 à n du PPI en défaut)

représentent les relations sous-adjacentes entre les différents indicateurs de performance.

La Figure 2.38 représente le PPI relatif au taux d'humidité du papier (KA4_XM1009_PPI) et ses 13 prédécesseurs directs.

Etape 1 :

L'humidité est détectée en défaut à partir de l'instant 13h16. En se basant sur le graphe de la Figure 2.38 et en appliquant les tests statistiques sur les différents prédécesseurs de l'humidité, l'analyse des causes aboutit à l'identification de deux causes candidates. Il s'agit des pressions des groupes de vapeur 4 et 8 (KA4_PC0654_LPI et KA4_PC0660_LPI) qui correspondent aux sommets 4 et 6 du graphe de la Figure (2.38). Le résultat des tests statistiques est présenté à la Figure 2.40. Ceci peut être représenté sur le graphe de causalité comme le montre la Figure 2.39.

La Table 2.6 regroupe les différents défauts détectés au niveau de l'humidité ainsi que les pressions des groupes de vapeur 4 et 8 et leurs instants de détection. A partir de la Figure 2.39 et la Table 2.6, il est possible de déterminer les vecteurs Q_{ij} relatifs aux niveaux 1 et 2 du graphe de causalité dans l'équation 2.5.

- Niveau 1 :

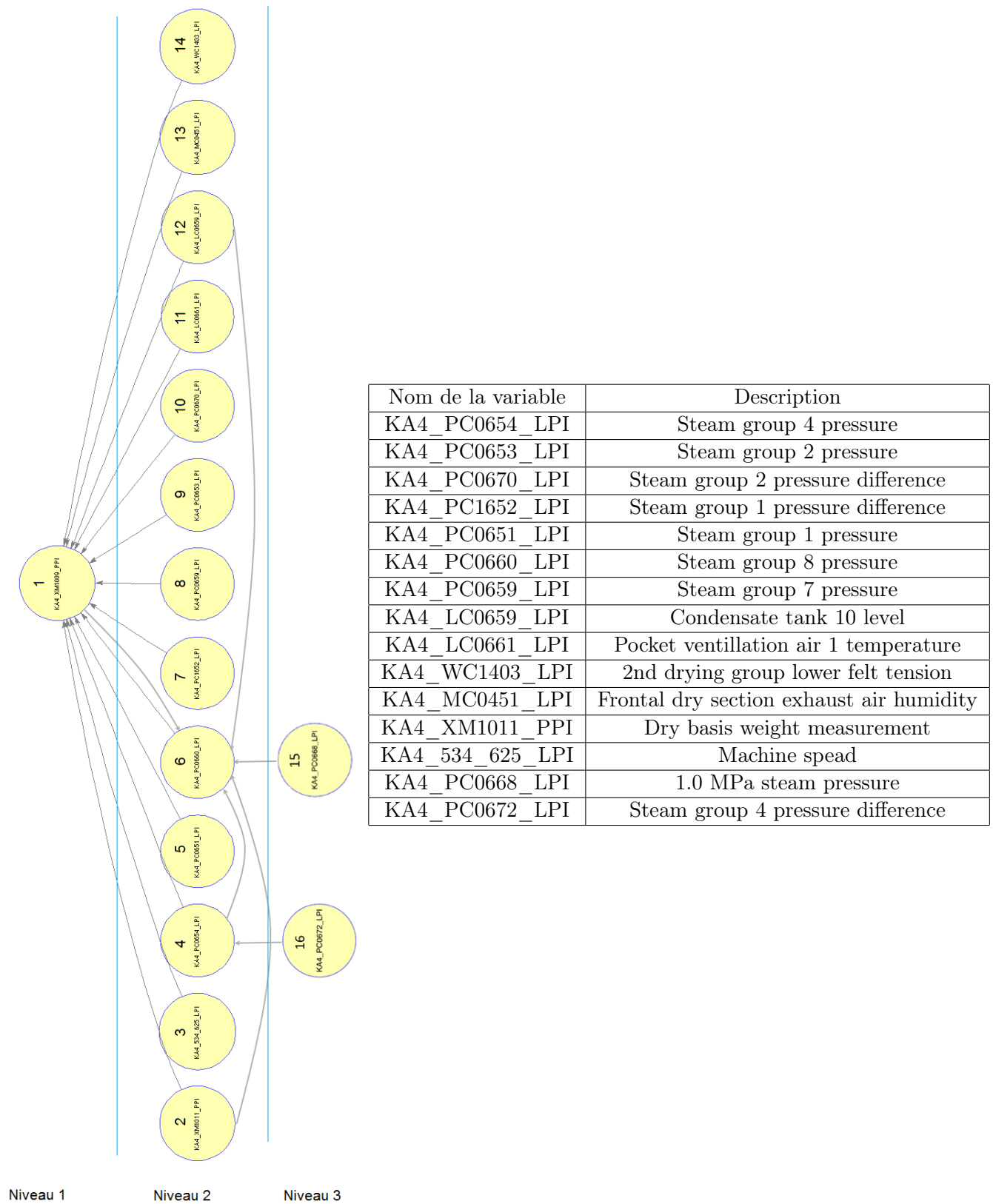
$$Q_{10} = [TD_1] \text{ avec } TD_1 \text{ relatif au sommet 1}$$

- Niveau 2 :

$$Q_{21} = [TD_4, TD_6] \text{ avec } TD_4 \text{ relatif au sommet 4 et } TD_6 \text{ relatif au sommet 6}$$

(2.53)

Le premier test sur le sommet 1 a identifié les deux boucles correspondantes aux sommets n° 4 et 6 en défaut. Nous nous intéressons donc aux prédécesseurs de ces deux sommets. Le Niveau 2 concerne l'utilisation de l'AdC sur la cause candidate prioritaire (la pression du groupe de vapeur 8 correspondante au sommet n° 6). Ce sommet devient le sommet cible.

FIGURE 2.38 – Modèle de causalité : $PPI_i \leftarrow LPI_j$ (Prédécesseurs directs du PPI_i)

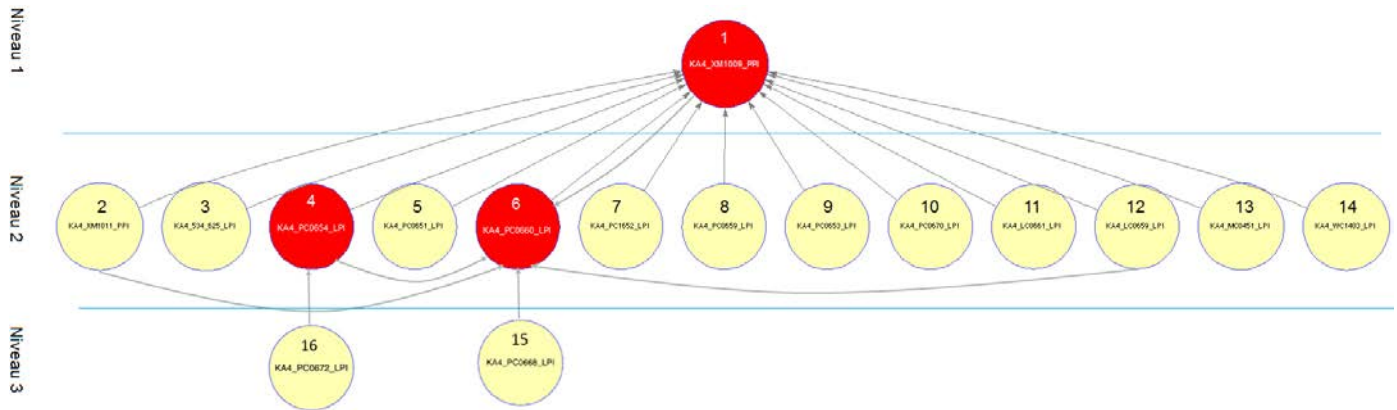
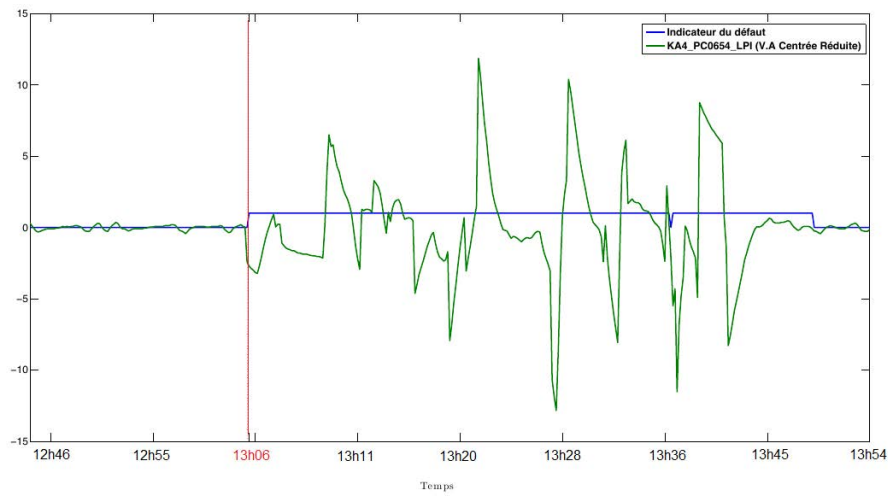
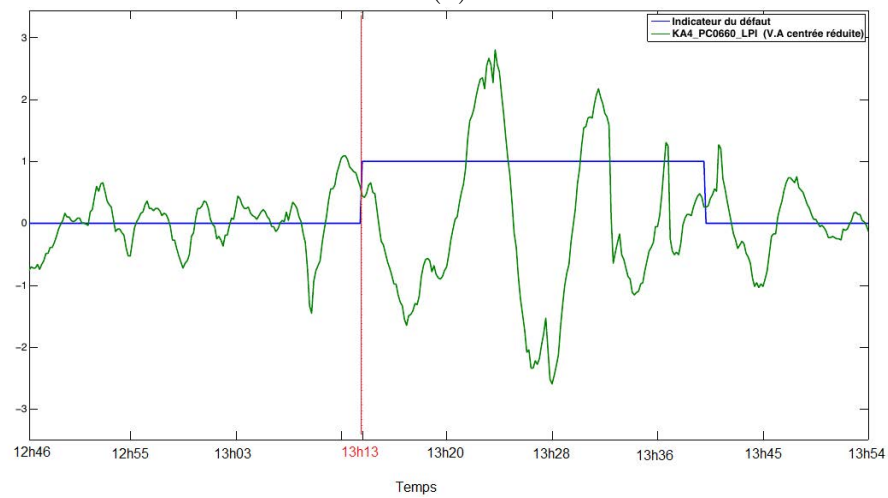


FIGURE 2.39 – Graphe de causalité représentant les défauts aux niveaux 1 et 2



(a)



(b)

FIGURE 2.40 – Dégradation de performance de la pression du groupe de vapeur 4 détectée à partir de 13h06 (a) et du groupe 8 détectée à partir de 13h13 (b)

	TD_k
KA4_XM1009_PPI (sommet n° 1) (Taux d'humidité)	13h16
KA4_PC0660_LPI (sommet n°6) (Pression du groupe de vapeur 8)	13h13
KA4_PC0654_LPI (sommet n°4) (Pression du groupe de vapeur 4)	13h06

TABLE 2.6 – Instants de détection des défauts pour l'humidité et les pressions des groupes de vapeur 4 et 8

En effet, ses prédécesseurs et leurs états sont représentés dans la Figure 2.41.

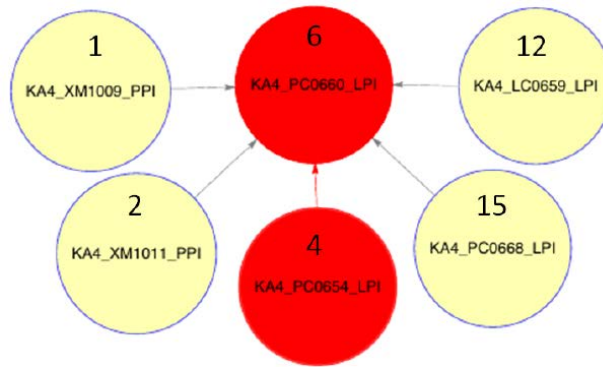


FIGURE 2.41 – Niveau 2 : Le prédécesseur du $KA4_PC0660_LPI$ qui correspond au sommet n° 6 est en défaut (sommet n° 4 en rouge)

De même, la procédure AdC est appliquée sur le sommet $KA4_PC0654_LPI$ (sommet n° 4) ($T_d=13h06$) en défaut. La Figure 2.42 représente le sommet $KA4_PC0654_LPI$ (sommet n° 4) en défaut et son prédécesseur en mode nominal.

Etape 2 :

Puisque le test effectué sur le niveau 2 n'identifie pas de nouveaux sommets en défaut, le graphe de la Figure 2.39 peut être simplifié en fonction des vecteurs Q_{ij} donnés dans l'Equation 2.5.4.2. Ce graphe simplifié est représenté dans la Figure 2.43.

Les pondérations des arcs de ce graphe correspondent à une distance temporelle de détection des défauts (voir section 2.5.4). Cette distance est obtenue à partir de la Table 2.6.

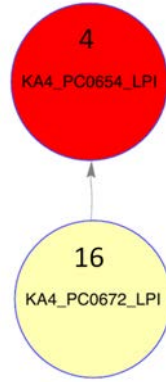


FIGURE 2.42 – Niveau 3 : Prédécesseur du *KA4_PC0654_LPI* (sommet n° 4) est dans l'état nominal (sommet n° 16 en jaune)

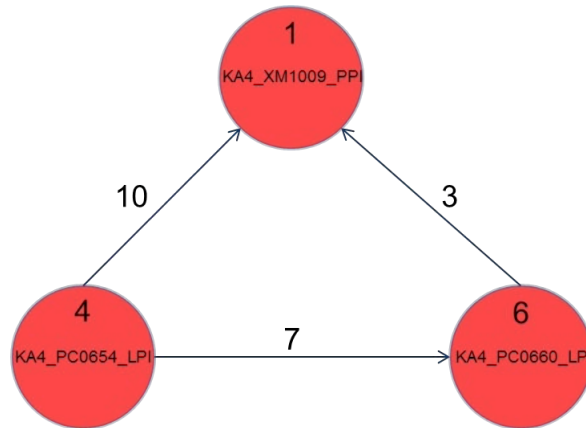


FIGURE 2.43 – Graphe de causalité simplifié

Ainsi, l'arc reliant la pression du groupe de vapeur 4 à l'humidité est pondéré par 10 (10 min) (13h16-13h06), l'arc reliant la pression du groupe de vapeur 8 à l'humidité est pondéré par 3 (3 min) (13h16-13h13) et l'arc reliant la pression du groupe de vapeur 4 à la pression du groupe de vapeur 8 est pondéré par 7 (7 min) (13h13-13h06). Ces pondérations permettent de prioriser une cause candidate par rapport à une autre.

Etape 3 :

Dans ce graphe, nous avons 3 chemins possibles associés aux pondérations suivantes :

$$s_6 \longrightarrow s_1 : 3; s_4 \longrightarrow s_1 : 10; s_4 \longrightarrow s_6 \longrightarrow s_1 : 10.$$

Etape 4 :

A partir des chemins donnés à l'étape 3, nous avons deux chemins avec la même pondération 10, le chemin indirect $s_4 \longrightarrow s_1 : 10$ et chemin direct $s_4 \longrightarrow s_6 \longrightarrow s_1$.

Dans ce cas, le chemin à privilégier est $s_4 \longrightarrow s_6 \longrightarrow s_1$ puisque ce dernier couvre l'arc $s_6 \longrightarrow s_1$ ayant la plus petite pondération.

Ainsi, la cause prioritaire du défaut au niveau de l'humidité est celle correspondante à la pression du groupe de vapeur 4. Le résultat obtenu dans cette analyse des causes est validé par le partenaire industriel. Cette validation est consignée dans un lot de travail WP10 (voir Chapitre 1).

2.6 Conclusion

Dans ce chapitre, nous nous sommes intéressés à la modélisation et au diagnostic des systèmes. Ce chapitre est principalement composé de trois parties essentielles qui correspondent aux méthodes de modélisation et au diagnostic.

La première partie est consacrée à la méthode basée sur le modèle. Pour cette méthode, nous considérons que les équations d'état du procédé sont disponibles. A partir de ces équations, il est possible de représenter la structure du système sous forme d'un graphe. Ce dernier représente les variables du système ainsi que les relations les liant. La représentation graphique de la structure du système est intuitive et, permet une analyse simple des propriétés du système. Dans ce chapitre, nous nous sommes intéressés à la propriété de diagnosticabilité. Cette propriété ainsi vérifiée permet le diagnostic des défauts d'un système et la synthèse d'observateurs. Les conditions graphiques de la diagnosticabilité sont présentées et illustrées par un exemple.

La deuxième partie concerne les méthodes basées sur les données à savoir l'inter-

corrélation, causalité de Granger et Transfert d'entropie. Ces trois méthodes permettent de déterminer les liens de causalité entre les différentes variables du système, leur sens et éventuellement leur force. En déterminant ces liens de causalité, il est possible d'obtenir un graphe de causalité du système qui permettra l'étude de ce dernier.

Une méthode combinée basée sur les données et le modèle est aussi présentée permettant ainsi l'obtention d'un modèle plus précis et plus proche du système réel.

La dernière partie de ce chapitre est consacrée au diagnostic. Nous avons introduit l'approche Top-Down et les dégradations statistiques des performances du système. Le diagnostic est ensuite étudié à travers des tests statistiques de détection de dégradation (test de Wald). Enfin, une analyse des causes est présentée et une application sur un procédé de fabrication de papier illustre la procédure d'analyse de(s) cause (s) introduite en amont.

Chapitre 3

Système tolérant aux défauts : Approche conjointe par gestion de références et analyse graphique

3.1 Introduction

De nos jours, les procédés industriels sont soumis à des contraintes de rendement et de sécurité croissantes, nécessitant de ce fait l'application de commandes toujours plus performantes. Ainsi les procédés industriels modernes sont souvent de grande dimension et ils deviennent de plus en plus vulnérables aux défauts qui les rendent incapables de respecter les contraintes de sécurité et de production. Cette évolution a conduit au développement des systèmes de surveillance et de tolérance aux défauts adaptés dont l'objectif est de fournir à tout moment des informations à l'opérateur sur le fonctionnement du système puis d'ajuster les paramètres de la commande ou les consignes afin de corriger les comportements anormaux observés. Plus précisément, la tolérance aux défauts permet de réduire, voire d'annuler l'effet des défauts ayant un impact inacceptable sur la mission, la sécurité (de l'être humain et du matériel), l'environnement et la rentabilité du système.

La problématique de la commande tolérante aux défauts est abordée dans plusieurs travaux de la littérature tels que Zhang and Jiang (2003); Theilliol et al. (1998); Staroswiecki and Hoblos (2004). L'objectif principal du diagnostic des défauts est de détecter, à partir des observations et de la connaissance du système, l'occurrence des défauts et de les localiser (Zwingelstein (1995); Pessel et al. (2007); Ligeza and Koscielny (2008)).

En présence d'un défaut, les performances du système se dégradent et, dans certains cas, la stabilité du système peut être affectée. L'accommodation aux défauts consiste à modifier la commande du système pour réagir contre l'effet d'un défaut en modifiant la loi de commande de façon à ce que la stabilité et les performances nominales du système soient maintenues. Pour ceci, on a besoin de mesurer la capacité du système à tolérer les défauts. De façon générale, l'analyse de la reconfigurabilité, dans le contexte de cette thèse, consiste à définir dans quelle mesure les performances du système, et en particulier ses boucles de régulation, peuvent assurer leurs fonctions requises en présence d'un défaut.

L'étude de la reconfigurabilité définit les possibilités et les limites de la stratégie de commande à utiliser quand le système est soumis à un défaut. Cette étude peut être effectuée dès la phase de conception du système. Dans les travaux de Wu et al. (2000); Theilliol et al. (2008), cette étude est réalisée lors de la phase de conception du système de commande. Meguetta et al. (2013) proposent de prendre en compte les contraintes de coût au niveau de la conception du système de commande tolérant aux défauts.

Dans ce chapitre, nous nous intéressons aux conditions pour que le système, après l'occurrence de défauts, soit capable de maintenir un niveau de performances exigé pendant le temps de mission du système.

La méthode proposée est basée sur l'analyse d'un modèle graphique préalablement obtenu représentant le système. Le modèle doit intégrer l'ensemble des boucles de régulations ainsi que leurs interactions. Nous allons introduire, dans un premier temps, une représentation topologique du système basée sur des outils d'algèbre usuels. Dans un second

temps, nous allons introduire un graphe de causalité orienté. Ce type de modèle graphique est analysé afin de statuer sur la reconfigurabilité dans le contexte proposé dans cette thèse.

3.2 Représentation topologique du système

3.2.1 Matrice d'adjacence du système

Le modèle de causalité d'un système peut être décrit par une matrice d'adjacence notée $\mathcal{A}$. Cette matrice est constituée de 1 et de 0 indiquant respectivement l'existence ou l'absence d'un lien entre les variables.

Une matrice dite "d'accessibilité", notée $\mathcal{R}$, peut être calculée à partir de la matrice d'adjacence $\mathcal{A}$. On dit que y est accessible par x si et seulement s'il existe un chemin de x vers y dans le graphe de causalité et qu' y est un descendant de x . Ainsi, tout chemin dans le graphe de causalité est représenté par un élément '1' dans la matrice d'accessibilité $\mathcal{R}$. Ci-après un exemple d'une matrice d'adjacence $\mathcal{A}$ où les paramètres A_{ij} correspondent à des éléments '1'.

$$\mathcal{A} = \begin{pmatrix} A_{11} & A_{12} & \cdots & A_{1n} \\ 0 & A_{22} & \cdots & \vdots \\ \vdots & \vdots & \ddots & \vdots \\ 0 & \cdots & \cdots & A_{nn} \end{pmatrix}$$

Considérons une matrice d'adjacence $\mathcal{A}$, avec des permutations simultanées de lignes et de colonnes de cette matrice. Chaque bloc de la diagonale représente une partie du système dans le sens d'ordre partiel. La partie avec le plus grand nombre ne peut pas atteindre celle avec un plus petit nombre. Ainsi, nous obtenons une nouvelle matrice d'adjacence notée $\mathcal{TAT}^T$.

La décomposition d'un système peut donc être effectuée à l'aide d'un outil matriciel en considérant la matrice d'adjacence $\mathcal{A}$ ou la matrice d'accessibilité $\mathcal{R}$.

L'inconvénient de ce type de représentation réside dans le fait qu'il n'est pas intuitif comparé à une représentation par un graphe orienté.

Nous considérons que le système global étudié est composé de plusieurs boucles de régulation présentant des structures différentes. Ainsi, selon le concept de la commandabilité, les variables de chaque boucle sont fortement connectées. Cet aspect est présent dans la matrice d'accessibilité $\mathcal{R}$ comme un bloc où tous les éléments sont égaux à '1'.

Exemple 3 *Considérons le système automatisé présenté par la Figure 3.1.*

La matrice d'adjacence correspondante est donnée par :

$$\mathcal{A} = \begin{pmatrix} 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 1 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \end{pmatrix}, \quad \mathcal{TAT}^T = \begin{pmatrix} 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \\ \hline 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & 1 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \end{pmatrix}$$

Dans cette matrice d'adjacence $\mathcal{A}$, les variables sont ordonnées comme suit : $L, F_1, F_2, e_{F_1}, e_{F_2}, e_L, sp_L, sp_{F_1}$ et sp_{F_2} . Les variables sp , e et X_i correspondent respectivement aux consignes, aux erreurs et aux sorties des boucles de régulation.

Après une permutation basique des lignes et des colonnes de cette matrice, les variables de la matrice d'adjacence $\mathcal{TAT}^T$ sont données comme suit : $sp_{F_2}, e_{F_2}, F_2, sp_L, e_L, sp_{F_1}, e_{F_1}, F_1$ et L .

La matrice $\mathcal{TAT}^T$ est principalement divisée en deux blocs : le premier bloc correspond au débit F_2 commandé par le débit $FC01$ et le deuxième bloc correspond à la boucle de régulation en cascade pour la commande du niveau L .

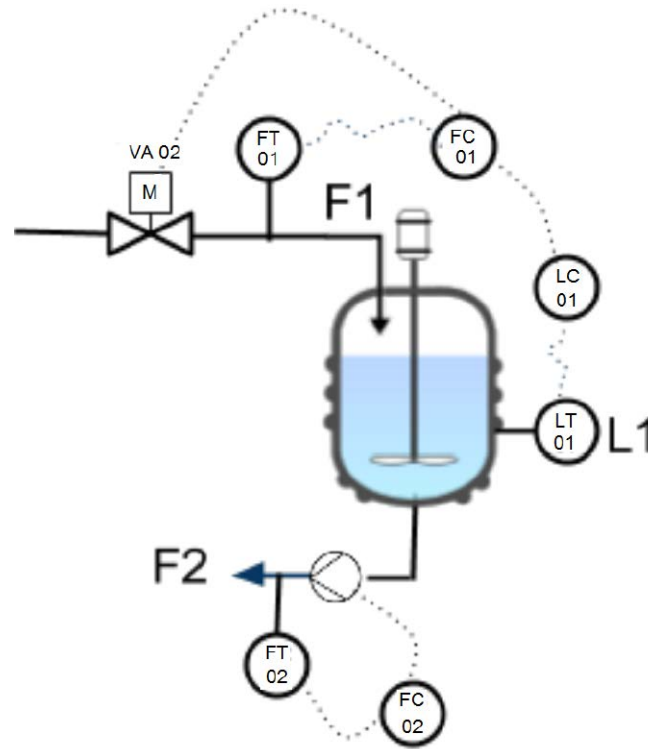


FIGURE 3.1 – Système automatisé

A partir de la matrice d'adjacence $\mathcal{TAT}^T$, la variable L est affectée par la boucle F_2 de commande (premier bloc) à travers le débit F_2 et également affectée par le débit F_1 de la boucle de commande du niveau (deuxième bloc).

Dans le but de représenter intuitivement la topologie du système, le graphe de causalité est associé à la matrice d'adjacence. Ce graphe de causalité est présenté dans la section 3.2.2. L'analyse de la solvabilité de la configuration dite "structurelle" est basée sur l'étude du graphe de causalité du système.

3.2.2 Graphe de causalité

Dans le graphe de causalité orienté, les variables mesurées du système sont représentées par des sommets et les relations cause/effet entre deux variables sont représentées par un arc orienté entre les sommets correspondants. Le sommet de la source correspond à la cause (ou le sommet parent), et le sommet de destination correspond à l'effet (ou le sommet enfant).

Le graphe peut être obtenu en utilisant des méthodes basées sur les données et / ou des méthodes basées sur le modèle. Toutefois, pour un système physique, le graphe orienté est établi en utilisant également la connaissance qualitative du système et le retour d'expérience. Les équations physiques associées au système ne sont généralement pas disponibles en raison de la complexité des phénomènes physiques décrivant le comportement et la structure du système étudié. Néanmoins, un ensemble d'informations importantes prélevées à partir des schémas P&ID est souvent utilisé en industrie afin de consolider les modèles obtenus.

Définition du graphe de causalité

Dans le même esprit que la définition du graphe orienté donnée en chapitre 2, notre objectif dans ce paragraphe est de présenter la modélisation de la topologie d'un système en tenant compte des différents types de ses variables mesurées. Pour une telle structure, on peut y associer de façon naturelle un graphe orienté noté $\mathcal{G}(\mathcal{V}, \mathcal{A})$ constitué d'un ensemble non-vide $\mathcal{V}$ de sommets et d'un ensemble fini $\mathcal{A}$ de paires ordonnées de sommets appelés arcs. La notation $\mathcal{G}(\mathcal{V}, \mathcal{A})$ signifie que $\mathcal{V}$ et $\mathcal{A}$ sont respectivement des ensembles de sommets et d'arcs dans le graphe $\mathcal{G}$.

L'ensemble de sommets $\mathcal{V}$ est défini par $\mathcal{V} = \mathcal{S} \cup \mathcal{E} \cup \mathcal{O}$ où $\mathcal{S}$, $\mathcal{E}$ et $\mathcal{O}$ correspondent respectivement à l'ensemble des consignes, l'ensemble des erreurs et l'ensemble des sorties des boucles de régulation.

L'ensemble $\mathcal{A} = \{(v_i, v_j) \mid v_i \in \mathcal{V} \text{ et } v_j \in \mathcal{V}\}$ correspond à l'ensemble des arcs qui représentent la relation de causalité entre les éléments v_i et v_j de $\mathcal{V}$.

Définitions et notations

Quelques définitions et notations utilisées dans la suite de ce chapitre sont rappelées ci-après.

- Un chemin p couvrant les sommets $v_{r_0}, \dots, v_{r_i}$ est noté $p = v_{r_0} \rightarrow v_{r_1}, \dots, \rightarrow v_{r_i}$ où

pour $j = 0, 1, \dots, i - 1$, nous avons $(v_{r_j}, v_{r_{j+1}}) \in \mathcal{A}$. Un chemin est dit “simple” s’il ne passe pas plus d’une fois par un même sommet.

- Un cycle est un chemin de la forme $v_{s_0} \rightarrow v_{s_1} \rightarrow \dots \rightarrow v_{s_i} \rightarrow v_{s_0}$, où les sommets $v_{s_0}, v_{s_1}, \dots, v_{s_i}$ sont distincts.
- $Succ(\{v_i\})$ est l’ensemble des successeurs du sommet v_i , *i.e.* l’ensemble des sommets v_j tels qu’il existe un chemin de v_i vers v_j . Un successeur d’un sommet v_i est dit “direct” s’il existe un arc (v_i, v_j) . $Succ_d(v_i)$ est l’ensemble des successeurs directs de v_i .
- $Pred(\{v_i\})$ est l’ensemble de prédécesseurs de v_i , *i.e.* l’ensemble de sommets v_j tel qu’il existe un chemin de v_j vers v_i . $Pred_d(v_i)$ est l’ensemble de prédécesseurs directs of v_i .
- $dist(v_i, v_j)$ représente la distance entre les deux sommets v_i et v_j du graphe. Cette distance est définie comme la longueur du plus court chemin entre ces deux sommets.
- Soit S un ensemble. La cardinalité de l’ensemble S est notée $card(S)$ et correspond au nombre d’éléments dans cet ensemble.
- Considérons deux sous-ensembles $S1$ et $S2$ de l’ensemble S . $S1 \setminus S2$ représente les éléments dans l’ensemble $S1$ qui ne sont pas dans l’ensemble $S2$.

3.3 Accommodation aux défauts par apprentissage itératif

Nous supposons que le système à étudier n’est défini que par son graphe de causalité associé à sa structure. Nous supposons également que seuls les défauts d’actionneurs sont considérés dans cette étude. L’approche intégrée de diagnostic et de tolérance au défaut se base sur les trois étapes macroscopiques suivantes :

- Dans un premier temps, nous effectuons la détection de défaut à travers la surveillance des KPI (cf. chapitre 2) ;

- Après la détection des défauts, l'analyse des causes est effectuée et le chemin de propagation du défaut est identifié (cf. chapitre 2) ;
- L'origine du défaut ainsi identifiée, l'accommodation itérative au défaut est effectuée en-ligne en réajustant des consignes exogènes sélectionnées à travers l'analyse du graphe de causalité pour la reconfiguration.

La Figure 3.2 représente le graphe de causalité correspondant à la matrice d'adjacence $\mathcal{A}$ donnée dans l'exemple 3. Ce graphe contient les variables de consignes sp_L , sp_{F_1} et sp_{F_2} , les variables d'erreurs e_L , e_{F_1} et e_{F_2} , et les variables de sorties des boucles des régulations L , F_1 et F_2 .

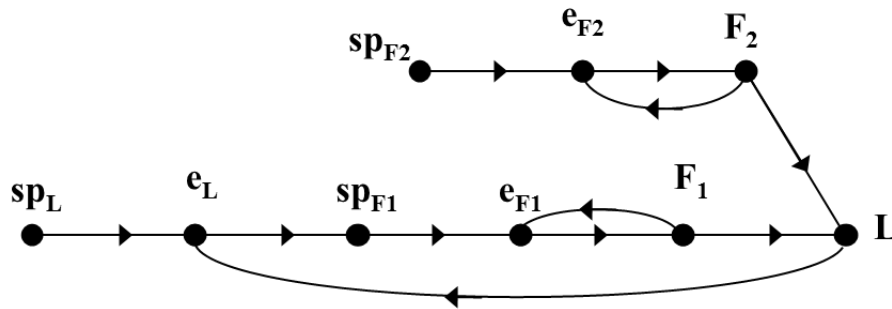


FIGURE 3.2 – Graphe de causalité associé à la matrice d'incidence $\mathcal{A}$ de l'exemple 3

Le problème de l'accommodation aux défauts, en considérant les défaillances des actionneurs, est traité dans plusieurs travaux de la littérature Tao et al. (2001, 2002); Tang et al. (2005); Tao et al. (2013) pour les systèmes définis par leur représentation d'état ou fonction de transfert. Cependant, pour des raisons de complexité de modélisation, ces représentations sont difficiles à obtenir dans le cas des systèmes complexes à grande dimension. Par conséquent, aucun modèle précis pour ce type de systèmes n'est généralement disponible.

Dans ce contexte, un modèle qualitatif du système est considéré. Ce modèle représente les relations de causalité entre les boucles de régulation constituant le système. Ce modèle est capable de représenter, de manière intuitive, les relations entre les indicateurs de performance de plusieurs sous-systèmes et des consignes des boucles de régulation. Cette représentation permet d'exhiber les consignes exogènes candidates au réajustement par le

gouverneur des références.

Pour les défauts d'actionneurs mineurs, les boucles fermées peuvent intrinsèquement compenser les effets de ces défauts. Cependant, les défauts conduisant à la saturation des actionneurs doivent être traités avec précaution. Ainsi, le gouverneur de références par apprentissage itératif est, potentiellement, une solution permettant de faire face à ce problème de saturation des actionneurs en ajustant les consignes exogènes sélectionnées de certaines boucles de régulation. Pour des raisons de simplicité, nous considérons seulement le cas où le système fonctionne en régime permanent. Sans perte de généralité, lorsque cette condition est vérifiée, nous supposons que $y \simeq y^*$ où y et y^* représentent la trajectoire de sortie et la trajectoire de référence que le système doit suivre. Dans ce qui suit, le signal de sortie exprime explicitement la sortie d'une boucle de régulation et la trajectoire de référence reflète l'objectif que le système doit atteindre. Ainsi, en occurrence d'un défaut majeur qui mène à la saturation d'un actionneur, le système ne parviendra pas à atteindre ses performances en termes de suivi de trajectoires (Figure 3.3).

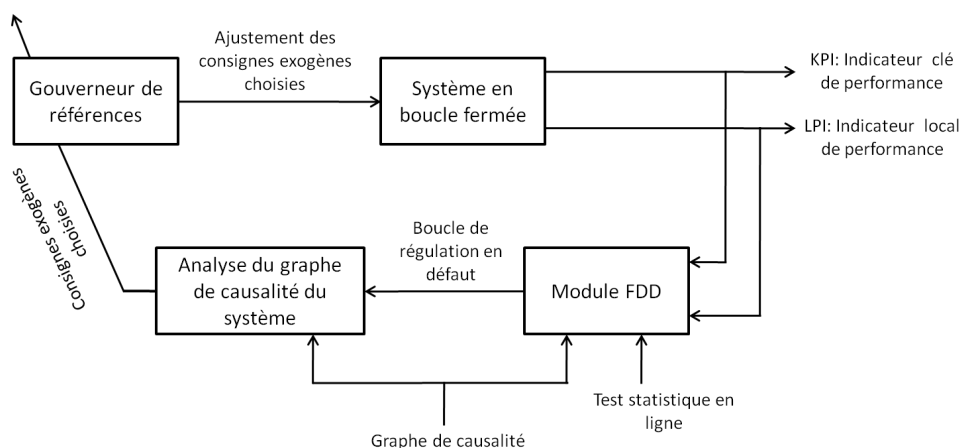


FIGURE 3.3 – Schéma de commande tolérante aux défauts adaptative

L'idée principale est de déterminer un ensemble de trajectoires de référence des consignes exogènes sélectionnées qui permettent au système, dans des conditions défectueuses, de fonctionner suivant des performances dégradées.

3.4 Synthèse d'un gouverneur de références

Dans cette section, nous nous intéressons à la formulation d'un indice de performance pour l'évaluation des performances du système sur différents niveaux.

3.4.1 Définition de l'indice de performance

L'indice de performance est assimilé à un *LPI* si l'évaluation de la performance concerne un sous-système de bas-niveau (boucle de régulation de niveau, de débit, etc).

À contrario, l'indice de performance est assimilé à un *PPI* si l'indice d'évaluation concerne des performances liées directement à la qualité du produit, à la consommation énergétique, à un taux de production, etc. Cette décomposition dépend fortement de l'analyse fonctionnelle du système (voir WP04, Chapitre 1). Le résultat de l'analyse associe à chaque niveau du système des indicateurs pour évaluer les performances économiques du système KPI (Business), procédé/moyen de production PPI (vue macroscopique du fonctionnement du procédé) et local (fonctionnement granulaire des sous-systèmes adjacents).

Nous considérons ici les vecteurs de mesures $y_{l,i}$ et $c_{l,i}$ qui représentent respectivement le signal de sortie et la consigne du $i^{\text{ème}}$ sous-système, l'indice de performance $P_{l,i}$ peut être défini comme suit :

$$P_{l,i}(k) = \left\| \frac{y_{l,i}(k) - c_{l,i}(k)}{c_{l,i}(k)} \right\|_2 \quad (3.1)$$

Avec :

$\| \cdot \|_2$ désigne une norme l_2 .

L'évaluation de l'indice de performance $P_{l,i}$ standardisée est effectué en-ligne et donc instantanée, c.à.d. $P_{l,i}$ est calculé à chaque instant k . Le test séquentiel du rapport de vraisemblance présenté dans le Chapitre 2, section 2.5.3 permet d'évaluer le bon fonctionnement du système.

Nous considérons qu'il y a deux régions de fonctionnement :

- 1) $\mathcal{R}_{l,i}^{Nom} \stackrel{def}{=} \{p_i \text{ tel que } 0 \leq P_{l,i} \leq P_{l,i}^{Nom}\}$
- 2) $\mathcal{R}_{l,i}^{Deg} \stackrel{def}{=} \{p_i \text{ tel que } P_{l,i}^{Nom} < P_{l,i} < P_{l,i}^{Deg}\}$

Ceci implique l'utilisation d'un test d'hypothèses simple afin de permettre l'identification de la région de fonctionnement à un instant donné. Un fonctionnement instable du système est inacceptable et requière un arrêt d'urgence. Une stratégie de tolérance au défaut par gestion de références n'est donc pas pertinente.

L'évaluation des régions de fonctionnement revient à la formulation de deux hypothèses $\mathcal{H}_0$ et $\mathcal{H}_1$ qui sont évalués à chaque instant par le test séquentiel (voir Chapitre 2, section 2.5).

3.4.2 Méthode du gradient pour l'ajustement d'une consigne

La gestion des consignes exogènes sélectionnées par un gouverneur de références doit permettre d'avoir, en fonctionnement dégradé, des performances les plus proches possibles des performances requises (spécifications en fonctionnement nominal).

Ce problème peut être résolu en minimisant la fonction quadratique suivante sous des contraintes sur le critère de performance $P_{l,j}$ et la consigne $c_j \neq c_i$:

$$J = \arg \min_{|c_j| < c_{lim} \wedge P_{l,j} \in \mathcal{R}_{l,j}^{Nom}} \mathbb{E}(P_{l,i}^2(k)) \quad (3.2)$$

La consigne $c_j(k)$ en mode dégradé notée $c_j^{Deg}(k)$ peut être générée d'une manière itérative en considérant l'équation récurrente 3.3.

$$c_j^{Deg}(k) = c_j^{Deg}(k - \tau_S) - \eta \nabla (J(c_j^{Deg}(k - \tau_S))) \quad (3.3)$$

où η est le pas de descente, τ_S est le temps de séjour permettant l'établissement du régime permanent après un changement de valeur de consigne et $\nabla J(c_j^{Deg}(k - \tau_S))$ est la dérivée numérique définie par l'équation suivante :

$$\nabla J(c_j^{Deg}(k - \tau_S)) = \frac{J(k) - J(k - \tau_S)}{c_j^{Deg}(k) - c_j^{Deg}(k - \tau_S)} (\text{Approximation tangente}) \quad (3.4)$$

Le cas d'étude proposé dans la sous-section 3.7 illustre les résultats de l'implémentation de cette méthode.

3.5 Analyse graphique pour la tolérance au défaut

3.5.1 Graphe associé des composantes

La structure de la boucle de régulation est représentée par un sous-graphe du graphe de causalité de l'ensemble du système. Ce sous-graphe peut ensuite être représenté par un seul sommet dans un nouveau graphe appelé "graphe associé des composantes".

Ce type d'abstraction permet de représenter la topologie du système global sous forme d'interaction entre ses boucles de régulation et les consignes exogènes considérées comme des variables de décision dans la stratégie de commande tolérante aux défauts active.

Dans le même esprit que la sous-section 3.2.2, le graphe associé des composantes est noté $\mathcal{G}_C(\mathcal{V}_C, \mathcal{A}_C)$. Il est constitué d'un ensemble de sommets $\mathcal{V}_C$ défini par $\mathcal{V}_C = \mathcal{SP} \cup \mathcal{C}$ qui correspond respectivement à un ensemble de sommets de consignes $\mathcal{SP}$ et un ensemble de sommets de composantes $\mathcal{C}$ correspondant aux structures de boucles de régulation.

L'ensemble $\mathcal{C}$ est composé de sommets de composantes c_i . En se basant sur le graphe orienté $\mathcal{G}(\mathcal{V}, \mathcal{E})$, ces sommets c_i sont associés à des sous-ensembles $\mathbf{c}_i$ définis de la manière suivante :

$$\mathbf{c}_i = \{v_l \mid v_l \text{ couvert par un chemin simple } \{e_i\} - \{v_j\}, \text{ avec } e_i \in \text{Succ}_d(v_j)\}$$

où $e_i \in \mathcal{E}$ et $v_j \in \mathcal{O}$ dans le graphe $\mathcal{G}(\mathcal{V}, \mathcal{E})$.

L'ensemble $\mathcal{SP}$ peut être défini par : $\mathcal{SP} = \mathcal{S} \setminus \{sp_i\}$ avec $sp_i \in \mathbf{c}_i$ avec $\mathcal{S}$ l'ensemble des sommets de consignes dans le graphe $\mathcal{G}(\mathcal{V}, \mathcal{E})$. L'ensemble $\mathcal{SP}$ peut également être défini par : $\mathcal{SP} = \mathcal{V}_C \setminus \mathcal{C}$.

L'ensemble des arcs $\mathcal{A}_C$ peut être défini par $\mathcal{A}_C = \{(v_i, v_j) \mid v_i \in \mathcal{V}_C \text{ et } v_j \in \mathcal{V}_C\}$. Il correspond à des arcs représentant les relations de causalité entre les éléments de $\mathcal{V}_C$.

À noter qu'il n'y a pas de relation de causalité d'une composante $c_i \in \mathcal{C}$ vers une consigne $sp_i \in \mathcal{SP}$, i.e. les arcs (c_i, sp_i) ne sont pas présents dans le graphe $\mathcal{G}_C(\mathcal{V}_C, \mathcal{A}_C)$, puisqu'une consigne est toujours considérée comme le point d'entrée (sommets source) d'une structure de boucle de régulation.

Remarque 2 *Quand une structure de boucle de régulation c_j inclut une autre boucle c_i i.e. $c_i \subseteq c_j$, dans ce cas, nous représentons dans le graphe associé des composantes $\mathcal{G}_C(\mathcal{V}_C, \mathcal{A}_C)$ la structure de boucle de régulation c_j qui inclut c_i .*

Contrairement au graphe de causalité $\mathcal{G}(\mathcal{V}, \mathcal{E})$ qui représente les consignes, les erreurs et les sorties des boucles de régulations, le graphe associé des composantes $\mathcal{G}_C(\mathcal{V}_C, \mathcal{A}_C)$ ne contient que deux types de sommets : les consignes exogènes et les boucles de régulation. Ceci diminue le nombre de sommets dans le graphe associé des composantes tout en gardant la structure du système nécessaire à l'étude de l'accommodation aux défauts. Cette approche permet de réduire la complexité du modèle.

Exemple 4 *Considérons le graphe de causalité de la Figure 3.2. Les sous-ensembles de sommets de ce graphe sont donnés par : $\mathcal{SP} = \{sp_L, sp_{F1}, sp_{F2}\}$, $\mathcal{E} = \{e_L, e_{F1}, e_{F2}\}$ et $\mathcal{O} = \{F_1, F_2, L\}$ correspondant respectivement aux sous-ensembles de consignes, d'erreurs et de sorties des boucles de régulation du système.*

Dans le graphe de causalité de la Figure 3.2, nous avons principalement trois structures de boucles de régulations :

- a- Celle qui couvre les sommets sp_{F_2} et F_2 ;
- b- Celle qui couvre les sommets sp_{F_1} et F_1 ;
- c- Celle qui couvre les sommets e_L , sp_{F_1} , e_{F_1} , F_1 et L .

A travers ces trois structures de boucles de régulation, nous remarquons que la structure de boucle de régulation **c** inclut la **b**. Dans ce cas et selon la Remarque 2, le graphe associé des composantes contient les deux composantes **a** et **c**.

A partir de la définition du graphe associé des composantes $\mathcal{G}_C(\mathcal{V}_C, \mathcal{A}_C)$, ce dernier comporte trois regroupements de boucles de régulation. Ainsi, le graphe orienté de la Figure 3.2 est associé au graphe associé des composantes donné par la Figure 3.4.

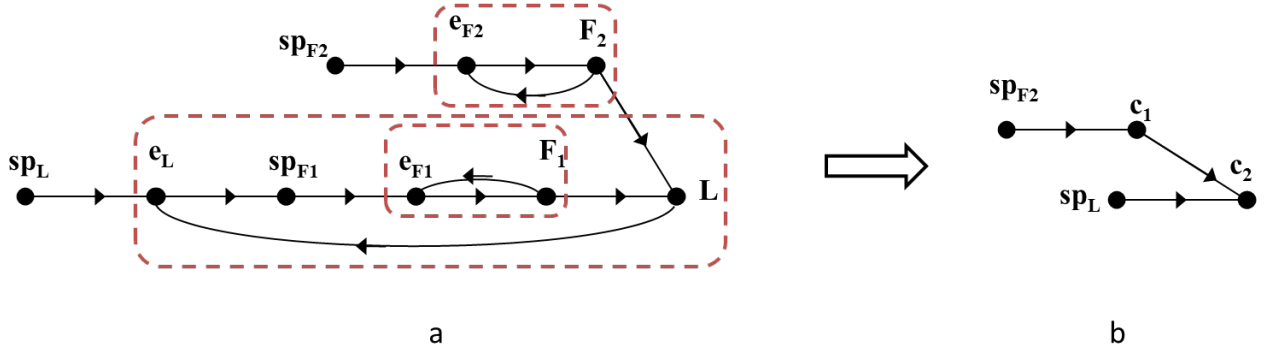


FIGURE 3.4 – (a) Graphe de causalité de la Figure 3.2 et (b) Graphe de composantes associé

3.5.2 Analyse d'atteignabilité

Dans l'analyse graphique, l'atteignabilité est une condition nécessaire dans l'étude de plusieurs propriétés. Une évaluation de cette propriété doit être effectuée en utilisant le graphe associé des composantes. Cette analyse se concentre principalement sur l'atteignabilité des composantes par les consignes exogènes. Ainsi, chaque composante c_i doit être atteignable par une ou plusieurs consignes sp_j , les consignes identifiées seront sélectionnées comme candidates au réajustement par le gouverneur des références.

Puisque les contraintes du système sont inconnues (limitations d'actionneurs, force de couplage entre variables, ...), la stratégie de commande tolérante aux défauts active est basée sur une commande par apprentissage itératif.

L'importance de la propriété d'atteignabilité pour les systèmes automatisés est soulignée dans plusieurs travaux pour l'analyse structurelle Boukhobza (2010); Conrard et al. (2011) et pour la tolérance aux défauts de la commande/diagnostic Staroswiecki and Hoblos (2004).

3.5.2.1 Classification des consignes

Pour l'élaboration d'une stratégie de commande tolérante au défaut, toutes les consignes ne sont pas candidates à la compensation d'un défaut sur une boucle de régulation cible. Pour cette raison, nous proposons une classification des consignes dans le graphe associé des composantes afin de déterminer, pour chaque composante c_i représentant une structure de boucle de régulation, les consignes exogènes pouvant être utilisées pour une accommodation au défaut au niveau de la composante cible en défaut.

Ainsi, les consignes sont classées en deux catégories : "utiles" et "inutiles" .

Pour toutes les composantes c_i du graphe $\mathcal{G}_c$, les consignes inutilisées appartenant à $\mathcal{SP}$ sont groupées dans un ensemble noté $\mathcal{V}_{ul}$ (ul : *useless*, inutilisées) et défini par ce qui suit :

$$\mathcal{V}_{ul} = \left\{ sp_i \in \mathcal{SP} \mid Pred(\{sp_i\}) \neq \emptyset \right\}.$$

$\mathcal{V}_{ul} \subseteq \mathcal{SP}$ est un sous-ensemble de consignes endogènes. Cela signifie qu'il y a une relation de précedence entre ces consignes et les autres variables du système.

Pour chaque composante c_i du graphe $\mathcal{G}_c$, les consignes utiles appartenant à $\mathcal{SP}$ sont groupées dans un ensemble noté $\mathcal{V}_{uf}$ (uf : *useful*, utiles) et défini par ce qui suit :

$$\mathcal{V}_{uf}(c_i) = \left\{ sp_j \in \mathcal{SP} \setminus \mathcal{V}_{ul} \mid \exists \{sp_j\} - \{c_i\} \text{ chemin simple et } Pred(\{sp_j\}) = \emptyset \right\}$$

$\mathcal{V}_{uf}(c_i)$ est un sous-ensemble qui contient toutes les consignes exogènes atteignant une composante cible c_i . Cela signifie que ce sous-ensemble contient des consignes exogènes candidates pour l'accommodation aux défauts au niveau de la composante cible c_i .

3.5.2.2 Composantes impactées par une sélection de consignes

En sélectionnant une consigne sp_t dans $\mathcal{V}_{uf}(c_i)$ pour agir sur la composante cible c_i , il est possible que cette consigne puisse atteindre d'autres composantes qui ne soient pas ciblées. Ces composantes impactées par la variation de la consigne $sp_t \in \mathcal{V}_{uf}(c_i)$ sont groupées dans un ensemble noté $\mathcal{C}_{imp}(sp_t, c_i)$ et dont les paramètres sont la consigne sp_t sélectionnée et la composante cible c_i . Cet ensemble est défini par :

$$\mathcal{C}_{imp}(sp_t, c_i) = \left\{ c_j \in \mathcal{C} \text{ couvert par un chemin simple } \{sp_t\} - \{c_i\} \right\}$$

La définition ci-dessus de l'ensemble $\mathcal{C}_{imp}(sp_t, c_i)$ des composantes impactées par la variation de la consigne sp_t peut être généralisée à un ensemble de consignes sélectionnées $\mathcal{V}_{SP} \subseteq \mathcal{SP}$. Ainsi, pour une composante cible c_i et un ensemble de consignes sélectionnées $\mathcal{V}_{SP}$, nous avons :

$$\mathcal{C}_{imp}(\mathcal{V}_{SP}, c_i) = \bigcup_{sp_t \in \mathcal{V}_{SP}} \mathcal{C}_{imp}(sp_t, c_i)$$

Ainsi, pour agir sur la composante cible c_i , la variation des consignes dans l'ensemble $\mathcal{V}_{SP}$ impacte les composantes dans l'ensemble $\mathcal{C}_{imp}(\mathcal{V}_{SP}, c_i)$. Nous proposons d'associer ces composantes impactées à une fonction "coût". Cette fonction coût peut être représentée par un nombre défini par :

$$\psi(\mathcal{C}_{imp}(\mathcal{V}_{SP}, c_i)) = \text{card}(\mathcal{C}_{imp}(\mathcal{V}_{SP}, c_i))$$

La valeur de ψ correspond au nombre de composantes impactées par la sélection des consignes de $\mathcal{V}_{SP}$ en tenant compte de la composante cible c_i .

Exemple 5 *Considérons le graphe de composantes de la Figure 3.5 avec une composante cible $c_i = c_1$.*

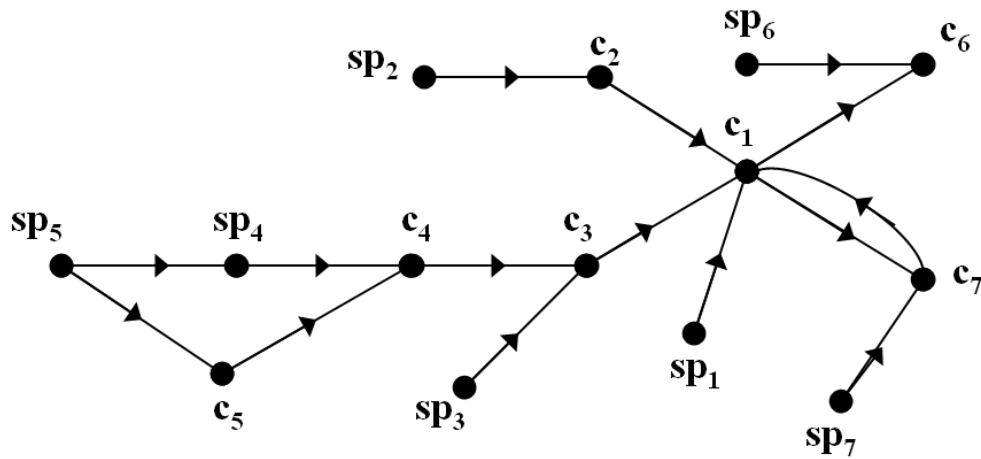


FIGURE 3.5 – Exemple d'un graphe associé des composantes

Pour la composante cible c_1 , nous avons : $\mathcal{V}_{uf}(c_1) = \{sp_1, sp_2, sp_3, sp_5, sp_7\}$ et $\mathcal{V}_{ul} = \{sp_4, sp_6\}$.

Les composantes impactées par la sélection de chaque ensemble de consignes sont résumées dans la Table 3.1.

Consignes	Composantes impactées	$\psi(C_{imp}(\mathcal{V}_{SP}, c_i))$	Chemins considérés
sp_1	$\mathcal{C}_{imp}(sp_1, c_1) = \{c_1\}$	1	$p_1 : sp_1 \rightarrow c_1$
sp_2	$\mathcal{C}_{imp}(sp_2, c_1) = \{c_1, c_2\}$	2	$p_2 : sp_2 \rightarrow c_2 \rightarrow c_1$
sp_3	$\mathcal{C}_{imp}(sp_3, c_1) = \{c_1, c_3\}$	2	$p_3 : sp_3 \rightarrow c_3 \rightarrow c_1$
sp_7	$\mathcal{C}_{imp}(sp_7, c_1) = \{c_1, c_7\}$	2	$p_4 : sp_7 \rightarrow c_7 \rightarrow c_1$
sp_5	$\mathcal{C}_{imp}(sp_5, c_1) = \{c_1, c_3, c_4, c_5\}$	4	$p_5 : sp_5 \rightarrow sp_4 \rightarrow c_4 \rightarrow c_3 \rightarrow c_1$ $p_6 : sp_5 \rightarrow c_5 \rightarrow c_4 \rightarrow c_3 \rightarrow c_1$

TABLE 3.1 – Composantes impactées par les consignes sélectionnées

Remarque 3 Les composantes c_6 et c_7 ne sont pas considérées comme des composantes impactées par la sélection des consignes $sp_i \in \mathcal{V}_{uf}(c_1)$ puisqu'il n'y a aucun chemin simple $\{sp_i\} - \{c_6\}$ ne couvrant pas l'arc (c_1, c_6) et, $\{sp_i\} - \{c_7\}$ ne couvrant pas l'arc (c_1, c_7) . Ainsi, en agissant sur la composante cible c_1 , les deux composantes c_6 et c_7 ne sont pas impactées.

3.5.3 Diagramme de Hasse

Afin de résoudre le problème d'accommodation aux défauts, nous proposons une méthode basée sur le diagramme de Hasse Pavan and Todeschini (2004).

Avant de définir le diagramme de Hasse utilisé pour la stratégie de choix des consignes exogènes, quelques définitions relatives à la théorie des ensembles sont requises.

3.5.3.1 Théorie des ensembles : Définitions

Définition 9 *Un ensemble partiellement ordonné (POSET : Partially ordered Set) est un ensemble P avec une relation binaire $R \subseteq P \times P$ qui vérifie les conditions suivantes pour $x \in P$ et $y \in P$:*

- *Réflexivité : $(x, x) \in R$ pour tout $x \in P$.*
- *Antisymétrie : $(x, y) \in R$ et $(y, x) \in R \Rightarrow x = y$.*
- *Transitivité : $(x, y) \in R$ et $(y, z) \in R \Rightarrow (x, z) \in R$.*

Définition 10 *Soit $P = \{A : A \subseteq X\}$ correspondant à tout ensemble X . Pour $x \in P$ et $y \in P$, nous avons $x \leq y$ si $x \subseteq y$.*

Dans ce cas, $(P, \leq)$ est un POSET pour tout sous-ensemble de X ordonné par inclusion.

3.5.3.2 Diagramme de Hasse

Un ensemble P partiellement ordonné peut être associé à un diagramme de Hasse. Ce diagramme est un graphe orienté acyclique composé de sommets et d'arcs Okada and Gen (1993); Pavan and Todeschini (2004). Le diagramme de Hasse est noté $\mathcal{H}(\mathcal{N}, \mathcal{L})$ où les deux ensembles $\mathcal{N}$ et $\mathcal{L}$ représentent respectivement l'ensemble des sommets associés à des sous-ensembles de X et l'ensemble des arcs qui représentent une relation ordonnée $(P, \leq)$ entre les sous-ensembles de X . Chaque arc (x_i, y_i) de l'ensemble $\mathcal{L}$ caractérise le lien entre les deux sommets $x_i \in P$ et $y_i \in P$. Il est à noter que $\forall x_i \in P$, x_i est associé à un sous-ensemble de X .

Le diagramme de Hasse peut également être considéré comme la représentation graphique de la table de vérité dont les éléments sont ceux de l'ensemble P . Ainsi, les

sommets du diagramme de Hasse correspondent aux différentes lignes de la table de vérité correspondant aux différentes combinaisons des éléments de P .

Exemple 6 *Considérons un ensemble $X = \{a, b, c\}$. L'ensemble P représente tous les sous-ensembles de X composés de trois éléments et ordonnés par inclusion. Ainsi, $P = \{\{\emptyset\}, \{a\}, \{b\}, \{c\}, \{a, b\}, \{a, c\}, \{b, c\}, \{a, b, c\}\}$.*

Le diagramme de Hasse $\mathcal{H}(\mathcal{N}, \mathcal{L})$ associé au POSET $(P, \leq)$ est donné à la Figure 3.6.

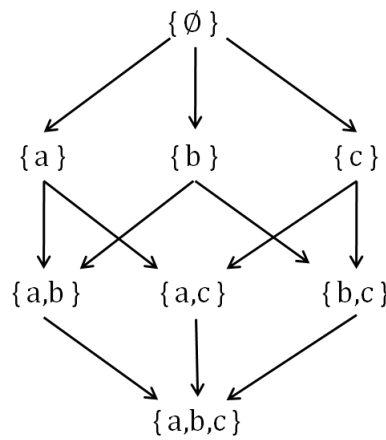


FIGURE 3.6 – Diagramme de Hasse $\mathcal{H}(\mathcal{N}, \mathcal{L})$ associé au POSET $(P, \leq)$

Remarque 4 *Noter que dans un diagramme de Hasse, il n'existe aucune relation de précedence entre les sommets d'un même niveau horizontal ainsi que d'autres paires de sommets telles que $\{a\}$ et $\{c, b\}$.*

Il est possible de faire le lien entre les sommets du diagramme de Hasse de la Figure 3.6 et la table de vérité correspondant aux différentes combinaisons de a , b et c . Cette table de vérité est donnée à la Table 3.2.

L'utilisation du diagramme de Hasse présente beaucoup d'avantages par rapport à l'utilisation de la table de vérité. Ces avantages peuvent être résumés dans :

- L'ordre partiel de l'ensemble P est mieux représenté avec le diagramme de Hasse ;
- À l'aide des chemins sur le diagramme de Hasse, les inclusions entre les sous-ensembles

Sommets de $\mathcal{H}(\mathcal{N}, \mathcal{L})$	Table de vérité de		
	a	b	c
$\{\emptyset\}$	0	0	0
$\{a\}$	1	0	0
$\{b\}$	0	1	0
$\{c\}$	0	0	1
$\{a, b\}$	1	1	0
$\{a, c\}$	1	0	1
$\{b, c\}$	0	1	1
$\{a, b, c\}$	1	1	1

TABLE 3.2 – Diagramme de Hasse/Table de vérité

sont représentées ;

- La représentation des combinaisons des éléments de P est intuitive ;
- Les combinaisons sont classées de telle sorte à ce que chaque ligne du diagramme de Hasse contienne les combinaisons de la même cardinalité.

3.5.3.3 Diagramme de Hasse pondéré

Un graphe pondéré est un graphe dont les arcs sont associés à des indices correspondant à des poids numériques. Ces poids peuvent correspondre à des nombres entiers ou réels représentant une caractéristique des arcs associés.

A partir du diagramme de Hasse $\mathcal{H}(\mathcal{N}, \mathcal{L})$ présenté dans la sous-section 3.5.3, nous pouvons définir un diagramme de Hasse pondéré. Ainsi, le diagramme de Hasse pondéré est noté $\mathcal{H}^w(\mathcal{H}, W)$. Il correspond à un diagramme de Hasse $\mathcal{H}$ et un ensemble de poids numériques W associés aux arcs du diagramme de Hasse $\mathcal{H}$. Un poids numérique $w_{ij} \in W$ est associé à l'arc liant un sommet $N_i \in \mathcal{N}$ à un autre sommet $N_j \in \mathcal{N}$.

Sachant que les sommets d'un diagramme de Hasse correspondent à des sélections de consignes, les poids w_{ij} dans un diagramme de Hasse pondéré sont définis par :

$$w_{ij} = \psi(\mathcal{C}_{imp}(N_i, c_i))$$

Un poids w_{ij} associé à un arc de N_i à N_j correspond donc au nombre de composantes impactées par la sélection de consignes associée au sommet N_i .

Exemple 7 *Considérons le graphe associé des composantes donné dans la Figure 3.7.*

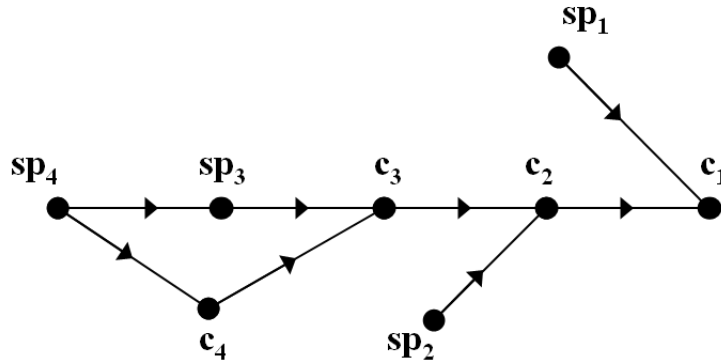


FIGURE 3.7 – Exemple d'un graphe associé des composantes

Scénario de défaut :

Considérons que le défaut survient au niveau de la boucle de régulation 2 associée à la composante c_2 . Ce défaut est propagé jusqu'à la boucle de régulation 1 associée à la composante c_1 . Cette boucle de régulation correspond donc à un indicateur de performance KPI dégradé.

Une stratégie de commande tolérante aux défauts est adoptée pour compenser l'effet du défaut sur la boucle de régulation 1 en ajustant les consignes sélectionnées sur les autres boucles de régulation.

Analyse d'atteignabilité :

Dans cet exemple, la composante cible est c_1 . L'analyse d'atteignabilité porte donc sur l'atteignabilité entre les consignes et la composante c_1 . Cette analyse est effectuée en plusieurs étapes.

La première étape consiste à obtenir l'ensemble des consignes utiles et inutiles :

- $\mathcal{V}_{ul} = \{sp_3\}$;
- $\mathcal{V}_{uf}(c_1) = \{sp_1, sp_2, sp_4\}$.

Après avoir défini les consignes utiles, les ensembles de composantes impactées par l'ajustement de ces consignes ainsi que la fonction coût représentée par ψ sont donnés par la Table 3.3 :

Consignes	$\mathcal{C}_{imp}(sp_i, c_1)$	$\psi(\mathcal{C}_{imp}(sp_i, c_1))$
sp_1	$\mathcal{C}_{imp}(sp_1, c_1) = \{c_1\}$	1
sp_2	$\mathcal{C}_{imp}(sp_2, c_1) = \{c_1, c_2\}$	2
sp_4	$\mathcal{C}_{imp}(sp_4, c_1) = \{c_1, c_2, c_3, c_4\}$	4

TABLE 3.3 – Composantes impactées par les consignes dans $\mathcal{V}_{uf}(c_1)$ et la fonction coût

Ce système correspond au diagramme de Hasse $\mathcal{H}(\mathcal{N}, \mathcal{L})$ donné à la Figure 3.8.

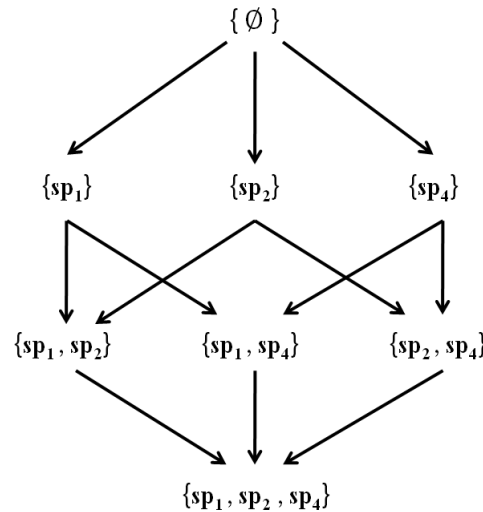


FIGURE 3.8 – Diagramme de Hasse correspondant à l'exemple 7

A partir de ce diagramme de Hasse, il est possible d'obtenir le diagramme de Hasse pondéré $\mathcal{H}^w(\mathcal{H}, W)$ en calculant les poids w_{ij} . Pour ce faire et en utilisant la Table 3.3, nous dressons la Table 3.4 représentant les composantes impactées par la sélection des consignes associées aux sommets de $\mathcal{H}(\mathcal{N}, \mathcal{L})$. Pour chaque sommet, les sélections peuvent être associées à des

sous-ensembles $\mathcal{V}_{SP}$.

$\mathcal{V}_{SP}$ associés aux sommets	$\mathcal{C}_{imp}(\mathcal{V}_{SP}, c_1)$	$\psi(\mathcal{C}_{imp}(\mathcal{V}_{SP}, c_1))$
$\{\emptyset\}$	$\{\emptyset\}$	0
$\{sp_1\}$	$\{c_1\}$	1
$\{sp_2\}$	$\{c_1, c_2\}$	2
$\{sp_4\}$	$\{c_1, c_2, c_3, c_4\}$	4
$\{sp_1, sp_2\}$	$\{c_1, c_2\}$	2
$\{sp_1, sp_4\}$	$\{c_1, c_2, c_3, c_4\}$	4
$\{sp_2, sp_4\}$	$\{c_1, c_2, c_3, c_4\}$	4
$\{sp_1, sp_2, sp_4\}$	$\{c_1, c_2, c_3, c_4\}$	4

TABLE 3.4 – Composantes impactées

A partir de la Table 3.4, chaque arc entrant dans un sommet N_i , qui correspond à une sélection $\mathcal{V}_{SP}$, est associé au poids $w_{ij} = \mathcal{C}_{imp}(\mathcal{V}_{SP}, c_1)$.

Le diagramme de Hasse pondéré $\mathcal{H}^w(\mathcal{H}, W)$ correspondant est donné par la Figure 3.9.

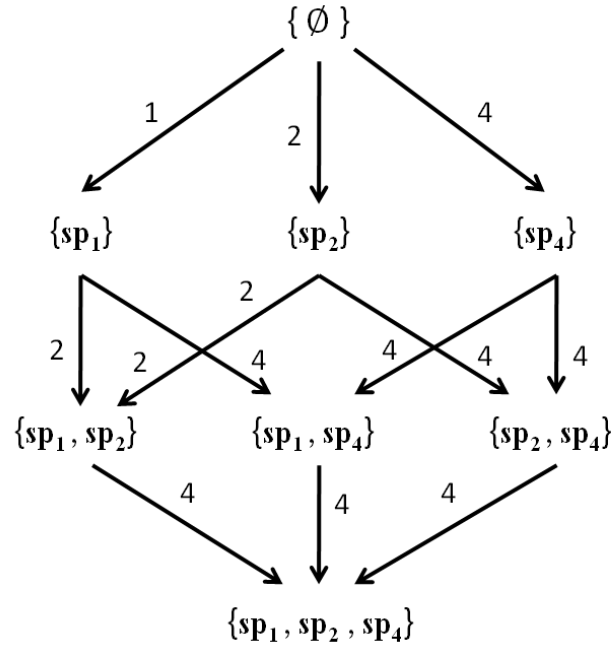


FIGURE 3.9 – Diagramme de Hasse pondéré $\mathcal{H}^w(\mathcal{H}, W)$

Ce diagramme de Hasse pondéré peut être décomposé en niveaux. Le niveau 0 est donc composé du sommet $\{\emptyset\}$, le niveau 1 est composé de $\{sp_1\}$, $\{sp_2\}$ et $\{sp_4\}$, le niveau 2 est composé de $\{sp_1, sp_2\}$, $\{sp_2, sp_4\}$ et $\{sp_1, sp_4\}$, et le niveau 3 est composé de $\{sp_1, sp_2, sp_4\}$.

3.6 Stratégie de la tolérance au défaut adoptée

Dans cette section, nous énonçons une stratégie conjointe d'analyse graphique et de gestion de références. Un cas d'étude est présenté dans la section 3.7.

Le diagramme de Hasse pondéré $\mathcal{H}^\omega$, obtenu dans la section précédente, permet d'avoir plusieurs combinaisons de consignes candidates. Dans ce diagramme, nous avons également l'information concernant les composantes impactées par la variation de chacune des combinaisons de consignes. La stratégie de sélection des consignes est donc basée sur le diagramme de Hasse pondéré $\mathcal{H}^\omega$.

Quand une boucle de régulation est en défaut et nécessite qu'on agisse dessus par l'intermédiaire d'une ou plusieurs consignes, il faut prévoir une stratégie de sélection parmi les consignes candidates en utilisant le diagramme de Hasse pondéré $\mathcal{H}^\omega$. Nous rappelons les 3 étapes de construction de ce diagramme.

Étape 1 : Identification de la composante cible

Cette première étape consiste à déterminer la composante en question (dont le PPI est dégradé). Ainsi, le point de départ est de fixer la composante cible c_i .

Étape 2 : Définition des consignes utiles

Après avoir identifié la composante cible, il faut déterminer la liste des m consignes utiles constituant l'ensemble $\mathcal{V}_{uf}$ défini dans la section 3.5.2.1. Les consignes exogènes $sp_j \in \mathcal{V}_{uf}$ sont des consignes candidates à la compensation du défaut impactant la composante cible c_i . Ceci est traduit par l'existence d'un chemin entre sp_j et c_i dans le graphe associé des composantes $\mathcal{G}$.

Étape 3 : Construction du diagramme de Hasse et diagramme de Hasse pondéré

Le diagramme de Hasse peut être construit à partir de l'ensemble des consignes utiles $sp_j \in \mathcal{V}_{uf}$. Ses sommets s_{ij} correspondent à toutes les combinaisons possibles entre les consignes sp_j pouvant agir sur la composante cible.

Le diagramme de Hasse ainsi construit peut être pondéré. La pondération des arcs du diagramme de Hasse dépend des composantes impactées $\mathcal{C}_{imp}$ par chacune des combinaisons de consignes (voir 3.5.2.2).

Après l'obtention du diagramme de Hasse pondéré, la stratégie de sélection peut être présentée. Cette dernière repose sur deux principaux traitements.

Le premier est basé sur un parcours du diagramme de Hasse pondéré de façon itérative. Ce parcours permet de sélectionner une combinaison parmi les différents sommets s_{ij} des différents niveaux du diagramme de Hasse pondéré en fonction des composantes de $\mathcal{C}_{imp}$.

Une fois qu'une combinaison de consignes est choisie (premier traitement), le deuxième traitement consiste à établir un ordre d'ajustement des consignes de la combinaison choisie.

1^{er} traitement : Parcours du diagramme de Hasse pondéré $\mathcal{H}^\omega$

Niveau 0

Ce premier traitement est effectué sur le diagramme de Hasse pondéré $\mathcal{H}^\omega$. Initialement, la composante cible c_i n'est pas en défaut. Dans ce cas, il n'y a pas besoin d'avoir une sélection de consignes à ajuster. Ce qui correspond au premier sommet $\{\emptyset\}$ du diagramme de Hasse pondéré $\mathcal{H}^\omega$. Ainsi, le premier traitement ne s'applique pas sur ce niveau 0 du diagramme mais à partir du niveau 1.

Quand il y a occurrence de défaut au niveau de la composante cible c_i , il faut quitter le niveau 0 du diagramme de Hasse pondéré (qui contient le sommet $\{\emptyset\}$). Les arcs

sortants de ce sommet $\{\emptyset\}$ mènent vers les m singletons $\{sp_j\}$ des consignes utiles candidates dans le niveau 1.

Niveau 1

Les consignes des m singletons $\{sp_j\}$ peuvent être ajustées dans un ordre précis. Leur classement dépend des poids w_{ij} des arcs associés au nombre de composantes impactées C_{imp} de chaque consigne sp_j . La consigne privilégiée est celle avec un arc entrant associé au nombre minimal de composantes impactées. Si les poids w_{ij} de deux arcs entrants dans deux sommets s_{ij} différents sont égaux, alors la technique d'échantillonnage aléatoire simple est nécessaire. Par exemple, sur le diagramme de Hasse pondéré de la Figure 3.9, la consigne privilégiée est sp_1 puisque son arc entrant correspond au poids minimal 1.

A chaque ajustement d'une consigne sp_j , un test est effectué par le gouverneur de références sur la composante cible c_i , et plus précisément le PPI correspondant, afin de vérifier si cette dernière n'est plus affectée par le défaut. Si le résultat du test est négatif et que la composante cible a besoin d'être ajustée à nouveau, une deuxième consigne est associée à la consigne choisie pour être ajustées. Un test est également effectué après le nouvel ajustement. Par exemple, sur le diagramme de Hasse pondéré de la Figure 3.9, si l'ajustement de la consigne sp_1 n'est pas suffisant, sp_1 et sp_2 sont toutes les deux ajustées et ainsi de suite.

Niveau i

De la même façon que le niveau 1, les combinaisons de consignes d'un niveau i sont choisies en fonction des poids des arcs entrants w_{ij} . Un test est ensuite effectué après chaque ajustement de consignes. Si les résultats du niveau i sont négatifs, alors les consignes des combinaisons du niveau $i + 1$ sont ajustées puis testées et ainsi de suite jusqu'à accommodation au défaut sur c_i ou jusqu'au dernier $n^{\text{ème}}$ niveau du diagramme de Hasse pondéré. Ce dernier niveau correspond à un seul sommet s_{ij} où toutes les consignes utiles

de $\mathcal{V}_{uf}$ sont ajustées.

2^{ème} traitement : Ordre d'ajustement de consignes dans une même combinaison

A travers le diagramme de Hasse pondéré, des combinaisons d'une ou plusieurs consignes sont ajustées. Quand une combinaison contient plus d'une consigne, un ordre d'ajustement de consigne est à prévoir. Cet ordre est obtenu à travers le graphe de composantes associé $\mathcal{G}_C$ (voir la section 3.5.1).

Dans une combinaison $\{sp_1, sp_2, \dots, sp_i\}$, la consigne privilégiée est celle qui atteint la composante cible c_i à travers le chemin le plus court dans $\mathcal{G}_C$. Par exemple, pour le graphe associé des composantes de la Figure 3.7, nous avons une composante cible $c_i = c_1$ et un ensemble de consignes utiles $\mathcal{V}_{uf} = \{sp_1, sp_2, sp_4\}$. Ces trois consignes peuvent être ajustées dans l'ordre suivant : sp_1 qui est lié à la cible par un chemin de longueur 1, puis sp_2 qui est lié à la cible par un chemin de longueur 2, puis sp_4 qui est lié à la cible par deux chemins de longueur 4.

3.7 Cas d'étude : Application à un système de chauffage-mixage

3.7.1 Présentation du système

L'objectif de l'utilisation d'un banc d'essai est de permettre la simulation du comportement d'un système de pointe ainsi que sa commande. Les bancs d'essai sont généralement utilisés pour des systèmes industriels.

Ainsi, nous considérons pour ce cas d'étude un système non linéaire de chauffage-mixage composé principalement de trois cuves. Ce système est conçu avec plusieurs boucles de régulation. L'objectif principal du système de commande correspondant consiste à maintenir un rapport constant entre des liquides dans les trois cuves ainsi que de maintenir

la température de ces liquides en respectant certaines exigences.

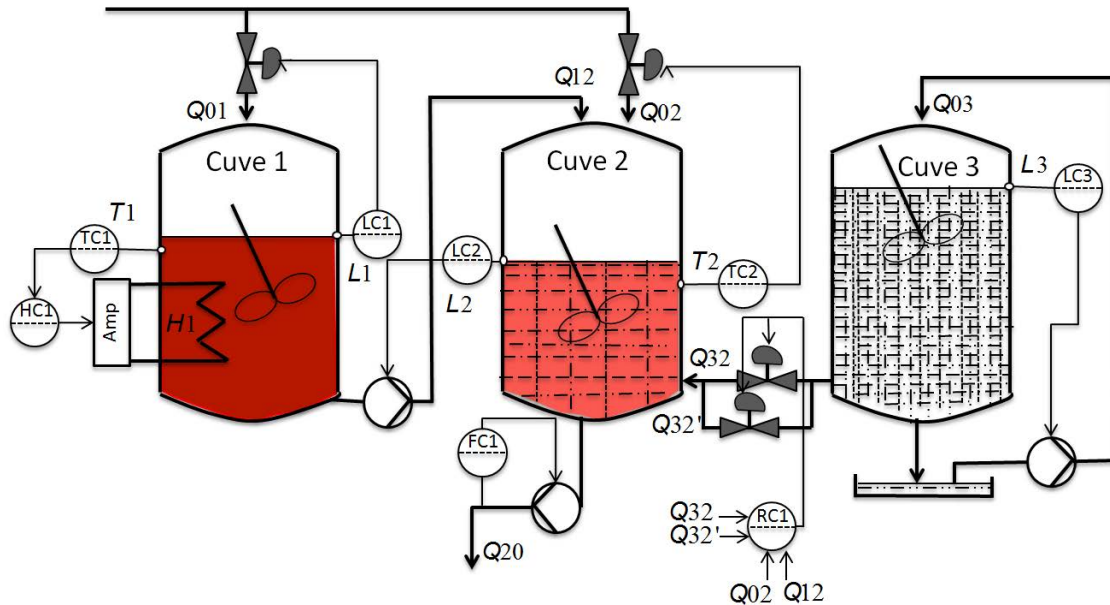


FIGURE 3.10 – Système de chauffage-mixage à trois cuves

Le schéma d'usine de ce système est donné à la Figure 3.10. Ce schéma représente la tuyauterie ainsi que l'instrumentation du système (P&ID). La cuve 1 est utilisée pour le chauffage. La cuve 2 est dédiée au mixage des liquides provenant des cuves 1 et 2. La Figure 3.10 représente également les différents niveaux, flux de liquides, chauffages et mixeurs. Les variables mesurables de ce système sont groupées dans la Table 3.5.

3.7.2 Représentation graphique du système

Ce système de chauffage-mixage à trois cuves et ses variables sont représentés par le graphe de causalité orienté donné à la Figure 3.11. Ce graphe de causalité représente les différents liens entre les consignes, les erreurs et les variables de sortie du système (débits, températures et niveaux).

A partir du graphe de causalité orienté associé au système, les composantes associées aux boucles de régulation peuvent être déterminées comme précisé dans la section 3.5.1. La

Variable mesurée	Unité de mesure	Désignation
L_1	m	Niveau de la cuve 1
L_2	m	Niveau de la cuve 2
L_3	m	Niveau de la cuve 3
K_{ratio}	-	Ratio entre les flux d'entrée des liquides des cuves 1 et 3
T_1	°C	Température dans la cuve 1
T_2	°C	Température dans la cuve 2
Q_{00}	l/m	Débit d'entrée de la pompe 1
Q_{01}	l/m	Débit d'entrée de la pompe 1 vers la cuve 1
Q_{02}	l/m	Débit d'entrée de la pompe 1 vers la cuve 2
Q_{12}	l/m	Débit d'entrée de la cuve 1 (liquide chauffé) vers la cuve 2
Q_{20}	l/m	Débit de sortie du produit final de la cuve 2
Q_{30}	l/m	Débit de sortie de la cuve 3
Q_{32}	l/m	Débit du liquide de la cuve 2 vers la cuve 3
$Q_{32'}$	l/m	Débit du liquide de la cuve 2 vers la cuve 3

TABLE 3.5 – Sorties des boucles de régulation du cas d'étude

Table 3.6 représente chacune des boucles de régulation de ce système ainsi que les variables qui y correspondent.

Composantes c_i	Variables correspondantes	Composantes $c_j \subset c_i$
c_1	$\{e_1, sp_{H_1}, e_2, H_1, T_1\}$	$\{e_2, H_1\}$
c_2	$\{e_3, sp_{Q_{02}}, e_4, Q_{02}, T_2\}$	$\{e_4, Q_{02}\}$
c_3	$\{e_5, sp_{Q_{01}}, e_6, Q_{01}, L_1\}$	$\{e_6, Q_{01}\}$
c_4	$\{e_7, sp_{Q_{12}}, e_8, Q_{12}, L_2\}$	$\{e_8, Q_{12}\}$
c_5	$\{e_9, sp_{Q_{03}}, e_{10}, Q_{03}, L_3\}$	$\{e_{10}, Q_{03}\}$
c_6	$\{e_{11}, sp_{Q_{32}/Q_{32'}}, e_{12}, Q_{32}/Q_{32'}, K_{ratio}\}$	$\{e_{12}, Q_{32}/Q_{32'}\}$
c_7	$\{e_{13}, Q_{30}\}$	—
c_8	$\{e_{14}, Q_{20}\}$	—

TABLE 3.6 – Composantes c_i du graphe à la Figure 3.12

Ainsi, le graphe de causalité orienté de la Figure 3.11 peut être simplifié en un graphe associé des composantes où seules les consignes sp_i et les composantes c_i correspondantes aux boucles de régulation sont représentées. Le schéma simplifié est donné à la Figure 3.12.

3.7.3 Identification du diagramme de Hasse pondéré

Nous considérons que la composante en défaut est c_1 . L'ensemble des consignes utiles définies dans la Section 3.5.2.1 est donné par :

$$\mathcal{V}_{uf} = \{sp_{T_2}, sp_{L_1}, sp_{L_2}, sp_{L_3}, sp_{Q_{20}}, sp_{Q_{30}}, sp_{K_{ratio}}\}.$$

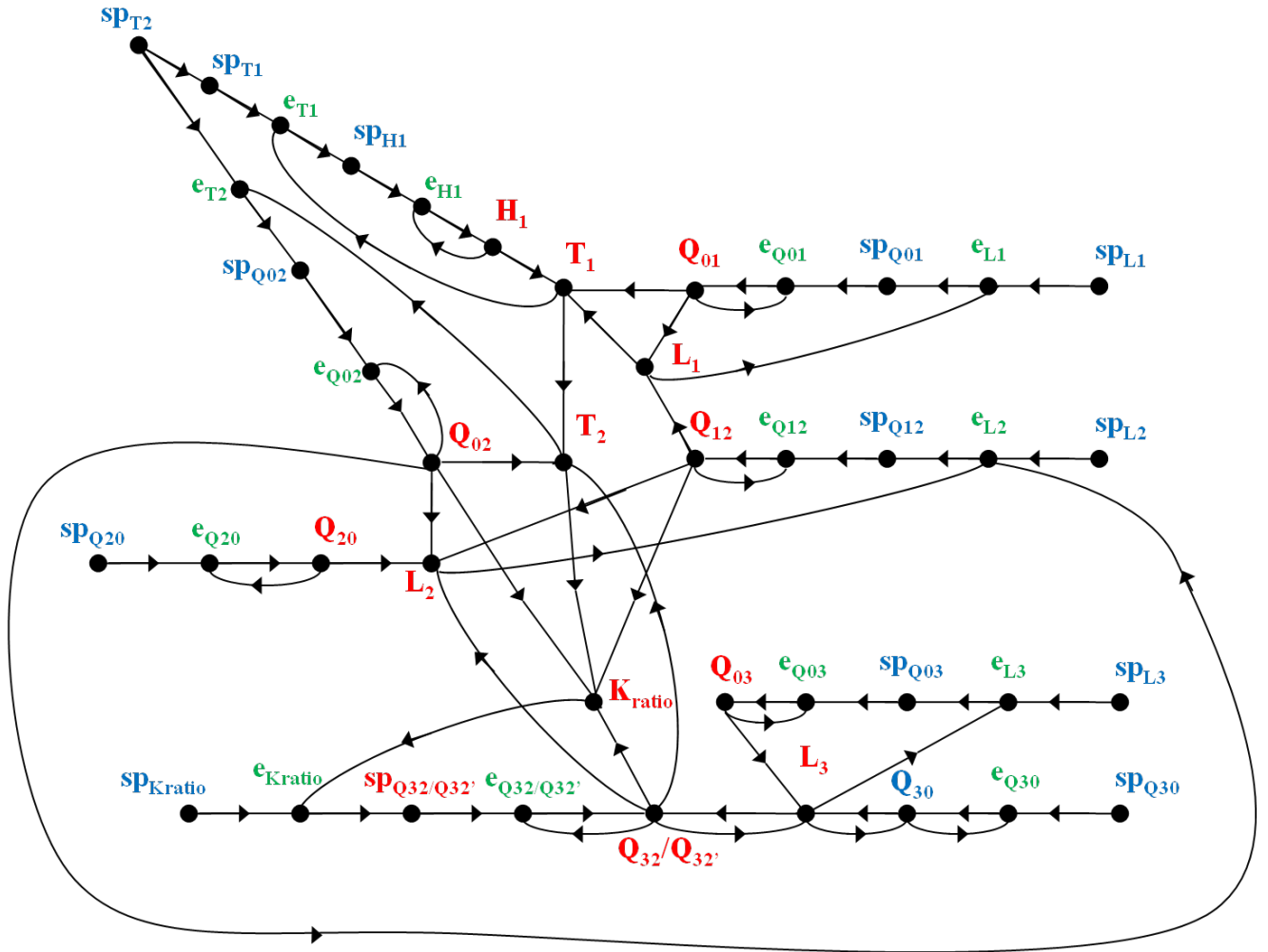


FIGURE 3.11 – Graphe de causalité associé au système de chauffage-mixage

Pour ce système, les 7 consignes de $\mathcal{V}_{uf}$ sont associées aux ensemble $\mathcal{C}_{imp}(sp_i, c_1)$ donnés dans la Table 3.7. A partir de cette Table, nous obtenons le diagramme de Hasse pondéré donné dans la Figure 3.13.

Consigne sp_i	$\mathcal{C}_{imp}(sp_i, c_1)$	$\psi(\mathcal{C}_{imp}(sp_i, c_1))$
sp_{T2}	$\{c_1, c_2\}$	2
sp_{L1}	$\{c_1, c_2, c_3\}$	3
sp_{L2}	$\{c_1, c_2, c_3, c_4, c_6\}$	5
sp_{Kratio}	$\{c_1, c_2, c_3, c_4, c_6\}$	5
sp_{L3}	$\{c_1, c_2, c_3, c_4, c_6, c_7\}$	6
sp_{Q20}	$\{c_1, c_2, c_3, c_4, c_5, c_6\}$	6
sp_{Q30}	$\{c_1, c_2, c_3, c_4, c_6, c_7, c_8\}$	7

TABLE 3.7 – Composantes impactées

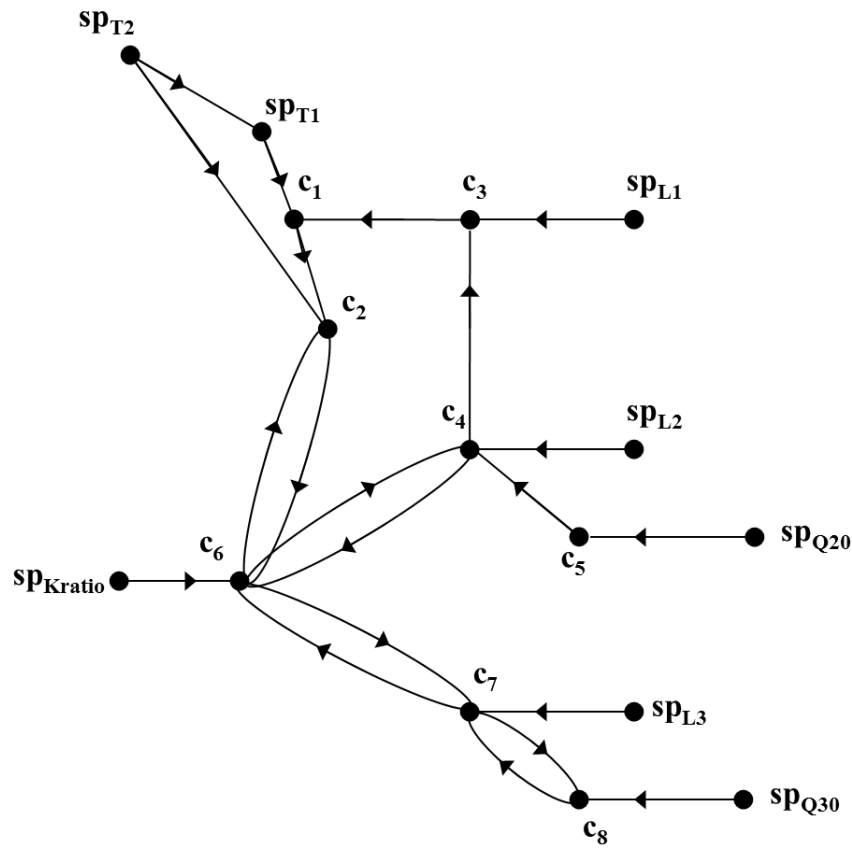


FIGURE 3.12 – Graphe associé des composantes au système de chauffage-mixage

3.7.4 Définitions des indicateurs de performance

1. Indicateurs de Performances Procédé (PPI) :

(a) Indicateurs de Performances Qualité (Température) :

$$\Delta T = T_2^{\text{SP}} - T_2;$$

$$Q_T = \begin{cases} 1, & |\Delta T| \leq 0.5 \\ 1 - \frac{|\Delta T|}{3.5}, & 0.5 < |\Delta T| \leq 3.5 \\ 0, & \text{sinon;} \end{cases}$$

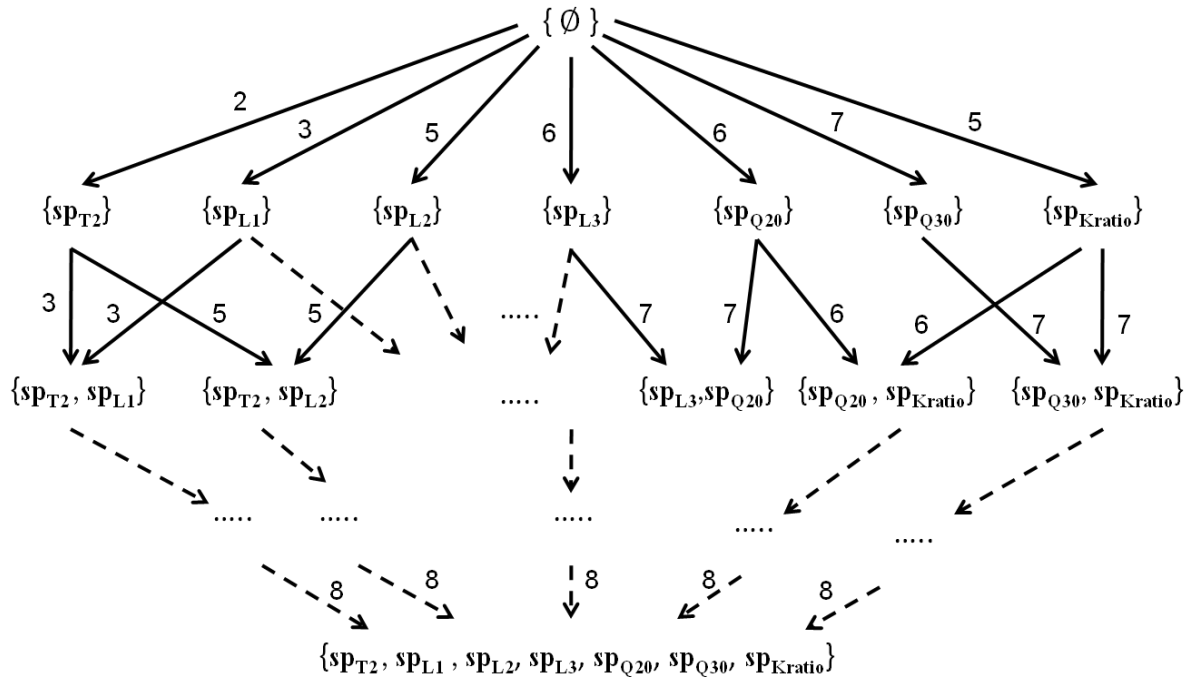


FIGURE 3.13 – Diagramme de Hasse pondéré associé au cas d'étude

(b) Indicateurs de Performances Qualité (Ratio) :

$$\Delta k_{\text{ratio}} = k_{\text{ratio}}^{\text{SP}} - k_{\text{ratio}};$$

$$Q_{k_{\text{ratio}}} = \begin{cases} 1 - \frac{|\Delta k_{\text{ratio}}|}{k_{\text{ratio}}^{\text{SP}}}, & |\Delta k_{\text{ratio}}| \leq k_{\text{ratio}}^{\text{SP}} \\ 0, & \text{sinon;} \end{cases}$$

(c) Indicateurs de Performances Production :

$$Q_f = \frac{Q_{20}}{Q_{20}^{\text{SP}}}$$

3.7.5 Modes de fonctionnement

1. Mode de fonctionnement nominal :

Les caractéristiques statistiques des différentes mesures du banc d'essai en mode de fonctionnement nominal sont représentées par μ_0 et σ_0 . Ces caractéristiques sont considérées comme étant des données d'apprentissage pour l'algorithme de détection

de défaut. Cet algorithme, présenté dans le Chapitre 2, section 2.5.3, est le test séquentiel de vraisemblance (en anglais SPRT) pour la vérification de deux hypothèses (binaire).

Le but de ce test est de détecter un saut de moyenne ou de variance et d'estimer le temps d'occurrence du saut.

En occurrence d'un défaut, on peut considérer intuitivement deux types de changement abrupts (sauts).

- Une déviation de la valeur moyenne nominale μ_0 vers une valeur moyenne μ_1 , avec un écart-type constant.
- Une augmentation de l'écart type de σ_0 à σ_1 , avec moyenne constante.

Les paramètres μ_0 et σ_0 de chacune des deux situations sont supposés connus. Un indice de performance est considéré comme étant une variable aléatoire. En mode de fonctionnement nominal, la variance de l'indice de performance est caractérisée par σ_0^2 , ainsi, une dégradation de l'indice de performance, dans le cas d'un saut de variance, se caractérise par une augmentation de variance σ_0 ($\sigma_1 > \sigma_0$).

- H_0 : V.A de variance σ_0^2
- H_1 : V.A de variance σ_1^2 ($\sigma_1 > \sigma_0$)

Une dégradation est détectée si $\sum_{i=1}^{\tau_d} (x_i - \mu_0)^2 = TH_1(\tau_d)$ et τ_d est le temps de détection de la dégradation et x_i est la $i^{\text{ème}}$ observation.

TH_1 est le seuil d'acceptation de l'hypothèse H_1 .

2. Mode de fonctionnement dégradé :

- Scénario de défaut :

Perte d'efficacité de 50 % de capacité de chauffage de la résistance H_1 .

Le défaut commence à partir de l'instant 2000.

▷ Graphe de causalité pour le diagnostic : Figure 3.14

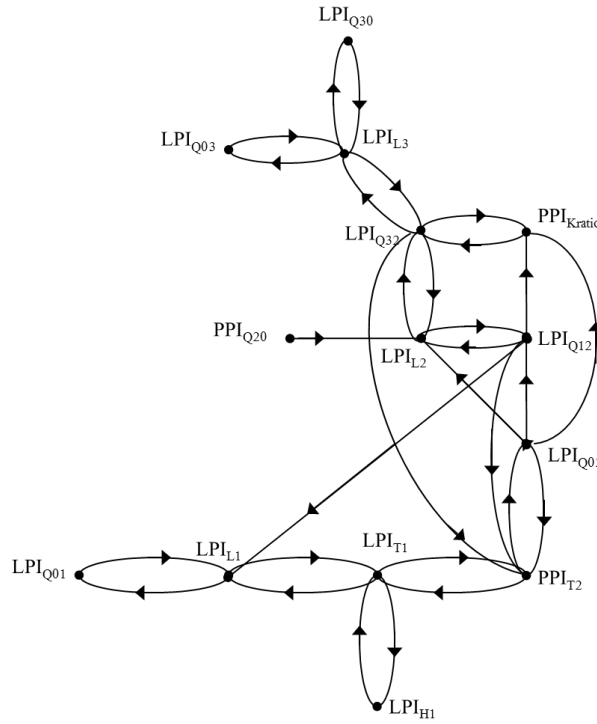


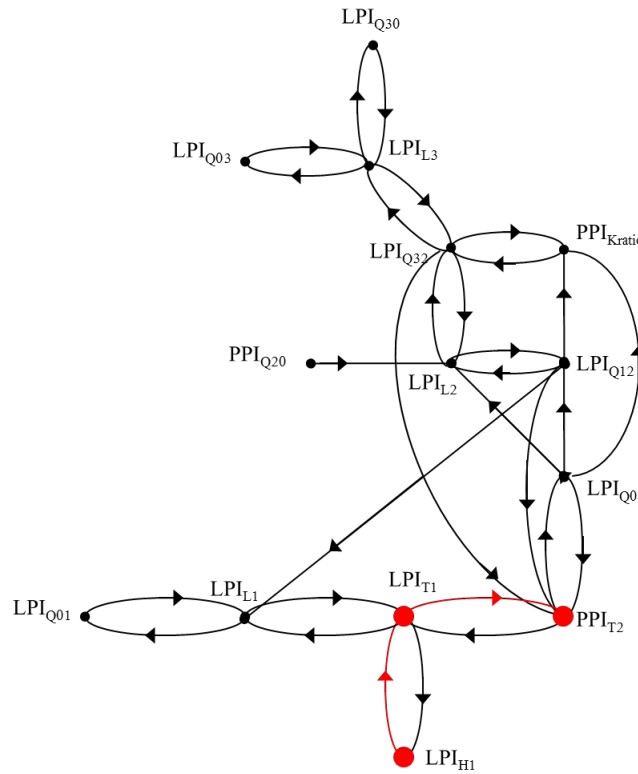
FIGURE 3.14 – Graphe orienté associé au modèle de causalité du benchmark

▷ Graphe de causalité pour le diagnostic illustrant le chemin de propagation du défaut : Figures 3.15.

3.7.6 Sélection et Ajustement des consignes

Pour des objectifs de simplification, nous reformulons le diagramme de Hasse dans la Figure 3.16 en considérons un ensemble réduit de consignes sp_{L1} , sp_{L2} , sp_{Q20} et sp_{Kratio} .

Remarque 5 La consigne sp_{T2} est écartée de l'analyse pour la raison suivante. En effet, le PPI_{T2} est défini comme étant l'erreur de poursuite qui dépend de la sortie de la température T_2 et la consigne sp_{T2} . Le choix de sp_{T2} est trivial, néanmoins son ajustement implique le changement des caractéristiques du produit et notamment sa température. Ceci implique le

FIGURE 3.15 – Chemin de propagation du défaut $\{LPI_{H1}\} - \{LPI_{T2}\}$

non-respect des exigences du cahier des charges (continuer à produire avec la même température de référence).

Ainsi, en se basant sur ce diagramme de Hasse, nous appliquons les premier et deuxième traitements de la stratégie de sélection des consignes exogènes.

Niveau 0

En ce niveau 0, aucun traitement n'est nécessaire puisque ce niveau correspond à la phase avant occurrence du défaut.

Niveau 1

En ce niveau 1, nous nous intéressons aux consignes appartenant aux singletons du diagramme de Hasse pondéré. Plus précisément, aux poids w_{ij} des arcs entrants dans les

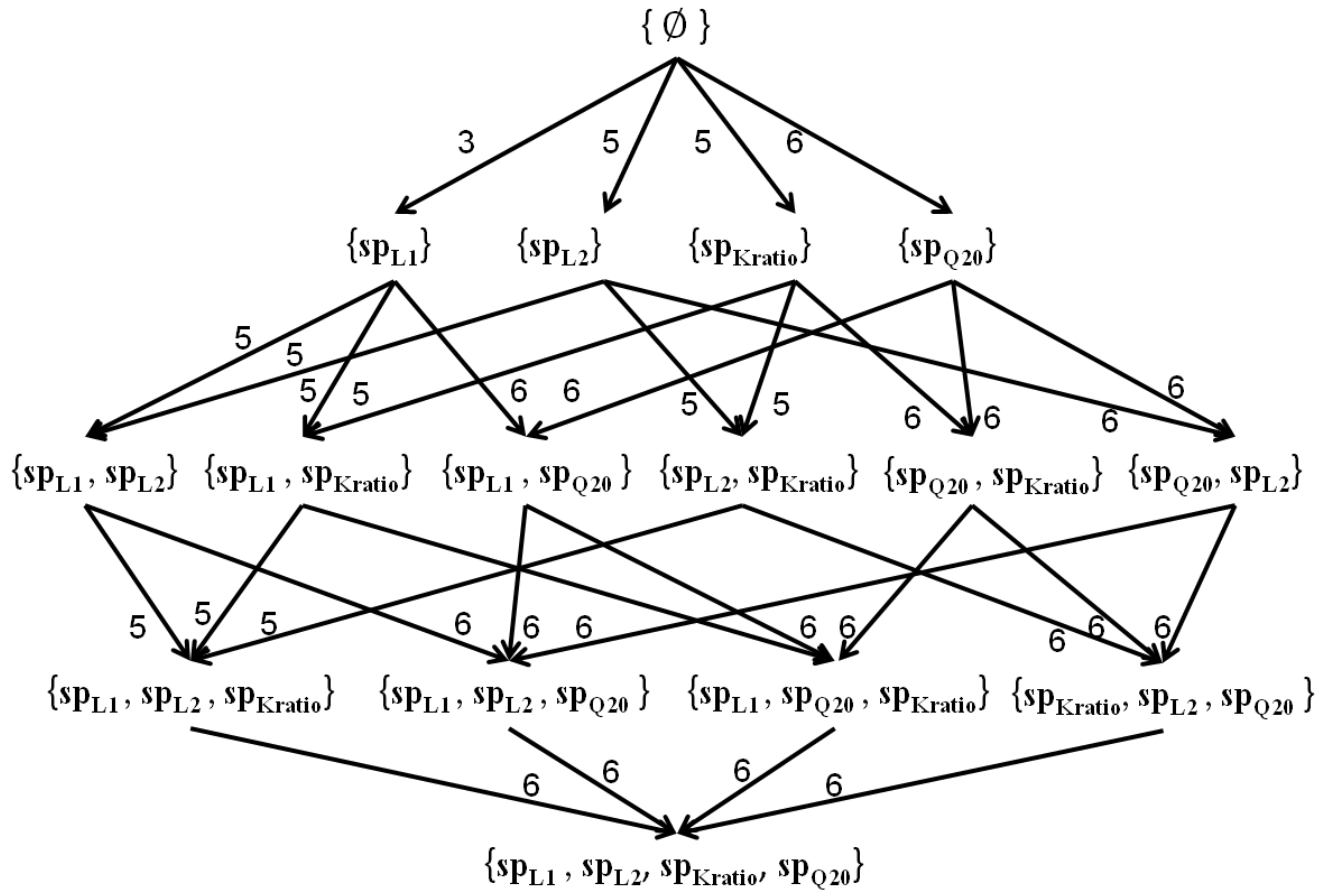


FIGURE 3.16 – Diagramme de Hasse pondéré simplifié

sommets associés à ces singletons.

Dans notre cas, l'ajustement des consignes du niveau 1 suit un ordre bien précis. La première consigne à être ajustée est sp_{L1} . Cette consigne correspond à un nombre minimal de composantes impactées c.à.d. 3 composantes.

Un test est ensuite effectué par le gouverneur de références. Le résultat de ce test est donné par la Figure 3.17. Nous remarquons que l'ajustement de la consigne sp_{L1} n'est pas suffisant puisque le LPI_{T1} associée à la composante c_1 ne permet pas une compensation du défaut sur la composante cible c_2 dont l'indicateur PPI_{T2} est associé. À travers l'analyse énoncée précédemment, sp_{L1} peut être associée à une deuxième consigne sp_{Q20} .

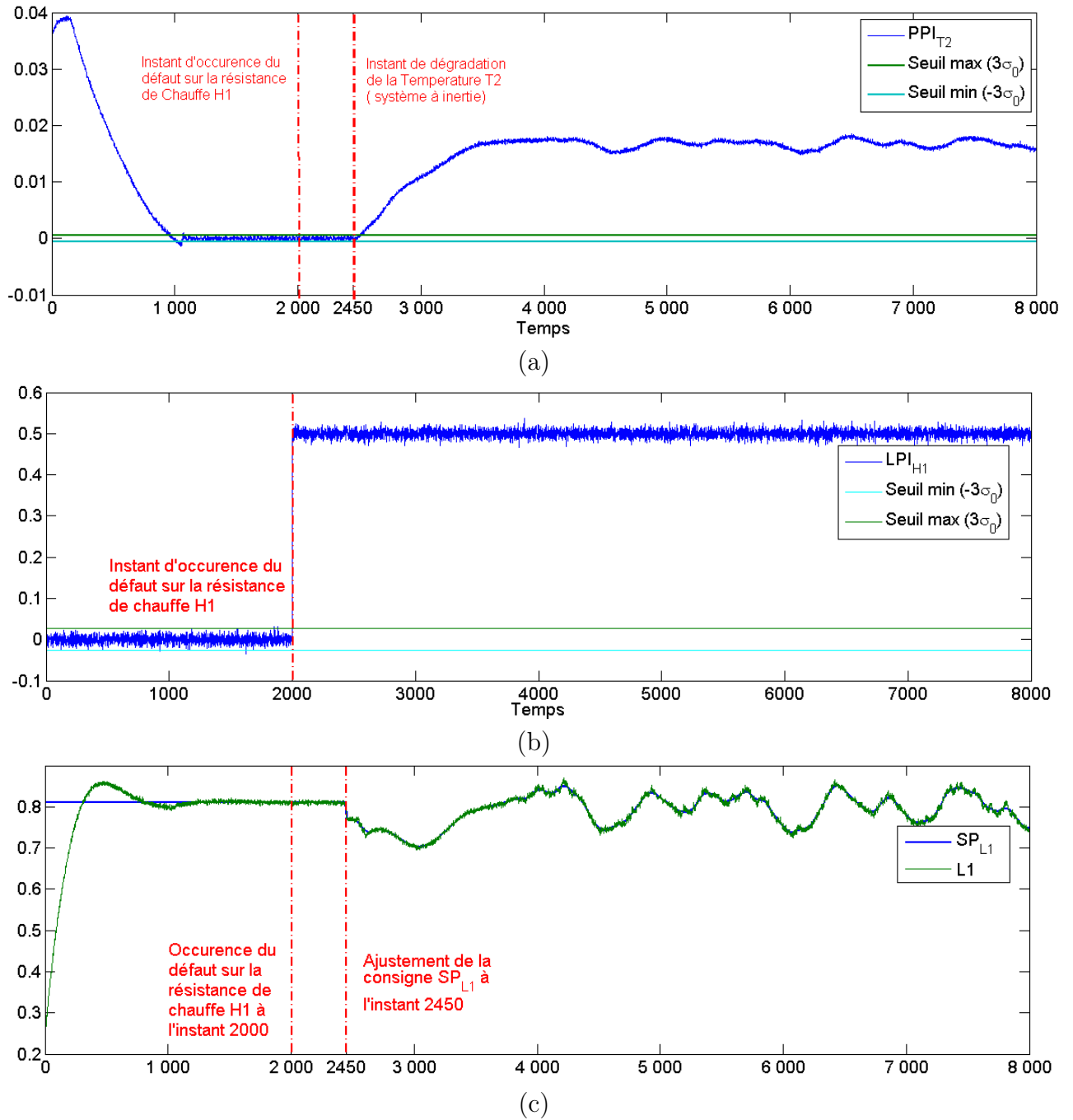


FIGURE 3.17 – Test par rapport à la variation de sp_{L1} : (a) Indicateur de performance Procédé T_2 , (b) Indicateur de performance Local H_1 , (c) Consigne sp_{L1} et la variable $L1$.

Niveau 2

En ce niveau 2, nous considérons par exemple la combinaison $\{sp_{L1}, sp_{Q20}\}$ puisqu'elle représente le nombre minimal de composantes impactées 5 de ce niveau. Une gestion de la combinaison est ensuite effectuée par le gouverneur de références. Le résultat de ce test est donné par la Figure 3.18.

Ces deux ajustements ont permis de revenir aux spécifications fixées par le cahier de charge. Aucun autre ajustement n'est donc nécessaire.

3.8 Conclusion

Ce chapitre a été consacré à l'étude des systèmes tolérants aux défauts. Chaque système peut être associé à une matrice d'adjacence contenant des éléments "1" et des éléments "0" qui représentent la causalité directe entre une paire de variables du système et son sens. A partir de cette matrice d'adjacence, il est possible de représenter les liens de causalité entre les différentes variables du système, c'est à dire les consignes, les erreurs et les sorties des boucles de régulation, par un graphe de causalité orienté représentant les différents liens entre les variables. Nous nous intéressons en particulier aux défauts survenant au niveau des boucles de régulation. Ainsi, le graphe de causalité orienté est simplifié pour ne représenter que les boucles de régulation ainsi que leurs consignes. Il s'agit du graphe associé des composantes. Notre étude est principalement basée sur cette représentation du système.

Afin de compenser l'effet d'un défaut dans une boucle de régulation, il est nécessaire d'agir sur cette dernière en ajustant une ou plusieurs consignes. Ces consignes doivent respecter la condition graphique d'atteignabilité. Cette condition nécessite l'existence d'un chemin dans le graphe associé des composantes qui lie la boucle de régulation en défaut et les consignes à ajuster. Ainsi, toutes les consignes pouvant agir sur une boucle de régulation en compensant l'effet du défaut sont regroupées dans un ensemble de consignes utiles.

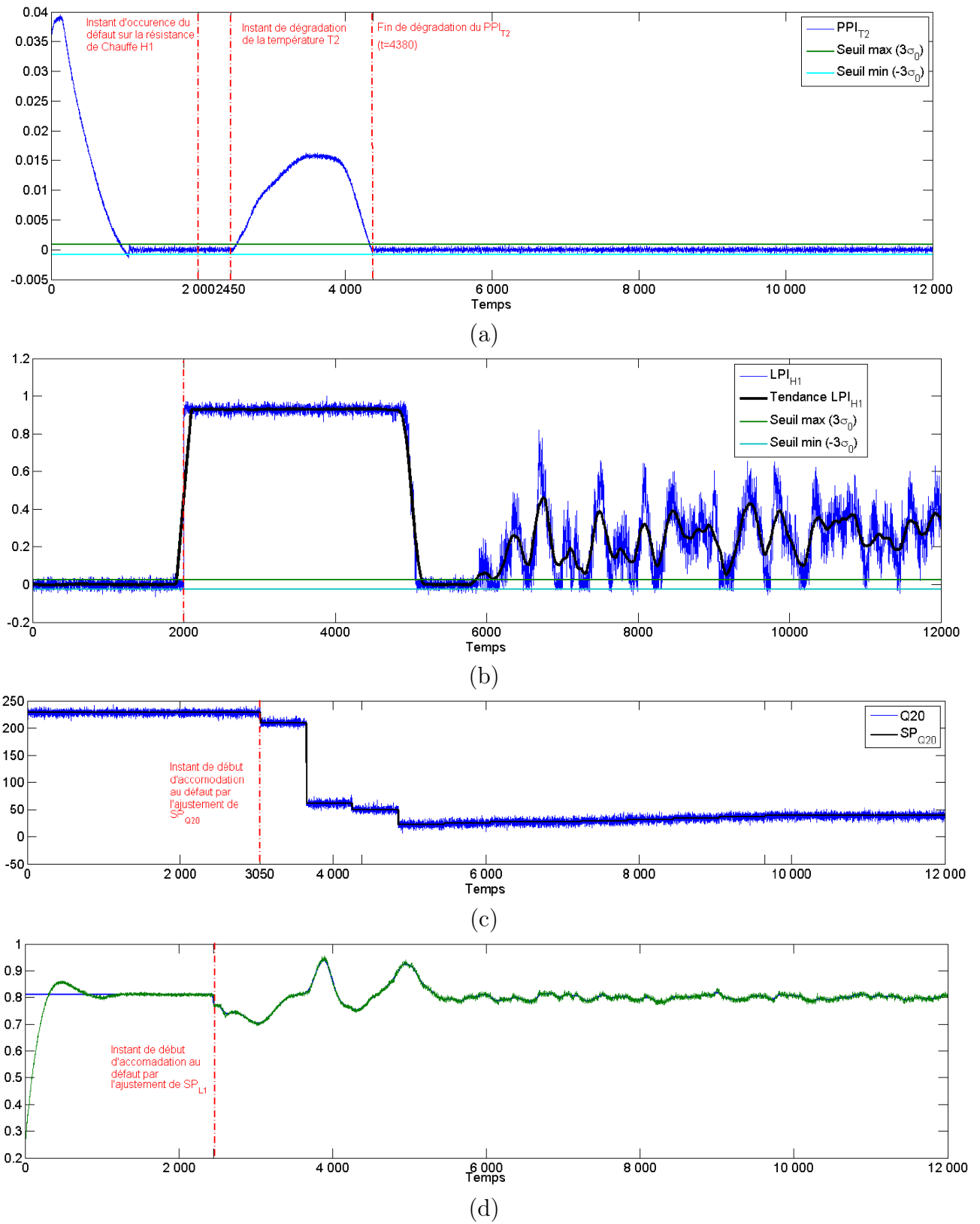


FIGURE 3.18 – Test par rapport à la variation de sp_{L1} et sp_{Kratio} : (a) Indicateur de performance Procédé T_2 , (b) Indicateur de performance Local H_1 , (c) Consigne $sp_{Q_{20}}$ et la variable Q_{20} , (d) Consigne sp_{L1} et la variable $L1$.

Après avoir déterminé les consignes utiles, un diagramme de Hasse peut être obtenu en indiquant les combinaisons possibles des consignes à ajuster. Ce diagramme de Hasse est ensuite pondéré où ses arcs sont associés à des indices représentant le nombre des boucles de régulation impactées par l'ajustement d'une ou plusieurs consignes. Le diagramme de Hasse pondéré obtenu est utilisé pour élaborer une stratégie de sélection de consignes exogènes. Cette stratégie est présentée sous forme d'étapes et de traitement effectués sur le diagramme de Hasse pondéré et le graphe associé des composantes. La méthode proposée est illustrée par un cas d'étude qui est un système de chauffage mixage à trois cuves.

Conclusion et perspectives

Dans le contexte économique actuel de l'industrie, une performance croissante est requise pour les systèmes industriels. Ces derniers doivent produire suivant des contraintes contractuelles et réglementaires strictes. Malheureusement, ces systèmes complexes sont sujets à l'occurrence des défauts qui peuvent causer des pertes économiques considérables ; l'analyse de la diagnosticabilité, le développement des approches de diagnostic et de tolérance aux défauts sont des éléments indispensables dans l'amélioration de la sûreté de fonctionnement du système.

Ainsi, le travail de thèse, présenté dans ce document, s'articule autour de 3 aspects, à savoir, la modélisation de causalité des systèmes, le diagnostic, l'analyse des causes et l'accommodation aux défauts afin que le système puisse accomplir sa mission. Le premier chapitre introduit le cadre de cette thèse. Le deuxième chapitre concerne la modélisation et le diagnostic des systèmes complexes. Le troisième et dernier chapitre est dédié à l'analyse des systèmes tolérant aux défauts.

Cette thèse est réalisée dans le cadre du projet européen PAPYRUS (*Plug And Play monitoring and control architecture for optimization of large scale production processes*) en collaboration avec des industriels et des universitaires européens (Allemagne et Finlande). Le premier chapitre présente l'architecture et le concept de la plateforme PAPYRUS. Plus particulièrement, le Système de Gestion d'Actifs du Procédé (PAM : Plant Asset Management) qui est utilisé principalement pour assurer des tâches de diagnostic et d'aide à la décision. Également, les différents objectifs du projet sont présentés en détails sous forme

de lots de travail (Work Packages) ainsi que les différents partenaires européens chargés de les accomplir. En dernière partie de ce chapitre, nous précisons les travaux de thèse réalisés dans le cadre de ce projet. Ces travaux concernent principalement le Work package 4 (WP4) et le Work package 7 (WP7) dédiés à la modélisation de causalité, au diagnostic et à la tolérance aux défauts des systèmes complexes.

Le chapitre 2 est dédié à la modélisation et le diagnostic des systèmes complexes. La modélisation est présentée à travers trois méthodes graphiques différentes. La première est celle basée sur un modèle structurel. À partir des équations d'état du système, nous pouvons représenter le système par un graphe orienté représentant la structure du système. Un exemple d'illustration est également fourni. La deuxième méthode est basée sur les données. Elle concerne trois approches de modélisations de la causalité. Il s'agit de l'inter-corrélation, la causalité de Granger et le transfert d'entropie. Un exemple d'illustration est fourni pour chacune de ces approches. La troisième méthode est basée sur l'utilisation combinée du modèle et des données. Elle permet de représenter le système sous forme d'un graphe de causalité à travers une matrice d'adjacence.

La propriété de diagnosticabilité est définie pour les systèmes représentés par le modèle structurel. Les conditions graphiques de cette propriété sont appliquées à un exemple d'illustration. Le diagnostic est analysé à travers un test statistique de détection de dégradation des performances du système, ainsi qu'une analyse des causes permettant de définir l'origine du défaut.

Dans le chapitre 3, nous nous intéressons à l'étude des systèmes tolérants aux défauts. Dans la première partie, nous considérons des systèmes sujets à des défauts et représentés par un graphe de causalité obtenu à partir de la matrice d'adjacence correspondante. Ce graphe est simplifié ensuite en un graphe de composantes associé où seules les boucles de régulation et les consignes du système sont représentées. Le graphe de composantes associé permet de déterminer les consignes pouvant agir sur la boucle de régulation en défaut.

En deuxième partie, une stratégie de sélection de ces consignes est proposée en utilisant le diagramme de Hasse puis un diagramme de Hasse pondéré par rapport à certaines contraintes graphiques. Un cas d'étude est fourni en dernière partie du chapitre concernant les résultats présentés pour l'analyse des systèmes tolérants aux défauts sur un système industriel de chauffage-mixage.

Les méthodes et algorithmes proposés dans cette thèse peuvent être adaptés à plusieurs applications dont l'amélioration de la sûreté de fonctionnement. Cette dernière représente un aspect pouvant améliorer les performances du système et limiter les pertes économiques. À titre d'exemple, nous citons les processus de rationalisation et de classification des alarmes permettant leur gestion en fonction de différents critères à savoir la non-conformité de la qualité des produits. Ces processus peuvent être améliorés en se basant sur l'approche conjointe de diagnostic présentée dans ce manuscrit. En effet, la priorisation des chemins de propagation des défauts (dans le graphe de causalité), l'isolation des sommets en défaut n'impactant pas les indicateurs de performance clé constituent une bonne approche dans la classification des alarmes.

Les outils proposés dans ce manuscrit permettront aussi une amélioration continue des outils d'analyse fonctionnelle/dysfonctionnelle existants, par exemple l'AMDEC qui est traditionnellement utilisés en industrie. En effet, la sévérité et la fréquence d'occurrence (paramètres AMDEC) de certains défauts affectant la qualité des produits peuvent être aisément réévaluées (amélioration continue des informations quantitatives des outils de maintenance) et ainsi permettre l'amélioration des actions de maintenance préventives et curatives.

Pour conclure, nous estimons que la modélisation, le diagnostic et la tolérance aux défauts sont des outils importants, nécessaires et très utiles dans le domaine de la surveillance des systèmes complexes à grande dimension. Ils assurent aux utilisateurs du système de production une continuité et qualité de service (QoS) rendant un système

robuste en présence d'un défaut.

Résumé

Le travail de thèse, qui s'inscrit dans le cadre du projet européen PAPYRUS (Plug and Play monitoring and control architecture for optimization of large scale production processes) du 7ème PCRD, a concerné tout d'abord la synthèse et la mise en œuvre d'une approche de modélisation, de diagnostic et de reconfiguration originale. Celle-ci se fonde sur la génération de graphes causaux permettant de modéliser en temps réel le comportement d'un système complexe dans un premier temps. La cible de cette première étude a été la papeterie Stora Enso d'Imatra en Finlande, qui était le procédé d'application du projet PAPYRUS. En suite logique à cette première partie, une approche permettant l'accommodation du système à certains défauts particuliers a été définie par l'ajustement des signaux de consigne de diverses boucles de régulation.

Le manuscrit est structuré en trois parties.

Dans la première, le projet européen PAPYRUS est présenté. Le rôle de chaque partenaire y est décrit au travers des différents « workpackages » et le travail de thèse y est positionné. La seconde partie de la thèse a pour objectif la génération d'un modèle utile au diagnostic en se fondant uniquement sur les différents signaux mesurés du système. Plus précisément, un modèle causal graphique est présenté par la mise en évidence des liens de causalité entre les différentes variables mesurées. Des analyses à base d'inter-corrélation, de transfert d'entropie et du test de causalité de Granger sont effectuées. Une approche de diagnostic fondée sur le modèle graphique ainsi obtenu est ensuite proposée en utilisant un test d'hypothèse séquentiel.

La dernière partie est dédiée au problème d'accommodation aux défauts. Le graphe utilisé pour établir le diagnostic du système est remanié afin de faire apparaître les différentes boucles de régulation du système. Une stratégie permettant la sélection de consignes influentes est alors proposée avec l'objectif d'ajuster ces dernières afin de compenser l'effet du défaut survenu.

Bibliographie

- Amadi-Echendu, J. E., Brown, K., Willett, R., and Mathew, J. (2010). *Definitions, Concepts and Scope of Engineering Asset Management*. Springer Science & Business Media.
- Baghli, M. (2006). A model-free characterization of causality. *Economics Letters*, 91(3) :380–388.
- Barnett, L. and Seth, A. K. (2014). The MVGC multivariate Granger causality toolbox : A new approach to Granger-causal inference. *Journal of Neuroscience Methods*, 223 :50–68.
- Basseville, M. and Nikiforov, I. V. (1993). *Detection of abrupt changes : theory and application*. Prentice Hall.
- Bauer, M. (2005). Data-driven methods for process analysis. *University College London, Diss.*
- Bauer, M., Cox, J., Caveness, M., Downs, J. J., and Thornhill, N. (2007). Finding the direction of disturbance propagation in a chemical process using transfer entropy. *IEEE Transactions on Control Systems Technology*, 15(1) :12–21.
- Bauer, M. and Thornhill, N. F. (2008). A practical method for identifying the propagation path of plant-wide disturbances. *Journal of Process Control*, 18(7–8) :707–719.
- Blanke, M., Kinnaert, M., Lunze, J., and Staroswiecki, M. (2006). Diagnosis and fault-tolerant control. *Springer-Verlag, Heidelberg*.

- Boukhobza, T. (2010). Partial state and input observability recovering by additional sensor implementation : a graph-theoretic approach. *International Journal of Systems Science*, 41(11) : 1281–1291.
- Brunet, J., Jaume, D., and Labarrere, M. (1990). *Détection et diagnostic de pannes : approche par modélisation*. Hermès.
- Burnham, K. and Anderson, D. (2002a). Model selection and multimodel inference : a practical information-theoretic approach. *2nd. Springer*.
- Burnham, K. P. and Anderson, D. R. (2002b). *Model Selection and Multimodel Inference : A Practical Information-Theoretic Approach*. Springer Science & Business Media.
- Burnham, K. P. and Anderson, D. R. (2004). Multimodel inference understanding AIC and BIC in model selection. *Sociological Methods & Research*, 33(2) :261–304.
- Campbell, J. D., Jardine, A. K. S., and McGlynn, J. (2010). *Asset Management Excellence : Optimizing Equipment Life-Cycle Decisions, Second Edition*. CRC Press.
- Cao, L. (1997). Practical method for determining the minimum embedding dimension of a scalar time series. *Physica A*, 110.
- Chávez, M., Martinerie, J., and Le Van Quyen, M. (2003). Statistical assessment of nonlinear causality : application to epileptic EEG signals. *Journal of Neuroscience Methods*, 124(2) :113–128.
- Chen, J. and Patton, R. J. (1999). Robust model-based fault diagnosis for dynamic systems. *The International Series on Asian Studies in Computer and Information Science*, 3.
- Chiang, L. H. and Braatz, R. D. (2003). Process monitoring using causal map and multivariate statistics : fault detection and identification. *Chemometrics and Intelligent Laboratory Systems*, 65(2) :159–178.
- Chu, T. and Glymour, C. (2006). Semi-parametric causal inference for nonlinear time series data. *Journal of Machine Learning Research*.

- Cocquempot, V. (2004). *Contribution à la surveillance des systèmes industriels complexes*. Habilitation à diriger des recherches, Université Des Sciences et Thechnique de Lille.
- Commault, C., Dion, J.-M., and Agha, S. Y. (2008). Structural analysis for the sensor location problem in fault detection and isolation. *Automatica*, 44(8) :2074–2080.
- Commault, C., Dion, J. M., Sename, O., and Motyeian, R. (2002). Observer-based fault detection and isolation for structured systems. *IEEE Transactions on Automatic Control*, 47(12) :2074–2079.
- Conrard, B., Cocquempot, V., and Mili, S. (2011). Fault-tolerant system design in multiple operating modes using a structural model. *In proceeding of the European Safety & Reliability Conference, Troyes, France*.
- Dion, J. M., Commault, C., and van der Woude, J. (2003). Generic properties and control of linear structured systems : a survey. *Automatica*, 39(7) : 1125–1144.
- Dustegor, D., Cocquempot, V., Staroswiecki, M., and Frisk, E. (2004). Isolabilité structurelle des défaillances. application à un modèle de vanne. *Journal Européen des Systèmes Automatisés*, 38(1-2) :103–124.
- Eichler, M. (2013). Causal inference with multiple time series : principles and problems. *Philosophical Transactions of the Royal Society of London A : Mathematical, Physical and Engineering Sciences*, 371(1997).
- Faghraoui, A., Kabadi, M., Sauter, D., Boukhobza, T., and Aubrun, C. (2014). Data-driven causality digraph modeling of large-scale complex system based on transfer entropy. In *2014 IEEE Conference on Control Applications (CCA)*, pages 705–710.
- Faghraoui, A., Kabadi, M.-G., Kosayyer, N., Morel, D., Sauter, D., and Aubrun, C. (2012). Soa-based platform implementing a structural modelling for large-scale system fault detection : application to a board machine. In *IEEE Multi-Conference on Systems and Control, MSC 2012*, page CDR0M, Dubrovnik, Croatie.

- Fleming, K. (1974). Reliability model for common mode failures in redundant systems. *GA-A-13284*.
- Forsell, U. and Ljung, L. (1999). Closed-loop identification revisited. *Automatica*, 35(7) :1215–1241.
- Franck, P. M. (1990). Fault diagnosis in dynamic systems using analytical and knowledge-based redundancy a survey and some new results. *Automatica*, 26(3) :459–474.
- Friedman, N., Nachman, I., and Peér, D. (1999). Learning bayesian network structure from massive datasets : the "sparse candidate" algorithm. In *Proceedings of the Fifteenth conference on Uncertainty in artificial intelligence*, UAI'99, pages 206–215, San Francisco, CA, USA. Morgan Kaufmann Publishers Inc.
- Friston, K., Moran, R., and Seth, A. K. (2013). Analysing connectivity with granger causality and dynamic causal modelling. *Current Opinion in Neurobiology*, 23(2) :172–178.
- Friston, K. J., Bastos, A. M., Oswal, A., van Wijk, B., Richter, C., and Litvak, V. (2014). Granger causality revisited. *NeuroImage*, 101 :796–808.
- GAO, D., WU, C., ZHANG, B., and MA, X. (2010). Signed directed graph and qualitative trend analysis based fault diagnosis in chemical industry. *Chinese Journal of Chemical Engineering*, 18(2) :265–276.
- Gilson, M. and Van den Hof, P. (2005). Instrumental variable methods for closed-loop system identification. *Automatica*, 41(2) :241–249.
- Granger, C. (1969). Investigating causal relations by econometric models and cross-spectral methods. *Econometrica*, 37(3) :424–438.
- Granger, C. (1988). Some recent development in a concept of causality. *Journal of Econometrics*, 39(1–2) :199–211.
- Heckerman, D. (1995). A tutorial on learning with bayesian networks. Technical report, Learning in Graphical Models.

- Heiman, G. W. (2010). Basic statistics for the behavioral sciences. *Cengage Learning*.
- Horch, A. (2000). Condition monitoring of control loops. *Tekniska hoegsk.*
- Jensen, F. (1996). *An Introduction to Bayesian Networks*. London.
- Johansson, R. (1993). System modeling and identification. *Prentice-Hall information and system sciences series*.
- Kabadi, M. G., Hamelin, F., Aubrun, C., and Boukhobza, T. (2012). Discrete mode observability analysis of switching structured linear systems with unknown input. a graphical approach. *In proceeding of 10th European Workshop on Advanced Control and Diagnosis, ACD Copenhagen, Denmark*.
- Kalisch, M., Bühlmann, P., and Chickering, M. (2005). Estimating high-dimensional directed acyclic graphs with the pc-algorithm. Technical report, Journal of Machine Learning Research.
- Lear, J. (1988). *Aristotle : The Desire to Understand*. Cambridge University Press.
- Lee, G., Song, S.-O., and Yoon, E. S. (2003). Multiple-fault diagnosis based on system decomposition and dynamic PLS. *Industrial & Engineering Chemistry Research*, 42(24) :6145–6154.
- Lee, W. B., Moh, S.-Y., and Choi, H.-J. (2012). Plant asset management today and tomorrow. In Mathew, J., Ma, L., Tan, A., Weijnen, M., and Lee, J., editors, *Engineering Asset Management and Infrastructure Sustainability*, pages 1–17. Springer London.
- Leung, D. and Romagnoli, J. (2002). An integration mechanism for multivariate knowledge-based fault diagnosis. *Journal of Process Control*, 12(1) :15–26.
- Ligeza, A. and Koscielny, J. (2008). A new approach to multiple fault diagnosis : A combination of diagnostic matrices, graphs, algebraic and rule-based models. the case of two-layer models. *International Journal of Applied Mathematics and Computer Science*, 18(4) :465–476.

- Lü, N. and Wang, X. (2007). Sdg-based hazop and fault diagnosis analysis to the inversion of synthetic ammonia. *Tsinghua Science & Technology*, 12(1) :30–37.
- Marschinski, R. and Kantz, H. (2002). Analysing the information flow between financial time series. *The European Physical Journal B - Condensed Matter and Complex Systems*, 30(2) :275–281.
- Maurya, M., Rengaswamy, R., and Venkatasubramanian, V. (2007). A signed directed graph and qualitative trend analysis-based framework for incipient fault diagnosis. *Chemical Engineering Research and Design*, 85(10) :1407–1422.
- Maurya, M. R., Rengaswamy, R., and Venkatasubramanian, V. (2006). A signed directed graph-based systematic framework for steady-state malfunction diagnosis inside control loops. *Chemical Engineering Science*, 61(6) :1790–1810.
- May, W. E. (1970). Knowledge of causality in hume and aquinas. *The Thomist*, 34 :254–288.
- Meek, C. (1996). *Graphical Models : Selecting Causal and Statistical Models*. PhD thesis, Carnegie Mellon University, United State of America.
- Meguetta, Z.-E., Conrard, B., and Bayart, M. (2013). Design of control architecture based search algorithm for fault tolerant control system. *Control and Fault-Tolerant Systems, SysTol, Nice, France*.
- Moneta, A. and Spirtes, P. (2006). Graphical models for the identification of causal structures in multivariate time series models. In *In Proceedings of the 9th Joint Conference on Information Sciences (JCIS), 1. Atlantis Press, Paris, France*, pages 1–4.
- Murthy, V. K. and Krishnamurthy, E. V. (2009). Multiset of agents in a network for simulation of complex systems. In Kyamakya, K., Halang, W. A., Unger, H., Chedjou, J. C., Rulkov, N. F., and Li, Z., editors, *Recent Advances in Nonlinear Dynamics and Synchronization*, number 254 in Studies in Computational Intelligence, pages 153–200. Springer Berlin Heidelberg.

- Naghoosi, E., Huang, B., Domlan, E., and Kadali, R. (2013). Information transfer methods in causality analysis of process variables with an industrial application. *Journal of Process Control*, 23(9) :1296–1305.
- Norvilas, A., Negiz, A., DeCicco, J., and Çinar, A. (2000). Intelligent process monitoring by interfacing knowledge-based systems and multivariate statistical monitoring. *Journal of Process Control*, 10(4) :341–350.
- Okada, S. and Gen, M. (1993). Order relation between intervals and its application to shortest path problem. *Computers & Industrial Engineering*, 25(1–4) :147–150.
- Oppenheim, A. V. (1997). *Signals and Systems*. Prentice Hall.
- Pavan, M. and Todeschini, R. (2004). New indices for analysing partial ranking diagrams. *Analytica Chimica Acta*, 515(1) :167–181.
- Pearl, J. (2000). *Causality : Models, Reasoning, and Inference*. Cambridge University Press.
- Pessel, N., Balmat, J., Lafont, F., and Bannal, J. (2007). An improved fault detection for the diagnosis. in *Proceeding of the 9th International Conference on Automatic Control, Modeling & Simulation*.
- Pukelsheim, F. (1994). The three sigma rule. *The American Statistician*, 48(2) :88–91.
- Ragwitz, M. and Kantz, H. (2002). Markov models from data by simple nonlinear time series predictors in delay embedding spaces. *Phys Rev E Stat Nonlin Soft Matter Phys*, 65.
- Ram Maurya, M., Rengaswamy, R., and Venkatasubramanian, V. (2004). Application of signed digraphs-based analysis for fault diagnosis of chemical process flowsheets. *Engineering Applications of Artificial Intelligence*, 17(5) :501–518.
- Rissanen, J. (1978). Modeling by shortest data description. *Automatica*, 14(5) :465–471.
- Rokhlin, D. B. (2015). Central limit theorem under uncertain linear transformations. *Statistics & Probability Letters*, 107 :191 – 198.

- Schreiber, T. (2000). Measuring information transfer. *Physical Review Letters*, 85(2) :461–464.
- Sensoy, A., Sobaci, C., Sensoy, S., and Alali, F. (2014). Effective transfer entropy approach to information flow between exchange rates and stock markets. *Chaos, Solitons & Fractals*, 68 :180–185.
- Seth, A. K. (2010). A MATLAB toolbox for Granger causal connectivity analysis. *Journal of Neuroscience Methods*, 186(2) :262–273.
- Shu, Y. and Zhao, J. (2013). Data-driven causal inference based on a modified transfer entropy. *Computers & Chemical Engineering*, 57 :173–180.
- Silva, R., Scheines, R., Glymour, C., and Spirtes, P. (2006). Learning the structure of linear latent variable models. *J. Mach. Learn. Res.*, 7 :191–246.
- Söderström, T. and Stoica, P. (2002). Instrumental variable methods for system identification. *Circuits, Systems and Signal Processing*, 21(1) :1–9.
- Staniek, M. and Lehnertz, K. (2009). Symbolic transfer entropy : inferring directionality in biosignals. *Biomedizinische Technik. Biomedical Engineering*, 54(6) :323–328.
- Staroswiecki, M. and Hoblos, G. (2004). Sensor network design for fault tolerant estimation. *In. J. Adapt. Control Signal Precess.*, 18 :55–72.
- Tang, X., Tao, G., and Joshi, S. M. (2005). Adaptive output feedback actuator failure compensation for a class of non-linear systems. *International Journal of Adaptive Control and Signal Processing*, 19(6) :419–444.
- Tao, G., Burkholder, J. O., and Guo, J. (2013). Adaptive state feedback actuator nonlinearity compensation for multivariable systems. *International Journal of Adaptive Control and Signal Processing*, 27(1-2) :82–107.
- Tao, G., Chen, S., and M. Joshi, S. (2002). An adaptive control scheme for systems with unknown actuator failures. *Automatica*, 38(6) :1027–1034.

- Tao, G., Joshi, S., and Ma, X. (2001). Adaptive state feedback and tracking control of systems with actuator failures. *IEEE Transactions on Automatic Control*, 46(1) :78–95.
- Theilliol, D., D.Sauter, and Ponsart, J. (1998). Fault tolerant control method for actuator and component faults. *In proceedings of IEEE Conference on Decesio and Control, CDC*.
- Theilliol, D., Zhang, Y., Ponsart, J.-C., and Aubrun, C. (2008). Actuator fault tolerant control system with re-configuring reference input design based on MPC. In *23rd IAR Workshop on Advanced Control and Diagnosis*, pages 275–280, Coventry, Royaume-Uni.
- Thoma, J. (1976). *Synthèse et simulation interdisciplinaire par bond graphs ou graphes de liaison*. Overlapping Tendencies in Operations Research Systems Theory and Cybernetics Part of the series Interdisciplinary Systems Research / Interdisziplinäre Systemforschung pp 555-560.
- Vedam, H. and Venkatasubramanian, V. (1999). PCA-SDG based process monitoring and fault diagnosis. *Control Engineering Practice*, 7(7) :903–917.
- Wald, A. (1945). Sequential tests of statistical hypotheses. *The Annals of Mathematical Statistics*, 16(2) :117–186.
- Wald, A. (2004). *Sequential Analysis*. Courier Dover Publications.
- Wallace, W. A. (1974). Aquinas on the temporal relation between cause and effect. *Review of Metaphysics*, 27(3) :569–584.
- Wheeler, D. J., Chambers, D. S., and 1917 (1992). *Understanding statistical process control*. SPC Press.
- Wu, N. E., Zhang, Y., and Zhou, K. (2000). Detection, estimation, and accommodation of loss of control effectiveness. *International Journal of Adaptive Control and Signal Processing*, 14(7) :775–795.

- Yang, F., Shah, S., and Xiao, D. (2012). Signed directed graph based modeling and its validation from process knowledge and process data. *International Journal of Applied Mathematics and Computer Science*, 22(1) :41–53.
- Yang, Y. (2005). Can the strengths of AIC and BIC be shared? a conflict between model identification and regression estimation. *Biometrika*, 92(4) :937–950.
- Zhang, Y. and Jiang, J. (2003). Fault tolerance control system design with explicit consideration of performance degradation. *IEEE Transactions on Aerospace and Electronic Systems*, 39(3) :838–848.
- Zhang, Y. and Jiang, J. (2006). Issues on integration of fault diagnosis and reconfigurable control in active fault-tolerant control. *Fault Detection, Supervision and Safety of Technical Processes*, 6(1) :1437–1448.
- Zwingelstein, G. (1995). Diagnostic des défaillances : Théorie et pratique pour les systèmes industriels. *Traité des nouvelles technologies Série diagnostic et maintenance*. Hermes .